

Notes on Diffy Qs

Differential Equations for Engineers

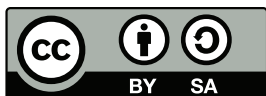
by Jiří Lebl

May 18, 2026
(version 6.11)

Typeset in L^AT_EX.

Edition 6 (version 6.11)

Copyright ©2008–2026 Jiří Lebl



This work is dual licensed under the Creative Commons Attribution-Noncommercial-Share Alike 4.0 International License and the Creative Commons Attribution-Share Alike 4.0 International License. To view a copy of these licenses, visit <https://creativecommons.org/licenses/by-nc-sa/4.0/> or <https://creativecommons.org/licenses/by-sa/4.0/> or send a letter to Creative Commons PO Box 1866, Mountain View, CA 94042, USA.

You can use, print, duplicate, and share this book as much as you want. You can base your own notes on it and reuse parts if you keep the license the same. You can assume the license is either CC-BY-NC-SA or CC-BY-SA, whichever is compatible with what you wish to do. Your derivative work must use at least one of the licenses. Derivative works must be prominently marked as such.

During the writing of this book, the author was in part supported by NSF grants DMS-0900885 and DMS-1362337.

The major version / edition number is raised only if there have been substantial changes. Edition number started at 5, that is, version 5.0, as it was not kept track of before.

See <https://www.jirka.org/diffyqs/> or <https://jirilebl.github.io/diffyqs/> for more information (including contact information).

The L^AT_EX source for the book is available for possible modification and customization at github: <https://github.com/jirilebl/diffyqs>

Contents

Introduction	7
0.1 Notes about these notes	7
0.2 Introduction to differential equations	10
0.3 Classification of differential equations	17
1 First-order equations	21
1.1 Integrals as solutions	21
1.2 Slope fields	27
1.3 Separable equations	33
1.4 Linear equations and the integrating factor	40
1.5 Substitution	46
1.6 Autonomous equations	51
1.7 Numerical methods: Euler's method	57
1.8 Exact equations	63
1.9 First-order linear PDEs	72
2 Higher-order linear ODEs	79
2.1 Second-order linear ODEs	79
2.2 Constant-coefficient second-order linear ODEs	84
2.3 Higher-order linear ODEs	90
2.4 Mechanical vibrations	96
2.5 Nonhomogeneous equations	104
2.6 Forced oscillations and resonance	111
3 Systems of ODEs	119
3.1 Introduction to systems of ODEs	119
3.2 Matrices and linear systems	127
3.3 Linear systems of ODEs	136
3.4 Eigenvalue method	140
3.5 Two-dimensional systems and their vector fields	147
3.6 Second-order systems and applications	152
3.7 Repeated eigenvalues	162
3.8 Matrix exponentials	169
3.9 Nonhomogeneous systems	177

4	Fourier series and PDEs	189
4.1	Boundary value problems	189
4.2	The trigonometric series	198
4.3	More on the Fourier series	208
4.4	Sine and cosine series	218
4.5	Applications of Fourier series	226
4.6	PDEs, separation of variables, and the heat equation	232
4.7	One-dimensional wave equation	243
4.8	D'Alembert solution of the wave equation	252
4.9	Steady state temperature and the Laplacian	258
4.10	Dirichlet problem in the circle and the Poisson kernel	264
5	More on eigenvalue problems	273
5.1	Sturm–Liouville problems	273
5.2	Higher-order eigenvalue problems	282
5.3	Steady periodic solutions	286
6	The Laplace transform	293
6.1	The Laplace transform	293
6.2	Transforms of derivatives and ODEs	300
6.3	Convolution	308
6.4	Dirac delta and impulse response	313
6.5	Solving PDEs with the Laplace transform	320
7	Power-series methods	327
7.1	Power series	327
7.2	Series solutions of linear second-order ODEs	335
7.3	Singular points and the method of Frobenius	342
8	Nonlinear systems	351
8.1	Linearization, critical points, and equilibria	351
8.2	Stability and classification of isolated critical points	357
8.3	Applications of nonlinear systems	364
8.4	Limit cycles	373
8.5	Chaos	378
A	Linear algebra	385
A.1	Vectors, mappings, and matrices	385
A.2	Matrix algebra	396
A.3	Elimination	407
A.4	Subspaces, dimension, and the kernel	421
A.5	Inner product and projections	428
A.6	Determinant	436

<i>CONTENTS</i>	5
B Table of Laplace Transforms	443
Further Reading	445
Solutions to Selected Exercises	447
Index	461

Introduction

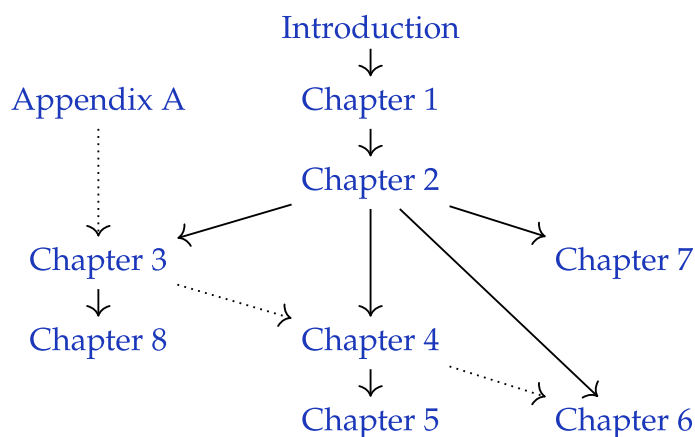
0.1 Notes about these notes

Note: A section for the instructor.

This book originated from my class notes for Math 286 at the [University of Illinois at Urbana-Champaign](#) (UIUC) in Fall 2008 and Spring 2009. It is a first course on differential equations for engineers. Using this book, I also taught Math 285 at UIUC, Math 20D at [University of California, San Diego](#) (UCSD), and Math 2233 and 4233 at [Oklahoma State University](#) (OSU). Normally these courses are taught with Edwards and Penney, *Differential Equations and Boundary Value Problems: Computing and Modeling* [EP], or Boyce and DiPrima's *Elementary Differential Equations and Boundary Value Problems* [BD], and this book aims to be more or less a drop-in replacement. Other books I used as sources of information and inspiration are E.L. Ince's classic (and inexpensive) *Ordinary Differential Equations* [I], Stanley Farlow's *Differential Equations and Their Applications* [F], now available from Dover, Berg and McGregor's *Elementary Partial Differential Equations* [BM], and William Trench's free book *Elementary Differential Equations with Boundary Value Problems* [T]. See the [Further Reading](#) chapter at the end of the book.

0.1.1 Organization

The organization of this book requires, to some degree, that the chapters be covered in order. Later chapters can be dropped. The dependence of the material covered is roughly:



There are a few references in chapters 4 and 5 to chapter 3 (some linear algebra), but these references are not essential and can be skimmed over, so chapter 3 can safely be dropped, while still covering chapters 4 and 5. Chapter 6 does not depend on chapter 4 except that the PDEs section 6.5 makes a few references to chapter 4, although it could, in theory, be covered separately. The more in-depth appendix A on linear algebra can replace the short review § 3.2 for a course that combines linear algebra and ODEs.

0.1.2 Typical courses

Several different courses can be run with the book. There are the two original one-semester courses at UIUC. Both cover ODEs as well as some PDEs. The first course runs at 4 hours a week (Math 286):

Introduction (0.2), chapter 1 (1.1–1.7), chapter 2, chapter 3, chapter 4 (4.1–4.9), chapter 5 (or 6 or 7 or 8).

The second course at UIUC runs 3 hours a week (Math 285):

Introduction (0.2), chapter 1 (1.1–1.7), chapter 2, chapter 4 (4.1–4.9), and maybe chapter 5, 6, or 7.

A semester-long course at 3 hours a week that does not cover either systems or PDEs will cover, beyond the introduction, chapter 1, chapter 2, chapter 6, and chapter 7, (with sections skipped as above). On the other hand, a typical course that covers systems will probably need to skip Laplace and power series and cover chapter 1, chapter 2, chapter 3, and chapter 8.

If sections need to be skipped in the beginning, a good core of the sections on single ODEs is: 0.2, 1.1–1.4, 1.6, 2.1, 2.2, 2.4–2.6.

The complete book can be covered at a reasonably fast pace at approximately 76 lectures (without appendix A) or 86 lectures (with appendix A replacing § 3.2). This is not accounting for exams, review, or time spent in a computer lab. A two-quarter or a two-semester course can be easily run with the material. For example (with some sections perhaps strategically skipped):

Semester 1: Introduction, chapter 1, chapter 2, chapter 6, chapter 7.

Semester 2: Chapter 3, chapter 8, chapter 4, chapter 5.

A combined course on ODEs with linear algebra can run as:

Introduction, chapter 1 (1.1–1.7), chapter 2, appendix A, chapter 3 (w/o § 3.2), possibly chapter 8.

The chapter on the Laplace transform (chapter 6), the chapter on Sturm–Liouville (chapter 5), the chapter on power series (chapter 7), and the chapter on nonlinear systems (chapter 8), are more or less interchangeable and can be treated as “topics”. If chapter 8 is covered, it may be best to place it right after chapter 3, and chapter 5 is best covered right after chapter 4. If time is short, the first two sections of chapter 7 make a reasonable self-contained unit.

0.1.3 Computer resources

The book's website <https://www.jirka.org/diffyqs/> or <https://jirilebl.github.io/diffyqs/> contains the following resources:

1. Interactive SAGE demos.
2. Online WeBWorK homeworks (using either your own WeBWorK installation or Edfinity) for most sections, customized for this book.
3. The PDFs of the figures used in this book.
4. YouTube videos and corresponding slides for some sections.

I used IODE (<https://publish.illinois.edu/iode-diffeq/>) in the UIUC courses. IODE is a free software package that works with Matlab (proprietary) or Octave (free software). The graphs in the book were made with the Genius software (see <https://www.jirka.org/genius.html>).

The $\text{\LaTeX}$ source of the book is also available for possible modification and customization at GitHub (<https://github.com/jirilebl/diffyqs>).

0.1.4 Acknowledgments

Firstly, I would like to acknowledge Rick Laugesen. I used his handwritten class notes the first time I taught Math 286. My organization of this book through chapter 5, and the choice of material covered, is heavily influenced by his notes. Many examples and computations are taken from his notes. I am also heavily indebted to Rick for all the advice he has given me, not just on teaching Math 286. For spotting errors and other suggestions, I would also like to acknowledge (in no particular order): John P. D'Angelo, Sean Raleigh, Jessica Robinson, Michael Angelini, Leonardo Gomes, Jeff Winegar, Ian Simon, Thomas Wicklund, Eliot Brenner, Sean Robinson, Jannett Susberry, Dana Al-Quadi, Cesar Alvarez, Cem Bagdatlioglu, Nathan Wong, Alison Shive, Shawn White, Wing Yip Ho, Joanne Shin, Gladys Cruz, Jonathan Gomez, Janelle Louie, Navid Froutan, Grace Victorine, Paul Pearson, Jared Teague, Ziad Adwan, Martin Weilandt, Sönmez Şahutoğlu, Pete Peterson, Thomas Gresham, Prentiss Hyde, Jai Welch, Simon Tse, Andrew Browning, James Choi, Dusty Grundmeier, John Marriott, Jim Kruidenier, Barry Conrad, Wesley Snider, Colton Koop, Sarah Morse, Erik Boczek, Asif Shakeel, Chris Peterson, Nicholas Hu, Paul Seeburger, Jonathan McCormick, David Leep, William Meisel, Shishir Agrawal, Tom Wan, Andres Valloud, Martin Irungu, Justin Corvino, Tai-Peng Tsai, Santiago Mendoza Reyes, Glen Pugh, Michael Tran, Heber Farnsworth, Tamás Zsoldos, Mark Mills, George Ballinger, and probably others I have forgotten. Finally, I would like to acknowledge NSF grants DMS-0900885 and DMS-1362337.

0.2 Introduction to differential equations

Note: more than 1 lecture, §1.1 in [EP], chapter 1 in [BD]

0.2.1 Differential equations

The laws of physics are generally written down as differential equations. Therefore, all of science and engineering uses differential equations to some degree. Differential equations are essential to understanding almost anything you will study in your science and engineering classes. You can think of mathematics as the language of science, and differential equations are one of the most important parts of this language as far as science and engineering are concerned. As an analogy, suppose all your classes from now on were given in Swahili. It would be important to first learn Swahili, or you would have a very tough time getting a good grade in your classes.

A *differential equation* is simply an equation that involves not only an unknown variable, but also its derivatives. You have seen and even solved many differential equations already in calculus perhaps without knowing about it. Here is an example you may not have seen:

$$\frac{dx}{dt} + x = 2 \cos t. \quad (1)$$

The unknown, x , is the *dependent variable* and t is the *independent variable*. That is, x is a function of t . Equation (1) is an example of a *first-order differential equation*, since it involves only the first derivative of the dependent variable. This equation arises from Newton's law of cooling where the ambient temperature oscillates with time.

0.2.2 Solutions of differential equations

Solving the differential equation (1) means finding x in terms of t . That is, we want to find a function of t , which we call x , such that when we plug x , t , and $\frac{dx}{dt}$ into (1), the equation holds; that is, the left-hand side equals the right-hand side. It is the same idea as it would be for a normal (algebraic) equation of just x and t . We claim that

$$x = x(t) = \cos t + \sin t$$

is a *solution*. How do we check? We simply plug x into equation (1)! First we need to compute $\frac{dx}{dt}$. We find that $\frac{dx}{dt} = -\sin t + \cos t$. Now let us compute the left-hand side of (1).

$$\frac{dx}{dt} + x = \underbrace{(-\sin t + \cos t)}_{\frac{dx}{dt}} + \underbrace{(\cos t + \sin t)}_x = 2 \cos t.$$

Yay! We got precisely the right-hand side. But there is more! We claim $x = \cos t + \sin t + e^{-t}$ is also a solution. Let us try,

$$\frac{dx}{dt} = -\sin t + \cos t - e^{-t}.$$

We plug into the left-hand side of (1)

$$\frac{dx}{dt} + x = \underbrace{(-\sin t + \cos t - e^{-t})}_{\frac{dx}{dt}} + \underbrace{(\cos t + \sin t + e^{-t})}_x = 2 \cos t.$$

And it works yet again!

So there can be many different solutions. For this equation, all solutions can be written in the form

$$x = \cos t + \sin t + Ce^{-t},$$

for some constant C . Different constants C will give different solutions, so there are really infinitely many possible solutions. See Figure 1 for the graph of a few of these solutions. We will see how we find these solutions in a later section.

Solving differential equations can be quite hard. There is no general method that solves every differential equation. We will generally focus on how to get exact formulas for solutions of certain differential equations, but we will also spend a little bit of time on getting approximate solutions. And we will spend some time on understanding the equations without solving them.

Most of this book is dedicated to *ordinary differential equations* or ODEs, that is, equations with only one independent variable, where derivatives are only with respect to this one variable. If there are several independent variables, we get *partial differential equations* or PDEs.

Even for ODEs, which are very well understood, it is not a simple question of turning a crank to get solutions. When you can find exact solutions, they are usually preferable to approximate solutions. It is important to understand how such solutions are found. Although in real applications you will leave many of the actual calculations to computers, you need to understand what they are doing. It is often necessary to simplify or transform your equations into something that a computer can understand and solve. You may even need to make certain assumptions and changes in your model to achieve this.

To be a successful engineer or scientist, you will be required to solve problems in your job that you have never seen before. It is important to learn problem solving techniques, so that you may apply those techniques to new problems. A common mistake is to expect to learn some prescription for solving all the problems you will encounter in your later career. This course is no exception.

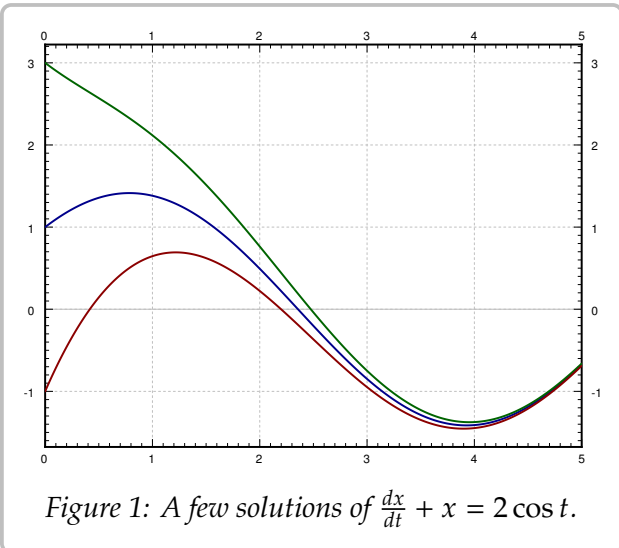
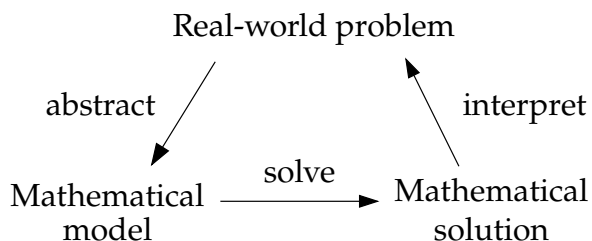


Figure 1: A few solutions of $\frac{dx}{dt} + x = 2 \cos t$.

0.2.3 Differential equations in practice

How do we use differential equations in science and engineering? First, we have some *real-world problem* we wish to understand. We make some simplifying assumptions and create a *mathematical model*. That is, we translate the real-world situation into a differential equation (or equations). Then we use mathematics to get some sort of a *mathematical solution*. There is still something left to do. We must interpret the results. We have to figure out what the mathematical solution says about the real-world problem we started with.



Formulating the mathematical model and interpreting the results is what you do in your physics and engineering classes. In this course, we focus mostly on the mathematical analysis. Sometimes we will work with simple real-world examples to have some intuition and motivation about what we are doing. Let us look at an example of the entire process.

Example 0.2.1: One of the most basic differential equations is the *exponential growth model*. Let P denote the population of some bacteria on a Petri dish. We assume that there is enough food and enough space. Then the rate of growth of the bacteria is proportional to the population—a larger population grows faster. Let t denote time (say in seconds) and P the population. Our model is

$$\frac{dP}{dt} = kP,$$

for some positive constant $k > 0$.

Suppose there are 100 bacteria at time 0 and 200 bacteria 10 seconds later. How many bacteria will there be 1 minute from time 0 (in 60 seconds)?

First we need to solve the equation. We claim that a solution is given by

$$P(t) = Ce^{kt},$$

where C is a constant. Let us try:

$$\frac{dP}{dt} = Cke^{kt} = kP.$$

And it really is a solution.

OK, now what? We do not know C , and we do not know k . But we know something. We know $P(0) = 100$, and we know $P(10) = 200$. We plug these conditions in and see what happens:

$$\begin{aligned} 100 &= P(0) = Ce^{k0} = C, \quad \text{so } C = 100, \\ 200 &= P(10) = Ce^{k10} = 100e^{k10}. \end{aligned}$$

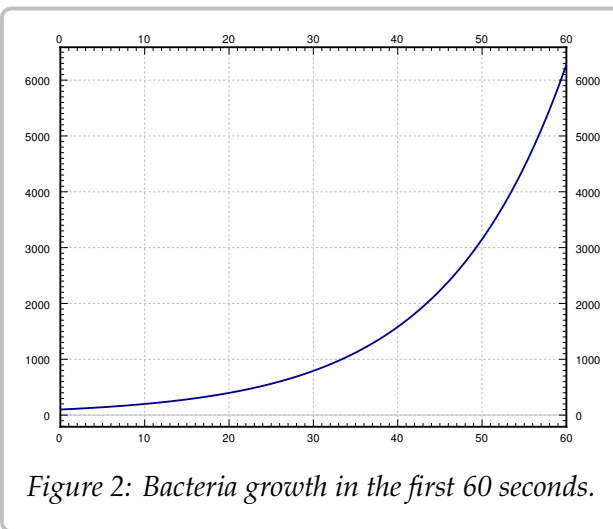


Figure 2: Bacteria growth in the first 60 seconds.

Therefore, $2 = e^{10k}$ or $k = \frac{\ln 2}{10} \approx 0.069$. So

$$P(t) = 100 e^{(\ln 2)t/10} \approx 100 e^{0.069t}.$$

At one minute, $t = 60$, the population is $P(60) = 6400$. See [Figure 2](#) on the preceding page.

Let us interpret the results. Does our solution mean that there must be exactly 6400 bacteria on the plate at 60 seconds? No! We made assumptions that might not be true exactly, just approximately. If our assumptions are reasonable, then there will be approximately 6400 bacteria. Also, in real life P is a discrete quantity, that is, a whole number. However, our model has no problem saying that for example at 61 seconds, $P(61) \approx 6859.35$.

Normally, the k in $\frac{dP}{dt} = kP$ is given, and we want to solve the equation for some *initial condition*. For example, the problem could be: Solve the equation $\frac{dP}{dt} = P$ (so $k = 1$) subject to $P(0) = 1000$ (the initial condition). Then the solution is (exercise)

$$P(t) = 1000 e^t.$$

We call $P(t) = Ce^t$ the *general solution*, as every solution of the equation can be written in this form for some constant C . We need an initial condition to find out what C is, in order to find the *particular solution* we are looking for, in this case, $P(t) = 1000 e^t$. Generally, when we say “particular solution,” we just mean some solution.

In real life, parameters such as k must often first be somehow computed or estimated. The example above shows how finding an analytic solution to the differential equation is useful in finding these parameters.

0.2.4 Four fundamental equations

A few equations appear often and it is useful to just memorize what their solutions are. Let us call them the four fundamental equations. Their solutions are reasonably easy to guess by recalling properties of exponentials, sines, and cosines. They are also simple to check, which is something that you should always do. No need to wonder if you remembered the solution correctly.

The first such equation is

$$\frac{dy}{dx} = ky,$$

for some constant $k > 0$. Here y is the dependent and x the independent variable. The general solution for this equation is

$$y(x) = Ce^{kx}.$$

We saw above that this function is a solution, although we used different variable names.

Next, we have

$$\frac{dy}{dx} = -ky,$$

for some constant $k > 0$. The general solution for this equation is

$$y(x) = Ce^{-kx}.$$

Exercise 0.2.1: Check that the y given is really a solution to the equation above.

Next, take the second-order differential equation

$$\frac{d^2y}{dx^2} = -k^2y,$$

for some constant $k > 0$. The general solution for this equation is

$$y(x) = C_1 \cos(kx) + C_2 \sin(kx).$$

Since the equation is a second-order differential equation, we have two constants in our general solution.

Exercise 0.2.2: Check that the y given is really a solution to the equation above.

Finally, consider the second-order differential equation

$$\frac{d^2y}{dx^2} = k^2y,$$

for some constant $k > 0$. The general solution for this equation is

$$y(x) = C_1 e^{kx} + C_2 e^{-kx},$$

or

$$y(x) = D_1 \cosh(kx) + D_2 \sinh(kx).$$

For those who do not know, cosh and sinh are defined by

$$\cosh x = \frac{e^x + e^{-x}}{2}, \quad \sinh x = \frac{e^x - e^{-x}}{2}.$$

They are called the *hyperbolic cosine* and *hyperbolic sine*. These functions are sometimes easier to work with than exponentials. They have some nice familiar properties such as $\cosh 0 = 1$, $\sinh 0 = 0$, and $\frac{d}{dx} \cosh x = \sinh x$ (no, that is not a typo) and $\frac{d}{dx} \sinh x = \cosh x$.

Exercise 0.2.3: Check that both forms of the y given are really solutions to the equation above.

Example 0.2.2: In equations of higher order, you get more constants you must solve for to get a particular solution and hence you need more initial conditions. The equation $\frac{d^2y}{dx^2} = 0$ has the general solution $y = C_1x + C_2$; simply integrate twice and don't forget about the constants of integration. Consider the initial conditions $y(0) = 2$ and $y'(0) = 3$. We plug in our general solution and solve for the constants:

$$2 = y(0) = C_1 \cdot 0 + C_2 = C_2, \quad 3 = y'(0) = C_1.$$

In other words, $y = 3x + 2$ is the particular solution we seek.

An interesting note about cosh: The graph of cosh is the exact shape of a hanging chain. This shape is called a *catenary*. Contrary to popular belief, this is not a parabola. If you invert the graph of cosh, it is also the ideal arch for supporting its weight—a parabola would not work. For example, per the National Park Service, the exact formula used by the architects of the Gateway Arch in Saint Louis is (a so-called weighted catenary)

$$y = A \left(\cosh \left(\frac{C}{L} x \right) - 1 \right), \quad \text{where } A = 68.7672, C = 3.0022, L = 299.2239.$$

0.2.5 Exercises

Exercise 0.2.4: Show that $x = e^{4t}$ is a solution to $x''' - 12x'' + 48x' - 64x = 0$.

Exercise 0.2.5: Show that $x = e^t$ is not a solution to $x''' - 12x'' + 48x' - 64x = 0$.

Exercise 0.2.6: Is $y = \sin t$ a solution to $\left(\frac{dy}{dt}\right)^2 = 1 - y^2$? Justify.

Exercise 0.2.7: Let $y'' + 2y' - 8y = 0$. Now try a solution of the form $y = e^{rx}$ for some (unknown) constant r . Is this a solution for some r ? If so, find all such r .

Exercise 0.2.8: Verify that $x = Ce^{-2t}$ is a solution to $x' = -2x$. Find C to solve for the initial condition $x(0) = 100$.

Exercise 0.2.9: Verify that $x = C_1e^{-t} + C_2e^{2t}$ is a solution to $x'' - x' - 2x = 0$. Find C_1 and C_2 to solve for the initial conditions $x(0) = 10$ and $x'(0) = 0$.

Exercise 0.2.10: Find a solution to $(x')^2 + x^2 = 4$ using your knowledge of derivatives of functions that you know from basic calculus.

Exercise 0.2.11: Solve:

a) $\frac{dA}{dt} = -10A, \quad A(0) = 5$

b) $\frac{dH}{dx} = 3H, \quad H(0) = 1$

c) $\frac{d^2y}{dx^2} = 4y, \quad y(0) = 0, \quad y'(0) = 1$

d) $\frac{d^2x}{dy^2} = -9x, \quad x(0) = 1, \quad x'(0) = 0$

Exercise 0.2.12: Is there a solution to $y' = y$, such that $y(0) = y(1)$?

Exercise 0.2.13: The population of city X was 100 thousand 20 years ago, and it was 120 thousand 10 years ago. Assuming constant growth, you can use the exponential population model (like for the bacteria). What do you estimate the population is now?

Exercise 0.2.14: Suppose that a football coach gets a salary of one million dollars now, and a raise of 10% every year (so exponential model, like population of bacteria). Let s be the salary in millions of dollars, and t is time in years.

a) What are $s(0)$ and $s(1)$?

b) Approximately how many years will it take for the salary to be 10 million?

c) Approximately how many years will it take for the salary to be 20 million?

d) Approximately how many years will it take for the salary to be 30 million?

Note: Exercises with numbers 101 and higher have solutions in the back of the book.

Exercise 0.2.101: Show that $x = e^{-2t}$ is a solution to $x'' + 4x' + 4x = 0$.

Exercise 0.2.102: Is $y = x^2$ a solution to $x^2y'' - 2y = 0$? Justify.

Exercise 0.2.103: Let $xy'' - y' = 0$. Try a solution of the form $y = x^r$. Is this a solution for some r ? If so, find all such r .

Exercise 0.2.104: Verify that $x = C_1 e^t + C_2$ is a solution to $x'' - x' = 0$. Find C_1 and C_2 so that x satisfies $x(0) = 10$ and $x'(0) = 100$.

Exercise 0.2.105: Solve $\frac{d\varphi}{ds} = 8\varphi$ and $\varphi(0) = -9$.

Exercise 0.2.106: Solve:

a) $\frac{dx}{dt} = -4x, \quad x(0) = 9$

b) $\frac{d^2x}{dt^2} = -4x, \quad x(0) = 1, \quad x'(0) = 2$

c) $\frac{dp}{dq} = 3p, \quad p(0) = 4$

d) $\frac{d^2T}{dx^2} = 4T, \quad T(0) = 0, \quad T'(0) = 6$

0.3 Classification of differential equations

Note: less than 1 lecture or left as reading, §1.3 in [BD]

There are many types of differential equations, and we classify them into different categories based on their properties. Let us quickly go over the most basic classification. We already saw the distinction between ordinary and partial differential equations:

- An *ordinary differential equation* (ODE) is an equation where the derivatives are taken with respect to only one variable. That is, there is only one independent variable.
- A *partial differential equation* (PDE) is an equation that depend on partial derivatives of several variables. That is, there are several independent variables.

Let us see some examples of ordinary differential equations:

$$\begin{aligned}\frac{dy}{dt} &= ky, & (\text{Exponential growth}) \\ \frac{dy}{dt} &= k(A - y), & (\text{Newton's law of cooling}) \\ m \frac{d^2x}{dt^2} + c \frac{dx}{dt} + kx &= f(t). & (\text{Mechanical vibrations})\end{aligned}$$

And of partial differential equations:

$$\begin{aligned}\frac{\partial y}{\partial t} + c \frac{\partial y}{\partial x} &= 0, & (\text{Transport equation}) \\ \frac{\partial u}{\partial t} &= \frac{\partial^2 u}{\partial x^2}, & (\text{Heat equation}) \\ \frac{\partial^2 u}{\partial t^2} &= \frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2}. & (\text{Wave equation in 2 dimensions})\end{aligned}$$

If there are several equations working together, we have a so-called *system of differential equations*. For example,

$$y' = x, \quad x' = y$$

is a system of ordinary differential equations. Maxwell's equations for electromagnetics,

$$\begin{aligned}\nabla \cdot \vec{D} &= \rho, & \nabla \cdot \vec{B} &= 0, \\ \nabla \times \vec{E} &= -\frac{\partial \vec{B}}{\partial t}, & \nabla \times \vec{H} &= \vec{J} + \frac{\partial \vec{D}}{\partial t},\end{aligned}$$

are a system of partial differential equations. The divergence operator $\nabla \cdot$ and the curl operator $\nabla \times$ can be written out in partial derivatives of the functions involved in the x , y , and z variables.

The next bit of information is the *order* of the equation (or system). The order is the order of the largest derivative that appears. If the highest derivative that appears is the first

derivative, the equation is of first order. If the highest derivative that appears is the second derivative, then the equation is of second order. For example, Newton's law of cooling above is a first-order equation, while the mechanical vibrations equation is a second-order equation. The equation governing transversal vibrations in a beam,

$$a^4 \frac{\partial^4 y}{\partial x^4} + \frac{\partial^2 y}{\partial t^2} = 0,$$

is a fourth-order partial differential equation. It is of fourth order as at least one derivative is the fourth derivative. It does not matter that the derivative in t is only of second order.

In the first chapter, we will start attacking first-order ordinary differential equations, that is, equations of the form $\frac{dy}{dx} = f(x, y)$. In general, lower-order equations are easier to work with and have simpler behavior, which is why we start with them.

We also distinguish how the dependent variables appear in the equation (or system). In particular, an equation is *linear* if the dependent variable (or variables) and their derivatives appear linearly, that is, only as first powers, they are not multiplied together, and no other functions of the dependent variables appear. The equation is a sum of terms, where each term is a function of the independent variables or a function of the independent variables multiplied by a dependent variable or one of its derivatives. An ordinary differential equation is linear if it can be put into the form

$$a_n(x) \frac{d^n y}{dx^n} + a_{n-1}(x) \frac{d^{n-1} y}{dx^{n-1}} + \cdots + a_1(x) \frac{dy}{dx} + a_0(x)y = b(x). \quad (2)$$

The functions $a_0, a_1, \dots, a_n$ are called the *coefficients*. The equation is allowed to depend arbitrarily on the independent variable. So

$$e^x \frac{d^2 y}{dx^2} + \sin(x) \frac{dy}{dx} + x^2 y = \frac{1}{x} \quad (3)$$

is a linear equation as y and its derivatives only appear linearly.

All the equations and systems given as examples above are linear. Linearity may not be immediately obvious for Maxwell's equations unless you write out the divergence and curl in terms of partial derivatives. If an equation is not linear, we say it is *nonlinear*. Let us see some example nonlinear equations. Burgers' equation,

$$\frac{\partial y}{\partial t} + y \frac{\partial y}{\partial x} = \nu \frac{\partial^2 y}{\partial x^2},$$

is a nonlinear second-order partial differential equation. It is nonlinear because y and $\frac{\partial y}{\partial x}$ are multiplied together. The equation

$$\frac{dx}{dt} = x^2 \quad (4)$$

is a nonlinear first-order ordinary differential equation as there is a second power of the dependent variable x . Another nonlinear ODE is the pendulum equation

$$\theta'' + \sin(\theta) = 0, \quad (5)$$

which is nonlinear as the dependent variable θ appears inside a sin function. Nonlinear equations are notoriously difficult to solve and their solutions may behave in strange and unexpected ways. Perhaps you have heard of chaos theory and butterflies in the Amazon causing hurricanes in the Atlantic, all due to nonlinear equations. So sometimes we study related linear equations, such as $\theta'' + \theta = 0$ for the pendulum, instead. The solutions to $\theta'' + \sin(\theta) = 0$ and $\theta'' + \theta = 0$ are different, but they are close for some short time, and the linear equation is far simpler to solve and analyze.

A linear equation is further called *homogeneous* if all terms depend on the dependent variable. That is, if no term is a function of the independent variables alone. Otherwise, the equation is called *nonhomogeneous* or *inhomogeneous*. For example, the exponential growth equation, the wave equation, or the transport equation above are homogeneous. The mechanical vibrations equation above is nonhomogeneous as long as $f(t)$ is not the zero function. Similarly, if the ambient temperature A is nonzero, Newton's law of cooling is nonhomogeneous. A homogeneous linear ODE can be put into the form

$$a_n(x) \frac{d^n y}{dx^n} + a_{n-1}(x) \frac{d^{n-1} y}{dx^{n-1}} + \cdots + a_1(x) \frac{dy}{dx} + a_0(x)y = 0.$$

Compare with (2) and notice there is no function $b(x)$.

If the coefficients of a linear equation are actually constant functions, then the equation is said to have *constant coefficients*. The coefficients are the functions multiplying the dependent variable(s) or one of its derivatives, not the function $b(x)$ standing alone. A constant-coefficient nonhomogeneous ODE is an equation of the form

$$a_n \frac{d^n y}{dx^n} + a_{n-1} \frac{d^{n-1} y}{dx^{n-1}} + \cdots + a_1 \frac{dy}{dx} + a_0 y = b(x),$$

where $a_0, a_1, \dots, a_n$ are all constants, but b may depend on the independent variable x . The mechanical vibrations equation is a constant-coefficient nonhomogeneous second-order ODE. The same nomenclature applies to PDEs, so the transport equation, heat equation and wave equation are all examples of constant-coefficient linear PDEs. A linear equation whose coefficients are not constants is sometimes called a *variable-coefficient* equation.

Finally, an equation (or system) is called *autonomous* if the equation does not depend on the independent variable. For autonomous ordinary differential equations, the independent variable is then thought of as time. An autonomous equation means an equation that does not change with time. For example, Newton's law of cooling is autonomous, so are the equations (4) and (5). On the other hand, mechanical vibrations (as long as $f(t)$ is nonconstant) or (3) are not autonomous. A general first-order autonomous ODE would have the form

$$\frac{dx}{dt} = f(x).$$

0.3.1 Exercises

Exercise 0.3.1: Classify the following equations. Are they ODEs or PDEs? Is it an equation or a system? What is the order? Is it linear or nonlinear, and if it is linear, is it homogeneous, constant-coefficient? If it is an ODE, is it autonomous?

$$a) \sin(t) \frac{d^2x}{dt^2} + \cos(t)x = t^2$$

$$b) \frac{\partial u}{\partial x} + 3 \frac{\partial u}{\partial y} = xy$$

$$c) y'' + 3y + 5x = 0, \quad x'' + x - y = 0$$

$$d) \frac{\partial^2 u}{\partial t^2} + u \frac{\partial^2 u}{\partial s^2} = 0$$

$$e) x'' + tx^2 = t$$

$$f) \frac{d^4x}{dt^4} = 0$$

Exercise 0.3.2: If $\vec{u} = (u_1, u_2, u_3)$ is a vector, we have the divergence $\nabla \cdot \vec{u} = \frac{\partial u_1}{\partial x} + \frac{\partial u_2}{\partial y} + \frac{\partial u_3}{\partial z}$ and the curl $\nabla \times \vec{u} = \left(\frac{\partial u_3}{\partial y} - \frac{\partial u_2}{\partial z}, \frac{\partial u_1}{\partial z} - \frac{\partial u_3}{\partial x}, \frac{\partial u_2}{\partial x} - \frac{\partial u_1}{\partial y} \right)$. Notice that the curl of a vector is still a vector. Write out Maxwell's equations in terms of partial derivatives and classify the system.

Exercise 0.3.3: Suppose F is a linear function of two variables, that is, $F(x, y) = ax + by$ for constants a and b . Suppose $a \neq 0$. Classify equations of the form $F(y', y) = 0$.

Exercise 0.3.4: Write down an explicit example of a third-order, linear, variable-coefficient (i.e. not constant-coefficient), nonautonomous, nonhomogeneous system of two ODEs such that every derivative that could appear, does appear.

Exercise 0.3.101: Classify the following equations. Are they ODEs or PDEs? Is it an equation or a system? What is the order? Is it linear or nonlinear, and if it is linear, is it homogeneous, constant-coefficient? If it is an ODE, is it autonomous?

$$a) \frac{\partial^2 v}{\partial x^2} + 3 \frac{\partial^2 v}{\partial y^2} = \sin(x)$$

$$b) \frac{dx}{dt} + \cos(t)x = t^2 + t + 1$$

$$c) \frac{d^7 F}{dx^7} = 3F(x)$$

$$d) y'' + 8y' = 1$$

$$e) x'' + tyx' = 0, \quad y'' + txy = 0$$

$$f) \frac{\partial u}{\partial t} = \frac{\partial^2 u}{\partial s^2} + u^2$$

Exercise 0.3.102: Write down the general zeroth-order linear ordinary differential equation. Write down the general solution.

Exercise 0.3.103: For which k is $\frac{dx}{dt} + x^k = t^{k+2}$ linear? Hint: There are two answers.

Chapter 1

First-order equations

1.1 Integrals as solutions

Note: 1 lecture (or less), §1.2 in [EP], covered in §1.2 and §2.1 in [BD]

A first-order ODE is an equation of the form

$$\frac{dy}{dx} = f(x, y),$$

or just

$$y' = f(x, y).$$

In general, there is no simple formula or procedure one can follow to find solutions. In the next few lectures, we will look at cases where solutions are not difficult to obtain. In this section, we consider the case when f is a function of x alone, that is, the equation is

$$y' = f(x). \tag{1.1}$$

We can integrate (antidifferentiate) both sides of the equation with respect to x :

$$\int y'(x) dx = \int f(x) dx + C,$$

that is,

$$y(x) = \int f(x) dx + C.$$

This $y(x)$ is actually the general solution. So to solve (1.1), we find some antiderivative of $f(x)$ and then we add an arbitrary constant to get the general solution.

Example 1.1.1: Find the general solution of $y' = 3x^2$.

Elementary calculus tells us that the general solution must be $y = x^3 + C$. Let us check by differentiating: $y' = 3x^2$. We got *precisely* our equation back.

Now is a good time to discuss a point about calculus notation and terminology. Calculus textbooks muddy the waters by talking about the integral as primarily the

so-called indefinite integral. The indefinite integral is really the *antiderivative* (in fact the whole one-parameter family of antiderivatives). There really exists only one integral and that is the definite integral. The only reason for the indefinite integral notation is that we can always write an antiderivative as a definite integral. That is, by the fundamental theorem of calculus, we can always write $\int f(x) dx + C$ as

$$\int_{x_0}^x f(t) dt + C.$$

Hence the terminology *to integrate* when we may really mean *to antidifferentiate*. Integration is just one way to compute the antiderivative (and it is a way that always works, see the following examples). Integration is defined as the area under the graph—it only happens to also compute antiderivatives. For the sake of consistency, we will keep using the indefinite integral notation when we want an antiderivative, and you should *always* think of the definite integral as a way to write it.

Normally, we also have an initial condition such as $y(x_0) = y_0$ for some two numbers x_0 and y_0 (x_0 is often 0, but not always). We can then write the solution as a definite integral in a nice way. Suppose our problem is $y' = f(x)$, $y(x_0) = y_0$. Then the solution is

$$y(x) = \int_{x_0}^x f(t) dt + y_0. \quad (1.2)$$

Let us check! We compute $y' = f(x)$ via the fundamental theorem of calculus, and by Jupiter, y is a solution. Is it the one satisfying the initial condition? Well, $y(x_0) = \int_{x_0}^{x_0} f(t) dt + y_0 = y_0$. It is!

Do note that the definite integral and the indefinite integral (antidifferentiation) are completely different beasts. The definite integral always evaluates to a number. Therefore, (1.2) is a formula we can plug into the calculator or a computer, and it will be happy to calculate specific values for us. We will easily be able to plot the solution and work with it just like with any other function. It is not so crucial to always find a closed form* for the antiderivative.

Example 1.1.2: Solve

$$y' = e^{-x^2}, \quad y(0) = 1.$$

By the preceding discussion, the solution must be

$$y(x) = \int_0^x e^{-t^2} dt + 1.$$

Here is a good way to make fun of your friends taking second-semester calculus. Tell them to find the closed-form solution. Ha ha ha (bad math joke). It is not possible (in closed form). There is absolutely nothing wrong with writing the solution as a definite integral. This particular integral is in fact very important in statistics.

*By closed form, we mean a formula in terms of known functions with no integrals or limits.

Using this method, we can also solve equations of the form

$$y' = f(y).$$

We write the equation in Leibniz notation:

$$\frac{dy}{dx} = f(y).$$

We use the inverse function theorem from calculus to switch the roles of x and y , that is, we treat x as a function of y and find

$$\frac{dx}{dy} = \frac{1}{f(y)}.$$

What we are doing seems like algebra with dx and dy . It is tempting to just do algebra with dx and dy as if they were numbers. And in this case it does work. Be careful, however, as this sort of hand-waving calculation can lead to trouble, especially when more than one independent variable is involved. At this point, we integrate with respect to y ,

$$x(y) = \int \frac{1}{f(y)} dy + C.$$

That is not quite the form we wanted the solution in, so we next try to solve for y .

Example 1.1.3: Previously, we guessed $y' = ky$ (for some $k > 0$) has the solution $y = Ce^{kx}$. We could have found the solution by integrating. First we note that $y = 0$ is a solution. Henceforth, we assume $y \neq 0$. We write

$$\frac{dx}{dy} = \frac{1}{ky}.$$

We integrate in y to obtain

$$x(y) = x = \frac{1}{k} \ln |y| + D,$$

where D is an arbitrary constant. Now we solve for y (actually for $|y|$).

$$|y| = e^{kx-kD} = e^{-kD} e^{kx}.$$

If we replace e^{-kD} with an arbitrary constant C , we can get rid of the absolute value bars (which we can do as D was arbitrary). In this way, we also incorporate the solution $y = 0$. We get the same general solution as we guessed before, $y = Ce^{kx}$.

Example 1.1.4: Find the general solution of $y' = y^2$.

First we note that $y = 0$ is a solution. We can now assume that $y \neq 0$. Write

$$\frac{dx}{dy} = \frac{1}{y^2}.$$

We integrate to get

$$x = \frac{-1}{y} + C.$$

We solve for $y = \frac{1}{C-x}$. So the general solution is

$$y = \frac{1}{C-x} \quad \text{together with} \quad y = 0.$$

Note the singularities of the solution. If, for example, $C = 1$, then the solution “blows up” as we approach $x = 1$. See [Figure 1.1](#). Generally, it is hard to tell from just looking at the equation itself how the solution is going to behave. The equation $y' = y^2$ is very nice and defined everywhere, but the solution is only defined on the interval $(-\infty, C)$ or (C, ∞) . Usually when this happens, we only consider the solution on one of these intervals and not both. For example, if we impose an initial condition $y(0) = 1$, then the solution is $y = \frac{1}{1-x}$, and we consider this solution only for x on the interval $(-\infty, 1)$, because our initial x (that is, 0) is in this interval. In the figure, it is the left side of the graph.

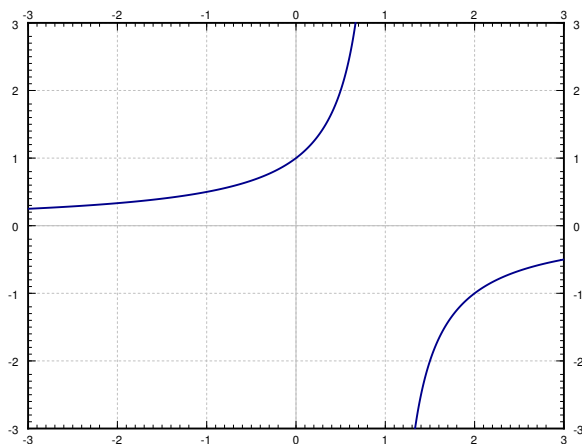


Figure 1.1: Plot of $y = \frac{1}{1-x}$.

On the other hand, if we impose the initial condition $y(2) = -1$ instead, then we again get the solution $y = \frac{1}{1-x}$. Now, however, it is the right-hand side of the graph, that is, it is defined on the interval $(1, \infty)$, as that is the interval containing 2, the initial x .

Classical problems leading to differential equations solvable by integration are problems dealing with velocity, acceleration, and distance. You have surely seen these problems before in your calculus class.

Example 1.1.5: Suppose a car drives at a speed of $e^{t/2}$ meters per second, where t is time in seconds. How far did the car get in 2 seconds (starting at $t = 0$)? How far in 10 seconds?

Let x denote the distance the car traveled. The equation is

$$x' = e^{t/2}.$$

We just integrate this equation to get that

$$x(t) = 2e^{t/2} + C.$$

We still need to figure out C . We know that when $t = 0$, then $x = 0$. That is, $x(0) = 0$. So

$$0 = x(0) = 2e^{0/2} + C = 2 + C.$$

Thus $C = -2$ and

$$x(t) = 2e^{t/2} - 2.$$

Now we just plug in to get where the car is at 2 and at 10 seconds. We obtain

$$x(2) = 2e^{2/2} - 2 \approx 3.44 \text{ meters}, \quad x(10) = 2e^{10/2} - 2 \approx 294 \text{ meters}.$$

Example 1.1.6: Suppose that the car accelerates at $t^2 \text{ m/s}^2$. At time $t = 0$ the car is at the 1-meter mark and is traveling at 10 m/s . Where is the car at time $t = 10$?

Well, this is actually a second-order problem. If x is the distance traveled, then x' is the velocity, and x'' is the acceleration. The equation with initial conditions is

$$x'' = t^2, \quad x(0) = 1, \quad x'(0) = 10.$$

We can add a new dependent variable v and declare that $x' = v$. Then we solve

$$v' = t^2, \quad v(0) = 10.$$

Once we find v , we solve $x' = v$, $x(0) = 1$ to find x . We leave the integration to the reader.

Exercise 1.1.1: Solve for v , then solve for x . Find $x(10)$ to answer the question.

1.1.1 Exercises

Exercise 1.1.2: Solve $\frac{dy}{dx} = x^2 + x$ for $y(1) = 3$.

Exercise 1.1.3: Solve $\frac{dy}{dx} = \sin(5x)$ for $y(0) = 2$.

Exercise 1.1.4: Solve $\frac{dy}{dx} = \frac{1}{x^2-1}$ for $y(0) = 0$.

Exercise 1.1.5: Solve $y' = y^3$ for $y(0) = 1$.

Exercise 1.1.6 (little harder): Solve $y' = (y-1)(y+1)$ for $y(0) = 3$.

Exercise 1.1.7: Solve $\frac{dy}{dx} = \frac{1}{y+1}$ for $y(0) = 0$.

Exercise 1.1.8 (harder): Solve $y'' = \sin x$ for $y(0) = 0$, $y'(0) = 2$.

Exercise 1.1.9: A spaceship is traveling at the speed $2t^2 + 1 \text{ km/s}$ (t is time in seconds). It is pointing directly away from earth and at time $t = 0$ it is 1000 kilometers from earth. How far from earth is it at one minute from time $t = 0$?

Exercise 1.1.10: Solve $\frac{dx}{dt} = \sin(t^2) + t$, $x(0) = 20$. It is OK to leave your answer as a definite integral.

Exercise 1.1.11: A dropped ball accelerates downwards at a constant rate 9.8 meters per second squared. Set up the differential equation for the height above ground h in meters. Then supposing $h(0) = 100$ meters, how long does it take for the ball to hit the ground?

Exercise 1.1.12: Find the general solution of $y' = e^x$, and then $y' = e^y$.

Exercise 1.1.101: Solve $\frac{dy}{dx} = e^x + x$ and $y(0) = 10$.

Exercise 1.1.102: Solve $x' = \frac{1}{x^2}$, $x(1) = 1$.

Exercise 1.1.103: Solve $x' = \frac{1}{\cos(x)}$, $x(0) = \frac{\pi}{4}$.

Exercise 1.1.104: Sid is in a car traveling at speed $10t + 70$ miles per hour away from Las Vegas, where t is in hours. At $t = 0$, Sid is 10 miles away from Vegas. How far from Vegas is Sid 2 hours later?

Exercise 1.1.105: Solve $y' = y^n$, $y(0) = 1$, where n is a positive integer. Hint: You have to consider different cases.

Exercise 1.1.106: The rate of change of the volume of a snowball that is melting is proportional to the surface area of the snowball. Suppose the snowball is perfectly spherical. The volume (in centimeters cubed) of a ball of radius r centimeters is $(4/3)\pi r^3$. The surface area is $4\pi r^2$. Set up the differential equation for how the radius r is changing. Then, suppose that at time $t = 0$ minutes, the radius is 10 centimeters. After 5 minutes, the radius is 8 centimeters. At what time t will the snowball be completely melted?

Exercise 1.1.107: Find the general solution to $y'''' = 0$. How many distinct constants do you need?

1.2 Slope fields

Note: 1 lecture, §1.3 in [EP], §1.1 in [BD]

As we said, the general first-order equation we are studying looks like

$$y' = f(x, y).$$

Frequently, we cannot simply solve these kinds of equations explicitly. It would be nice if we could at least figure out the shape and behavior of the solutions, or find approximate solutions.

1.2.1 Slope fields

The equation $y' = f(x, y)$ gives you a slope at each point in the (x, y) -plane. And this is the slope a solution $y(x)$ would have at x if its value was y . In other words, $f(x, y)$ is the slope of a solution whose graph runs through the point (x, y) . At a point (x, y) , we draw a short line with the slope $f(x, y)$. For example, if $f(x, y) = xy$, then at point $(2, 1.5)$ we draw a short line of slope $xy = 2 \times 1.5 = 3$. If $y(x)$ is a solution and $y(2) = 1.5$, then the equation mandates that $y'(2) = 3$. See Figure 1.2.

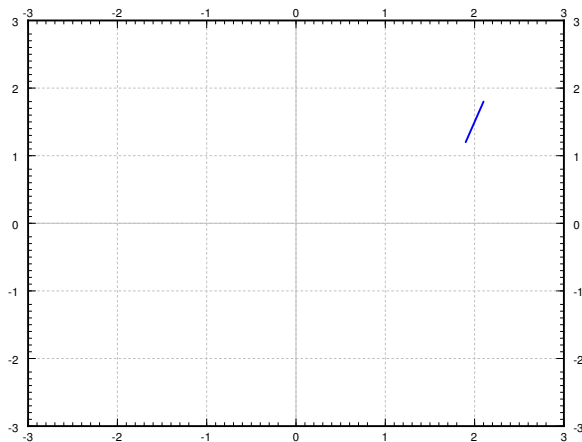
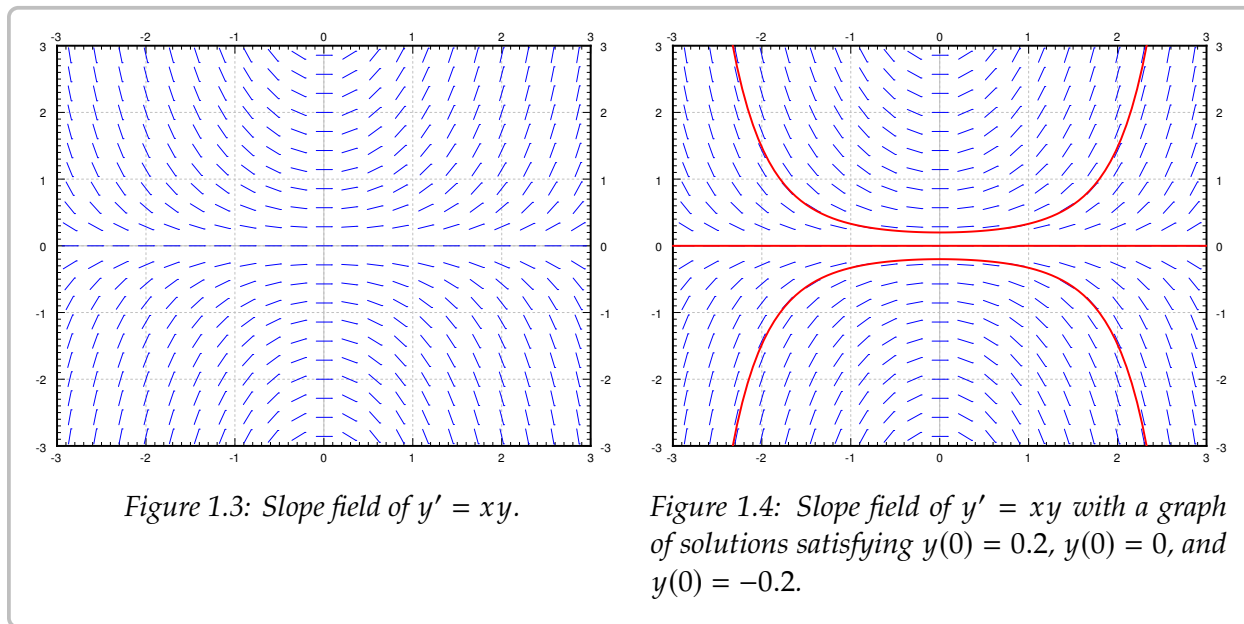


Figure 1.2: The slope $y' = xy$ at $(2, 1.5)$.

To get an idea of how solutions behave, we draw such lines at lots of points in the plane, not just the point $(2, 1.5)$. We would ideally want to see the slope at every point, but that is just not possible. Usually we pick a grid of points fine enough so that it shows the behavior, but not too fine so that we can still recognize the individual lines. We call this picture the *slope field* of the equation. See Figure 1.3 on the following page for the slope field of the equation $y' = xy$. In practice, one does not do this by hand; a computer can do the drawing.

Suppose we are given a specific initial condition $y(x_0) = y_0$. A solution, that is, the graph of the solution, would be a curve that follows the slopes we drew. For a few sample solutions, see Figure 1.4. It is easy to roughly sketch (or at least imagine) possible solutions in the slope field, just from looking at the slope field itself. You simply sketch a line that roughly fits the little line segments and goes through your initial condition.



By looking at the slope field, we get a lot of information about the behavior of solutions without having to solve the equation. For example, in Figure 1.4 we see what the solutions do when the initial conditions are $y(0) > 0$, $y(0) = 0$, and $y(0) < 0$. A small change in the initial condition causes quite different behavior. We see this behavior just from the slope field and imagining what solutions ought to do.

We see a different behavior for the equation $y' = -y$. For the slope field and a few solutions, see Figure 1.5 on the next page. If we think of moving from left to right (perhaps x is time and time is usually increasing), then we see that no matter what $y(0)$ is, all solutions tend to zero as x tends to infinity. Again that behavior is clear from simply looking at the slope field itself.

1.2.2 Existence and uniqueness

It seems that if we want a solution starting at some point, we could just follow the slope field. With that in mind, we wish to ask two fundamental questions about the problem

$$y' = f(x, y), \quad y(x_0) = y_0.$$

- (i) Does a solution *exist*?
- (ii) Is the solution *unique* (if it exists)?

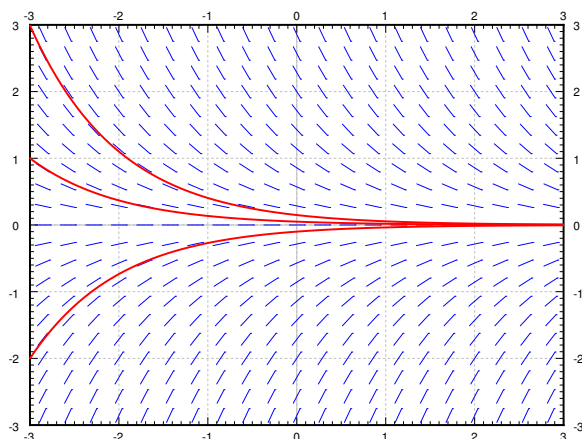


Figure 1.5: Slope field of $y' = -y$ with a graph of a few solutions.

What do you think is the answer? The answer to both questions seems to be yes, does it not? Well, it really is yes most of the time. But there are cases when the answer to either question can be no.

Since the equations we encounter in applications come from real life situations, it seems logical that a solution always exists. It also has to be unique if we believe our universe is deterministic. If the solution does not exist, or if it is not unique, we have probably not devised the correct model. Hence, it is good to know when things go wrong and why.

Example 1.2.1: Attempt to solve:

$$y' = \frac{1}{x}, \quad y(0) = 0.$$

Integrate to find the general solution $y = \ln|x| + C$. The solution does not exist at $x = 0$. See Figure 1.6 on the following page. You may say one can see the division by zero a mile away, but the equation may have been written as the seemingly harmless $xy' = 1$.

Example 1.2.2: Solve:

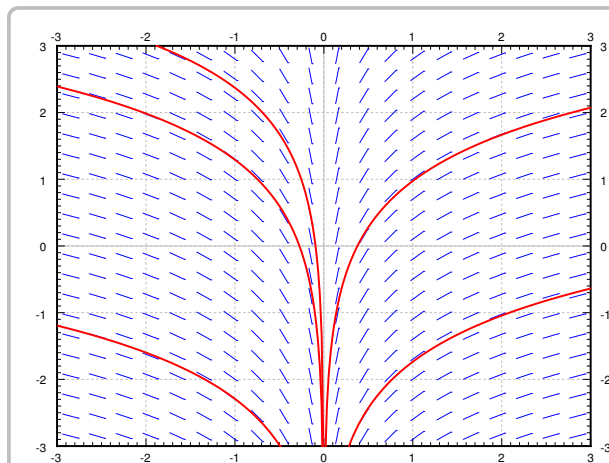
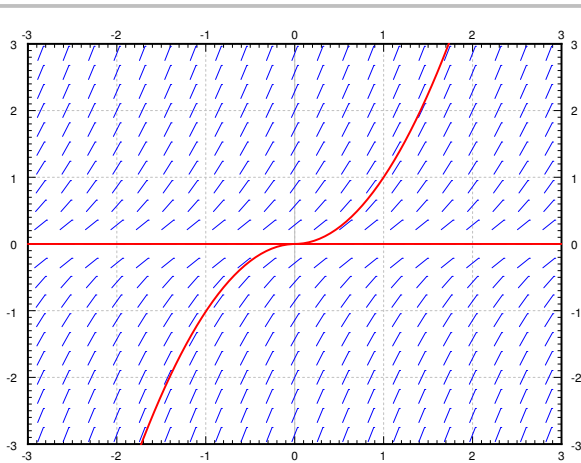
$$y' = 2\sqrt{|y|}, \quad y(0) = 0.$$

See Figure 1.7 on the next page. Note that $y = 0$ is a solution. But another solution is the function

$$y(x) = \begin{cases} x^2 & \text{if } x \geq 0, \\ -x^2 & \text{if } x < 0. \end{cases}$$

It is hard to tell by staring at the slope field that the solution is not unique. Is there any hope? Of course there is. We have the following theorem, known as Picard's theorem*.

*Named after the French mathematician [Charles Émile Picard](#) (1856–1941).

Figure 1.6: Slope field of $y' = 1/x$.Figure 1.7: Slope field of $y' = 2\sqrt{|y|}$ with two solutions satisfying $y(0) = 0$.

Theorem 1.2.1 (Picard's theorem on existence and uniqueness). *If $f(x, y)$ is continuous (as a function of two variables) and $\frac{\partial f}{\partial y}$ exists and is continuous near some (x_0, y_0) , then a solution to*

$$y' = f(x, y), \quad y(x_0) = y_0,$$

exists (at least for x in some small interval) and is unique.

Note that the problems $y' = 1/x$, $y(0) = 0$ and $y' = 2\sqrt{|y|}$, $y(0) = 0$ do not satisfy the hypothesis of the theorem. Even if we can use the theorem, we ought to be careful about this existence business. It is quite possible that the solution only exists for a short while.

Example 1.2.3: For some constant A , solve:

$$y' = y^2, \quad y(0) = A.$$

We know how to solve this equation. First assume that $A \neq 0$, so y is not equal to zero at least for some x near 0. So $x' = 1/y^2$, so $x = -1/y + C$, so $y = \frac{1}{C-x}$. If $y(0) = A$, then $C = 1/A$ so

$$y = \frac{1}{1/A - x}.$$

If $A = 0$, then $y = 0$ is a solution.

For instance, when $A = 1$ the solution “blows up” at $x = 1$. Hence, the solution does not exist for all x even if the equation itself is nice everywhere—it only exists in the interval $(-\infty, 1)$. The equation $y' = y^2$ certainly looks nice.

For most of this course, we will be interested in equations where existence and uniqueness hold, and in fact hold “globally” unlike for the equation $y' = y^2$.

1.2.3 Exercises

Exercise 1.2.1: Sketch the slope field for $y' = e^{x-y}$. How do the solutions behave as x grows? Can you guess a particular solution by looking at the slope field?

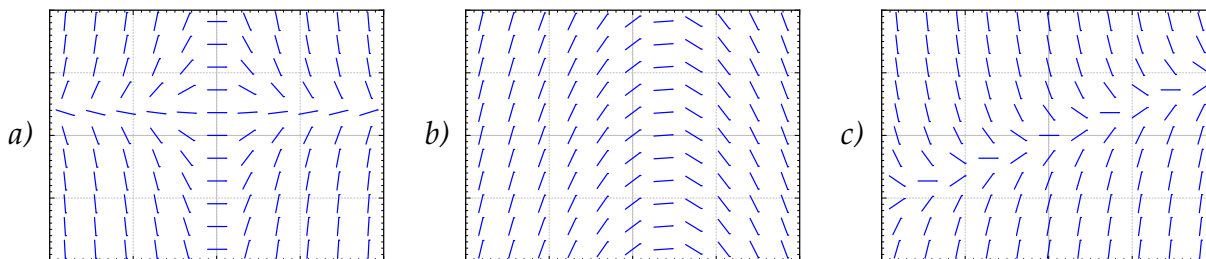
Exercise 1.2.2: Sketch the slope field for $y' = x^2$.

Exercise 1.2.3: Sketch the slope field for $y' = y^2$.

Exercise 1.2.4: Is it possible to solve the equation $y' = \frac{xy}{\cos x}$ for $y(0) = 1$? Justify.

Exercise 1.2.5: Is it possible to solve the equation $y' = y\sqrt{|x|}$ for $y(0) = 0$? Is the solution unique? Justify.

Exercise 1.2.6: Match equations $y' = 1 - x$, $y' = x - 2y$, $y' = x(1 - y)$ to slope fields. Justify.



Exercise 1.2.7 (challenging): Take $y' = f(x, y)$, $y(0) = 0$, where $f(x, y) > 1$ for all x and y . If the solution exists for all x , can you say what happens to $y(x)$ as x goes to positive infinity? Explain.

Exercise 1.2.8 (challenging): Take $(y - x)y' = 0$, $y(0) = 0$.

- Find two distinct solutions.
- Explain why this does not violate Picard's theorem.

Exercise 1.2.9: Suppose $y' = f(x, y)$. What will the slope field look like, explain and sketch an example, if you know the following about $f(x, y)$:

- f does not depend on y .
- f does not depend on x .
- $f(t, t) = 0$ for any number t .
- $f(x, 0) = 0$ and $f(x, 1) = 1$ for all x .

Exercise 1.2.10: Find a solution to $y' = |y|$, $y(0) = 0$. Does Picard's theorem apply?

Exercise 1.2.11: Take an equation $y' = (y - 2x)g(x, y) + 2$ for some function $g(x, y)$. Can you solve the problem for the initial condition $y(0) = 0$, and if so what is the solution?

Exercise 1.2.12 (challenging): Suppose $y' = f(x, y)$ is such that $f(x, 1) = 0$ for every x , f is continuous and $\frac{\partial f}{\partial y}$ exists and is continuous for every x and y .

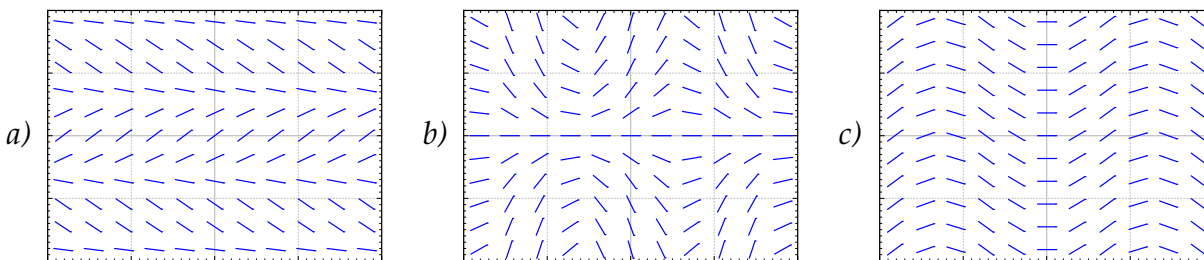
- Guess a solution given the initial condition $y(0) = 1$.
- Can graphs of two solutions of the equation for different initial conditions ever intersect?
- Given $y(0) = 0$, what can you say about the solution. In particular, can $y(x) > 1$ for any x ? Can $y(x) = 1$ for any x ? Why or why not?

Exercise 1.2.101: Sketch the slope field of $y' = y^3$. Can you visually find the solution that satisfies $y(0) = 0$?

Exercise 1.2.102: Is it possible to solve $y' = xy$ for $y(0) = 0$? Is the solution unique?

Exercise 1.2.103: Is it possible to solve $y' = \frac{x}{x^2-1}$ for $y(1) = 0$?

Exercise 1.2.104: Match equations $y' = \sin x$, $y' = \cos y$, $y' = y \cos(x)$ to slope fields. Justify.



Exercise 1.2.105 (tricky): Suppose

$$f(y) = \begin{cases} 0 & \text{if } y > 0, \\ 1 & \text{if } y \leq 0. \end{cases}$$

Does $y' = f(y)$, $y(0) = 0$ have a continuously differentiable solution? Does Picard apply? Why, or why not?

Exercise 1.2.106: Consider an equation of the form $y' = f(x)$ for some continuous function f , and an initial condition $y(x_0) = y_0$. Does a solution exist for all x ? Why or why not?

1.3 Separable equations

Note: 1 lecture, §1.4 in [EP], §2.2 in [BD]

When a differential equation is of the form $y' = f(x)$, we integrate: $y = \int f(x) dx + C$. Unfortunately, simply integrating no longer works for the general form of the equation $y' = f(x, y)$. Integrating both sides yields the rather unhelpful expression

$$y = \int f(x, y) dx + C.$$

Notice the dependence on y in the integral.

1.3.1 Separable equations

We say a differential equation is *separable* if we can write it as

$$y' = f(x)g(y),$$

for some functions $f(x)$ and $g(y)$. In Leibniz notation, the equation is

$$\frac{dy}{dx} = f(x)g(y),$$

which we rewrite as

$$\frac{dy}{g(y)} = f(x) dx.$$

Both sides now look like something we can integrate. We obtain

$$\int \frac{dy}{g(y)} = \int f(x) dx + C.$$

If we can find closed-form expressions for these two integrals, we can, perhaps, solve for y .

Example 1.3.1: Take the equation

$$y' = xy.$$

Note that $y = 0$ is a solution. We will remember that fact and assume $y \neq 0$ from now on, so that we can divide by y . Write the equation as $\frac{dy}{dx} = xy$ or $\frac{dy}{y} = x dx$. Then

$$\int \frac{dy}{y} = \int x dx + C.$$

We compute the antiderivatives to get

$$\ln|y| = \frac{x^2}{2} + C,$$

or

$$|y| = e^{\frac{x^2}{2} + C} = e^{\frac{x^2}{2}} e^C = D e^{\frac{x^2}{2}},$$

where $D > 0$ is some constant. Because $y = 0$ is also a solution and because of the absolute value, we can write:

$$y = D e^{\frac{x^2}{2}},$$

for any number D (including zero or negative).

We check:

$$y' = D x e^{\frac{x^2}{2}} = x \left(D e^{\frac{x^2}{2}} \right) = x y.$$

Yay!

You may be worried that we integrated in two different variables. We seemingly did a different operation to each side. Perhaps we should be a little bit more careful and work through this method more rigorously. Consider

$$\frac{dy}{dx} = f(x)g(y).$$

We rewrite the equation as follows. Note that $y = y(x)$ is a function of x and so is $\frac{dy}{dx}$!

$$\frac{1}{g(y)} \frac{dy}{dx} = f(x).$$

We integrate both sides with respect to x :

$$\int \frac{1}{g(y)} \frac{dy}{dx} dx = \int f(x) dx + C.$$

We use the change of variables formula (substitution) on the left-hand side:

$$\int \frac{1}{g(y)} dy = \int f(x) dx + C.$$

And we are done.

1.3.2 Implicit solutions

We sometimes get stuck even if we can do the integration. Consider the separable equation

$$y' = \frac{xy}{y^2 + 1}.$$

We separate variables,

$$\frac{y^2 + 1}{y} dy = \left(y + \frac{1}{y} \right) dy = x dx.$$

We integrate to get

$$\frac{y^2}{2} + \ln |y| = \frac{x^2}{2} + C,$$

or perhaps the less intimidating expression (where $D = 2C$)

$$y^2 + 2 \ln |y| = x^2 + D.$$

It is not easy to find the solution explicitly—it is hard to solve for y . We, therefore, leave the solution in this form and call it an *implicit solution*. It is still easy to check that an implicit solution satisfies the differential equation. In this case, we differentiate with respect to x , and remember that y is a function of x , to get

$$2yy' + 2\frac{1}{y}y' = y' \left(2y + \frac{2}{y} \right) = 2x.$$

That is usually called “implicit differentiation” in calculus, but it is just the chain rule. Anyway, now multiply both sides by y and divide by $2(y^2 + 1)$ and you will get exactly the differential equation. We leave this computation to the reader.

If you have an implicit solution, and you want to compute values for y , you might have to be tricky. You might get multiple solutions y for each x , so you have to pick one. Sometimes you can graph x as a function of y , and then turn your paper to see a graph. Sometimes you have to do more.

Computers are also good at some of these tricks. More advanced mathematical software usually has some way of plotting solutions to implicit equations. For example, for $D = 0$, if you plot all the points (x, y) that are solutions to $y^2 + 2 \ln |y| = x^2$, you find the two curves in [Figure 1.8](#) on the following page. This is not quite a graph of a function. For each x , there are two choices of y . To find a function, you have to pick one of these two curves. You pick the one that satisfies your initial condition if you have one. For instance, the top curve satisfies the condition $y(1) = 1$. So for each D , we really got two solutions. As you can see, computing values from an implicit solution can be somewhat tricky. But sometimes, an implicit solution is the best we can do.

The equation above also has the solution $y = 0$. So the general solution is

$$y^2 + 2 \ln |y| = x^2 + D \quad \text{and} \quad y = 0.$$

Sometimes these extra solutions that came up due to division by zero, such as $y = 0$ above, are called *singular solutions*.

We still solve for the constants using initial conditions in implicit solutions as usual. For example, consider the initial condition $y(0) = 1$. We can discount the solution $y = 0$. In the implicit solution, let us plug in $x = 0$ and $y = 1$ to get $1^2 + 2 \ln |1| = 0^2 + D$, in other words, $D = 1$. So the solution for this initial condition is $y^2 + 2 \ln |y| = x^2 + 1$, but only the part of this implicit solution that is positive as that is what our initial condition requires.

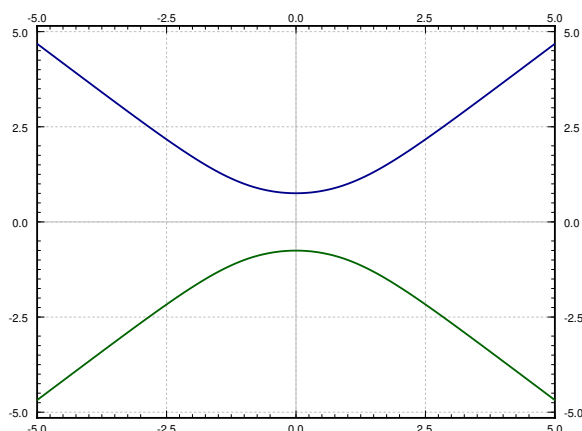


Figure 1.8: The implicit solution $y^2 + 2 \ln |y| = x^2$ to $y' = \frac{xy}{y^2 + 1}$.

1.3.3 More examples of separable equations

Example 1.3.2: Solve $x^2 y' = 1 - x^2 + y^2 - x^2 y^2$, $y(1) = 0$.

Factor the right-hand side

$$x^2 y' = (1 - x^2)(1 + y^2).$$

Separate variables, integrate, and solve for y :

$$\begin{aligned} \frac{y'}{1 + y^2} &= \frac{1 - x^2}{x^2}, \\ \frac{y'}{1 + y^2} &= \frac{1}{x^2} - 1, \\ \arctan(y) &= \frac{-1}{x} - x + C, \\ y &= \tan\left(\frac{-1}{x} - x + C\right). \end{aligned}$$

Solve for the initial condition, $0 = \tan(-2 + C)$ to get $C = 2$ (or $C = 2 + \pi$, or $C = 2 + 2\pi$, etc.). The particular solution we seek is, therefore,

$$y = \tan\left(\frac{-1}{x} - x + 2\right).$$

Example 1.3.3: Bob made a cup of coffee, and Bob likes to drink coffee only once it reaches 60 degrees Celsius and will not burn him. Initially, right after it was poured in the cup, Bob measured the temperature and the coffee was 89 degrees Celsius. One minute later, Bob measured the coffee again and it had 85 degrees. The temperature of the room (the ambient temperature) is 22 degrees. When should Bob start drinking?

Let T be the temperature of the coffee in degrees Celsius, let A be the ambient (room) temperature, also in degrees Celsius, and let t denote the time in minutes, with $t = 0$ denoting the time when Bob made the coffee. Newton's law of cooling states that the rate at which the temperature of the coffee is changing is proportional to the difference between the ambient temperature and the temperature of the coffee. That is,

$$\frac{dT}{dt} = k(A - T),$$

for some positive constant k . For our setup $A = 22$, $T(0) = 89$, $T(1) = 85$. Let C and D denote arbitrary constants, and notice that $T - A > 0$ since the coffee is clearly warmer than the room. We separate variables and integrate.

$$\begin{aligned}\frac{1}{T - A} \frac{dT}{dt} &= -k, \\ \ln(T - A) &= -kt + C, \\ T - A &= De^{-kt}, \\ T &= A + De^{-kt}.\end{aligned}$$

That is, $T = 22 + De^{-kt}$. We plug in the first condition: $89 = T(0) = 22 + D$, and hence $D = 67$. So $T = 22 + 67e^{-kt}$. The second condition says $85 = T(1) = 22 + 67e^{-k}$. Solving for k , we get $k = -\ln \frac{85-22}{67} \approx 0.0616$. We solve for the time t that gives a temperature of 60 degrees. Namely, we solve

$$60 = 22 + 67e^{-0.0616t}$$

to get $t = -\frac{\ln \frac{60-22}{67}}{0.0616} \approx 9.21$ minutes. So Bob can begin to drink the coffee at just over 9 minutes from the time Bob made it. That is probably about the amount of time it took us to calculate how long it would take. See [Figure 1.9](#) on the next page.

Example 1.3.4: Find the general solution to $y' = \frac{-xy^2}{3}$ (including any singular solutions).

First, note that $y = 0$ is a solution (a singular solution). Next, assume $y \neq 0$ and separate:

$$\begin{aligned}\frac{-3}{y^2} y' &= x, \\ \frac{3}{y} &= \frac{x^2}{2} + C, \\ y &= \frac{3}{x^2/2 + C} = \frac{6}{x^2 + 2C}.\end{aligned}$$

The general solution is

$$y = \frac{6}{x^2 + 2C} \quad \text{and} \quad y = 0.$$

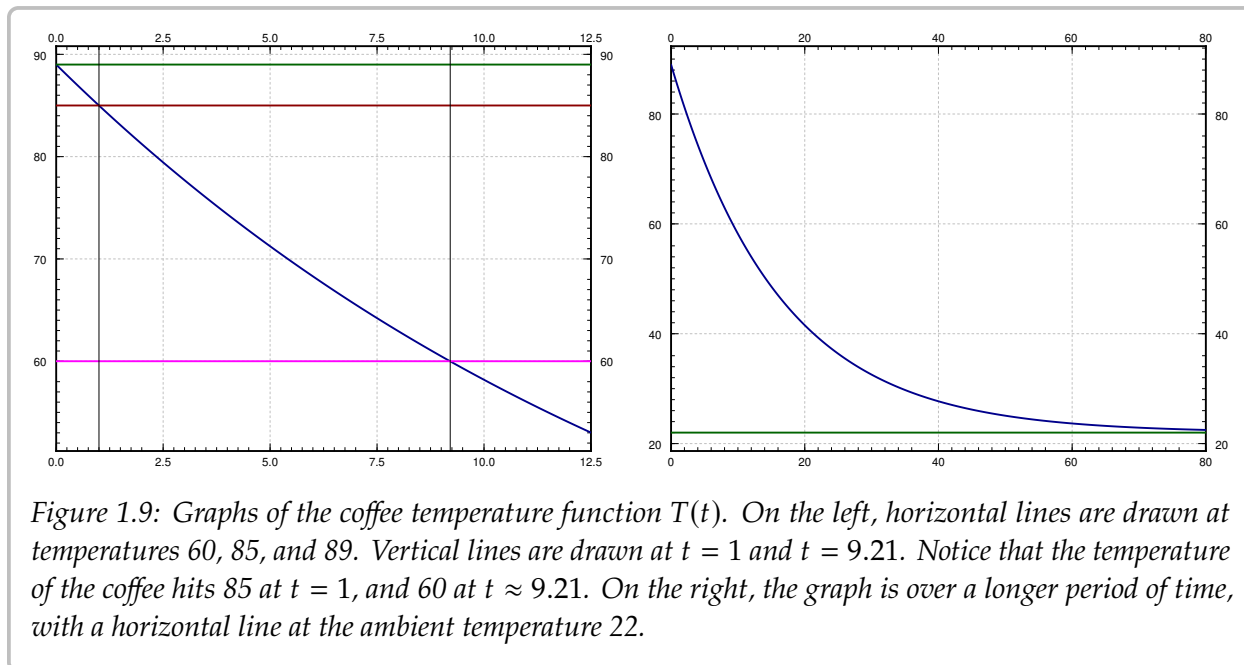


Figure 1.9: Graphs of the coffee temperature function $T(t)$. On the left, horizontal lines are drawn at temperatures 60, 85, and 89. Vertical lines are drawn at $t = 1$ and $t = 9.21$. Notice that the temperature of the coffee hits 85 at $t = 1$, and 60 at $t \approx 9.21$. On the right, the graph is over a longer period of time, with a horizontal line at the ambient temperature 22.

1.3.4 Exercises

Exercise 1.3.1: Solve $y' = \frac{x}{y}$.

Exercise 1.3.2: Solve $y' = x^2 y$.

Exercise 1.3.3: Solve $\frac{dx}{dt} = (x^2 - 1)t$, for $x(0) = 0$.

Exercise 1.3.4: Solve $\frac{dx}{dt} = x \sin(t)$, for $x(0) = 1$.

Exercise 1.3.5: Solve $\frac{dy}{dx} = xy + x + y + 1$. Hint: Factor the right-hand side.

Exercise 1.3.6: Solve $xy' = y + 2x^2 y$, where $y(1) = 1$.

Exercise 1.3.7: Solve $\frac{dy}{dx} = \frac{y^2 + 1}{x^2 + 1}$, for $y(0) = 1$.

Exercise 1.3.8: Find an implicit solution for $\frac{dy}{dx} = \frac{x^2 + 1}{y^2 + 1}$, for $y(0) = 1$.

Exercise 1.3.9: Find an explicit solution for $y' = xe^{-y}$, $y(0) = 1$.

Exercise 1.3.10: Find an explicit solution for $xy' = e^{-y}$, for $y(1) = 1$.

Exercise 1.3.11: Find an explicit solution for $y' = ye^{-x^2}$, $y(0) = 1$. It is alright to leave a definite integral in your answer.

Exercise 1.3.12: Suppose a cup of coffee is at 100 degrees Celsius at time $t = 0$, it is at 70 degrees at $t = 10$ minutes, and it is at 50 degrees at $t = 20$ minutes. Compute the ambient temperature.

Exercise 1.3.101: Solve $y' = 2xy$.

Exercise 1.3.102: Solve $x' = 3xt^2 - 3t^2$, $x(0) = 2$.

Exercise 1.3.103: Find an implicit solution for $x' = \frac{1}{3x^2 + 1}$, $x(0) = 1$.

Exercise 1.3.104: Find an explicit solution to $xy' = y^2$, $y(1) = 1$.

Exercise 1.3.105: Find an implicit solution to $y' = \frac{\sin(x)}{\cos(y)}$.

Exercise 1.3.106: Take [Example 1.3.3](#) with the same numbers: 89 degrees at $t = 0$, 85 degrees at $t = 1$, and ambient temperature of 22 degrees. Suppose these temperatures were measured with a precision of ± 0.5 degrees. Given this imprecision, the time it takes the coffee to cool to (exactly) 60 degrees is also only known in a certain range. Find this range. Hint: Think about what kind of error makes the cooling time longer and what shorter.

Exercise 1.3.107: A population x of rabbits on an island is modeled by $x' = x - (1/1000)x^2$, where the independent variable t is time in months. At time $t = 0$, there are 40 rabbits on the island.

- Find the solution to the equation with the initial condition.
- How many rabbits are on the island in 1 month, 5 months, 10 months, 15 months (round to the nearest integer)?

1.4 Linear equations and the integrating factor

Note: 1 lecture, §1.5 in [EP], §2.1 in [BD]

Linear equations are one of the most important types of equations we will learn to solve. In fact, the majority of the course is about linear equations. In this section, we focus on the *first-order linear equation*. A first-order equation is linear if we can put it into the form:

$$y' + p(x)y = f(x). \quad (1.3)$$

The word “linear” means linear in y and y' : No higher powers or functions of y or y' appear. The dependence on x can be more complicated.

Solutions of linear equations have nice properties. For example, the solution exists wherever $p(x)$ and $f(x)$ are continuous, and it has essentially the same regularity (read: it is just as nice). Most importantly for us right now, there is a method for solving linear first-order equations.

The trick is to rewrite the left-hand side of (1.3) as a derivative of a product of y with another function. To this end, we wish to find a function $r(x)$ such that

$$r(x)y' + r(x)p(x)y = \frac{d}{dx} [r(x)y].$$

This is the left-hand side of (1.3) multiplied by $r(x)$. If we multiply (1.3) by $r(x)$, we obtain

$$\frac{d}{dx} [r(x)y] = r(x)f(x).$$

We can now integrate both sides, which we can do as the right-hand side does not depend on y and the left-hand side is written as a derivative of a function. After the integration, we solve for y by dividing by $r(x)$. The function $r(x)$ is called the *integrating factor* and the method is called the *integrating factor method*.

We are looking for a function $r(x)$, such that if we differentiate it, we get the same function back multiplied by $p(x)$. That seems like a job for the exponential function! Let

$$r(x) = e^{\int p(x) dx},$$

where $\int p(x) dx$ is any antiderivative of $p(x)$. We compute:

$$\begin{aligned} y' + p(x)y &= f(x), \\ e^{\int p(x) dx} y' + e^{\int p(x) dx} p(x)y &= e^{\int p(x) dx} f(x), \\ \frac{d}{dx} \left[e^{\int p(x) dx} y \right] &= e^{\int p(x) dx} f(x), \\ e^{\int p(x) dx} y &= \int e^{\int p(x) dx} f(x) dx + C, \\ y &= e^{-\int p(x) dx} \left(\int e^{\int p(x) dx} f(x) dx + C \right). \end{aligned}$$

Of course, to get a closed-form formula for y , we need to be able to find a closed-form formula for the integrals appearing above.

Example 1.4.1: Solve

$$y' + 2xy = e^{x-x^2}, \quad y(0) = -1.$$

First note that $p(x) = 2x$ and $f(x) = e^{x-x^2}$. The integrating factor is $r(x) = e^{\int p(x) dx} = e^{x^2}$. We multiply both sides of the equation by $r(x)$ to get

$$\begin{aligned} e^{x^2} y' + 2xe^{x^2} y &= e^{x^2} e^{x-x^2}, \\ \frac{d}{dx} [e^{x^2} y] &= e^x. \end{aligned}$$

We integrate

$$\begin{aligned} e^{x^2} y &= e^x + C, \\ y &= e^{x-x^2} + Ce^{-x^2}. \end{aligned}$$

Next, we solve for the initial condition $-1 = y(0) = 1 + C$, so $C = -2$. The solution is

$$y = e^{x-x^2} - 2e^{-x^2}.$$

Note that we do not care which antiderivative we take when computing $e^{\int p(x) dx}$. You can always add a constant of integration, but those constants will not matter in the end.

Exercise 1.4.1: Try it! Add a constant of integration to the integral in the integrating factor and show that the solution you get in the end is the same as what we got above.

Advice: Do not try to remember the formula for y itself—that is way too hard. It is easier to remember the process and repeat it.

Since we cannot always evaluate the integrals in closed form, it is useful to know how to write the solution in definite integral form. A definite integral is something that you can plug into a computer or a calculator. Suppose we are given

$$y' + p(x)y = f(x), \quad y(x_0) = y_0.$$

Look at the solution and write the integrals as definite integrals.

$$y(x) = e^{-\int_{x_0}^x p(s) ds} \left(\int_{x_0}^x e^{\int_{x_0}^t p(s) ds} f(t) dt + y_0 \right). \quad (1.4)$$

You should be careful to properly use dummy variables here. If you now plug such a formula into a computer or a calculator, it will be happy to give you numerical answers.

Exercise 1.4.2: Check that $y(x_0) = y_0$ in formula (1.4).

Exercise 1.4.3: Write the solution of the following problem as a definite integral, but try to simplify as far as you can. You will not be able to find the solution in closed form.

$$y' + y = e^{x^2-x}, \quad y(0) = 10.$$

Remark 1.4.1: Before we move on, we should note some interesting properties of linear equations. First, for the linear initial value problem $y' + p(x)y = f(x)$, $y(x_0) = y_0$, there is an explicit formula (1.4) for the solution. Second, it follows from the formula (1.4) that if $p(x)$ and $f(x)$ are continuous on some interval (a, b) , then the solution $y(x)$ exists and is differentiable on (a, b) . Compare with the simple nonlinear example we have seen previously, $y' = y^2$, and compare to Theorem 1.2.1.

Example 1.4.2: We get to a common simple application of linear equations. Real-life applications of this type of problem include computing the concentration of chemicals in bodies of water (rivers and lakes).

A 100-liter tank contains 10 kilograms of salt dissolved in 60 liters of water. A solution of water and salt (brine) with a concentration of 0.1 kilograms per liter is flowing in at the rate of 5 liters a minute. The solution in the tank is well stirred and flows out at a rate of 3 liters a minute. How much salt is in the tank when the tank is full?

To solve this problem, we need to find the differential equation for this setup. Let x denote the kilograms of salt in the tank, and let t denote the time in minutes. For a small change Δt in time, the change in x (denoted Δx) is approximately

$$\Delta x \approx (\text{rate in} \times \text{concentration in})\Delta t - (\text{rate out} \times \text{concentration out})\Delta t.$$

Dividing through by Δt and taking the limit $\Delta t \rightarrow 0$, we see that

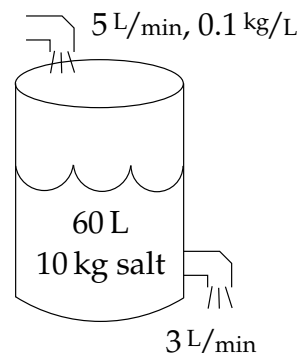
$$\frac{dx}{dt} = (\text{rate in} \times \text{concentration in}) - (\text{rate out} \times \text{concentration out}).$$

In our example,

$$\begin{aligned} \text{rate in} &= 5, \\ \text{concentration in} &= 0.1, \\ \text{rate out} &= 3, \\ \text{volume} &= \text{initial volume} + (\text{rate in} - \text{rate out})t = 60 + (5 - 3)t, \\ \text{concentration out} &= \frac{x}{\text{volume}} = \frac{x}{60 + (5 - 3)t}. \end{aligned}$$

Our differential equation is

$$\frac{dx}{dt} = (5 \times 0.1) - \left(3 \frac{x}{60 + 2t}\right).$$



In the form (1.3), it is

$$\frac{dx}{dt} + \frac{3}{60 + 2t}x = 0.5.$$

Let us solve. The integrating factor is

$$r(t) = \exp\left(\int \frac{3}{60 + 2t} dt\right) = \exp\left(\frac{3}{2} \ln(60 + 2t)\right) = (60 + 2t)^{3/2}.$$

We multiply both sides of the equation by $r(t)$ and simplify,

$$\begin{aligned} (60 + 2t)^{3/2} \frac{dx}{dt} + (60 + 2t)^{3/2} \frac{3}{60 + 2t} x &= 0.5(60 + 2t)^{3/2}, \\ \frac{d}{dt} \left[(60 + 2t)^{3/2} x \right] &= 0.5(60 + 2t)^{3/2}. \end{aligned}$$

We integrate and solve,

$$\begin{aligned} (60 + 2t)^{3/2} x &= \int 0.5(60 + 2t)^{3/2} dt + C, \\ x &= (60 + 2t)^{-3/2} \int \frac{(60 + 2t)^{3/2}}{2} dt + C(60 + 2t)^{-3/2}, \\ x &= (60 + 2t)^{-3/2} \frac{1}{10} (60 + 2t)^{5/2} + C(60 + 2t)^{-3/2}, \\ x &= \frac{60 + 2t}{10} + C(60 + 2t)^{-3/2}. \end{aligned}$$

To find C , note that at $t = 0$, we have $x = 10$. That is,

$$10 = x(0) = \frac{60}{10} + C(60)^{-3/2} = 6 + C(60)^{-3/2},$$

or

$$C = 4(60^{3/2}) \approx 1859.03.$$

As 5 liters per minute are flowing in and 3 liters per minute are flowing out, the volume is increasing by 2 liters a minute. So the tank is full when $60 + 2t = 100$, or when $t = 20$. We want the value of x when the tank is full, that is, we compute $x(20)$:

$$\begin{aligned} x(20) &= \frac{60 + 40}{10} + C(60 + 40)^{-3/2} \\ &\approx 10 + 1859.03(100)^{-3/2} \approx 11.86. \end{aligned}$$

There are 11.86 kg of salt in the tank when it is full. See Figure 1.10 for the graph of $x(t)$. The concentration when the tank is full is approximately $11.86/100 = 0.1186$ kg/liter, and we started with $1/6$ or approximately 0.1667 kg/liter.

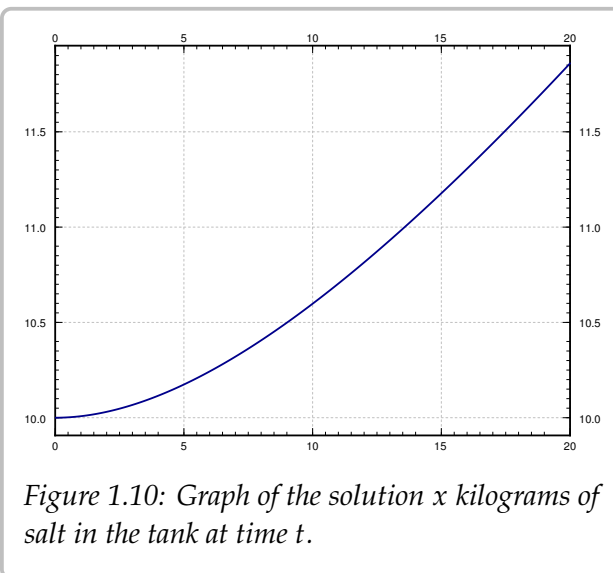


Figure 1.10: Graph of the solution x kilograms of salt in the tank at time t .

1.4.1 Exercises

In the exercises, feel free to leave the answer as a definite integral if a closed-form solution cannot be found. If you can find a closed-form solution, you should give that.

Exercise 1.4.4: Solve $y' + xy = x$.

Exercise 1.4.5: Solve $y' + 6y = e^x$.

Exercise 1.4.6: Solve $y' + 3x^2y = \sin(x)e^{-x^3}$, with $y(0) = 1$.

Exercise 1.4.7: Solve $y' + \cos(x)y = \cos(x)$.

Exercise 1.4.8: Solve $\frac{1}{x^2+1}y' + xy = 3$, with $y(0) = 0$.

Exercise 1.4.9: Suppose there are two lakes located on a stream. Clean water flows into the first lake, then the water from the first lake flows into the second lake, and then water from the second lake flows further downstream. The in and out flow from each lake is 500 liters per hour. The first lake contains 100 thousand liters of water and the second lake contains 200 thousand liters of water. A truck with 500 kg of toxic substance crashes into the first lake. Assume that the water is being continually mixed perfectly by the stream.

- Find the concentration of toxic substance as a function of time in both lakes.
- When will the concentration in the first lake be below 0.001 kg per liter?
- When will the concentration in the second lake be maximal?

Exercise 1.4.10: Newton's law of cooling states that $\frac{dx}{dt} = -k(x - A)$ where x is the temperature, t is time, A is the ambient temperature, and $k > 0$ is a constant. Suppose that $A = A_0 \cos(\omega t)$ for some constants A_0 and ω . That is, the ambient temperature oscillates (for example night and day temperatures).

- Find the general solution.
- In the long term, will the initial conditions make much of a difference? Why or why not?

Exercise 1.4.11: Initially 5 grams of salt are dissolved in 20 liters of water. Brine with a salt concentration of 2 grams per liter is added at a rate of 3 liters a minute. The tank is mixed well and is drained at 3 liters a minute. How long does the process have to continue until there are 20 grams of salt in the tank?

Exercise 1.4.12: Initially a tank contains 10 liters of pure water. Brine of unknown (but constant) concentration of salt is flowing in at 1 liter per minute. The water is mixed well and drained at 1 liter per minute. In 20 minutes there are 15 grams of salt in the tank. What is the concentration of salt in the incoming brine?

Exercise 1.4.101: Solve $y' + 3x^2y = x^2$.

Exercise 1.4.102: Solve $y' + 2 \sin(2x)y = 2 \sin(2x)$, $y(\pi/2) = 3$.

Exercise 1.4.103: Suppose a water tank is being pumped out at 3 L/min . The water tank starts at 10 L of clean water. Water with a toxic substance is flowing into the tank at 2 L/min , with concentration $20t \text{ g/L}$ at time t . When the tank is half empty, how many grams of toxic substance are in the tank (assuming perfect mixing)?

Exercise 1.4.104: There are bacteria on a plate and a toxic substance is being added that slows down the rate of growth of the bacteria. That is, suppose that $\frac{dP}{dt} = (2 - 0.1t)P$. If $P(0) = 1000$, find the population at $t = 5$.

Exercise 1.4.105: A cylindrical water tank has water flowing in at I cubic meters per second. Let A be the area of the cross section of the tank in square meters. Suppose water is flowing out from the bottom of the tank at a rate proportional to the height of the water level. Set up the differential equation for h , the height of the water, introducing and naming constants that you need. You should also give the units for your constants.

1.5 Substitution

Note: 1 lecture, can safely be skipped, §1.6 in [EP], not in [BD]

Just as when solving integrals, one method to try is to change variables to end up with a simpler equation to solve.

1.5.1 Substitution

The equation

$$y' = (x - y + 1)^2$$

is neither separable nor linear. What can we do? How about trying to change variables, so that in the new variables the equation is simpler? We use another variable v , which we treat as a function of x . We try

$$v = x - y + 1.$$

We need to figure out y' in terms of v' , v and x . We differentiate (in x) to obtain $v' = 1 - y'$. So $y' = 1 - v'$. We plug this into the equation to get

$$1 - v' = v^2.$$

In other words, $v' = 1 - v^2$. We know how to solve such an equation by separating variables:

$$\frac{1}{1 - v^2} dv = dx.$$

So

$$\frac{1}{2} \ln \left| \frac{v+1}{v-1} \right| = x + C, \quad \text{or} \quad \left| \frac{v+1}{v-1} \right| = e^{2x+2C}, \quad \text{or} \quad \frac{v+1}{v-1} = De^{2x},$$

for some constant D . Note that $v = 1$ and $v = -1$ are also solutions.

Now we need to “unsubstitute” to obtain

$$\frac{x - y + 2}{x - y} = De^{2x},$$

and also the two solutions $x - y + 1 = 1$ or $y = x$, and $x - y + 1 = -1$ or $y = x + 2$. We solve the first equation for y :

$$\begin{aligned} x - y + 2 &= (x - y)De^{2x}, \\ x - y + 2 &= Dxe^{2x} - yDe^{2x}, \\ -y + yDe^{2x} &= Dxe^{2x} - x - 2, \\ y(-1 + De^{2x}) &= Dxe^{2x} - x - 2, \\ y &= \frac{Dxe^{2x} - x - 2}{De^{2x} - 1}. \end{aligned}$$

Note that $D = 0$ gives $y = x + 2$, but no value of D gives the solution $y = x$.

Substitution in differential equations is applied in much the same way that it is applied in calculus. You guess. Several different substitutions might work. There are some general patterns to look for. We summarize a few of these in a table.

When you see	Try substituting
yy'	$v = y^2$
y^2y'	$v = y^3$
$(\cos y)y'$	$v = \sin y$
$(\sin y)y'$	$v = \cos y$
$e^y y'$	$v = e^y$

Usually, you try to substitute in the “most complicated” part of the equation, with the hope of simplifying it. The table above is just a rule of thumb. You might have to modify your guesses. If a substitution does not work (it does not make the equation any simpler), try a different one.

1.5.2 Bernoulli equations

There are some forms of equations where there is a general rule for substitution that always works. One such example is the so-called *Bernoulli equation**:

$$y' + p(x)y = q(x)y^n.$$

This equation looks a lot like a linear equation except for the y^n . If $n = 0$ or $n = 1$, then the equation is linear and we can solve it. Otherwise, the substitution $v = y^{1-n}$ transforms the Bernoulli equation into a linear equation. Note that n need not be an integer.

Example 1.5.1: Solve

$$xy' + y(x + 1) + xy^5 = 0, \quad y(1) = 1.$$

Divide the equation by x and rearrange to see it is a Bernoulli equation, $y' + \frac{x+1}{x}y = -y^5$, that is, $p(x) = \frac{x+1}{x}$ and $q(x) = -1$. We substitute

$$v = y^{1-5} = y^{-4}, \quad v' = -4y^{-5}y'.$$

In other words, $(-1/4)y^5v' = y'$. So

$$\begin{aligned} xy' + y(x + 1) + xy^5 &= 0, \\ \frac{-xy^5}{4}v' + y(x + 1) + xy^5 &= 0, \end{aligned}$$

*There are several things called Bernoulli equations; this is just one of them. The Bernoullis were a prominent Swiss family of mathematicians. These particular equations are named for [Jacob Bernoulli](#) (1654–1705).

$$\begin{aligned}\frac{-x}{4}v' + y^{-4}(x+1) + x &= 0, \\ \frac{-x}{4}v' + v(x+1) + x &= 0,\end{aligned}$$

and finally

$$v' - \frac{4(x+1)}{x}v = 4.$$

The equation is now linear. We can use the integrating factor method. In particular, we use formula (1.4). We assume that $x > 0$ so $|x| = x$. This assumption is OK, as our initial condition is at $x = 1 > 0$. Let us compute the integrating factor. Here $p(s)$ from formula (1.4) is $\frac{-4(s+1)}{s}$.

$$\begin{aligned}e^{\int_1^x p(s) ds} &= \exp\left(\int_1^x \frac{-4(s+1)}{s} ds\right) = e^{-4x-4\ln(x)+4} = e^{-4x+4}x^{-4} = \frac{e^{-4x+4}}{x^4}, \\ e^{-\int_1^x p(s) ds} &= e^{4x+4\ln(x)-4} = e^{4x-4}x^4.\end{aligned}$$

We now plug in to (1.4)

$$\begin{aligned}v(x) &= e^{-\int_1^x p(s) ds} \left(\int_1^x e^{\int_1^t p(s) ds} 4 dt + 1 \right) \\ &= e^{4x-4}x^4 \left(\int_1^x 4 \frac{e^{-4t+4}}{t^4} dt + 1 \right).\end{aligned}$$

The integral in this expression is not possible to find in closed form. As we said before, it is perfectly fine to have a definite integral in our solution. Now “unsubstitute”

$$\begin{aligned}y^{-4} &= e^{4x-4}x^4 \left(4 \int_1^x \frac{e^{-4t+4}}{t^4} dt + 1 \right), \\ y &= \frac{e^{-x+1}}{x \left(4 \int_1^x \frac{e^{-4t+4}}{t^4} dt + 1 \right)^{1/4}}.\end{aligned}$$

1.5.3 Homogeneous equations

Using substitution, we can also solve the so-called *homogeneous equations*. Suppose that we can write the differential equation as

$$y' = F\left(\frac{y}{x}\right).$$

Here we try the substitutions

$$v = \frac{y}{x} \quad \text{and therefore} \quad y' = v + xv'.$$

We note that the equation is transformed into

$$v + xv' = F(v) \quad \text{or} \quad xv' = F(v) - v \quad \text{or} \quad \frac{v'}{F(v) - v} = \frac{1}{x}.$$

Hence an implicit solution is

$$\int \frac{1}{F(v) - v} dv = \ln|x| + C.$$

Clearly this solution does not work when $x = 0$ (we would, after all, divide by zero in y/x). So we will either assume $x > 0$ or $x < 0$ depending on the initial condition.

Example 1.5.2: Solve

$$x^2 y' = y^2 + xy, \quad y(1) = 1.$$

We put the equation into the form $y' = (y/x)^2 + y/x$, that is, $F(v) = v^2 + v$. As the initial condition is for a positive x value, we will assume $x > 0$. We substitute $v = y/x$ to get the separable equation

$$xv' = v^2 + v - v = v^2,$$

which has a solution

$$\begin{aligned} \int \frac{1}{v^2} dv &= \ln|x| + C, \\ \frac{-1}{v} &= \ln x + C, \\ v &= \frac{-1}{\ln x + C}. \end{aligned}$$

We unsubstute

$$\frac{y}{x} = \frac{-1}{\ln x + C}, \quad \text{or} \quad y = \frac{-x}{\ln x + C}.$$

We want $y(1) = 1$, so

$$1 = y(1) = \frac{-1}{\ln 1 + C} = \frac{-1}{C}.$$

Thus $C = -1$ and the solution we are looking for is

$$y = \frac{-x}{\ln x - 1}.$$

1.5.4 Exercises

Hint: Answers need not always be in closed form.

Exercise 1.5.1: Solve $y' + y(x^2 - 1) + xy^6 = 0$, with $y(1) = 1$.

Exercise 1.5.2: Solve $2yy' + 1 = y^2 + x$, with $y(0) = 1$.

Exercise 1.5.3: Solve $y' + xy = y^4$, with $y(0) = 1$.

Exercise 1.5.4: Solve $yy' + x = \sqrt{x^2 + y^2}$.

Exercise 1.5.5: Solve $y' = (x + y - 1)^2$.

Exercise 1.5.6: Solve $y' = \frac{x^2 - y^2}{xy}$, with $y(1) = 2$.

Exercise 1.5.101: Solve $xy' + y + y^2 = 0$, $y(1) = 2$.

Exercise 1.5.102: Solve $xy' + y + x = 0$, $y(1) = 1$.

Exercise 1.5.103: Solve $y^2y' = y^3 - 3x$, $y(0) = 2$.

Exercise 1.5.104: Solve $2yy' = e^{y^2 - x^2} + 2x$.

1.6 Autonomous equations

Note: 1 lecture, §2.2 in [EP], §2.5 in [BD]

Consider problems of the form

$$\frac{dx}{dt} = f(x),$$

where the derivative of solutions depends only on x (the dependent variable). Such equations are called *autonomous equations*. If we think of t as time, the naming comes from the fact that the equation is independent of time.

We return to the cooling coffee problem (Example 1.3.3). Newton's law of cooling says

$$\frac{dx}{dt} = k(A - x),$$

where x is the temperature, t is time, k is some positive constant, and A is the ambient temperature. See Figure 1.11 for an example with $k = 0.3$ and $A = 5$.

Note the solution $x = A$ (in the figure $x = 5$). We call these constant solutions the *equilibrium solutions*. The points on the x -axis where $f(x) = 0$ are called *critical points* or *equilibria*. The point $x = A$ is a critical point. In fact, each critical point corresponds to an equilibrium solution. Note also, by looking at the graph, that the solution $x = A$ is “stable” in that small perturbations in x do not lead to substantially different solutions as t grows. If we change the initial condition a little bit, then as $t \rightarrow \infty$, we still get $x(t) \rightarrow A$. We call such a critical point *stable*. In this simple example, all solutions in fact go to A as $t \rightarrow \infty$. If a critical point is not stable, we say it is *unstable*.

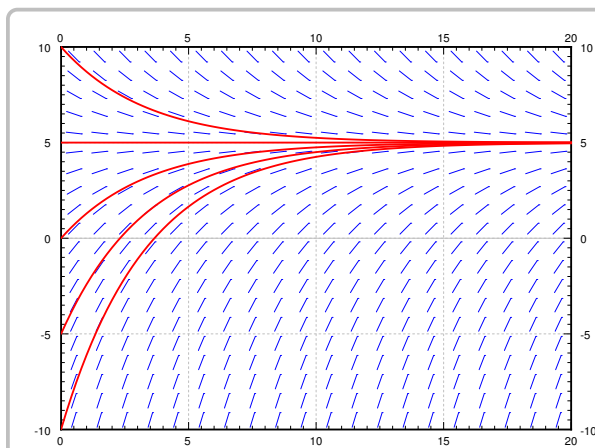


Figure 1.11: The slope field and some solutions of $x' = 0.3(5 - x)$.

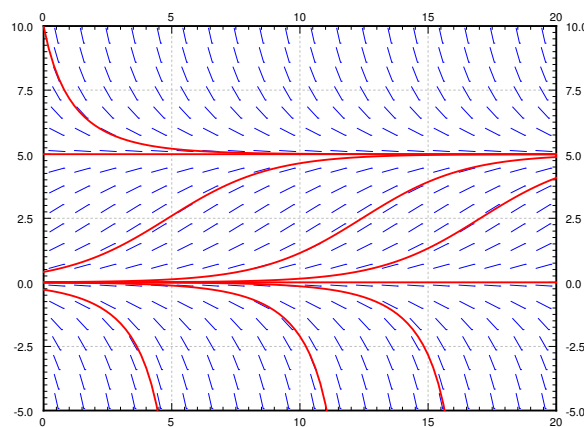


Figure 1.12: The slope field and some solutions of $x' = 0.1x(5 - x)$.

Consider now the *logistic equation*

$$\frac{dx}{dt} = kx(M - x),$$

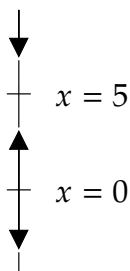
for some positive k and M . This equation is commonly used to model population if we know the limiting population M , that is, the maximum sustainable population. The logistic equation leads to less catastrophic predictions about world population than $x' = kx$. In the real world, there is no such thing as negative population, but we will still consider negative x for the purposes of the math.

See Figure 1.12 on the preceding page for an example, $x' = 0.1x(5 - x)$. There are two critical points, $x = 0$ and $x = 5$. The critical point at $x = 5$ is stable, while the critical point at $x = 0$ is unstable. It is not necessary to find the exact solutions to understand their long-term behavior, that is, behavior as time goes to infinity. From the slope field above of $x' = 0.1x(5 - x)$, we see that

$$\lim_{t \rightarrow \infty} x(t) = \begin{cases} 5 & \text{if } x(0) > 0, \\ 0 & \text{if } x(0) = 0, \\ \text{DNE or } -\infty & \text{if } x(0) < 0. \end{cases}$$

Here DNE means “does not exist.” From just looking at the slope field, we cannot quite decide what happens if $x(0) < 0$. It could be that the solution does not exist for t all the way to ∞ . Think of the equation $x' = x^2$. We have seen that solutions only exist for some finite period of time. Same can happen here. In our example equation above it turns out that the solution does not exist for all time, but to see that we would have to solve the equation. In any case, the solution does go to $-\infty$, but it may get there rather quickly.

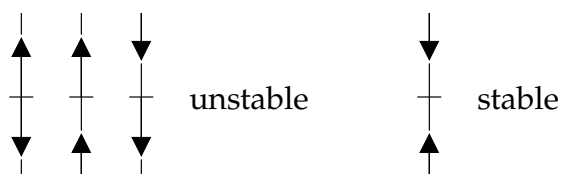
If we are interested only in the long-term behavior of the solution, we would be doing unnecessary work if we solved the equation exactly. We could draw the slope field, but it is easier to just look at the *phase diagram* or *phase portrait*, which is a simple way to visualize the behavior of autonomous equations. In this case there is one dependent variable, the x . We draw the x -axis, we mark all the critical points, and then we draw arrows in between. Since x is the dependent variable, we draw the axis vertically, as it appears in the slope field diagrams above. If $f(x) > 0$, we draw an up arrow. If $f(x) < 0$, we draw a down arrow. To figure out the sign, we could just plug in some x between the critical points because $f(x)$ will have the same sign at all x between two critical points as long as $f(x)$ is continuous. For example, $f(6) = -0.6 < 0$, so $f(x) < 0$ for $x > 5$, and the arrow above $x = 5$ is a down arrow. Next, $f(1) = 0.4 > 0$, so $f(x) > 0$ whenever $0 < x < 5$, and the arrow points up. Finally, $f(-1) = -0.6 < 0$, so $f(x) < 0$ when $x < 0$, and the arrow points down.



Armed with the phase diagram, it is easy to sketch the solutions approximately: As time t moves from left to right, the graph of a solution goes up if the arrow is up, and it goes down if the arrow is down.

Exercise 1.6.1: Try sketching a few solutions simply from looking at the phase diagram. Check with the preceding graphs if you are getting the same type of curves.

Once we draw the phase diagram, we classify critical points as stable or unstable*:



Since any mathematical model we cook up will only be an approximation to the real world, unstable points are generally bad news. We remark that you can figure out the arrows by plotting the graph $y = f(x)$. For that graph, note that x is the independent variable and will be on the horizontal axis.

Let us think about the logistic equation with harvesting. Suppose an alien race really likes to eat humans. They keep a planet with humans and harvest the humans at a rate of h million humans per year. Suppose x is the number of humans in millions on the planet and t is time in years. Let M be the limiting population when no harvesting is done. The number $k > 0$ is a constant depending on how fast humans multiply. Our equation becomes

$$\frac{dx}{dt} = kx(M - x) - h.$$

We expand the right-hand side and set it to zero.

$$kx(M - x) - h = -kx^2 + kMx - h = 0.$$

Solving for the critical points, let us call them A and B , we get

$$A = \frac{kM + \sqrt{(kM)^2 - 4hk}}{2k}, \quad B = \frac{kM - \sqrt{(kM)^2 - 4hk}}{2k}.$$

Exercise 1.6.2: Sketch a phase diagram for different possibilities. Note that these possibilities are $A > B$, or $A = B$, or A and B both complex (i.e. no real solutions). Hint: Fix some simple k and M and then vary h .

For example, let $M = 8$ and $k = 0.1$. When $h = 1$, then A and B are distinct and positive. See Figure 1.13 on the next page for the slope field. As long as the population starts above B , which is approximately 1.55 million, then the population will not die out, it will tend towards $A \approx 6.45$ million. If ever a catastrophe happens and the population drops below B , humans will die out, and the fast food restaurant serving them will go out of business.

*Unstable points with one of the arrows pointing towards the critical point are sometimes called *semistable*.

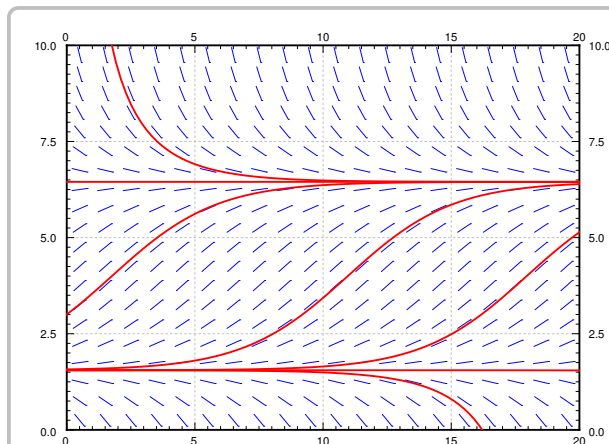


Figure 1.13: The slope field and some solutions of $x' = 0.1x(8 - x) - 1$.

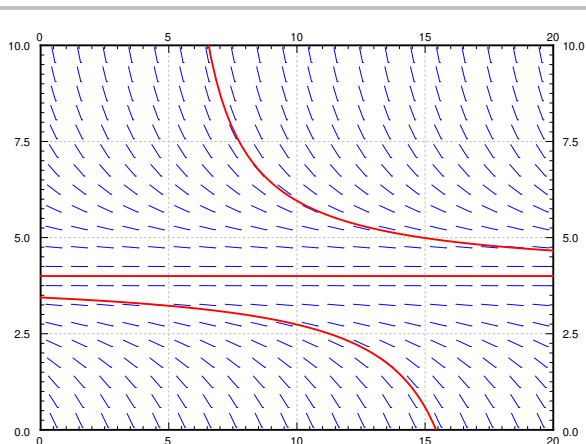


Figure 1.14: The slope field and some solutions of $x' = 0.1x(8 - x) - 1.6$.

When $h = 1.6$, then $A = B = 4$. There is only one critical point and it is unstable*. When the population starts above 4 million, it will tend towards 4 million. However, if it ever drops below 4 million, perhaps due to a worse than normal hurricane season one year, then humans will die out on the planet. This scenario is not one that we (as the human fast food proprietor) want to be in. A small perturbation of the equilibrium state and we are out of business. There is no room for error. See Figure 1.14.

Finally, if we are harvesting at 2 million humans per year, there are no critical points. The population will always plummet towards zero, no matter how well stocked the planet starts. See Figure 1.15.

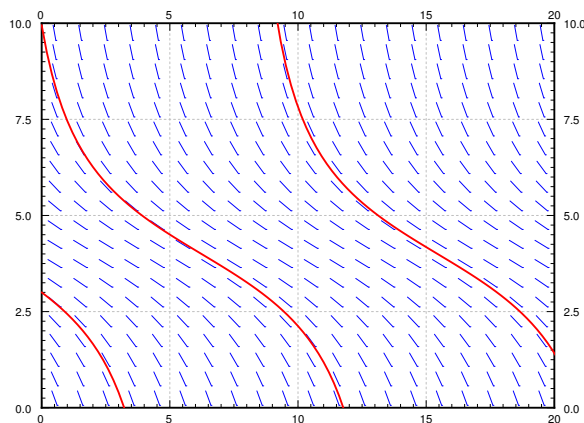


Figure 1.15: The slope field and some solutions of $x' = 0.1x(8 - x) - 2$.

*More precisely, it is semistable.

1.6.1 Exercises

Exercise 1.6.3: Consider $x' = x^2$.

- Draw the phase diagram, find the critical points, and mark them stable or unstable.
- Sketch typical solutions of the equation.
- Find $\lim_{t \rightarrow \infty} x(t)$ for the solution with the initial condition $x(0) = -1$.

Exercise 1.6.4: Consider $x' = \sin x$.

- Draw the phase diagram for $-4\pi \leq x \leq 4\pi$. On this interval mark the critical points stable or unstable.
- Sketch typical solutions of the equation.
- Find $\lim_{t \rightarrow \infty} x(t)$ for the solution with the initial condition $x(0) = 1$.

Exercise 1.6.5: Suppose $f(x)$ is positive for $0 < x < 1$, it is zero when $x = 0$ and $x = 1$, and it is negative for all other x .

- Draw the phase diagram for $x' = f(x)$, find the critical points, and mark them stable or unstable.
- Sketch typical solutions of the equation.
- Find $\lim_{t \rightarrow \infty} x(t)$ for the solution with the initial condition $x(0) = 0.5$.

Exercise 1.6.6: Start with the logistic equation $\frac{dx}{dt} = kx(M - x)$. Suppose we modify our harvesting. We will only harvest an amount proportional to the current population. In other words, we harvest hx per unit of time for some $h > 0$ (similar to earlier example with h replaced with hx).

- Construct the differential equation.
- Show that if $kM > h$, then the equation is still logistic.
- What happens when $kM < h$?

Exercise 1.6.7: A disease is spreading through the country. Let x be the number of people infected. Let the constant S be the number of people susceptible to infection. The infection rate $\frac{dx}{dt}$ is proportional to the product of already infected people, x , and the number of susceptible but uninfected people, $S - x$.

- Write down the differential equation.
- Supposing $x(0) > 0$, that is, some people are infected at time $t = 0$, what is $\lim_{t \rightarrow \infty} x(t)$.
- Does the solution to part b) agree with your intuition? Why or why not?

Exercise 1.6.101: Let $x' = (x - 1)(x - 2)x^2$.

- a) Sketch the phase diagram and find critical points.
- b) Classify the critical points.
- c) If $x(0) = 0.5$, then find $\lim_{t \rightarrow \infty} x(t)$.

Exercise 1.6.102: Let $x' = e^{-x}$.

- a) Find and classify all critical points.
- b) Find $\lim_{t \rightarrow \infty} x(t)$ given any initial condition.

Exercise 1.6.103: Assume that a population of fish in a lake satisfies $\frac{dx}{dt} = kx(M - x)$. Now suppose that fish are continually added at A fish per unit of time.

- a) Find the differential equation for x .
- b) What is the new limiting population?

Exercise 1.6.104: Suppose $\frac{dx}{dt} = (x - \alpha)(x - \beta)$ for two numbers $\alpha < \beta$.

- a) Find the critical points, and classify them.

For b), c), d), find $\lim_{t \rightarrow \infty} x(t)$ based on the phase diagram.

- b) $x(0) < \alpha$,
- c) $\alpha < x(0) < \beta$,
- d) $\beta < x(0)$.

1.7 Numerical methods: Euler's method

Note: 1 lecture, can safely be skipped, §2.4 in [EP], §8.1 in [BD]

Unless $f(x, y)$ is of a special form, it is generally very hard if not impossible to get a nice formula for the solution of the problem

$$y' = f(x, y), \quad y(x_0) = y_0.$$

If the equation can be solved in closed form, we should do that. But what if we have an equation that cannot be solved in closed form? What if we want to find the value of the solution at some particular x ? Or perhaps we want to produce a graph of the solution to inspect the behavior. In this section we will learn about the basics of numerical approximation of solutions.

The simplest method for approximating a solution is *Euler's method**. It works as follows: Take x_0 and y_0 and compute the slope $k = f(x_0, y_0)$. The slope is the change in y per unit change in x . Follow the line for an interval of length h on the x -axis. We call h the *step size*. Hence if $y = y_0$ at x_0 , then we say that y_1 , the approximate value of y at $x_1 = x_0 + h$, is $y_1 = y_0 + hk$. Rinse, repeat! Let $k = f(x_1, y_1)$, and then compute $x_2 = x_1 + h$ and $y_2 = y_1 + hk$. Now compute x_3 and y_3 using x_2 and y_2 , etc. Consider the equation $y' = y^2/3$, $y(0) = 1$, and $h = 1$. Then $x_0 = 0$ and $y_0 = 1$. We compute

$$x_1 = x_0 + h = 0 + 1 = 1, \quad y_1 = y_0 + h f(x_0, y_0) = 1 + 1 \cdot 1/3 = 4/3 \approx 1.333,$$

$$x_2 = x_1 + h = 1 + 1 = 2, \quad y_2 = y_1 + h f(x_1, y_1) = 4/3 + 1 \cdot \frac{(4/3)^2}{3} = 52/27 \approx 1.926.$$

We draw an approximate graph of the solution by connecting the points (x_0, y_0) , (x_1, y_1) , (x_2, y_2) , $\dots$. See Figure 1.16 on the following page for the first two steps of the method.

We can continue along: For any $i = 0, 1, 2, 3, \dots$, we compute

$$x_{i+1} = x_i + h, \quad y_{i+1} = y_i + h f(x_i, y_i).$$

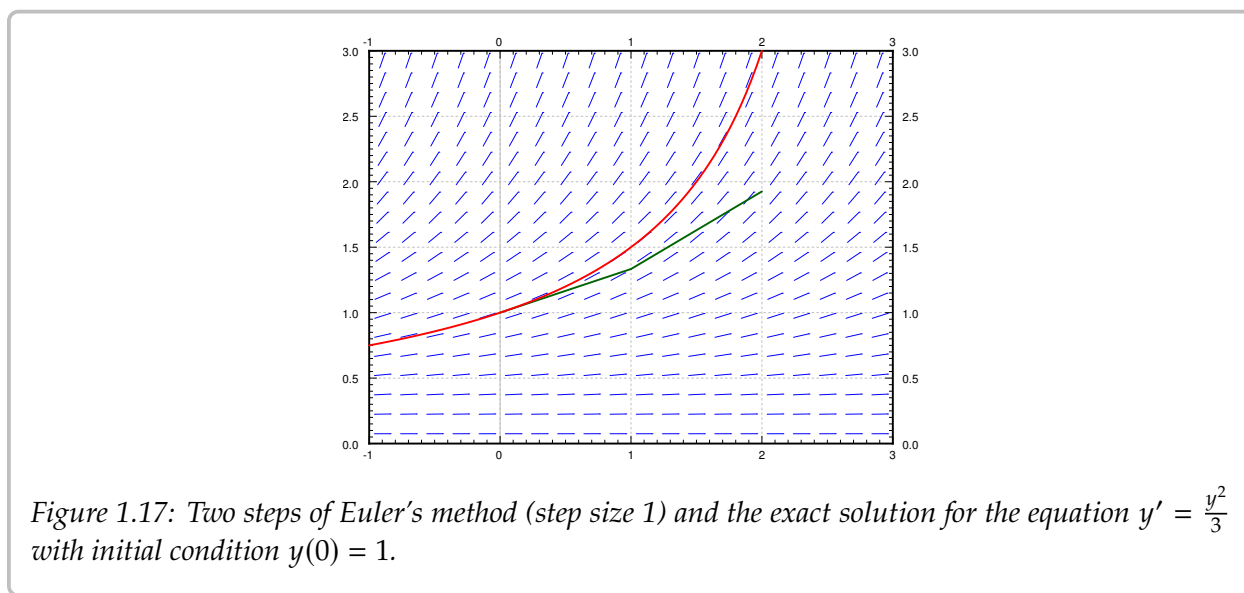
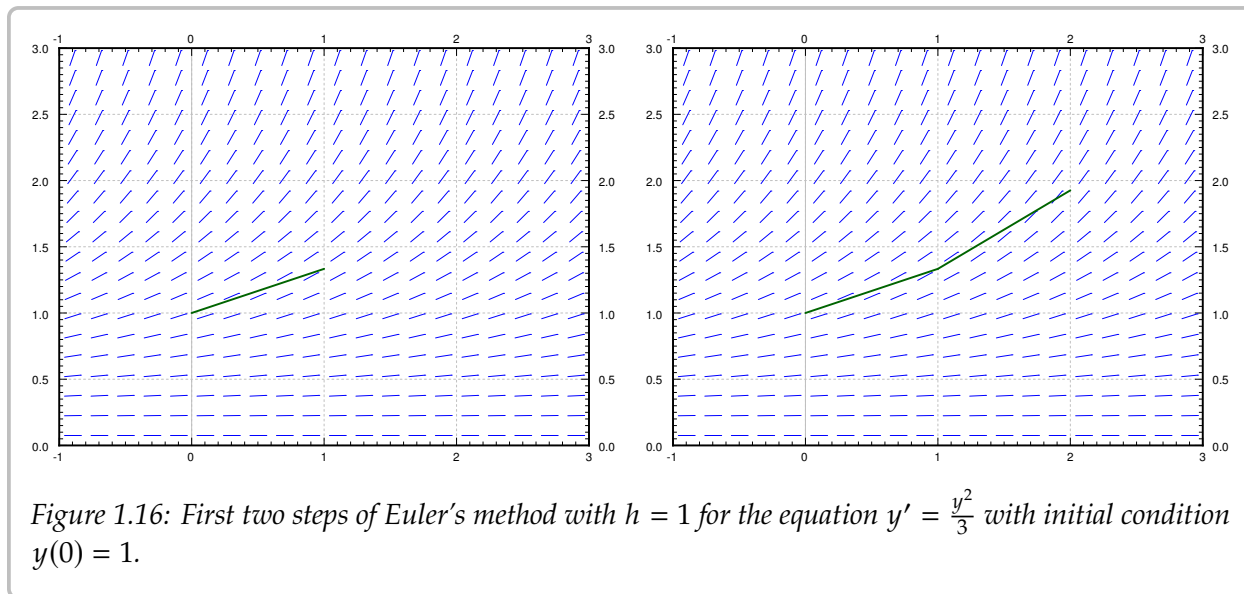
The line segments we get are an approximate graph of the solution. In particular, y_i is the approximation of $y(x_i)$. Generally it is not exactly the solution. See Figure 1.17 on the next page for the plot of the real solution and the approximation.

We continue with the example $y' = y^2/3$, $y(0) = 1$. We have approximated $y(2)$ with Euler's method with step size 1. We conclude that $y(2)$ is approximately y_2 , that is, $y(2) \approx 1.926$. The real answer is 3. We are off by roughly 1.074. The difference between the actual solution and the approximate solution is called the error. We usually talk about the size of the error and we do not care much about its sign.

$$\text{Error} = |\text{Actual } y - \text{Approximate } y|.$$

We usually do not know the error exactly. If we knew the error exactly, we would know the actual solution $\dots$ so what is the point of doing the approximation?

*Named after the Swiss mathematician [Leonhard Paul Euler](#) (1707–1783). The correct pronunciation of the name sounds more like “oiler.”



Exercise 1.7.1: Solve the problem $y' = y^2/3$, $y(0) = 1$ exactly and show that $y(2) = 3$.

Let us now halve the step size to $h = 0.5$ to, hopefully, improve the approximation. With half the step size, we do twice as many steps. Here, to find $y(2)$, we do 4 steps with $h = 0.5$ as then $2 = x_4$. Computing y_4 , we find $y(2) \approx 2.209$. An error of about 0.791. [Table 1.1](#) on the facing page gives the values computed, for a few more halvings.

Note that except for the first few times, each time we halve the h , the error approximately halves. Halving of the error is a general feature of Euler's method as it is a *first-order method*. A simple improvement of the Euler method (see the exercises) produces a second-order method. A second-order method reduces the error to approximately one quarter every time we halve the step size. The order being “second” means the squaring in $1/4 = 1/2 \times 1/2 = (1/2)^2$.

h	Approximate $y(2)$	Error	$\frac{\text{Error}}{\text{Previous error}}$	Steps done
1	1.92593	1.07407		2
0.5	2.20861	0.79139	0.73681	4
0.25	2.47250	0.52751	0.66656	8
0.125	2.68034	0.31966	0.60599	16
0.0625	2.82040	0.17960	0.56184	32
0.03125	2.90412	0.09588	0.53385	64
0.015625	2.95035	0.04965	0.51779	128
0.0078125	2.97472	0.02528	0.50913	256

Table 1.1: Euler's method approximation of $y(2)$ for $y' = y^2/3$, $y(0) = 1$.

To get the error to be less than 0.1, we did 64 steps. To get it below 0.01, we would have to halve another three or four times, doing 512 or 1024 steps. The improved Euler method from the exercises should quarter the error every time we halve the step size, so we would have to do (approximately) half as many “halvings” to get the same error. This reduction can be a big deal. With 10 halvings (starting at $h = 1$) we have 1024 steps, whereas with 5 halvings we only have to do 32 steps, assuming that the error was comparable to start with. A computer may not care about this difference for a problem this simple, but suppose each step would take a second to compute (the function may be substantially more difficult to compute than $y^2/3$). Then the difference is 32 seconds versus about 17 minutes. We are not being altogether fair; a second-order method would probably double the time to do each step. Even so, it is 1 minute versus 17 minutes. Next, suppose that we have to repeat such a calculation for different parameters a thousand times. You get the idea.

In practice, we do not know how large the error is! How do we know what is the right step size? One possibility is that we keep halving the step size, and if we are lucky, we can estimate the error from a few of these calculations and the assumption that the error halves each time (if we are using standard Euler).

Exercise 1.7.2: In the table above, suppose you do not know the error. Take the approximate values of the function in the last two lines, assume that the error halves. Can you estimate the error in the last line from this information? Does it (approximately) agree with the table? Now do it for the first two rows. Does this agree with the table? Feel free to assume that the real value of $y(2)$ is bigger than any of the approximate values, but you are not allowed to know the real value.

Let us talk a little bit more about the example $y' = y^2/3$, $y(0) = 1$. Suppose that instead of $y(2)$, we wish to find $y(3)$. Table 1.2 on the next page lists the results of this effort for successive halvings of h . What is going on here? Well, you should solve the equation exactly and you will notice that the solution does not exist at $x = 3$. In fact, the solution goes to infinity when you approach $x = 3$.

h	Approximate $y(3)$
1	3.16232
0.5	4.54329
0.25	6.86079
0.125	10.80321
0.0625	17.59893
0.03125	29.46004
0.015625	50.40121
0.0078125	87.75769

Table 1.2: Attempts to use Euler's to approximate $y(3)$ for $y' = y^2/3$, $y(0) = 1$.

Another case where things go bad is if the solution oscillates wildly near some point. The solution may exist at all points, but even a much better numerical method than Euler would need an insanely small step size to approximate the solution with reasonable precision. And computers might not be able to easily handle such a small step size.

In real applications, we would not use a simple method such as Euler's. The simplest method one would probably use in a real application is the standard Runge–Kutta method (see exercises). That is a fourth-order method, meaning that if we halve the step size, the error gets multiplied by $1/16$ (it is fourth-order as $1/16 = 1/2 \times 1/2 \times 1/2 \times 1/2$).

Choosing the right method to use and the right step size can be very tricky. There are several competing factors to consider.

- Computational time: Each step takes computer time. Even if the function f is simple to compute, we do it many times over. Large step size (fewer steps) means faster computation, but perhaps not the right precision.
- Roundoff errors: Computers only compute with a certain number of significant digits. Errors introduced by rounding numbers during computations become noticeable when the step size becomes too small relative to the quantities we are working with. So reducing step size too much may make errors worse. There is a certain optimum step size where the precision increases as we approach it, but then starts getting worse as we make our step size smaller still. The trouble is that this optimum may be hard to find.
- Stability: Certain equations may be numerically unstable. What may happen is that the numbers never seem to stabilize no matter how many times we halve the step size. We may need a ridiculously small step size, which may not be practical due to roundoff errors or computational time considerations. Such problems are sometimes called *stiff*. In the worst case, the numerical computations might be giving us bogus

numbers that look like a correct answer. Just because the numbers seem to have stabilized after successive halving, does not mean that we must have the right answer.

We have seen just the beginnings of the challenges that appear in real applications. Numerical approximation of solutions to differential equations is an active research area for engineers and mathematicians. For example, the general-purpose method used for the ODE solver in Matlab and Octave (as of this writing) is a method that appeared in the literature only in the 1980s.

1.7.1 Exercises

Exercise 1.7.3: Consider $\frac{dx}{dt} = (2t - x)^2$, $x(0) = 2$. Use Euler's method with step size $h = 0.5$ to approximate $x(1)$.

Exercise 1.7.4: Consider $\frac{dx}{dt} = t - x$, $x(0) = 1$.

- Use Euler's method with step sizes $h = 1, 1/2, 1/4, 1/8$ to approximate $x(1)$.
- Solve the equation exactly.
- Describe what happens to the errors for each h you used. That is, find the factor by which the error changed each time you halved the step size.

Exercise 1.7.5: Approximate the value of e by looking at the initial value problem $y' = y$ with $y(0) = 1$ and approximating $y(1)$ using Euler's method with a step size of 0.2.

Exercise 1.7.6: Example of numerical instability: Take $y' = -5y$, $y(0) = 1$. We know that the solution should decay to zero as x grows. Using Euler's method, start with $h = 1$ and compute y_1, y_2, y_3, y_4 to try to approximate $y(4)$. What happened? Now halve the step size. Keep halving the step size and approximating $y(4)$ until the numbers you are getting start to stabilize (that is, until they start going towards zero). Note: You might want to use a calculator.

The simplest method used in practice is the *Runge-Kutta method*. Consider $\frac{dy}{dx} = f(x, y)$, $y(x_0) = y_0$, and a step size h . Everything is the same as in Euler's method, except the computation of y_{i+1} and x_{i+1} . That is, in each step we compute slopes k_1, k_2, k_3 , and k_4 , and then we compute the next x_{i+1} and y_{i+1} :

$$\begin{aligned} k_1 &= f(x_i, y_i), & k_2 &= f(x_i + h/2, y_i + k_1(h/2)), \\ k_3 &= f(x_i + h/2, y_i + k_2(h/2)), & k_4 &= f(x_i + h, y_i + k_3h), \\ x_{i+1} &= x_i + h, & y_{i+1} &= y_i + \frac{k_1 + 2k_2 + 2k_3 + k_4}{6} h. \end{aligned}$$

Exercise 1.7.7: Consider $\frac{dy}{dx} = yx^2$, $y(0) = 1$.

- Approximate $y(1)$ using Runge–Kutta (see above) with step sizes $h = 1$ and $h = 1/2$.
- Approximate $y(1)$ using Euler's method with $h = 1$ and $h = 1/2$.
- Solve exactly, find the exact value of $y(1)$, and compare with the approximations.

Exercise 1.7.101: Let $x' = \sin(xt)$, and $x(0) = 1$. Approximate $x(1)$ using Euler's method with step sizes 1, 0.5, 0.25. Use a calculator and compute up to 4 decimal digits.

Exercise 1.7.102: Let $x' = 2t$, and $x(0) = 0$.

- Approximate $x(4)$ using Euler's method with step sizes 4, 2, and 1.
- Solve exactly, and compute the errors.
- Compute the factor by which the errors changed.

Exercise 1.7.103: Let $x' = xe^{x+1}$, and $x(0) = 0$.

- Approximate $x(4)$ using Euler's method with step sizes 4, 2, and 1.
- Guess an exact solution based on part a) and compute the errors.

There is a simple way to improve Euler's method to make it a second-order method by doing just one extra step. Consider $\frac{dy}{dx} = f(x, y)$, $y(x_0) = y_0$, and a step size h . What we do is to pretend we compute the next step as in Euler, that is, we start with (x_i, y_i) , we compute a slope $k_1 = f(x_i, y_i)$, and then look at the point $(x_i + h, y_i + k_1h)$. Instead of letting our new point be $(x_i + h, y_i + k_1h)$, we compute the slope at that point, call it k_2 , and then take the average of k_1 and k_2 , hoping that the average is going to be closer to the actual slope on the interval from x_i to $x_i + h$. And we are correct: If we halve the step, the error should get multiplied by $(1/2)^2 = 1/4$. To summarize, the setup is the same as for regular Euler, except the computation of y_{i+1} and x_{i+1} . At each step, we compute the new slopes k_1 and k_2 and then the next y_{i+1} and x_{i+1} :

$$\begin{aligned} k_1 &= f(x_i, y_i), & k_2 &= f(x_i + h, y_i + k_1h), \\ x_{i+1} &= x_i + h, & y_{i+1} &= y_i + \frac{k_1 + k_2}{2} h. \end{aligned}$$

Exercise 1.7.104: Consider $\frac{dy}{dx} = x + y$, $y(0) = 1$.

- Approximate $y(1)$ using the improved Euler's method (see above) with step sizes $h = 1/4$ and $h = 1/8$.
- Approximate $y(1)$ using Euler's method with $h = 1/4$ and $h = 1/8$.
- Solve exactly, find the exact value of $y(1)$.
- Compute the errors, and the factors by which the errors changed.

1.8 Exact equations

Note: 1–2 lectures, can safely be skipped, §1.6 in [EP], §2.6 in [BD]

A type of equation that comes up quite often in physics and engineering is an *exact equation*. Suppose $F(x, y)$ is a function of two variables, which we call the *potential function*. The naming should suggest potential energy, or electric potential. Exact equations and potential functions appear when there is a conservation law at play, such as conservation of energy. Let us make up a simple example. Consider

$$F(x, y) = x^2 + y^2.$$

We are interested in the lines of constant potential*: curves given by $F(x, y) = C$, for some constant C . In our example, the curves $x^2 + y^2 = C$ are circles. See Figure 1.18.

We take the *total derivative* of F :

$$dF = \frac{\partial F}{\partial x} dx + \frac{\partial F}{\partial y} dy.$$

For convenience, we will use the notation $F_x = \frac{\partial F}{\partial x}$ and $F_y = \frac{\partial F}{\partial y}$. In our example,

$$dF = 2x dx + 2y dy.$$

If we apply the total derivative to the equation $F(x, y) = C$, we find the differential equation $dF = 0$. This differential equation has the form

$$M dx + N dy = 0, \quad \text{or} \quad M + N \frac{dy}{dx} = 0.$$

An equation of this form is called *exact* if it was obtained as $dF = 0$ for some potential function F . In our simple example, we obtain the equation

$$2x dx + 2y dy = 0, \quad \text{or} \quad 2x + 2y \frac{dy}{dx} = 0.$$

Since we obtained this equation by differentiating $x^2 + y^2 = C$, the equation is exact. We often wish to solve for y in terms of x . In our example,

$$y = \pm \sqrt{C - x^2}.$$

In terms of multivariable calculus, at each point (x, y) in the plane, $\vec{v} = (M, N)$ is a vector, that is, a direction and a magnitude. As M and N are functions of (x, y) , we have a

*Accordingly, these curves are sometimes called *equipotential curves*.

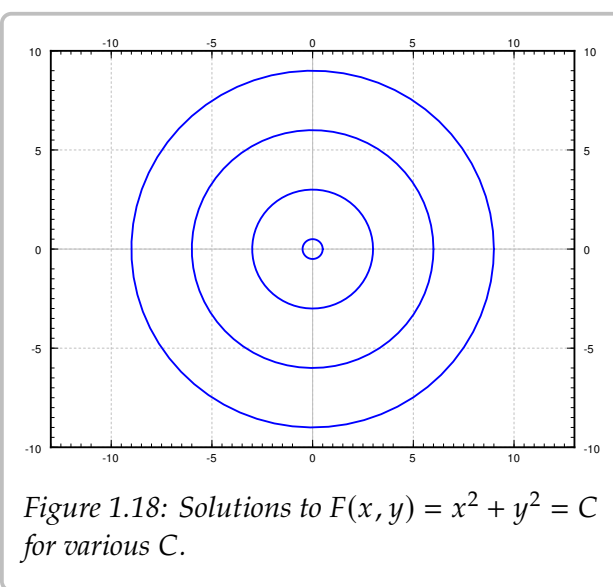


Figure 1.18: Solutions to $F(x, y) = x^2 + y^2 = C$ for various C .

vector field. A vector field $\vec{v}$ that comes from an exact equation is a so-called *conservative vector field*, that is, a vector field that comes with a potential function $F(x, y)$, such that

$$\vec{v} = \left(\frac{\partial F}{\partial x}, \frac{\partial F}{\partial y} \right).$$

Let γ be a path in the plane starting at (x_1, y_1) and ending at (x_2, y_2) . If we think of $\vec{v}$ as force, then the work required to move along γ is the path integral

$$\int_{\gamma} \vec{v}(\vec{r}) \cdot d\vec{r} = \int_{\gamma} M dx + N dy = F(x_2, y_2) - F(x_1, y_1).$$

In other words, the work done depends only on the endpoints—where we start and where we end. For example, suppose F is gravitational potential (physicists use the convention that $-F$ is the potential, but the idea is the same). The derivative of F given by $\vec{v}$ is the gravitational force. What we are saying is that the work required to move a heavy box from the ground floor to the roof only depends on the change in potential energy. The work done is the same no matter what path we took; if we took the stairs or the elevator. Although if we took the elevator, the elevator did the work for us. Movement along the curves $F(x, y) = C$ requires no work done, such as the heavy box sliding along without accelerating or braking on a perfectly flat roof on a cart with incredibly well-oiled wheels.

An exact equation is a conservative vector field, and the implicit solution of this equation is the potential function.

1.8.1 Solving exact equations

Now you, the reader, should ask: Where did we solve a differential equation? Well, in applications we generally know M and N , but we do not know F . That is, we may have just started with $2x + 2y \frac{dy}{dx} = 0$, or perhaps even

$$x + y \frac{dy}{dx} = 0.$$

It is up to us to find some potential F that works. Many different F will work; adding a constant to F does not change the equation. Once we have a potential function F , the equation $F(x, y(x)) = C$ gives an implicit solution of the ODE.

Example 1.8.1: Let us find the general solution to $2x + 2y \frac{dy}{dx} = 0$. Forget we know F .

If we know that this is an exact equation, we start looking for a potential function F . We have $M = 2x$ and $N = 2y$. If F exists, it must be such that $F_x(x, y) = 2x$. Integrate in the x variable to find

$$F(x, y) = x^2 + A(y), \tag{1.5}$$

for some function $A(y)$. The function A is the “constant of integration,” though it is only constant as far as x is concerned, and may still depend on y . Now differentiate (1.5) in y and set it equal to N , which is what F_y is supposed to be:

$$2y = F_y(x, y) = A'(y).$$

Integrating, we find $A(y) = y^2$. We could add a constant of integration if we wanted to, but there is no need. We found $F(x, y) = x^2 + y^2$. Next for a constant C , we solve

$$F(x, y(x)) = C$$

for y in terms of x . In this case, we obtain $y = \pm\sqrt{C - x^2}$ as we did before.

Exercise 1.8.1: Why did we not need to add a constant of integration when integrating $A'(y) = 2y$? Add a constant of integration, say 3, and see what F you get. What is the difference from what we got above, and why does it not matter?

The procedure, once we know that the equation is exact, is:

- (i) Integrate $F_x = M$ in x resulting in $F(x, y) = \text{something} + A(y)$.
- (ii) Differentiate this F in y and set that equal to N . Then find $A(y)$ by integration.

The procedure can also be done by first integrating in y and then differentiating in x . Pretty easy, huh? Let's try this again.

Example 1.8.2: Consider now $2x + y + xy \frac{dy}{dx} = 0$.

OK, so $M = 2x + y$ and $N = xy$. We try to proceed as before. Suppose F exists. Then $F_x(x, y) = 2x + y$. We integrate:

$$F(x, y) = x^2 + xy + A(y)$$

for some function $A(y)$. Differentiate in y and set equal to N :

$$N = xy = F_y(x, y) = x + A'(y).$$

But there is no way to satisfy this requirement! The function xy cannot be written as x plus a function of y . The equation is not exact; no potential function F exists.

Is there an easier way to check for exactness, other than failing in trying to find F ? Yes there is. Suppose $M = F_x$ and $N = F_y$. As long as the second derivatives are continuous,

$$\frac{\partial M}{\partial y} = \frac{\partial^2 F}{\partial y \partial x} = \frac{\partial^2 F}{\partial x \partial y} = \frac{\partial N}{\partial x}.$$

If $M_y \neq N_x$, the equation is not exact. For example, in [Example 1.8.2](#), where $M = 2x + y$ and $N = xy$, the equation is not exact. Indeed, $M_y = 1$ and $N_x = y$ are not equal.

What is interesting is that this idea works in reverse too. If $M_y = N_x$, then a potential exists, at least near an initial point. This result is called the Poincaré Lemma*.

Theorem 1.8.1 (Poincaré Lemma). If M and N are continuously differentiable functions of (x, y) , and $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$, then near any point there is a function $F(x, y)$ such that $M = \frac{\partial F}{\partial x}$ and $N = \frac{\partial F}{\partial y}$.

*Named for the French polymath [Jules Henri Poincaré](#) (1854–1912).

Example 1.8.3: Solve

$$\frac{dy}{dx} = \frac{-2x - y}{x - 1}, \quad y(0) = 1.$$

Write the equation as

$$(2x + y) + (x - 1)\frac{dy}{dx} = 0,$$

so $M = 2x + y$, $N = x - 1$, and

$$M_y = 1 = N_x.$$

The Poincaré Lemma applies, so let us look for F . Integrating M in x , we find

$$F(x, y) = x^2 + xy + A(y).$$

Differentiating in y and setting to N , we find

$$x - 1 = x + A'(y).$$

So $A'(y) = -1$, and $A(y) = -y$ will work. We obtain $F(x, y) = x^2 + xy - y$, so the implicit solution is $x^2 + xy - y = C$. First we find C . As $y(0) = 1$, we have $F(0, 1) = C$. Therefore, $0^2 + 0 \times 1 - 1 = C$, so $C = -1$. Now we solve $x^2 + xy - y = -1$ for y to get

$$y = \frac{-x^2 - 1}{x - 1}.$$

Example 1.8.4: Solve

$$\frac{-y}{x^2 + y^2} dx + \frac{x}{x^2 + y^2} dy = 0, \quad y(1) = 2.$$

We leave to the reader to check that $M_y = N_x$, so the Poincaré Lemma applies.

This vector field (M, N) is not conservative (the equation is not exact) if considered as a vector field of the entire plane minus the origin. The problem is that if the curve γ is a circle around the origin, say starting at $(1, 0)$ and ending at $(1, 0)$ going counterclockwise, then if F existed, we would expect

$$0 = F(1, 0) - F(1, 0) = \int_{\gamma} F_x dx + F_y dy = \int_{\gamma} \frac{-y}{x^2 + y^2} dx + \frac{x}{x^2 + y^2} dy = 2\pi.$$

That is nonsense! We leave the computation of the path integral to the interested reader, or you can consult your multivariable calculus textbook. So there is no potential function F defined everywhere outside the origin $(0, 0)$.

The Poincaré Lemma does not guarantee such an F anyway. It only guarantees a potential function F locally—only in some region containing the initial point. As $y(1) = 2$, we start at the point $(1, 2)$. Considering $x > 0$ and integrating M in x or N in y , we find

$$F(x, y) = \arctan(y/x).$$

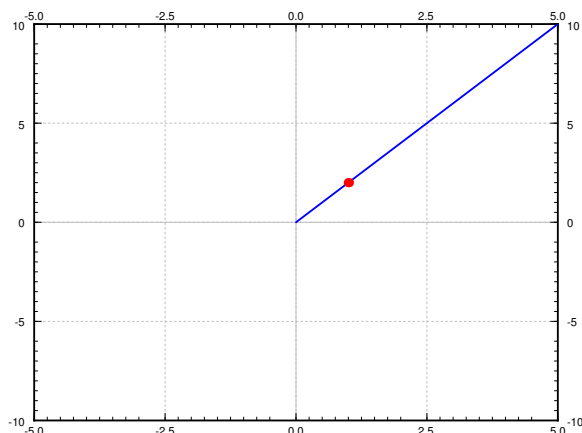


Figure 1.19: Solution to $\frac{-y}{x^2+y^2}dx + \frac{x}{x^2+y^2}dy = 0$, $y(1) = 2$, with initial point marked.

This F is defined on the region $x > 0$ and so the equation is exact on that region. The implicit solution is $\arctan(y/x) = C$. Solving, $y = \tan(C)x$. That is, the solution is a straight line. Solving $y(1) = 2$ gives us that $\tan(C) = 2$, and so $y = 2x$ is the desired solution. See Figure 1.19, and note again that the solution only exists for $x > 0$.

If the region where M and N are defined and where $M_y = N_x$ has “no holes,” such as, for instance, the entire xy -plane, then the equation is exact and an F defined on that region can be found. For example, in Example 1.8.3, $M = 2x + y$ and $N = x - 1$ are defined in the entire plane and $M_y = 1 = N_x$, so a global F exists there. Similarly, in Example 1.8.4 the half-plane where $x > 0$ also has “no holes” and a potential function exists on that region. So in such cases, checking $M_y = N_x$ is good enough to check for exactness.

Example 1.8.5: Solve

$$x^2 + y^2 + 2y(x + 1)\frac{dy}{dx} = 0.$$

The reader should check that this equation is exact. Let $M = x^2 + y^2$ and $N = 2y(x + 1)$. We follow the procedure for exact equations

$$F(x, y) = \frac{1}{3}x^3 + xy^2 + A(y),$$

and

$$2y(x + 1) = 2xy + A'(y).$$

Therefore $A'(y) = 2y$ or $A(y) = y^2$ and $F(x, y) = \frac{1}{3}x^3 + xy^2 + y^2$. We try to solve $F(x, y) = C$. We easily solve for y^2 and then just take the square root:

$$y^2 = \frac{C - (1/3)x^3}{x + 1}, \quad \text{so} \quad y = \pm \sqrt{\frac{C - (1/3)x^3}{x + 1}}.$$

When $x = -1$, the term in front of $\frac{dy}{dx}$ is zero, and our explicit solution is not valid. The given equation has no solution (for $y(x)$) near $x = -1$, but the equation $(x^2 + y^2) dx + 2y(x+1) dy = 0$ does have a solution $x = -1$. In fact, one could solve for x in terms of y for any initial condition. The solution is messy, so we leave it as $\frac{1}{3}x^3 + xy^2 + y^2 = C$.

1.8.2 Integrating factors

Sometimes an equation $M dx + N dy = 0$ is not exact, but it can be made exact by multiplying with a function $u(x, y)$. That is, perhaps for some nonzero function $u(x, y)$,

$$u(x, y)M(x, y) dx + u(x, y)N(x, y) dy = 0$$

is exact. Any solution to this new equation is also a solution to $M dx + N dy = 0$.

In fact, a linear equation

$$\frac{dy}{dx} + p(x)y = f(x), \quad \text{or} \quad (p(x)y - f(x)) dx + dy = 0,$$

is always such an equation. Let $r(x) = e^{\int p(x) dx}$ be the integrating factor for the linear equation. Multiply the equation by $r(x)$ and write it in the form of $M + N \frac{dy}{dx} = 0$.

$$r(x)p(x)y - r(x)f(x) + r(x)\frac{dy}{dx} = 0.$$

Then $M = r(x)p(x)y - r(x)f(x)$, so $M_y = r(x)p(x)$, while $N = r(x)$, so $N_x = r'(x) = r(x)p(x)$. In other words, we have an exact equation. Integrating factors for linear functions are just a special case of integrating factors for exact equations.

But how do we find the integrating factor u ? Well, for $uM dx + uN dy = 0$ to be exact, the function u should satisfy

$$\frac{\partial}{\partial y} [uM] = u_y M + u M_y = \frac{\partial}{\partial x} [uN] = u_x N + u N_x.$$

Therefore,

$$(M_y - N_x)u = u_x N - u_y M.$$

At first it may seem we replaced one differential equation by another. Even worse, the new equation is a PDE. True, but all hope is not lost.

A strategy that often works is to look for a u that is a function of x alone, or a function of y alone. After all, that is what worked for linear equations. If u is a function of x alone, that is, $u(x)$, then we write $u'(x)$ instead of u_x , and u_y is just zero. Then

$$\frac{M_y - N_x}{N} u = u'.$$

In particular, $\frac{M_y - N_x}{N}$ ought to be a function of x alone (not depend on y). If so, then we have a linear equation

$$u' - \frac{M_y - N_x}{N}u = 0.$$

Letting $P(x) = \frac{M_y - N_x}{N}$, we solve the linear equation to find $u(x) = Ce^{\int P(x) dx}$. The constant in the solution is not relevant, we need any nonzero solution, so we take $C = 1$. Then $u(x) = e^{\int P(x) dx}$ is the integrating factor making the equation exact.

Similarly, we could try a function of the form $u(y)$. Then

$$\frac{M_y - N_x}{M}u = -u'.$$

In particular, $\frac{M_y - N_x}{M}$ ought to be a function of y alone. If so, we have a linear equation

$$u' + \frac{M_y - N_x}{M}u = 0.$$

Letting $Q(y) = \frac{M_y - N_x}{M}$, we find $u(y) = Ce^{-\int Q(y) dy}$. We take $C = 1$. So $u(y) = e^{-\int Q(y) dy}$ is the integrating factor.

Example 1.8.6: Solve

$$\frac{x^2 + y^2}{x + 1} + 2y \frac{dy}{dx} = 0.$$

Let $M = \frac{x^2 + y^2}{x + 1}$ and $N = 2y$. Compute

$$M_y - N_x = \frac{2y}{x + 1} - 0 = \frac{2y}{x + 1}.$$

As this is not zero, the equation is not exact. We notice

$$P(x) = \frac{M_y - N_x}{N} = \frac{2y}{x + 1} \frac{1}{2y} = \frac{1}{x + 1}$$

is a function of x alone. We compute the integrating factor

$$e^{\int P(x) dx} = e^{\ln(x+1)} = x + 1.$$

We multiply our given equation by $(x + 1)$ to obtain

$$x^2 + y^2 + 2y(x + 1) \frac{dy}{dx} = 0,$$

which is an exact equation that we solved in [Example 1.8.5](#). The solution was

$$y = \pm \sqrt{\frac{C - (1/3)x^3}{x + 1}}.$$

Example 1.8.7: Solve

$$y^2 + (xy + 1)\frac{dy}{dx} = 0.$$

First compute

$$M_y - N_x = 2y - y = y.$$

As this is not zero, the equation is not exact. We observe

$$Q(y) = \frac{M_y - N_x}{M} = \frac{y}{y^2} = \frac{1}{y}$$

is a function of y alone. We compute the integrating factor

$$e^{-\int Q(y) dy} = e^{-\ln y} = \frac{1}{y}.$$

Therefore, we look at the exact equation

$$y + \frac{xy + 1}{y} \frac{dy}{dx} = 0.$$

The reader should double check that this equation is exact. We follow the procedure for exact equations

$$F(x, y) = xy + A(y),$$

and

$$\frac{xy + 1}{y} = x + \frac{1}{y} = x + A'(y).$$

Consequently, $A'(y) = 1/y$ or $A(y) = \ln|y|$. Thus $F(x, y) = xy + \ln|y|$. It is not possible to solve $F(x, y) = C$ for y in terms of elementary functions, so let us be content with the implicit solution:

$$xy + \ln|y| = C.$$

We are looking for the general solution, and we divided by y above. We should check what happens when $y = 0$, as the equation itself makes perfect sense in that case. We plug in $y = 0$ to find the equation is satisfied. So $y = 0$ is also a solution.

1.8.3 Exercises

Exercise 1.8.2: Solve the following exact equations, implicit general solutions will suffice:

a) $(2xy + x^2) dx + (x^2 + y^2 + 1) dy = 0$

b) $x^5 + y^5 \frac{dy}{dx} = 0$

c) $e^x + y^3 + 3xy^2 \frac{dy}{dx} = 0$

d) $(x + y) \cos(x) + \sin(x) + \sin(x)y' = 0$

Exercise 1.8.3: Find the integrating factors for the following equations making them into exact equations:

a) $e^{xy} dx + \frac{y}{x} e^{xy} dy = 0$

b) $\frac{e^x + y^3}{y^2} dx + 3x dy = 0$

c) $4(y^2 + x) dx + \frac{2x + 2y^2}{y} dy = 0$

d) $2 \sin(y) dx + x \cos(y) dy = 0$

Exercise 1.8.4: Suppose you have an equation of the form: $f(x) + g(y) \frac{dy}{dx} = 0$.

a) Show it is exact.

b) Find the form of the potential function in terms of f and g .

Exercise 1.8.5: Suppose that we have the equation $f(x) dx - dy = 0$.

a) Is this equation exact?

b) Find the general solution using a definite integral.

Exercise 1.8.6: Find the potential function $F(x, y)$ of the exact equation $\frac{1+xy}{x} dx + \left(\frac{1}{y} + x\right) dy = 0$ in two different ways.

a) Integrate M in terms of x and then differentiate in y and set to N .

b) Integrate N in terms of y and then differentiate in x and set to M .

Exercise 1.8.7: A function $u(x, y)$ is said to be a harmonic function if $u_{xx} + u_{yy} = 0$.

a) Show that if u is harmonic, $-u_y dx + u_x dy = 0$ is an exact equation. So there exists (at least locally) the so-called harmonic conjugate function $v(x, y)$ such that $v_x = -u_y$ and $v_y = u_x$.

Verify that the following u are harmonic and find the corresponding harmonic conjugates v :

b) $u = 2xy$

c) $u = e^x \cos y$

d) $u = x^3 - 3xy^2$

Exercise 1.8.101: Solve the following exact equations, implicit general solutions will suffice:

a) $\cos(x) + ye^{xy} + xe^{xy} y' = 0$

b) $(2x + y) dx + (x - 4y) dy = 0$

c) $e^x + e^y \frac{dy}{dx} = 0$

d) $(3x^2 + 3y) dx + (3y^2 + 3x) dy = 0$

Exercise 1.8.102: Find the integrating factor for the following equations making them into exact equations:

a) $\frac{1}{y} dx + 3y dy = 0$

b) $dx - e^{-x-y} dy = 0$

c) $\left(\frac{\cos(x)}{y^2} + \frac{1}{y}\right) dx + \frac{x}{y^2} dy = 0$

d) $\left(2y + \frac{y^2}{x}\right) dx + (2y + x) dy = 0$

Exercise 1.8.103:

a) Show that every separable equation $y' = f(x)g(y)$ can be written as an exact equation, and verify that it is indeed exact.

b) Rewrite $y' = xy$ as an exact equation, solve it, and verify that the solution is the same as it was in [Example 1.3.1](#).

1.9 First-order linear PDEs

Note: 1 lecture, can safely be skipped

We have only considered ODEs so far, so let us solve a linear first-order PDE. Consider

$$a(x, t)u_x + b(x, t)u_t + c(x, t)u = g(x, t), \quad u(x, 0) = f(x), \quad -\infty < x < \infty, \quad t > 0,$$

where $u(x, t)$ is a function of x and t . We again use the notation $u_x = \frac{\partial u}{\partial x}$ and $u_t = \frac{\partial u}{\partial t}$ for convenience. The *initial condition* $u(x, 0) = f(x)$ is now a function of x rather than just a number. In these problems, it is useful to think of x as position and t as time. The equation describes the evolution of a function of x as time goes on. Below, the coefficients a, b, c , and the function g are mostly going to be constant or zero. The method we describe works with nonconstant coefficients, although the computations may get difficult quickly.

We will use the *method of characteristics*. The idea is to find lines along which the equation is an ODE that we then solve. We will see this technique again for second-order PDEs when we encounter the wave equation in § 4.8.

Example 1.9.1: Consider

$$\alpha u_x + u_t = 0, \quad u(x, 0) = f(x),$$

where α is a constant. This particular equation, $\alpha u_x + u_t = 0$, is called the *transport equation*.

The data will propagate along curves called characteristics. The idea is to change to the so-called *characteristic coordinates*, which we will call (ξ, s) . If we change to these coordinates, the equation simplifies. The change of variables for this equation is

$$\xi = x - \alpha t, \quad s = t.$$

Let us see what the equation becomes. Remember the chain rule in several variables.

$$\begin{aligned} u_x &= u_\xi \xi_x + u_s s_x = u_\xi, \\ u_t &= u_\xi \xi_t + u_s s_t = -\alpha u_\xi + u_s. \end{aligned}$$

The equation in the coordinates ξ and s becomes

$$\underbrace{\alpha (u_\xi)}_{u_x} + \underbrace{(-\alpha u_\xi + u_s)}_{u_t} = 0,$$

or in other words,

$$u_s = 0.$$

Treating ξ as simply a parameter, we have obtained the ODE $\frac{du}{ds} = 0$. That is trivial to solve.

The solution is a function that does not depend on s (but it does depend on ξ). That is, there is some function A such that

$$u = A(\xi) = A(x - \alpha t).$$

The initial condition says that

$$f(x) = u(x, 0) = A(x - \alpha 0) = A(x),$$

so $A = f$. In other words,

$$u(x, t) = f(x - \alpha t).$$

Everything is simply moving to the right* with velocity α as t increases. The curves given by the equation

$$\xi = \text{constant}$$

are called the *characteristic curves*. See Figure 1.20. In this case, the solution does not change along the characteristic as $\frac{du}{ds} = 0$.

In the (x, t) coordinates, the characteristic curves satisfy $t = \frac{1}{\alpha}(x - \xi)$, and are in fact lines. The slope of characteristic lines is $\frac{1}{\alpha}$, and for each different ξ , we get a different characteristic line.

We see why $\alpha u_x + u_t = 0$ is called the transport equation: Everything travels at some constant speed. This behavior is called *convection*. An example application is material being moved by a river where the material does not diffuse and is simply carried along. In this setup, x is the position along the river, t is the time, and $u(x, t)$ the concentration of the material at position x and time t . See Figure 1.21 for an example.

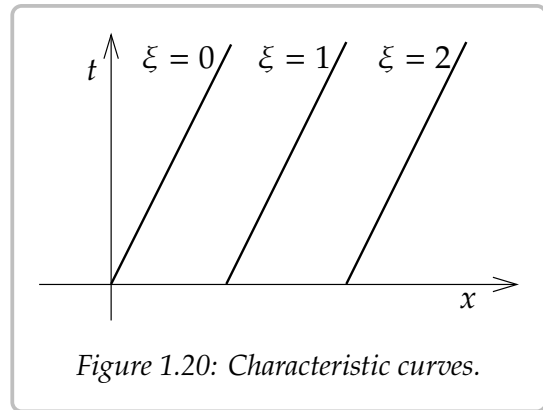


Figure 1.20: Characteristic curves.

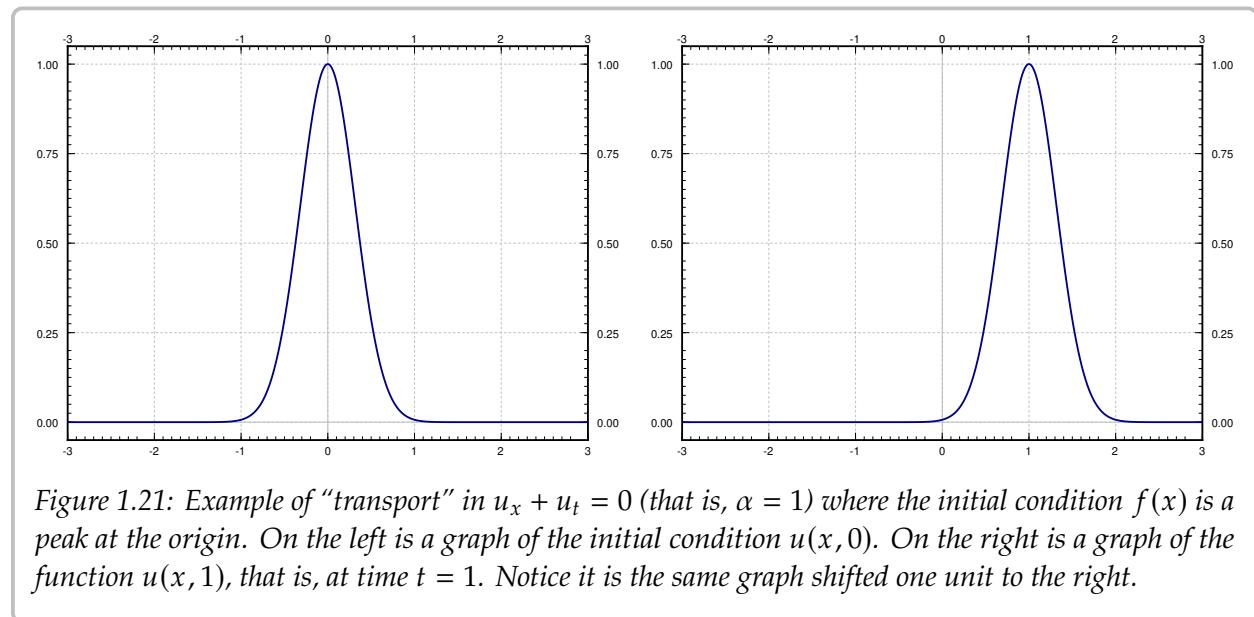


Figure 1.21: Example of "transport" in $u_x + u_t = 0$ (that is, $\alpha = 1$) where the initial condition $f(x)$ is a peak at the origin. On the left is a graph of the initial condition $u(x, 0)$. On the right is a graph of the function $u(x, 1)$, that is, at time $t = 1$. Notice it is the same graph shifted one unit to the right.

*That is assuming $\alpha > 0$. If $\alpha < 0$, then we are moving to the left.

We use a similar idea in the more general case:

$$au_x + bu_t + cu = g, \quad u(x, 0) = f(x).$$

We change coordinates to the characteristic coordinates, which we call (ξ, s) . These are coordinates where $au_x + bu_t$ becomes differentiation in the s variable.

Along the characteristic curves (where ξ is constant), we get a new ODE in the s variable. In the transport equation, we got the simple $\frac{du}{ds} = 0$. In general, we get the linear equation

$$\frac{du}{ds} + cu = g. \quad (1.6)$$

We think of everything as a function of ξ and s , although we are thinking of ξ as a parameter rather than an independent variable. So the equation is an ODE. It is a linear ODE that we can solve using the integrating factor.

To find the characteristics, think of a curve given parametrically $(x(s), t(s))$. We try to have the curve satisfy

$$\frac{dx}{ds} = a, \quad \frac{dt}{ds} = b.$$

Why? Because when we think of x and t as functions of s , we find, using the chain rule,

$$\frac{du}{ds} + cu = \underbrace{\left(u_x \frac{dx}{ds} + u_t \frac{dt}{ds} \right)}_{\frac{du}{ds}} + cu = au_x + bu_t + cu = g.$$

So we get the ODE (1.6), which then describes the value of the solution u of the PDE along this characteristic curve. It is convenient to make sure that $s = 0$ corresponds to $t = 0$, that is, $t(0) = 0$. It will also be convenient for $x(0) = \xi$. See Figure 1.22.

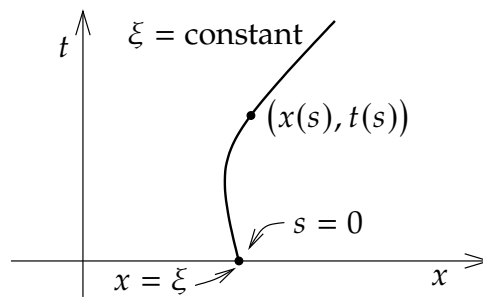


Figure 1.22: General characteristic curve.

Example 1.9.2: Consider

$$u_x + u_t + u = x, \quad u(x, 0) = e^{-x^2}.$$

We find the characteristics, that is, the curves given by

$$\frac{dx}{ds} = 1, \quad \frac{dt}{ds} = 1.$$

So

$$x = s + c_1, \quad t = s + c_2,$$

for some c_1 and c_2 . At $s = 0$, we want $x = \xi$ and $t = 0$. So we let $c_1 = \xi$ and $c_2 = 0$:

$$x = s + \xi, \quad t = s.$$

The ODE is $\frac{du}{ds} + u = x$, and $x = s + \xi$. So, the ODE to solve along the characteristic is

$$\frac{du}{ds} + u = s + \xi.$$

The general solution of this equation, treating ξ as a parameter, is $u = Ce^{-s} + s + \xi - 1$, for some C , which can depend on ξ . At $s = 0$, our initial condition is that u is $e^{-\xi^2}$, since at $s = 0$, we have $x = \xi$. Given this initial condition, we find $C = e^{-\xi^2} - \xi + 1$. So,

$$\begin{aligned} u &= (e^{-\xi^2} - \xi + 1)e^{-s} + s + \xi - 1 \\ &= e^{-\xi^2-s} + (1 - \xi)e^{-s} + s + \xi - 1. \end{aligned}$$

Substitute $\xi = x - t$ and $s = t$ to find u in terms of x and t :

$$\begin{aligned} u &= e^{-\xi^2-s} + (1 - \xi)e^{-s} + s + \xi - 1 \\ &= e^{-(x-t)^2-t} + (1 - x + t)e^{-t} + x - 1. \end{aligned}$$

See [Figure 1.23](#) on the next page for a plot of $u(x, t)$ as a function of two variables.

When the coefficients are not constants, the characteristic curves are not going to be straight lines anymore.

Example 1.9.3: Consider the following variable-coefficient equation:

$$xu_x + u_t + 2u = 0, \quad u(x, 0) = \cos(x).$$

We find the characteristics, that is, the curves given by

$$\frac{dx}{ds} = x, \quad \frac{dt}{ds} = 1.$$

So

$$x = c_1 e^s, \quad t = s + c_2.$$

At $s = 0$, we wish to get $x = \xi$ and $t = 0$ as before. So

$$x = \xi e^s, \quad t = s.$$

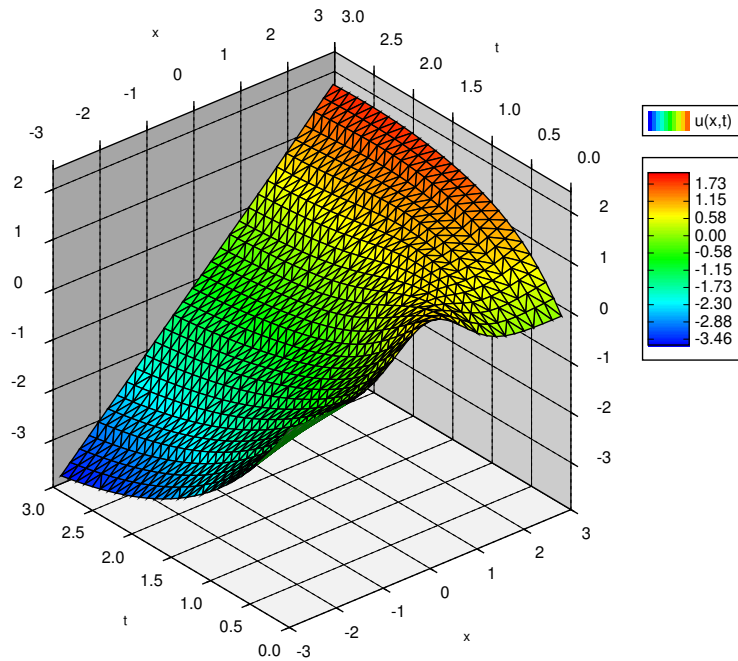


Figure 1.23: Plot of the solution $u(x, t)$ to $u_x + u_t + u = x$, $u(x, 0) = e^{-x^2}$.

OK, the ODE we need to solve is

$$\frac{du}{ds} + 2u = 0.$$

This is for a fixed ξ . We find $u = Ce^{-2s}$. At $s = 0$, we want u to be $\cos(\xi)$, so that is our initial condition for the ODE. Hence, $C = \cos(\xi)$. We solve for ξ and s to find $\xi = xe^{-t}$ and $s = t$. Consequently,

$$u = e^{-2s} \cos(\xi) = e^{-2t} \cos(xe^{-t}).$$

We make a few closing remarks. One thing to keep in mind is that we would get into trouble if the coefficient in front of u_t , that is, the b , is ever zero. Let us consider a quick example of what can go wrong:

$$u_x + u = 0, \quad u(x, 0) = \sin(x).$$

This problem has no solution. If we had a solution, it would imply that $u_x(x, 0) = \cos(x)$, but $u_x(x, 0) + u(x, 0) = \cos(x) + \sin(x) \neq 0$. The problem is that the characteristic curve is now the line $t = 0$, and the solution is already provided on that line! As b ought to then be nonzero, it is convenient to ensure that b is positive by multiplying the equation by -1 if necessary, so that a positive s means a positive t .

Another remark is that if a or b in the equation are not constants, the computations can quickly get out of hand, as the expressions for the characteristic coordinates become messy

and then solving the ODE becomes even messier. In the examples above, b was always 1, meaning we got $s = t$ in the characteristic coordinates. If b is not constant, your expression for s will be more complicated.

Finding the characteristic coordinates is really a system of ODEs in general if a depends on t or if b depends on x . In that case, we would need techniques of systems of ODEs to solve, see [chapter 3](#) or [chapter 8](#). In general, if a and b are not linear functions or constants, finding closed-form expressions for the characteristic coordinates may be impossible.

Finally, the method of characteristics applies to nonlinear first-order PDEs as well. In the nonlinear case, the characteristics depend not only on the differential equation, but also on the initial data. This leads to not only more difficult computations, but also the formation of singularities where the solution breaks down at a certain point in time. An example application where first-order nonlinear PDEs come up is traffic flow theory, and you have probably experienced the formation of singularities: traffic jams. But we digress.

1.9.1 Exercises

Exercise 1.9.1: Solve

- | | |
|--|--|
| a) $9u_x + u_t = 0, \quad u(x, 0) = \sin(x),$ | b) $-8u_x + u_t = 0, \quad u(x, 0) = \sin(x),$ |
| c) $\pi u_x + u_t = 0, \quad u(x, 0) = \sin(x),$ | d) $\pi u_x + u_t + u = 0, \quad u(x, 0) = \sin(x).$ |

Exercise 1.9.2: Solve $3u_x + u_t = 1, \quad u(x, 0) = x^2$.

Exercise 1.9.3: Solve $3u_x + u_t = x, \quad u(x, 0) = e^x$.

Exercise 1.9.4: Solve $u_x + u_t + xu = 0, \quad u(x, 0) = \cos(x)$.

Exercise 1.9.5:

- a) Find the characteristic coordinates for the following equations:
- | | |
|--|---|
| 1) $u_x + u_t + u = 1, \quad u(x, 0) = \cos(x),$ | 2) $2u_x + 2u_t + 2u = 2, \quad u(x, 0) = \cos(x).$ |
|--|---|
- b) Solve the two equations using the coordinates.
- c) Explain why you got the same solution, although the characteristic coordinates you found were different.

Exercise 1.9.6: Solve $x^2 u_x + (1 + x^2) u_t + e^x u = 0, \quad u(x, 0) = 0$. Hint: Think a little out of the box.

Exercise 1.9.101: Solve

- | | |
|--|---|
| a) $-5u_x + u_t = 0, \quad u(x, 0) = \frac{1}{1+x^2},$ | b) $2u_x + u_t = 0, \quad u(x, 0) = \cos(x).$ |
|--|---|

Exercise 1.9.102: Solve $u_x + u_t + tu = 0, \quad u(x, 0) = \cos(x)$.

Exercise 1.9.103: Solve $u_x + u_t = 5, \quad u(x, 0) = x$.

Chapter 2

Higher-order linear ODEs

2.1 Second-order linear ODEs

Note: 1 lecture, reduction of order optional, first part of §3.1 in [EP], parts of §3.1 and §3.2 in [BD]

Consider the general *second-order linear differential equation*

$$A(x)y'' + B(x)y' + C(x)y = F(x).$$

We often divide through by $A(x)$ to get

$$y'' + p(x)y' + q(x)y = f(x), \tag{2.1}$$

where $p(x) = B(x)/A(x)$, $q(x) = C(x)/A(x)$, and $f(x) = F(x)/A(x)$. The word *linear* means that the equation contains no powers or functions of y , y' , and y'' .

In the special case when $f(x) = 0$, we have a so-called *homogeneous* equation

$$y'' + p(x)y' + q(x)y = 0. \tag{2.2}$$

We have already seen some second-order linear homogeneous equations.

$y'' + k^2y = 0$	Two solutions are: $y_1 = \cos(kx), \quad y_2 = \sin(kx).$
$y'' - k^2y = 0$	Two solutions are: $y_1 = e^{kx}, \quad y_2 = e^{-kx}.$

If we know two solutions of a linear homogeneous equation, we know many more of them.

Theorem 2.1.1 (Superposition). *Suppose y_1 and y_2 are two solutions of the homogeneous equation (2.2). Then*

$$y(x) = C_1y_1(x) + C_2y_2(x),$$

also solves (2.2) for arbitrary constants C_1 and C_2 .

That is, we can add solutions together and multiply them by constants to obtain new and different solutions. We call the expression $C_1y_1 + C_2y_2$ a *linear combination* of y_1 and y_2 . Let us prove this theorem; the proof is very enlightening and illustrates how linear equations work.

Proof: Let $y = C_1y_1 + C_2y_2$. Then

$$\begin{aligned} y'' + py' + qy &= (C_1y_1 + C_2y_2)'' + p(C_1y_1 + C_2y_2)' + q(C_1y_1 + C_2y_2) \\ &= C_1y_1'' + C_2y_2'' + C_1py_1' + C_2py_2' + C_1qy_1 + C_2qy_2 \\ &= C_1(y_1'' + py_1' + qy_1) + C_2(y_2'' + py_2' + qy_2) \\ &= C_1 \cdot 0 + C_2 \cdot 0 = 0. \quad \square \end{aligned}$$

The proof becomes even simpler to state if we use the operator notation. An *operator* is an object that eats functions and spits out functions (kind of like what a function is, but a function eats numbers and spits out numbers). Define the operator L by

$$Ly = y'' + py' + qy.$$

The differential equation now becomes $Ly = 0$. The operator (and the equation) L being *linear* means that $L(C_1y_1 + C_2y_2) = C_1Ly_1 + C_2Ly_2$. It is almost as if we were “multiplying” by L . The proof above becomes

$$Ly = L(C_1y_1 + C_2y_2) = C_1Ly_1 + C_2Ly_2 = C_1 \cdot 0 + C_2 \cdot 0 = 0.$$

Two different solutions to the second equation $y'' - k^2y = 0$ are $y_1 = \cosh(kx)$ and $y_2 = \sinh(kx)$. Recalling the definition of $\sinh$ and $\cosh$, we note that these are solutions by superposition as they are linear combinations of the two exponential solutions: $\cosh(kx) = \frac{e^{kx} + e^{-kx}}{2} = (1/2)e^{kx} + (1/2)e^{-kx}$ and $\sinh(kx) = \frac{e^{kx} - e^{-kx}}{2} = (1/2)e^{kx} + (-1/2)e^{-kx}$.

The functions $\sinh$ and $\cosh$ are sometimes more convenient to use than the exponential. Let us review some of their properties:

$$\begin{aligned} \cosh 0 &= 1, & \sinh 0 &= 0, \\ \frac{d}{dx} [\cosh x] &= \sinh x, & \frac{d}{dx} [\sinh x] &= \cosh x, \\ \cosh^2 x - \sinh^2 x &= 1. \end{aligned}$$

Exercise 2.1.1: Derive these properties using the definitions of $\sinh$ and $\cosh$ in terms of exponentials.

Linear equations have nice and simple answers to the existence and uniqueness question.

Theorem 2.1.2 (Existence and uniqueness). Suppose p, q, f are continuous functions on some interval I , a is a number in I , and b_0, b_1 are constants. Then the equation

$$y'' + p(x)y' + q(x)y = f(x),$$

has exactly one solution $y(x)$ defined on the interval I satisfying the initial conditions

$$y(a) = b_0, \quad y'(a) = b_1.$$

For example, the equation $y'' + k^2y = 0$ with $y(0) = b_0$ and $y'(0) = b_1$ has the solution

$$y(x) = b_0 \cos(kx) + \frac{b_1}{k} \sin(kx).$$

The equation $y'' - k^2y = 0$ with $y(0) = b_0$ and $y'(0) = b_1$ has the solution

$$y(x) = b_0 \cosh(kx) + \frac{b_1}{k} \sinh(kx).$$

Using cosh and sinh in this solution allows us to solve for the initial conditions in a cleaner way than if we had used the exponentials.

The initial conditions for a second-order ODE consist of two equations. Common sense tells us that if we have two arbitrary constants and two equations, then we should be able to solve for the constants and find a solution to the differential equation satisfying the initial conditions.

Question: Suppose we find two different solutions y_1 and y_2 to the homogeneous equation (2.2). Can every solution be written (using superposition) in the form $y = C_1y_1 + C_2y_2$?

The answer is affirmative! Provided that y_1 and y_2 are different enough in the following sense. We say y_1 and y_2 are *linearly independent* if one is not a constant multiple of the other.

Theorem 2.1.3. *Let p, q be continuous functions. Let y_1 and y_2 be two linearly independent solutions to the homogeneous equation (2.2). Then every other solution is of the form*

$$y = C_1y_1 + C_2y_2.$$

That is, $y = C_1y_1 + C_2y_2$ is the general solution.

For example, we found the solutions $y_1 = \cos x$ and $y_2 = \sin x$ for the equation $y'' + y = 0$. It is not hard to see that sine and cosine are not constant multiples of each other. Indeed, if $\sin x = A \cos x$ for some constant A , plugging in $x = 0$ would imply $A = 0$. But then $\sin x = 0$ for all x , which is preposterous. So y_1 and y_2 are linearly independent. Hence,

$$y = C_1 \cos x + C_2 \sin x$$

is the general solution to $y'' + y = 0$.

For two functions, checking linear independence is rather simple. Here is another example. Consider $y'' - 2x^{-2}y = 0$. Then $y_1 = x^2$ and $y_2 = 1/x$ are solutions. To see that they are linearly independent, suppose one is a multiple of the other: $y_1 = Ay_2$, we just have to show that A cannot be a constant. In this case, we have $A = y_1/y_2 = x^3$, which is most decidedly not a constant. So $y = C_1x^2 + C_21/x$ is the general solution.

If you have one nonzero solution to a second-order linear homogeneous equation, then you can find another one. This is the *reduction of order method*. The idea is that if we

somehow found y_1 as a solution of $y'' + p(x)y' + q(x)y = 0$, then we try a second solution of the form $y_2(x) = y_1(x)v(x)$. We just need to find v . We plug y_2 into the equation:

$$\begin{aligned} 0 = y_2'' + p(x)y_2' + q(x)y_2 &= \underbrace{y_1''v + 2y_1'y' + y_1v''}_{y_2''} + \underbrace{p(x)(y_1'v + y_1v')}_{y_2'} + \underbrace{q(x)y_1v}_{y_2} \\ &= y_1v'' + (2y_1' + p(x)y_1)v' + \underbrace{(y_1'' + p(x)y_1' + q(x)y_1)}_{\rightarrow 0}v. \end{aligned}$$

In other words, $y_1v'' + (2y_1' + p(x)y_1)v' = 0$. Using $w = v'$, we have the first-order linear equation $y_1w' + (2y_1' + p(x)y_1)w = 0$. After solving this equation for w (integrating factor), we find v by antidifferentiating w . We then form y_2 by computing y_1v . For example, suppose we somehow know $y_1 = x$ is a solution to $y'' + x^{-1}y' - x^{-2}y = 0$. The equation for w is then $xw' + 3w = 0$. We find a solution, $w = Cx^{-3}$, and we find an antiderivative $v = \frac{-C}{2x^2}$. Hence $y_2 = y_1v = \frac{-C}{2x}$. Any C works and so $C = -2$ makes $y_2 = 1/x$. Thus, the general solution is $y = C_1x + C_21/x$.

Since we have a formula for the solution to the first-order linear equation, we can write a formula for y_2 :

$$y_2(x) = y_1(x) \int \frac{e^{-\int p(x) dx}}{(y_1(x))^2} dx$$

However, it is much easier to remember that we just need to try $y_2(x) = y_1(x)v(x)$ and find $v(x)$ as we did above. The technique works for higher-order equations too: You get to reduce the order by one for each solution you find. So it is better to remember how to do it rather than a specific formula.

We will study the solution of nonhomogeneous equations in § 2.5. We will first focus on finding general solutions to homogeneous equations.

2.1.1 Exercises

Exercise 2.1.2: Show that $y = e^x$ and $y = e^{2x}$ are linearly independent.

Exercise 2.1.3: Take $y'' + 5y = 10x + 5$. Find (guess!) a solution.

Exercise 2.1.4: Prove the superposition principle for nonhomogeneous equations. Suppose that y_1 is a solution to $Ly_1 = f(x)$ and y_2 is a solution to $Ly_2 = g(x)$ (same linear operator L). Show that $y = y_1 + y_2$ solves $Ly = f(x) + g(x)$.

Exercise 2.1.5: For the equation $x^2y'' - xy' = 0$, find two solutions, show that they are linearly independent and find the general solution. Hint: Try $y = x^r$.

Equations of the form $ax^2y'' + bxy' + cy = 0$ are called *Euler's equations* or *Cauchy–Euler equations*. They are solved by trying $y = x^r$ and solving for r (assume $x \geq 0$ for simplicity).

Exercise 2.1.6: Suppose that $(b - a)^2 - 4ac > 0$.

a) Find a formula for the general solution of Euler's equation (see above) $ax^2y'' + bxy' + cy = 0$.
Hint: Try $y = x^r$ and find a formula for r .

b) What happens when $(b - a)^2 - 4ac = 0$ or $(b - a)^2 - 4ac < 0$?

We will revisit the case when $(b - a)^2 - 4ac < 0$ later.

Exercise 2.1.7: Same equation as in [Exercise 2.1.6](#). Suppose $(b - a)^2 - 4ac = 0$. Find a formula for the general solution of $ax^2y'' + bxy' + cy = 0$. Hint: Try $y = x^r \ln x$ for the second solution.

Exercise 2.1.8 (reduction of order): Suppose y_1 is a solution to $y'' + p(x)y' + q(x)y = 0$. By directly plugging into the equation, show that

$$y_2(x) = y_1(x) \int \frac{e^{-\int p(x) dx}}{(y_1(x))^2} dx$$

is also a solution.

Exercise 2.1.9 (Chebyshev's equation of order 1): Take $(1 - x^2)y'' - xy' + y = 0$.

- Show that $y = x$ is a solution.
- Use reduction of order to find a second linearly independent solution.
- Write down the general solution.

Exercise 2.1.10 (Hermite's equation of order 2): Take $y'' - 2xy' + 4y = 0$.

- Show that $y = 1 - 2x^2$ is a solution.
- Use reduction of order to find a second linearly independent solution. (It's OK to leave a definite integral in the formula.)
- Write down the general solution.

Exercise 2.1.101: Are $\sin(x)$ and e^x linearly independent? Justify.

Exercise 2.1.102: Are e^x and e^{x+2} linearly independent? Justify.

Exercise 2.1.103: Guess a solution to $y'' + y' + y = 5$.

Exercise 2.1.104: Find the general solution to $xy'' + y' = 0$. Hint: It is a first-order ODE in y' .

Exercise 2.1.105: Write down an equation (guess) for which we have the solutions e^x and e^{2x} . Hint: Try an equation of the form $y'' + Ay' + By = 0$ for constants A and B , plug in both e^x and e^{2x} and solve for A and B .

2.2 Constant-coefficient second-order linear ODEs

Note: more than 1 lecture, second part of §3.1 in [EP], §3.1 in [BD]

2.2.1 Solving constant-coefficient equations

Consider the problem

$$y'' - 6y' + 8y = 0, \quad y(0) = -2, \quad y'(0) = 6.$$

This is a second-order linear homogeneous equation with constant coefficients. *Constant coefficients* means that the functions in front of y'' , y' , and y are constants; they do not depend on x .

To guess a solution, think of a function that stays essentially the same when we differentiate it, so that we can take the function and its derivatives, add some multiples of these together, and end up with zero. Yes, we are talking about the exponential.

We try* a solution of the form $y = e^{rx}$. Then $y' = re^{rx}$ and $y'' = r^2e^{rx}$. Plug in to get

$$\begin{aligned} y'' - 6y' + 8y &= 0, \\ \underbrace{r^2e^{rx}}_{y''} - 6 \underbrace{re^{rx}}_{y'} + 8 \underbrace{e^{rx}}_y &= 0, \\ r^2 - 6r + 8 &= 0 \quad (\text{divide through by } e^{rx}), \\ (r - 2)(r - 4) &= 0. \end{aligned}$$

Hence, if $r = 2$ or $r = 4$, then e^{rx} is a solution. So let $y_1 = e^{2x}$ and $y_2 = e^{4x}$.

Exercise 2.2.1: Check that y_1 and y_2 are solutions.

The functions e^{2x} and e^{4x} are linearly independent. If they were not linearly independent, we could write $e^{4x} = Ce^{2x}$ for some constant C , implying that $e^{2x} = C$ for all x , which is clearly not possible. Hence, we can write the general solution as

$$y = C_1e^{2x} + C_2e^{4x}.$$

We need to solve for C_1 and C_2 . To apply the initial conditions, we first find $y' = 2C_1e^{2x} + 4C_2e^{4x}$. We plug $x = 0$ into y and y' and solve.

$$\begin{aligned} -2 &= y(0) = C_1 + C_2, \\ 6 &= y'(0) = 2C_1 + 4C_2. \end{aligned}$$

*Making an educated guess with some parameters to solve for is such a central technique in differential equations that people sometimes use a fancy name for such a guess: *ansatz*, German for “initial placement of a tool at a work piece.” Yes, the Germans have a word for that.

Either apply some matrix algebra, or just solve these by high school math. For example, divide the second equation by 2 to obtain $3 = C_1 + 2C_2$, then subtract the two equations to get $5 = C_2$. Then $C_1 = -7$ as $-2 = C_1 + 5$. Hence, the solution we are looking for is

$$y = -7e^{2x} + 5e^{4x}.$$

We generalize this example into a method. Suppose that we have an equation

$$ay'' + by' + cy = 0, \quad (2.3)$$

where a, b, c are constants. Try the solution $y = e^{rx}$ to obtain

$$ar^2e^{rx} + bre^{rx} + ce^{rx} = 0.$$

Divide by e^{rx} to obtain the so-called *characteristic equation* of the ODE:

$$ar^2 + br + c = 0.$$

Solve for the r by using the quadratic formula:

$$r_1, r_2 = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}.$$

Suppose that $b^2 - 4ac \geq 0$ for now so that r_1 and r_2 are real. So e^{r_1x} and e^{r_2x} are solutions. There is still a difficulty if $r_1 = r_2$, but it is not hard to overcome.

Theorem 2.2.1. *Suppose that r_1 and r_2 are the roots of the characteristic equation.*

(i) *If r_1 and r_2 are distinct and real (when $b^2 - 4ac > 0$), then (2.3) has the general solution*

$$y = C_1e^{r_1x} + C_2e^{r_2x}.$$

(ii) *If $r_1 = r_2$ (happens when $b^2 - 4ac = 0$), then (2.3) has the general solution*

$$y = (C_1 + C_2x)e^{r_1x}.$$

Example 2.2.1: Solve

$$y'' - k^2y = 0.$$

The characteristic equation is $r^2 - k^2 = 0$ or $(r - k)(r + k) = 0$. Consequently, e^{-kx} and e^{kx} are the two linearly independent solutions, and the general solution is

$$y = C_1e^{kx} + C_2e^{-kx}.$$

Since $\cosh s = \frac{e^s + e^{-s}}{2}$ and $\sinh s = \frac{e^s - e^{-s}}{2}$, we can also write the general solution as

$$y = D_1 \cosh(kx) + D_2 \sinh(kx).$$

Example 2.2.2: Find the general solution of

$$y'' - 8y' + 16y = 0.$$

The characteristic equation is $r^2 - 8r + 16 = (r - 4)^2 = 0$. The equation has a double root $r_1 = r_2 = 4$. The general solution is, therefore,

$$y = (C_1 + C_2x)e^{4x} = C_1e^{4x} + C_2xe^{4x}.$$

Exercise 2.2.2: Check that e^{4x} and xe^{4x} are linearly independent.

It is good to check your work. That e^{4x} solves the equation is clear. Let us check that xe^{4x} solves the equation. Compute $y' = e^{4x} + 4xe^{4x}$ and $y'' = 8e^{4x} + 16xe^{4x}$. Plug in,

$$y'' - 8y' + 16y = 8e^{4x} + 16xe^{4x} - 8(e^{4x} + 4xe^{4x}) + 16xe^{4x} = 0.$$

In some sense, a double root rarely happens. If coefficients are picked randomly, a double root is unlikely. There are, however, some real-world problems where a double root does happen naturally (e.g. critically damped mass-spring system, as we will see).

Here is a short argument for why the solution xe^{rx} works for a double root. A double root is a limit case of two distinct but very close roots. Note that $\frac{e^{r_2x} - e^{r_1x}}{r_2 - r_1}$ is a solution when the roots are distinct. When we take the limit as r_1 goes to r_2 , we are really taking the derivative of e^{rx} using r as the variable. Therefore, the limit is xe^{rx} , and hence this is a solution in the double root case. We remark that in some numerical computations, two very close roots may lead to numerical instability while a double root will not.

2.2.2 Complex numbers and Euler's formula

A polynomial may have complex roots. The equation $r^2 + 1 = 0$ has no real roots, but it does have two complex roots. Here we review some properties of complex numbers.

Complex numbers may seem a strange concept, perhaps due to the terminology. There is nothing imaginary or complicated about complex numbers. A complex number is simply a pair of real numbers, (a, b) , that is, it is a point in the plane. We add complex numbers in the straightforward way: $(a, b) + (c, d) = (a + c, b + d)$. We define multiplication by

$$(a, b) \times (c, d) \stackrel{\text{def}}{=} (ac - bd, ad + bc).$$

It turns out that with this multiplication rule, all the standard properties of arithmetic hold. Further, and most importantly $(0, 1) \times (0, 1) = (-1, 0)$.

Generally we write (a, b) as $a + ib$, and we treat i as if it were an unknown. When b is zero, then $(a, 0)$ is just the number a . We do arithmetic with complex numbers just as we would with polynomials. The property we just mentioned becomes $i^2 = -1$. So whenever we see i^2 , we replace it by -1 . For example,

$$(2 + 3i)(4i) - 5i = (2 \times 4)i + (3 \times 4)i^2 - 5i = 8i + 12(-1) - 5i = -12 + 3i.$$

The numbers i and $-i$ are the two roots of $r^2 + 1 = 0$. Some engineers use the letter j instead of i for the square root of -1 . We follow the mathematicians' convention and use i .

Exercise 2.2.3: Make sure you understand (that you can justify) the following identities:

- a) $i^2 = -1, i^3 = -i, i^4 = 1,$ b) $\frac{1}{i} = -i,$
c) $(3 - 7i)(-2 - 9i) = \dots = -69 - 13i,$ d) $(3 - 2i)(3 + 2i) = 3^2 - (2i)^2 = 3^2 + 2^2 = 13,$
e) $\frac{1}{3-2i} = \frac{1}{3-2i} \frac{3+2i}{3+2i} = \frac{3+2i}{13} = \frac{3}{13} + \frac{2}{13}i.$

We also define the exponential e^{a+ib} of a complex number by writing down the Taylor series and plugging in the complex number. Most properties of the exponential follow from its Taylor series, and so these properties hold for the complex exponential too. For example, we have the property: $e^{x+y} = e^x e^y$. In particular, $e^{a+ib} = e^a e^{ib}$. Hence, if we can compute e^{ib} , we can compute e^{a+ib} . For e^{ib} , we use the so-called *Euler's formula*.

Theorem 2.2.2 (Euler's formula).

$$e^{i\theta} = \cos \theta + i \sin \theta \quad \text{and} \quad e^{-i\theta} = \cos \theta - i \sin \theta.$$

In other words, $e^{a+ib} = e^a (\cos(b) + i \sin(b)) = e^a \cos(b) + i e^a \sin(b)$.

Exercise 2.2.4: Using Euler's formula, check the identities:

$$\cos \theta = \frac{e^{i\theta} + e^{-i\theta}}{2} \quad \text{and} \quad \sin \theta = \frac{e^{i\theta} - e^{-i\theta}}{2i}.$$

Exercise 2.2.5 (Double angle identities): Start with $e^{i(2\theta)} = (e^{i\theta})^2$. Use Euler on each side to deduce:

$$\cos(2\theta) = \cos^2 \theta - \sin^2 \theta \quad \text{and} \quad \sin(2\theta) = 2 \sin \theta \cos \theta.$$

For a complex number $a + ib$, we call a the *real part* and b the *imaginary part* of the number. Often the following notation is used:

$$\operatorname{Re}(a + ib) = a \quad \text{and} \quad \operatorname{Im}(a + ib) = b.$$

2.2.3 Complex roots

Suppose the differential equation $ay'' + by' + cy = 0$ has the characteristic equation $ar^2 + br + c = 0$ that has complex roots. By the quadratic formula, the roots are $\frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$. These roots are complex if $b^2 - 4ac < 0$. In this case, we write the roots as

$$r_1, r_2 = \frac{-b}{2a} \pm i \frac{\sqrt{4ac - b^2}}{2a}.$$

As you can see, we get a pair of roots of the form $\alpha \pm i\beta$. We could still write the solution as

$$y = C_1 e^{(\alpha+i\beta)x} + C_2 e^{(\alpha-i\beta)x}.$$

However, the exponential is now complex-valued. We need to allow C_1 and C_2 to be complex numbers to obtain a real-valued solution (which is what we are after). While there is nothing particularly wrong with this approach, it can make calculations harder and it is generally preferred to find two real-valued solutions.

Euler's formula comes to the rescue. Let

$$y_1 = e^{(\alpha+i\beta)x} \quad \text{and} \quad y_2 = e^{(\alpha-i\beta)x}.$$

Then

$$y_1 = e^{\alpha x} \cos(\beta x) + ie^{\alpha x} \sin(\beta x),$$

$$y_2 = e^{\alpha x} \cos(\beta x) - ie^{\alpha x} \sin(\beta x).$$

Linear combinations of solutions are also solutions. Hence,

$$y_3 = \frac{y_1 + y_2}{2} = e^{\alpha x} \cos(\beta x),$$

$$y_4 = \frac{y_1 - y_2}{2i} = e^{\alpha x} \sin(\beta x),$$

are also solutions. It is not hard to see that y_3 and y_4 are linearly independent (not multiples of each other). So the general solution can be written in terms of y_3 and y_4 . And as they are real-valued, no complex numbers need to be used for the arbitrary constants in the general solution. We summarize what we found as a theorem.

Theorem 2.2.3. *Take the equation*

$$ay'' + by' + cy = 0.$$

If the characteristic equation has the roots $\alpha \pm i\beta$ (when $b^2 - 4ac < 0$), then the general solution is

$$y = C_1 e^{\alpha x} \cos(\beta x) + C_2 e^{\alpha x} \sin(\beta x).$$

Example 2.2.3: Find the general solution of $y'' + k^2 y = 0$, for a constant $k > 0$.

The characteristic equation is $r^2 + k^2 = 0$, and the roots are $r = \pm ik$. By the theorem, the general solution is

$$y = C_1 \cos(kx) + C_2 \sin(kx).$$

Example 2.2.4: Find the solution of $y'' - 6y' + 13y = 0$, $y(0) = 0$, $y'(0) = 10$.

The characteristic equation is $r^2 - 6r + 13 = 0$. Completing the square, we get $(r - 3)^2 + 2^2 = 0$ and hence the roots are $r = 3 \pm 2i$. Per the theorem, the general solution is

$$y = C_1 e^{3x} \cos(2x) + C_2 e^{3x} \sin(2x).$$

To find the solution satisfying the initial conditions, we first plug in zero to get

$$0 = y(0) = C_1 e^0 \cos 0 + C_2 e^0 \sin 0 = C_1.$$

Hence, $C_1 = 0$ and $y = C_2 e^{3x} \sin(2x)$. We differentiate,

$$y' = 3C_2 e^{3x} \sin(2x) + 2C_2 e^{3x} \cos(2x).$$

We again plug in the initial condition and obtain $10 = y'(0) = 2C_2$, or $C_2 = 5$. The solution we are seeking is

$$y = 5e^{3x} \sin(2x).$$

2.2.4 Exercises

Exercise 2.2.6: Find the general solution of $2y'' + 2y' - 4y = 0$.

Exercise 2.2.7: Find the general solution of $y'' + 9y' - 10y = 0$.

Exercise 2.2.8: Solve $y'' - 8y' + 16y = 0$ for $y(0) = 2$, $y'(0) = 0$.

Exercise 2.2.9: Solve $y'' + 9y' = 0$ for $y(0) = 1$, $y'(0) = 1$.

Exercise 2.2.10: Find the general solution of $2y'' + 50y = 0$.

Exercise 2.2.11: Find the general solution of $y'' + 6y' + 13y = 0$.

Exercise 2.2.12: Find the general solution of $y'' = 0$ using the methods of this section.

Exercise 2.2.13: The method of this section applies to equations of other orders than two. We will see higher orders later. Solve the first-order equation $2y' + 3y = 0$ using the methods of this section.

Exercise 2.2.14: Let us revisit the Cauchy–Euler equations of [Exercise 2.1.6](#) on page 82. Suppose now that $(b - a)^2 - 4ac < 0$. Find a formula for the general solution of $ax^2y'' + bxy' + cy = 0$. Hint: Note that $x^r = e^{r \ln x}$.

Exercise 2.2.15: Find the solution to $y'' - (2\alpha)y' + \alpha^2y = 0$, $y(0) = a$, $y'(0) = b$, where α , a , and b are real numbers.

Exercise 2.2.16: Construct an equation such that $y = C_1e^{-2x} \cos(3x) + C_2e^{-2x} \sin(3x)$ is the general solution.

Exercise 2.2.101: Find the general solution to $y'' + 4y' + 2y = 0$.

Exercise 2.2.102: Find the general solution to $y'' - 6y' + 9y = 0$.

Exercise 2.2.103: Find the solution to $2y'' + y' + y = 0$, $y(0) = 1$, $y'(0) = -2$.

Exercise 2.2.104: Find the solution to $2y'' + y' - 3y = 0$, $y(0) = a$, $y'(0) = b$.

Exercise 2.2.105: Find the solution to $z''(t) = -2z'(t) - 2z(t)$, $z(0) = 2$, $z'(0) = -2$.

Exercise 2.2.106: Find the solution to $y'' - (\alpha + \beta)y' + \alpha\beta y = 0$, $y(0) = a$, $y'(0) = b$, where α , β , a , and b are real numbers, and $\alpha \neq \beta$.

Exercise 2.2.107: Construct an equation such that $y = C_1e^{3x} + C_2e^{-2x}$ is the general solution.

2.3 Higher-order linear ODEs

Note: somewhat more than 1 lecture, §3.2 and §3.3 in [EP], §4.1 and §4.2 in [BD]

We briefly study higher-order equations. Equations appearing in applications tend to be second-order. Higher-order equations do appear from time to time, but generally the world around us is “second-order.”

The basic results about linear ODEs of higher order are essentially the same as for second-order equations, with 2 replaced by n . The important concept of linear independence is somewhat more complicated when more than two functions are involved. For higher-order constant-coefficient ODEs, the methods developed are also somewhat harder to apply, but we will not dwell on these complications. It is also possible to use the methods for systems of linear equations from [chapter 3](#) to solve higher-order constant-coefficient equations.

Consider a general homogeneous linear equation

$$y^{(n)} + p_{n-1}(x)y^{(n-1)} + \cdots + p_1(x)y' + p_0(x)y = 0. \quad (2.4)$$

Theorem 2.3.1 (Superposition). *If $y_1, y_2, \dots, y_n$ are solutions of the homogeneous equation (2.4), then*

$$y(x) = C_1y_1(x) + C_2y_2(x) + \cdots + C_ny_n(x)$$

also solves (2.4) for arbitrary constants $C_1, C_2, \dots, C_n$.

That is, a *linear combination* of solutions to (2.4) is a solution to (2.4). There is also the existence and uniqueness theorem for linear equations, including nonhomogeneous ones.

Theorem 2.3.2 (Existence and uniqueness). *Suppose $p_0, p_1, \dots, p_{n-1}$, and f are continuous functions on some interval I , a is a number in I , and $b_0, b_1, \dots, b_{n-1}$ are constants. Then the equation*

$$y^{(n)} + p_{n-1}(x)y^{(n-1)} + \cdots + p_1(x)y' + p_0(x)y = f(x)$$

has exactly one solution $y(x)$ defined on the same interval I satisfying the initial conditions

$$y(a) = b_0, \quad y'(a) = b_1, \quad \dots, \quad y^{(n-1)}(a) = b_{n-1}.$$

2.3.1 Linear independence

When we had two functions y_1 and y_2 , we said they were linearly independent if one was not a multiple of the other. Same idea holds for n functions, although in this case it is easier to state as follows. The functions $y_1, y_2, \dots, y_n$ are *linearly independent* if the equation

$$c_1y_1 + c_2y_2 + \cdots + c_ny_n = 0$$

has only the trivial solution $c_1 = c_2 = \cdots = c_n = 0$, where the equation must hold for all x . If we can solve the equation with some constants $c_1, c_2, \dots, c_n$, where for example $c_1 \neq 0$, then we can solve for y_1 as a linear combination of the others. If the functions are not linearly independent, they are *linearly dependent*.

Example 2.3.1: Verify that e^x, e^{2x}, e^{3x} are linearly independent.

Let us give several ways to show this fact. Many textbooks (including [EP] and [F]) introduce Wronskians, but it is difficult to see why they work and they are not really necessary here.

Consider

$$c_1 e^x + c_2 e^{2x} + c_3 e^{3x} = 0.$$

We use rules of exponentials and write $z = e^x$. Hence $z^2 = e^{2x}$ and $z^3 = e^{3x}$. Then we have

$$c_1 z + c_2 z^2 + c_3 z^3 = 0.$$

The left-hand side is a third degree polynomial in z . It is either identically zero, or it has at most 3 zeros. Therefore, it is identically zero, $c_1 = c_2 = c_3 = 0$, and the functions are linearly independent.

Let us try another way. As before we write

$$c_1 e^x + c_2 e^{2x} + c_3 e^{3x} = 0.$$

This equation has to hold for all x . We divide through by e^{3x} to get

$$c_1 e^{-2x} + c_2 e^{-x} + c_3 = 0.$$

As the equation is true for all x , let $x \rightarrow \infty$. After taking the limit we see that $c_3 = 0$. Hence our equation becomes

$$c_1 e^x + c_2 e^{2x} = 0.$$

Rinse, repeat!

How about yet another way? We again write

$$c_1 e^x + c_2 e^{2x} + c_3 e^{3x} = 0.$$

We can evaluate the equation and its derivatives at different values of x to obtain equations for c_1, c_2 , and c_3 . Let us first divide by e^x for simplicity.

$$c_1 + c_2 e^x + c_3 e^{2x} = 0.$$

We set $x = 0$ to get the equation $c_1 + c_2 + c_3 = 0$. Now differentiate both sides

$$c_2 e^x + 2c_3 e^{2x} = 0.$$

We set $x = 0$ to get $c_2 + 2c_3 = 0$. We divide by e^x again and differentiate to get $2c_3 e^x = 0$. It is clear that c_3 is zero. Then c_2 must be zero as $c_2 = -2c_3$, and c_1 must be zero because $c_1 + c_2 + c_3 = 0$.

There is no one best way to do it. All of these methods are perfectly valid. The important thing is to understand why the functions are linearly independent.

Example 2.3.2: On the other hand, the functions e^x , e^{-x} , and $\cosh x$ are linearly dependent. Simply apply definition of the hyperbolic cosine:

$$\cosh x = \frac{e^x + e^{-x}}{2} \quad \text{or} \quad 2 \cosh x - e^x - e^{-x} = 0.$$

Once we have enough linearly independent solutions, we have the general solution to the homogeneous equation, just as we did for second-order equations.

Theorem 2.3.3. If $y_1, y_2, \dots, y_n$ are linearly independent solutions of the homogeneous equation (2.4), then the general solution to (2.4) can be written as

$$y(x) = C_1 y_1(x) + C_2 y_2(x) + \dots + C_n y_n(x).$$

2.3.2 Constant-coefficient higher-order ODEs

When we have a higher-order constant-coefficient homogeneous linear equation, the song and dance is exactly the same as it was for second order. We just need to find more solutions. If the equation is n^{th} -order, we need to find n linearly independent solutions. It is best seen by example.

Example 2.3.3: Find the general solution to

$$y''' - 3y'' - y' + 3y = 0. \quad (2.5)$$

Try: $y = e^{rx}$. We plug in and get

$$\underbrace{r^3 e^{rx}}_{y'''} - 3 \underbrace{r^2 e^{rx}}_{y''} - \underbrace{r e^{rx}}_{y'} + 3 \underbrace{e^{rx}}_y = 0.$$

We divide through by e^{rx} . Then

$$r^3 - 3r^2 - r + 3 = 0.$$

The trick now is to find the roots. There is a formula for the roots of degree 3 and 4 polynomials, but it is very complicated. There is no formula for higher degree polynomials. That does not mean that the roots do not exist. There are always n roots for an n^{th} degree polynomial. They may be repeated and they may be complex. Computers are pretty good at finding roots approximately for reasonable size polynomials.

A good place to start is to plot the polynomial and check where it is zero. We can also simply try plugging in. We just start plugging in numbers $r = -2, -1, 0, 1, 2, \dots$ and see if we get a hit (we can also try complex numbers). Even if we do not get a hit, we may get an indication of where the root is. For example, we plug $r = -2$ into our polynomial and get -15 ; we plug in $r = 0$ and get 3 . That means there is a root between $r = -2$ and $r = 0$, because the sign changed. If we find one root, say r_1 , then we know $(r - r_1)$ is a factor of our polynomial. Polynomial long division can then be used.

A good strategy is to begin with $r = 0, 1$, or -1 . These are easy to compute. Our polynomial has two such roots, $r_1 = -1$ and $r_2 = 1$. There should be 3 roots and the last root is reasonably easy to find. The constant term in a monic* polynomial such as this is the

*The word monic means that the coefficient of the top degree r^d , in our case r^3 , is 1.

product of the negations of all the roots because $r^3 - 3r^2 - r + 3 = (r - r_1)(r - r_2)(r - r_3)$. So

$$3 = (-r_1)(-r_2)(-r_3) = (1)(-1)(-r_3) = r_3.$$

You should check that $r_3 = 3$ really is a root. Hence e^{-x} , e^x and e^{3x} are solutions to (2.5). They are linearly independent as can easily be checked, and there are 3 of them, which happens to be exactly the number we need. So the general solution is

$$y = C_1 e^{-x} + C_2 e^x + C_3 e^{3x}.$$

Suppose we were given some initial conditions $y(0) = 1$, $y'(0) = 2$, and $y''(0) = 3$. Then

$$1 = y(0) = C_1 + C_2 + C_3,$$

$$2 = y'(0) = -C_1 + C_2 + 3C_3,$$

$$3 = y''(0) = C_1 + C_2 + 9C_3.$$

It is possible to find the solution by high school algebra, but it would be a pain. The sensible way to solve a system of equations such as this is to use matrix algebra, see § 3.2 or appendix A. For now we note that the solution is $C_1 = -1/4$, $C_2 = 1$, and $C_3 = 1/4$. The specific solution to the ODE is

$$y = \frac{-1}{4} e^{-x} + e^x + \frac{1}{4} e^{3x}.$$

Next, suppose that we have real roots, but they are repeated. Let us say we have a root r repeated k times. In the spirit of the second-order solution, and for the same reasons, we have the solutions

$$e^{rx}, \quad x e^{rx}, \quad x^2 e^{rx}, \quad \dots, \quad x^{k-1} e^{rx}.$$

We take a linear combination of these solutions to find the general solution.

Example 2.3.4: Solve

$$y^{(4)} - 3y''' + 3y'' - y' = 0.$$

We note that the characteristic equation is

$$r^4 - 3r^3 + 3r^2 - r = 0.$$

By inspection, we note that $r^4 - 3r^3 + 3r^2 - r = r(r - 1)^3$. Hence the roots given with multiplicity are $r = 0, 1, 1, 1$. Thus the general solution is

$$y = \underbrace{(C_1 + C_2 x + C_3 x^2) e^x}_{\text{terms coming from } r=1} + \underbrace{C_4}_{\text{from } r=0}.$$

The case of complex roots is similar to second-order equations. Complex roots always come in pairs $r = \alpha \pm i\beta$. Suppose we have two such complex roots, each repeated k times. The corresponding solution is

$$(C_0 + C_1 x + \dots + C_{k-1} x^{k-1}) e^{\alpha x} \cos(\beta x) + (D_0 + D_1 x + \dots + D_{k-1} x^{k-1}) e^{\alpha x} \sin(\beta x),$$

where $C_0, \dots, C_{k-1}, D_0, \dots, D_{k-1}$ are arbitrary constants.

Example 2.3.5: Solve

$$y^{(4)} - 4y''' + 8y'' - 8y' + 4y = 0.$$

The characteristic equation is

$$r^4 - 4r^3 + 8r^2 - 8r + 4 = 0,$$

$$(r^2 - 2r + 2)^2 = 0,$$

$$((r - 1)^2 + 1)^2 = 0.$$

Hence the roots are $1 \pm i$, both with multiplicity 2. Hence the general solution to the ODE is

$$y = (C_1 + C_2x)e^x \cos x + (C_3 + C_4x)e^x \sin x.$$

The way we solved the characteristic equation above is really by guessing or by inspection. It is not so easy in general. We could also have asked a computer or an advanced calculator for the roots.

2.3.3 Exercises

Exercise 2.3.1: Find the general solution for $y''' - y'' + y' - y = 0$.

Exercise 2.3.2: Find the general solution for $y^{(4)} - 5y''' + 6y'' = 0$.

Exercise 2.3.3: Find the general solution for $y''' + 2y'' + 2y' = 0$.

Exercise 2.3.4: Suppose the characteristic equation for an ODE is $(r - 1)^2(r - 2)^2 = 0$.

- Find such a differential equation.
- Find its general solution.

Exercise 2.3.5: Suppose that a fourth-order equation has a solution $y = 2e^{4x}x \cos x$.

- Find such an equation.
- Find the initial conditions that the given solution satisfies.

Exercise 2.3.6: Find the general solution for the equation of [Exercise 2.3.5](#).

Exercise 2.3.7: Let $f(x) = e^x - \cos x$, $g(x) = e^x + \cos x$, and $h(x) = \cos x$. Are $f(x)$, $g(x)$, and $h(x)$ linearly independent? If so, show it, if not, find a linear combination that works.

Exercise 2.3.8: Let $f(x) = 0$, $g(x) = \cos x$, and $h(x) = \sin x$. Are $f(x)$, $g(x)$, and $h(x)$ linearly independent? If so, show it, if not, find a linear combination that works.

Exercise 2.3.9: Are x , x^2 , and x^4 linearly independent? If so, show it, if not, find a linear combination that gives 0.

Exercise 2.3.10: Are e^x , xe^x , and x^2e^x linearly independent? If so, show it, if not, find a linear combination that gives 0.

Exercise 2.3.11: Find an equation such that $y = xe^{-2x} \sin(3x)$ is a solution.

Exercise 2.3.101: Find the general solution of $y^{(5)} - y^{(4)} = 0$.

Exercise 2.3.102: Suppose that the characteristic equation of a third-order differential equation has roots $\pm 2i$ and 3.

- a) What is the characteristic equation?
- b) Find the corresponding differential equation.
- c) Find the general solution.

Exercise 2.3.103: Solve $1001y''' + 3.2y'' + \pi y' - \sqrt{4}y = 0$, $y(0) = 0$, $y'(0) = 0$, $y''(0) = 0$.

Exercise 2.3.104: Are e^x , e^{x+1} , e^{2x} , $\sin(x)$ linearly independent? If so, show it, if not find a linear combination that gives 0.

Exercise 2.3.105: Are $\sin(x)$, x , $x \sin(x)$ linearly independent? If so, show it, if not find a linear combination that gives 0.

Exercise 2.3.106: Find an equation such that $y = \cos(x)$, $y = \sin(x)$, $y = e^x$ are solutions.

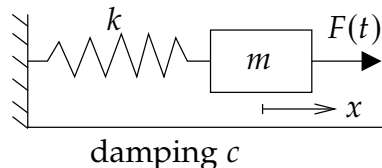
2.4 Mechanical vibrations

Note: 2 lectures, §3.4 in [EP], §3.7 in [BD]

Let us look at some applications of linear second-order constant-coefficient equations.

2.4.1 Some examples

Our first example is a mass on a spring. Suppose we have a mass $m > 0$ (in kilograms) connected by a spring with spring constant $k > 0$ (in newtons per meter) to a fixed wall. There may be some external force $F(t)$ (in newtons) acting on the mass. Here, t is time in seconds. Finally, there is some friction measured by $c \geq 0$ (in newton-seconds per meter) as the mass slides along the floor (or perhaps a damper is connected).



Let x be the displacement of the mass in meters ($x = 0$ is the rest position), with x growing to the right (away from the wall). The force exerted by the spring is proportional to the compression of the spring by Hooke's law. Therefore, it is kx in the negative direction. Similarly, the force exerted by friction is proportional to the velocity of the mass. Newton's second law says that force equals mass times acceleration. Hence, $mx'' = F(t) - cx' - kx$ or

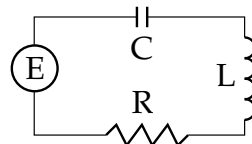
$$mx'' + cx' + kx = F(t).$$

The equation is a linear second-order constant-coefficient ODE. We say the motion is

- (i) *forced* if $F \not\equiv 0$ (if F is not identically zero),
- (ii) *unforced* or *free* if $F \equiv 0$ (if F is identically zero),
- (iii) *damped* if $c > 0$, and
- (iv) *undamped* if $c = 0$.

This equation appears in lots of applications even if it does not seem like it at first. Many real-world scenarios can be simplified to a mass on a spring. For example, a bungee jump is essentially a mass and spring system (you are the mass). It would be good if someone did the math before you jump off the bridge, right? Let us give two other examples.

Here is an example for electrical engineers. Consider the pictured RLC circuit. There is a resistor with a resistance of R ohms, an inductor with an inductance of L henries, and a capacitor with a capacitance of C farads. There is also an electric source (such as a battery) with voltage of $E(t)$ volts at time t (measured in seconds). Let $Q(t)$ be the charge (in coulombs) on the capacitor and $I(t)$ be the current (in amperes) in the circuit. The



relation between the two is $Q' = I$. By elementary principles, we find $LI' + RI + Q/C = E$. We differentiate to get

$$LI''(t) + RI'(t) + \frac{1}{C}I(t) = E'(t).$$

This is a nonhomogeneous second-order constant-coefficient linear equation. As L, R , and C are all positive, this circuit behaves just like the mass and spring system. Position of the mass is replaced by current. Mass is replaced by inductance, damping is replaced by resistance, and the spring constant is replaced by one over the capacitance. The change in voltage becomes the forcing function—for constant voltage this is an unforced motion.

Our next example behaves like a mass and spring system only approximately. Suppose a mass of m kilograms hangs on a pendulum of length L meters. We seek an equation for the angle $\theta(t)$ (in radians). Let g be the acceleration due to gravity in meters per second squared. Let us derive an equation for θ using Newton's second law: Force equals mass times acceleration. The acceleration is $L\theta''$ and mass is m . So $mL\theta''$ has to be equal to the tangential component of the force given by the gravity, which is $mg \sin \theta$ in the opposite direction. So $mL\theta'' = -mg \sin \theta$. The m curiously cancels from the equation and we end up with the equation

$$\theta'' + \frac{g}{L} \sin \theta = 0.$$

Now we make our approximation. For small θ we have that approximately $\sin \theta \approx \theta$. This can be seen by looking at the graph. In Figure 2.1, we can see that for approximately $-0.5 < \theta < 0.5$ (in radians) the graphs of $\sin \theta$ and θ are almost the same.

Therefore, when the swings are small, θ is small, and we can model the behavior by the simpler linear equation

$$\theta'' + \frac{g}{L} \theta = 0.$$

The errors from this approximation build up. After a long time, the state of the real-world system might be substantially different from our solution. Also, we will see that in a mass-spring system, the amplitude is independent of the period. This is not true for a pendulum. Nevertheless, for reasonably short periods of time and small swings (that is, only small angles θ), the approximation is reasonably good.

In real-world problems it is often necessary to make these types of simplifications. We must understand both the mathematics and the physics of the situation to see if the simplification is valid in the context of the questions we are trying to answer.

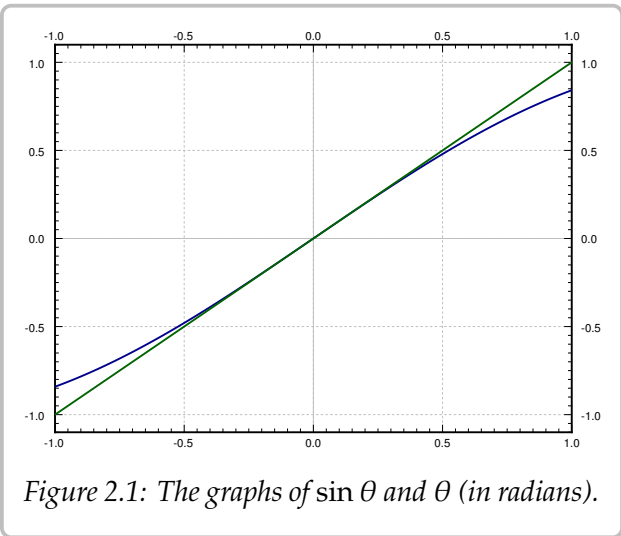
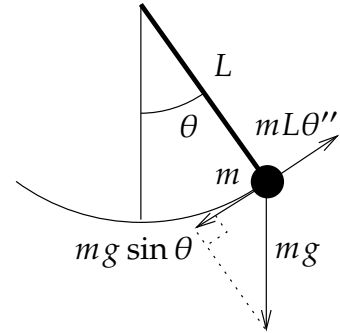


Figure 2.1: The graphs of $\sin \theta$ and θ (in radians).

2.4.2 Free undamped motion

In this section we only consider free or unforced motion, as we do not know yet how to solve nonhomogeneous equations. We start with undamped motion where $c = 0$. The equation is

$$mx'' + kx = 0.$$

We divide by m and let $\omega_0 = \sqrt{k/m}$ to rewrite the equation as

$$x'' + \omega_0^2 x = 0.$$

The general solution to this equation is

$$x(t) = A \cos(\omega_0 t) + B \sin(\omega_0 t).$$

By a trigonometric identity

$$A \cos(\omega_0 t) + B \sin(\omega_0 t) = C \cos(\omega_0 t - \gamma),$$

for two different constants C and γ . One finds that $A = C \cos \gamma$ and $B = C \sin \gamma$, and therefore it is not hard to compute that $C = \sqrt{A^2 + B^2}$ and $\tan \gamma = B/A$. Therefore, we let C and γ be our arbitrary constants and write $x(t) = C \cos(\omega_0 t - \gamma)$.

Exercise 2.4.1: Justify the identity $A \cos(\omega_0 t) + B \sin(\omega_0 t) = C \cos(\omega_0 t - \gamma)$ and verify the equations $C = \sqrt{A^2 + B^2}$ and $\tan \gamma = \frac{B}{A}$. Hint: Start with $\cos(\alpha - \beta) = \cos(\alpha)\cos(\beta) + \sin(\alpha)\sin(\beta)$ and multiply by C . Then what should α and β be?

While it is easier to use the first form with A and B to solve for the initial conditions, the second form is more natural. The constants C and γ have nice physical interpretation. Write the solution as

$$x(t) = C \cos(\omega_0 t - \gamma).$$

This is a pure-frequency oscillation (a sine wave). The *amplitude* is C , ω_0 is the (angular) *frequency*, and γ is the so-called *phase shift*. The phase shift just shifts the graph left or right. We call ω_0 the *natural (angular) frequency*. This entire setup is called *simple harmonic motion*.

Let us pause to explain the word *angular* before the word *frequency*. The units of ω_0 are radians per unit time, not cycles per unit time as is the usual measure of frequency. Because one cycle is 2π radians, the usual frequency is given by $\frac{\omega_0}{2\pi}$. It is simply a matter of where we put the constant 2π , and that is a matter of taste.

The *period* of the motion is one over the frequency (in cycles per unit time) and hence it is $\frac{2\pi}{\omega_0}$. That is the amount of time it takes to complete one full cycle.

Example 2.4.1: Suppose that $m = 2$ kg and $k = 8$ N/m. The whole mass and spring setup is sitting on a truck that was traveling at 1 m/s. The truck crashes and hence stops. The mass was held in place 0.5 meters forward from the rest position. During the crash the mass gets loose. That is, the mass is now moving forward at 1 m/s, while the other end of the spring is held in place. The mass therefore starts oscillating. What is the

frequency of the resulting oscillation? What is the amplitude? The units are the mks units (meters-kilograms-seconds).

The setup means that the mass was at half a meter in the positive direction during the crash and relative to the wall the spring is mounted to, the mass was moving forward (in the positive direction) at 1 m/s. This gives the initial conditions. So the equation with initial conditions is

$$2x'' + 8x = 0, \quad x(0) = 0.5, \quad x'(0) = 1.$$

We compute $\omega_0 = \sqrt{k/m} = \sqrt{4} = 2$. Hence the angular frequency is 2. The usual frequency in Hertz (cycles per second) is $2/2\pi = 1/\pi \approx 0.318$.

The general solution is

$$x(t) = A \cos(2t) + B \sin(2t).$$

Letting $x(0) = 0.5$ means $A = 0.5$. Then $x'(t) = -2(0.5) \sin(2t) + 2B \cos(2t)$. Letting $x'(0) = 1$, we get $B = 0.5$. Therefore, the amplitude is $C = \sqrt{A^2 + B^2} = \sqrt{0.25 + 0.25} = \sqrt{0.5} \approx 0.707$. The solution is

$$x(t) = 0.5 \cos(2t) + 0.5 \sin(2t).$$

A plot of $x(t)$ is shown in Figure 2.2.

In general, for free undamped motion, a solution of the form

$$x(t) = A \cos(\omega_0 t) + B \sin(\omega_0 t),$$

corresponds to the initial conditions $x(0) = A$ and $x'(0) = \omega_0 B$. It is easy to find A and B from the initial conditions. The amplitude and the phase shift are computed from A and B . In the example, we found the amplitude C . How about the phase shift γ ? We know $\tan \gamma = B/A = 1$. The arctangent of 1 is $\pi/4$ or approximately 0.785. We must check if $\pi/4$ gives the correct quadrant—if the angle from the positive x -axis is in the same quadrant as the point (A, B) —if not, we add π . That is because $A = C \cos \gamma$ and $B = C \sin \gamma$. Both A and B are positive, so γ must be in the first quadrant. As $\pi/4$ radians is in the first quadrant, $\gamma = \pi/4$.

Note: Many calculators and computer software have not only the `atan` function for arctangent, but also what is sometimes called `atan2`. This function takes two arguments, B and A , and returns a γ in the correct quadrant for you.

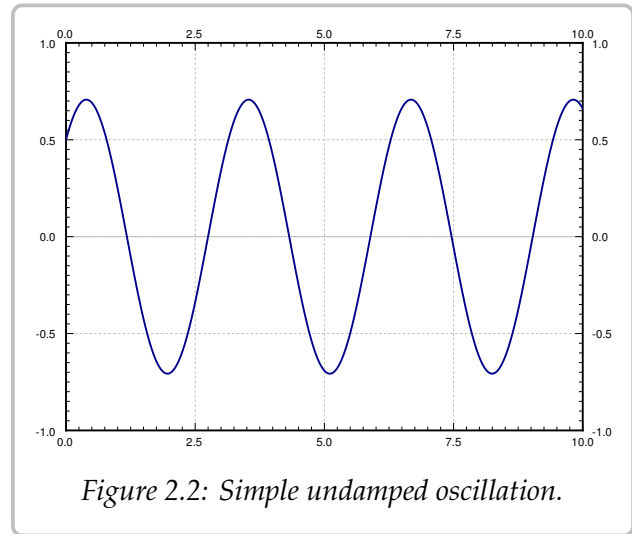


Figure 2.2: Simple undamped oscillation.

2.4.3 Free damped motion

Let us now focus on damped motion. We rewrite the equation

$$mx'' + cx' + kx = 0,$$

as

$$x'' + 2px' + \omega_0^2 x = 0,$$

where

$$\omega_0 = \sqrt{\frac{k}{m}}, \quad p = \frac{c}{2m}.$$

The characteristic equation is

$$r^2 + 2pr + \omega_0^2 = 0.$$

Using the quadratic formula, we get that the roots are

$$r = -p \pm \sqrt{p^2 - \omega_0^2}.$$

The form of the solution depends on whether we get complex or real roots. We get real roots if and only if the following number is nonnegative:

$$p^2 - \omega_0^2 = \left(\frac{c}{2m}\right)^2 - \frac{k}{m} = \frac{c^2 - 4km}{4m^2}.$$

The sign of $p^2 - \omega_0^2$ is the same as the sign of $c^2 - 4km$. Thus we get real roots if and only if $c^2 - 4km$ is nonnegative, or in other words if $c^2 \geq 4km$.

Overdamping

When $c^2 - 4km > 0$, the system is *overdamped*. In this case, there are two distinct real roots r_1 and r_2 . As $\sqrt{p^2 - \omega_0^2}$ is always less than p , the expression for the roots $-p \pm \sqrt{p^2 - \omega_0^2}$ is negative either way. So both roots are negative.

The solution is

$$x(t) = C_1 e^{r_1 t} + C_2 e^{r_2 t}.$$

Since r_1, r_2 are negative, $x(t) \rightarrow 0$ as $t \rightarrow \infty$. Thus the mass will tend towards the rest position as time goes to infinity. For a few sample plots for different initial conditions, see [Figure 2.3](#).

No oscillation happens. In fact, the graph crosses the t -axis at most once. To see why, we try to solve $0 = C_1 e^{r_1 t} + C_2 e^{r_2 t}$. Therefore, $C_1 e^{r_1 t} = -C_2 e^{r_2 t}$ and using laws of exponents we obtain

$$\frac{-C_1}{C_2} = e^{(r_2 - r_1)t}.$$

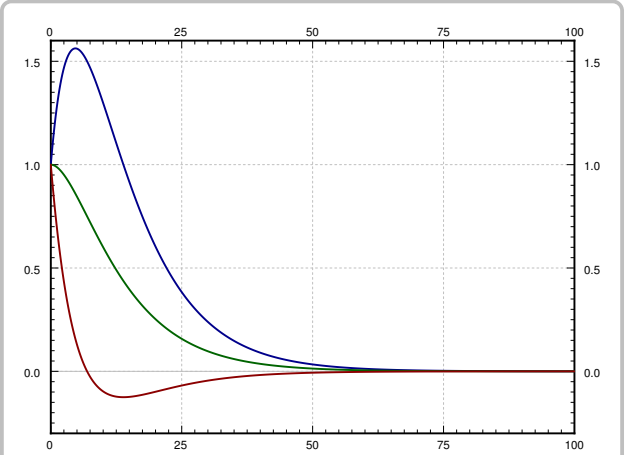


Figure 2.3: Overdamped motion for several different initial conditions.

This equation has at most one solution $t \geq 0$; the graph of $e^{(r_2-r_1)t}$ is either always increasing or always decreasing. For some initial conditions the graph never crosses the t -axis, as is evident from the sample graphs.

Example 2.4.2: Suppose the mass is released from rest. That is, $x(0) = x_0$ and $x'(0) = 0$. Then

$$x(t) = \frac{x_0}{r_1 - r_2} (r_1 e^{r_2 t} - r_2 e^{r_1 t}).$$

It is not hard to see that this satisfies the initial conditions.

Critical damping

When $c^2 - 4km = 0$, the system is *critically damped*. In this case, there is one root of multiplicity 2 and this root is $-p$. Our solution is

$$x(t) = C_1 e^{-pt} + C_2 t e^{-pt}.$$

The behavior of a critically damped system is very similar to an overdamped system. After all, a critically damped system is, in some sense, a limit of overdamped systems. Since these equations are really only an approximation to the real world, in reality, we are never critically damped; it is a place we can only reach in theory. We are always a little bit underdamped or a little bit overdamped. It is better not to dwell on critical damping.

Underdamping

When $c^2 - 4km < 0$, the system is *underdamped*. In this case, the roots are complex.

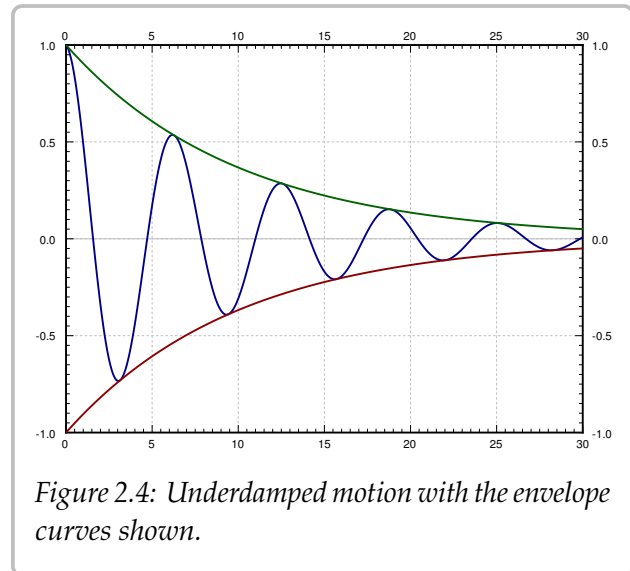
$$\begin{aligned} r &= -p \pm \sqrt{p^2 - \omega_0^2} \\ &= -p \pm \sqrt{-1} \sqrt{\omega_0^2 - p^2} \\ &= -p \pm i\omega_1, \end{aligned}$$

where $\omega_1 = \sqrt{\omega_0^2 - p^2}$. Our solution is

$$x(t) = e^{-pt} (A \cos(\omega_1 t) + B \sin(\omega_1 t)),$$

or

$$x(t) = C e^{-pt} \cos(\omega_1 t - \gamma).$$



An example plot is given in Figure 2.4. Note that we still have that $x(t) \rightarrow 0$ as $t \rightarrow \infty$.

The figure also shows the *envelope curves* Ce^{-pt} and $-Ce^{-pt}$. The solution is the oscillating line between the two envelope curves. The envelope curves give the maximum amplitude of the oscillation at any given point in time. For example, if you are bungee jumping, you are really interested in computing the envelope curve so as not to hit the concrete with

your head. The phase shift γ shifts the oscillation left or right, but within the envelope curves (the envelope curves do not change if γ changes).

Notice that the angular *pseudo-frequency** ω_1 becomes smaller when the damping c (and hence p) becomes larger. This makes sense. First, when we change the damping just a little bit, we do not expect the behavior of the solution to change dramatically. Second, if we keep making c larger, then at some point the solution should start looking like the solution for critical damping or overdamping, where no oscillation happens. As c gets larger and c^2 approaches $4km$, we find that ω_1 approaches 0. On the other hand, when c gets smaller, ω_1 approaches ω_0 (ω_1 is always smaller than ω_0), and the solution looks more and more like the steady periodic motion of the undamped case. The envelope curves become flatter and flatter as c (and hence p) goes to 0.

2.4.4 Exercises

Exercise 2.4.2: Consider a mass and spring system with a mass $m = 2$, spring constant $k = 3$, and damping constant $c = 1$.

- Set up and find the general solution of the system.
- Is the system underdamped, overdamped, or critically damped?
- If the system is not critically damped, find a c that makes the system critically damped.

Exercise 2.4.3: Do [Exercise 2.4.2](#) for $m = 3$, $k = 12$, and $c = 12$.

Exercise 2.4.4: Using the mks units (meters-kilograms-seconds), suppose you have a spring with spring constant 4 N/m . You want to use it to weigh items. Assume no friction. You place the mass on the spring and put it in motion.

- You count and find that the frequency is 0.8 Hz (cycles per second). What is the mass?
- Find a formula for the mass m given the frequency ω in Hz .

Exercise 2.4.5: Suppose we add possible friction to [Exercise 2.4.4](#). Further, suppose you do not know the spring constant, but you have two reference weights 1 kg and 2 kg to calibrate your setup. You put each in motion on your spring and measure the frequency. For the 1 kg weight you measured 1.1 Hz , for the 2 kg weight you measured 0.8 Hz .

- Find k (spring constant) and c (damping constant).
- Find a formula for the mass in terms of the frequency in Hz . Note that there may be more than one possible mass for a given frequency.
- For an unknown object you measured 0.2 Hz , what is the mass of the object? Suppose that you know that the mass of the unknown object is more than a kilogram.

*We do not call ω_1 a frequency since the solution is not really a periodic function.

Exercise 2.4.6: Suppose you wish to measure the friction a mass of 0.1 kg experiences as it slides along a floor (you wish to find c). You have a spring with spring constant $k = 5 \text{ N/m}$. You take the spring, you attach it to the mass and fix it to a wall. Then you pull on the spring and let the mass go. You find that the mass oscillates with frequency 1 Hz. What is the friction?

Exercise 2.4.101: A mass of 2 kilograms is on a spring with spring constant k newtons per meter with no damping. Suppose the system is at rest and at time $t = 0$ the mass is kicked and starts traveling at 2 meters per second. How large does k have to be to so that the mass does not go further than 3 meters from the rest position?

Exercise 2.4.102: Suppose we have an RLC circuit with a resistor of 100 milliohms (0.1 ohms), inductor of inductance of 50 millihenries (0.05 henries), and a capacitor of 5 farads, with constant voltage.

- a) Set up the ODE equation for the current I .
- b) Find the general solution.
- c) Solve for $I(0) = 10$ and $I'(0) = 0$.

Exercise 2.4.103: A 5000 kg railcar hits a bumper (a spring) at 1 m/s , and the spring compresses by 0.1 m. Assume no damping.

- a) Find k .
- b) How far does the spring compress when a 10000 kg railcar hits the spring at the same speed?
- c) If the spring would break if it compresses further than 0.3 m, what is the maximum mass of a railcar that can hit it at 1 m/s ?
- d) What is the maximum mass of a railcar that can hit the spring without it breaking at 2 m/s ?

Exercise 2.4.104: A mass of m kg is on a spring with $k = 3 \text{ N/m}$ and $c = 2 \text{ Ns/m}$. Find the mass m_0 for which there is critical damping. If $m < m_0$, does the system oscillate or not, that is, is it underdamped or overdamped?

2.5 Nonhomogeneous equations

Note: 2 lectures, §3.5 in [EP], §3.5 and §3.6 in [BD]

2.5.1 Solving nonhomogeneous equations

We have solved linear constant-coefficient homogeneous equations. What about nonhomogeneous linear ODEs? For example, the equations for forced mechanical vibrations. That is, consider an equation such as

$$y'' + 5y' + 6y = 2x + 1. \quad (2.6)$$

We will write $Ly = 2x + 1$ when the exact form of the operator is not important. We solve (2.6) in the following manner. First, we find the general solution y_c to the *associated homogeneous equation*

$$y'' + 5y' + 6y = 0. \quad (2.7)$$

We call y_c the *complementary solution*. Next, we find a single *particular solution* y_p to (2.6) in some way. Then

$$y = y_c + y_p$$

is the general solution to (2.6). We have $Ly_c = 0$ and $Ly_p = 2x + 1$. As L is a *linear operator*, we verify that y is a solution: $Ly = L(y_c + y_p) = Ly_c + Ly_p = 0 + (2x + 1)$.

Let us see why we obtain the *general* solution. Let y_p and $\tilde{y}_p$ be two different particular solutions to (2.6). Write the difference as $w = y_p - \tilde{y}_p$. Then plug w into the left-hand side of the equation to get

$$w'' + 5w' + 6w = (y_p'' + 5y_p' + 6y_p) - (\tilde{y}_p'' + 5\tilde{y}_p' + 6\tilde{y}_p) = (2x + 1) - (2x + 1) = 0.$$

With the operator notation, the calculation becomes even clearer. As L is a linear operator,

$$Lw = L(y_p - \tilde{y}_p) = Ly_p - L\tilde{y}_p = (2x + 1) - (2x + 1) = 0.$$

So $w = y_p - \tilde{y}_p$ is a solution to (2.7), that is, $Lw = 0$. Any two solutions of (2.6) differ by a solution to the homogeneous equation (2.7). The solution $y = y_c + y_p$ includes *all* solutions to (2.6), since y_c is the general solution to the associated homogeneous equation.

Theorem 2.5.1. *Let $Ly = f(x)$ be a linear ODE (not necessarily constant-coefficient). Let y_c be the complementary solution (the general solution to the associated homogeneous equation $Ly = 0$) and let y_p be any particular solution to $Ly = f(x)$. Then the general solution to $Ly = f(x)$ is*

$$y = y_c + y_p.$$

The moral of the story is that we can find the particular solution in any old way. If we find a different particular solution (by a different method or by simply guessing), then we still get the same general solution. The formula may look different, and the constants we have to choose to satisfy the initial conditions may be different, but it is the same solution.

2.5.2 Undetermined coefficients

The trick is to somehow, in a smart way, guess one particular solution to (2.6). Note that $2x + 1$ is a polynomial, and the left-hand side of the equation will be a polynomial if we let y be a polynomial of the same degree. Let us try

$$y_p = Ax + B.$$

We plug y_p into the left-hand side to obtain

$$\begin{aligned} y_p'' + 5y_p' + 6y_p &= (Ax + B)'' + 5(Ax + B)' + 6(Ax + B) \\ &= 0 + 5A + 6Ax + 6B = 6Ax + (5A + 6B). \end{aligned}$$

So $6Ax + (5A + 6B) = 2x + 1$. Therefore, $A = 1/3$ and $B = -1/9$. That means $y_p = \frac{1}{3}x - \frac{1}{9} = \frac{3x-1}{9}$. Solving the complementary problem (exercise!), we get

$$y_c = C_1 e^{-2x} + C_2 e^{-3x}.$$

Hence the general solution to (2.6) is

$$y = C_1 e^{-2x} + C_2 e^{-3x} + \frac{3x-1}{9}.$$

Now suppose we are further given some initial conditions. For example, $y(0) = 0$ and $y'(0) = 1/3$. First find $y' = -2C_1 e^{-2x} - 3C_2 e^{-3x} + 1/3$. Then

$$0 = y(0) = C_1 + C_2 - \frac{1}{9}, \quad \frac{1}{3} = y'(0) = -2C_1 - 3C_2 + \frac{1}{3}.$$

We solve to get $C_1 = 1/3$ and $C_2 = -2/9$. The particular solution we want is

$$y = \frac{1}{3}e^{-2x} - \frac{2}{9}e^{-3x} + \frac{3x-1}{9} = \frac{3e^{-2x} - 2e^{-3x} + 3x - 1}{9}.$$

Exercise 2.5.1: Check that y really solves the equation (2.6) and the given initial conditions.

Note: A common mistake is to solve for constants using the initial conditions with y_c and only add the particular solution y_p after that. That will *not* work. You need to first compute $y = y_c + y_p$ and *only then* solve for the constants using the initial conditions.

Another important remark is that you should not forget the lower degree terms, even if they do not appear on the right-hand side. If the equation is $Ly = x^3 + 1$, you must try $y_p = Ax^3 + Bx^2 + Cx + D$, even though there is no x^2 nor x on the right-hand side of the equation. It is a general polynomial of degree 3 that must be tried.

A right-hand side consisting of exponentials, sines, and cosines can be handled similarly. For example,

$$y'' + 2y' + 2y = \cos(2x).$$

Let us find some y_p . We start by guessing the solution includes some multiple of $\cos(2x)$. We may have to also add a multiple of $\sin(2x)$ to our guess since derivatives of cosine are sines. We try

$$y_p = A \cos(2x) + B \sin(2x).$$

We plug y_p into the equation and we get

$$\underbrace{-4A \cos(2x) - 4B \sin(2x)}_{y_p''} + 2 \underbrace{(-2A \sin(2x) + 2B \cos(2x))}_{y_p'} + 2 \underbrace{(A \cos(2x) + B \sin(2x))}_{y_p} = \cos(2x),$$

or

$$(-4A + 4B + 2A) \cos(2x) + (-4B - 4A + 2B) \sin(2x) = \cos(2x).$$

The left-hand side must equal the right-hand side. Namely, $-4A + 4B + 2A = 1$ and $-4B - 4A + 2B = 0$. So $-2A + 4B = 1$ and $2A + B = 0$ and hence $A = -1/10$ and $B = 1/5$. So

$$y_p = A \cos(2x) + B \sin(2x) = \frac{-\cos(2x) + 2 \sin(2x)}{10}.$$

Similarly, if the right-hand side contains exponentials, we try exponentials. If

$$Ly = e^{3x},$$

we try $y_p = Ae^{3x}$ as our guess and try to solve for A .

When the right-hand side is a multiple of sines, cosines, exponentials, and polynomials, we can use the product rule for differentiation to come up with a guess. We need to guess a form for y_p such that Ly_p is of the same form, and has all the terms needed for the right-hand side. For example,

$$Ly = (1 + 3x^2) e^{-x} \cos(\pi x).$$

For this equation, we guess

$$y_p = (A + Bx + Cx^2) e^{-x} \cos(\pi x) + (D + Ex + Fx^2) e^{-x} \sin(\pi x).$$

We plug in and then hopefully get equations that we can solve for A, B, C, D, E , and F . As you can see this can make for a very long and tedious calculation very quickly. C'est la vie!

There is one hiccup in all this. It could be that our guess actually solves the associated homogeneous equation. That is, consider

$$y'' - 9y = e^{3x}.$$

We would love to guess $y_p = Ae^{3x}$, but if we plug this into the left-hand side of the equation, we get

$$y_p'' - 9y_p = 9Ae^{3x} - 9Ae^{3x} = 0 \neq e^{3x}.$$

There is no way to choose A to make the left-hand side e^{3x} . The trick is to multiply our guess by x to get rid of duplication with the complementary solution. First, we find y_c (solution to $Ly = 0$),

$$y_c = C_1 e^{-3x} + C_2 e^{3x}.$$

We note that the e^{3x} term is a duplicate with our desired guess. We modify our guess to $y_p = A x e^{3x}$ so that there is no duplication anymore. Let us try: $y_p' = A e^{3x} + 3A x e^{3x}$ and $y_p'' = 6A e^{3x} + 9A x e^{3x}$, so

$$y_p'' - 9y_p = 6A e^{3x} + 9A x e^{3x} - 9A x e^{3x} = 6A e^{3x}.$$

Thus $6A e^{3x}$ is supposed to equal e^{3x} . Hence, $6A = 1$ and so $A = 1/6$. We can now write the general solution as

$$y = y_c + y_p = C_1 e^{-3x} + C_2 e^{3x} + \frac{1}{6} x e^{3x}.$$

It is possible that multiplying by x does not get rid of all duplication. Consider,

$$y'' - 6y' + 9y = e^{3x}.$$

The complementary solution is $y_c = C_1 e^{3x} + C_2 x e^{3x}$. Guessing $y_p = A x e^{3x}$ does not get us anywhere. We want to guess $y_p = A x^2 e^{3x}$. Basically, we multiply our guess by x until all duplication is gone. *But no more!* Multiplying too many times will not work.

Also make sure to look for duplication in the entire expression you try. For the equation $y'' - y = x e^x$, where $y_c = C_1 e^x + C_2 e^{-x}$, it is not enough that $x e^x$ does not appear in y_c . Trying $y_p = A e^x + B x e^x$ will not work as e^x appears in y_c , you must try $y_p = A x e^x + B x^2 e^x$.

Finally, what if the right-hand side has several terms? For instance,

$$Ly = e^{2x} + \cos x.$$

In this case, we find u that solves $Lu = e^{2x}$ and v that solves $Lv = \cos x$ (that is, do each term separately). If we set $y_p = u + v$, then we find our desired particular solution. That is because L is linear; we have $Ly_p = L(u + v) = Lu + Lv = e^{2x} + \cos x$.

2.5.3 Variation of parameters

The method of undetermined coefficients works for many basic problems that crop up. But it does not work all the time. It only works when the right-hand side of the equation $Ly = f(x)$ has finitely many linearly independent derivatives, so that we can write a guess that consists of them all. Some equations are a bit tougher. Consider

$$y'' + y = \tan x.$$

Each new derivative of $\tan x$ looks completely different and cannot be written as a linear combination of the previous derivatives. If we start differentiating $\tan x$, we get:

$$\begin{aligned} \sec^2 x, \quad 2 \sec^2 x \tan x, \quad 4 \sec^2 x \tan^2 x + 2 \sec^4 x, \\ 8 \sec^2 x \tan^3 x + 16 \sec^4 x \tan x, \quad 16 \sec^2 x \tan^4 x + 88 \sec^4 x \tan^2 x + 16 \sec^6 x, \quad \dots \end{aligned}$$

This equation calls for a different method. We present the method of *variation of parameters*, which handles any equation of the form $Ly = f(x)$, provided we can solve certain integrals. For simplicity, we restrict ourselves to second-order constant-coefficient equations, but the method works for higher-order equations just as well (the computations become more tedious). The method also works for equations with nonconstant coefficients, provided we can solve the associated homogeneous equation.

The details below will work for any equation of the form $y'' + p(x)y' + q(x)y = f(x)$, but perhaps it is best to explain the method with a specific example. Consider the equation

$$Ly = y'' + y = \tan x.$$

First we find the complementary solution (solution to $Ly_c = 0$). We get $y_c = C_1y_1 + C_2y_2$, where $y_1 = \cos x$ and $y_2 = \sin x$. To find a particular solution to the nonhomogeneous equation, the trick to this method is to try

$$y_p = y = u_1y_1 + u_2y_2,$$

where u_1 and u_2 are *functions* and not constants. We are trying to satisfy $Ly = \tan x$. That gives us one condition on the functions u_1 and u_2 . Compute (note the product rule!)

$$y' = (u_1'y_1 + u_2'y_2) + (u_1y_1' + u_2y_2').$$

We can still impose one more condition at our discretion to simplify computations (we have two unknown functions, so we should be allowed two conditions). We require that $u_1'y_1 + u_2'y_2 = 0$ (the first term above). This makes computing the second derivative easier:

$$\begin{aligned} y' &= u_1y_1' + u_2y_2', \\ y'' &= (u_1'y_1' + u_2'y_2') + (u_1y_1'' + u_2y_2''). \end{aligned}$$

Since y_1 and y_2 are solutions to $y'' + y = 0$, we find $y_1'' = -y_1$ and $y_2'' = -y_2$. (If the equation were the more general $y'' + p(x)y' + q(x)y = 0$, we would have $y_i'' = -p(x)y_i' - q(x)y_i$.) So

$$y'' = (u_1'y_1' + u_2'y_2') - (u_1y_1 + u_2y_2).$$

We have $u_1y_1 + u_2y_2 = y$ and so

$$y'' = (u_1'y_1' + u_2'y_2') - y,$$

and hence

$$y'' + y = Ly = u_1'y_1' + u_2'y_2'.$$

For y to satisfy $Ly = f(x)$, we must have $f(x) = u_1'y_1' + u_2'y_2'$.

What we need to solve are the two equations (conditions) we imposed on u_1 and u_2 :

$$\begin{aligned} u_1'y_1 + u_2'y_2 &= 0, \\ u_1'y_1' + u_2'y_2' &= f(x). \end{aligned}$$

We get these same equations for any $Ly = f(x)$, where $Ly = y'' + p(x)y' + q(x)y$. We solve for u'_1 and u'_2 in terms of $f(x)$, y_1 , and y_2 . There is a general formula for the solution we could just plug into, but instead of memorizing that, it is easier to simply solve it as we do below. In our case, $y_1 = \cos x$, $y_2 = \sin x$, and $f(x) = \tan x$, so the two equations are

$$\begin{aligned}u'_1 \cos x + u'_2 \sin x &= 0, \\ -u'_1 \sin x + u'_2 \cos x &= \tan x.\end{aligned}$$

Multiply the first equation by $\sin x$ and the second by $\cos x$:

$$\begin{aligned}u'_1 \cos x \sin x + u'_2 \sin^2 x &= 0, \\ -u'_1 \sin x \cos x + u'_2 \cos^2 x &= \tan x \cos x = \sin x.\end{aligned}$$

Add the two equations to eliminate u'_1 , solve for u'_2 , and then solve for u'_1 :

$$\begin{aligned}u'_2(\sin^2 x + \cos^2 x) &= \sin x, \\ u'_2 &= \sin x, \\ u'_1 &= \frac{-\sin^2 x}{\cos x} = \cos x - \sec x.\end{aligned}$$

We integrate u'_1 and u'_2 to get u_1 and u_2 .

$$\begin{aligned}u_1 &= \int u'_1 dx = \int (\cos x - \sec x) dx = \sin x - \ln|\sec x + \tan x|, \\ u_2 &= \int u'_2 dx = \int \sin x dx = -\cos x.\end{aligned}$$

We are looking for a particular solution, so we forget about the constants of integration. So our particular solution is

$$\begin{aligned}y_p &= u_1 y_1 + u_2 y_2 = \cos x \sin x - \cos x \ln|\sec x + \tan x| - \cos x \sin x \\ &= -\cos x \ln|\sec x + \tan x|.\end{aligned}$$

The general solution to $y'' + y = \tan x$ is, therefore,

$$y = C_1 \cos x + C_2 \sin x - \cos x \ln|\sec x + \tan x|.$$

So the general idea for any $y'' + py' + qy = f(x)$ is to first find solutions y_1, y_2 to $y'' + py' + qy = 0$. Then solve the two boxed equations for u'_1 and u'_2 , that is, solve $u'_1 y_1 + u'_2 y_2 = 0$ and $u'_1 y'_1 + u'_2 y'_2 = f(x)$. Integrate u'_1 and u'_2 to find u_1 and u_2 , and plug those into $y = u_1 y_1 + u_2 y_2$ to find the particular solution. We remark that if y'' has some coefficient that is not 1, that is, if the equation is $ay'' + by' + cy = f(x)$, you must first divide the equation (and hence f) by a .

2.5.4 Exercises

Exercise 2.5.2: Find a particular solution of $y'' - y' - 6y = e^{2x}$.

Exercise 2.5.3: Find a particular solution of $y'' - 4y' + 4y = e^{2x}$.

Exercise 2.5.4: Solve the initial value problem $y'' + 9y = \cos(3x) + \sin(3x)$ for $y(0) = 2$, $y'(0) = 1$.

Exercise 2.5.5: Set up the form of the particular solution but do not solve for the coefficients for $y^{(4)} - 2y''' + y'' = e^x$.

Exercise 2.5.6: Set up the form of the particular solution but do not solve for the coefficients for $y^{(4)} - 2y''' + y'' = e^x + x + \sin x$.

Exercise 2.5.7:

- Using variation of parameters, find a particular solution of $y'' - 2y' + y = e^x$.
- Find a particular solution using undetermined coefficients.
- Are the two solutions you found the same? See also [Exercise 2.5.10](#).

Exercise 2.5.8: Find a particular solution of $y'' - 2y' + y = \sin(x^2)$. It is OK to leave the answer as a definite integral.

Exercise 2.5.9: For an arbitrary constant c find a particular solution to $y'' - y = e^{cx}$. Hint: Make sure to handle every possible real c .

Exercise 2.5.10:

- Using variation of parameters, find a particular solution of $y'' - y = e^x$.
- Find a particular solution using undetermined coefficients.
- Are the two solutions you found the same? What is going on?

Exercise 2.5.11: Find a polynomial $P(x)$, so that $y = 2x^2 + 3x + 4$ solves $y'' + 5y' + y = P(x)$.

Exercise 2.5.101: Find a particular solution to $y'' - y' + y = 2\sin(3x)$.

Exercise 2.5.102:

- Find a particular solution to $y'' + 2y = e^x + x^3$.
- Find the general solution.

Exercise 2.5.103: Solve $y'' + 2y' + y = x^2$, $y(0) = 1$, $y'(0) = 2$.

Exercise 2.5.104: Use variation of parameters to find a particular solution of $y'' - y = \frac{1}{e^x + e^{-x}}$.

Exercise 2.5.105: For an arbitrary constant c find the general solution to $y'' - 2y = \sin(x + c)$.

Exercise 2.5.106: Undetermined coefficients can sometimes be used to guess a particular solution to other equations than those with constant coefficients. Find a polynomial $y(x)$ that solves $y' + xy = x^3 + 2x^2 + 5x + 2$.

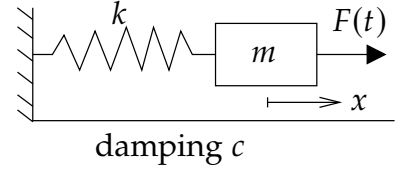
Note: Not every right-hand side will allow a polynomial solution, for example, $y' + xy = 1$ does not, but a technique based on undetermined coefficients does work, see [chapter 7](#).

2.6 Forced oscillations and resonance

Note: 2 lectures, §3.6 in [EP], §3.8 in [BD]

Let us return back to the example of a mass on a spring. We examine the case of forced oscillations, which we did not yet handle. That is, we consider the equation

$$mx'' + cx' + kx = F(t),$$



for some nonzero $F(t)$. The setup is again: x is position, m is mass, c is friction, k is the spring constant, and $F(t)$ is an external force acting on the mass.

We are interested in periodic forcing, such as noncentered rotating parts, or perhaps loud sounds, or other sources of periodic force. Using Fourier series, see [chapter 4](#), we note that we can understand all periodic functions by considering $F(t) = F_0 \cos(\omega t)$ (or sine instead of cosine, the calculations are essentially the same), so we focus on this simple case.

2.6.1 Undamped forced motion and resonance

First, let us consider undamped ($c = 0$) motion. We have the equation

$$mx'' + kx = F_0 \cos(\omega t).$$

This equation has the complementary solution (solution to the associated homogeneous equation)

$$x_c = C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t),$$

where $\omega_0 = \sqrt{k/m}$ is the *natural frequency* (angular). It is the frequency at which the system “wants to oscillate” without external interference.

Suppose that $\omega_0 \neq \omega$. We solve using the method of undetermined coefficients. We try the solution $x_p = A \cos(\omega t)$ and solve for A . We do not need a sine in our trial solution as after plugging in we only have cosines. If you include a sine, it is fine; you will find that its coefficient is zero (I could not find a second rhyme). We plug into the equation and solve for A to find

$$x_p = \frac{F_0}{m(\omega_0^2 - \omega^2)} \cos(\omega t).$$

We leave it as an exercise to do the algebra required.

The general solution is

$$x = C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t) + \frac{F_0}{m(\omega_0^2 - \omega^2)} \cos(\omega t).$$

Written another way

$$x = C \cos(\omega_0 t - \gamma) + \frac{F_0}{m(\omega_0^2 - \omega^2)} \cos(\omega t).$$

The solution is a superposition of two (shifted) cosine waves at different frequencies.

Example 2.6.1: Take

$$0.5x'' + 8x = 10 \cos(\pi t), \quad x(0) = 0, \quad x'(0) = 0.$$

Let us compute. First we read off the parameters: $\omega = \pi$, $\omega_0 = \sqrt{8/0.5} = 4$, $F_0 = 10$, $m = 0.5$. The general solution is

$$x = C_1 \cos(4t) + C_2 \sin(4t) + \frac{20}{16 - \pi^2} \cos(\pi t).$$

Solve for C_1 and C_2 using the initial conditions: $C_1 = \frac{-20}{16 - \pi^2}$ and $C_2 = 0$. Hence

$$x = \frac{20}{16 - \pi^2} (\cos(\pi t) - \cos(4t)).$$

Do notice the “beating” behavior in Figure 2.5. To analyze it, we use the trigonometric identity

$$2 \sin\left(\frac{A - B}{2}\right) \sin\left(\frac{A + B}{2}\right) = \cos B - \cos A$$

to get

$$x = \frac{20}{16 - \pi^2} \left(2 \sin\left(\frac{4 - \pi}{2}t\right) \sin\left(\frac{4 + \pi}{2}t\right) \right).$$

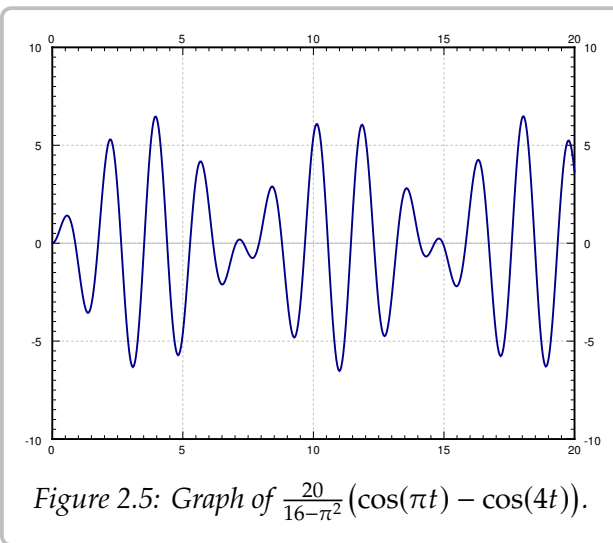


Figure 2.5: Graph of $\frac{20}{16 - \pi^2} (\cos(\pi t) - \cos(4t))$.

The function x is a high-frequency wave modulated by a low-frequency wave.

Now suppose $\omega_0 = \omega$. We notice that $\cos(\omega t)$ solves the associated homogeneous equation. Hence, we cannot try the solution $A \cos(\omega t)$ with the method of undetermined coefficients. Therefore, we try $x_p = At \cos(\omega t) + Bt \sin(\omega t)$. This time we do need the sine term, since the second derivative of $t \cos(\omega t)$ contains sines. We write the equation

$$x'' + \omega^2 x = \frac{F_0}{m} \cos(\omega t).$$

Plugging x_p into the left-hand side, we get

$$2B\omega \cos(\omega t) - 2A\omega \sin(\omega t) = \frac{F_0}{m} \cos(\omega t).$$

Hence $A = 0$ and $B = \frac{F_0}{2m\omega}$. Our particular solution is $\frac{F_0}{2m\omega} t \sin(\omega t)$ and the general solution is

$$x = C_1 \cos(\omega t) + C_2 \sin(\omega t) + \frac{F_0}{2m\omega} t \sin(\omega t).$$

The important term is the last one (the particular solution we found). This term grows without bound as $t \rightarrow \infty$. In fact it oscillates between $\frac{F_0 t}{2m\omega}$ and $-\frac{F_0 t}{2m\omega}$. The first two terms only oscillate between $\pm \sqrt{C_1^2 + C_2^2}$, which becomes smaller and smaller in proportion to the oscillations of the last term as t gets larger. In Figure 2.6 on the facing page, we see the graph with $C_1 = C_2 = 0$, $F_0 = 2$, $m = 1$, $\omega = \pi$.

By forcing the system at just the right frequency, we produce very wild oscillations. This kind of behavior is called *resonance* or perhaps *pure resonance*. Sometimes resonance is desired. For example, remember when as a kid you could start swinging by just moving back and forth on the swing seat in the “correct frequency”? You were trying to achieve resonance. The force of each one of your moves was small, but after a while it produced large swings.

On the other hand, resonance can be destructive. In an earthquake, some buildings collapse while others may be relatively undamaged. This is due to different buildings having different resonance frequencies. So figuring out the resonance frequency can be very important.

A common (but wrong) example of the destructive force of resonance is the Tacoma Narrows bridge failure. It turns out there was a different phenomenon at play*.

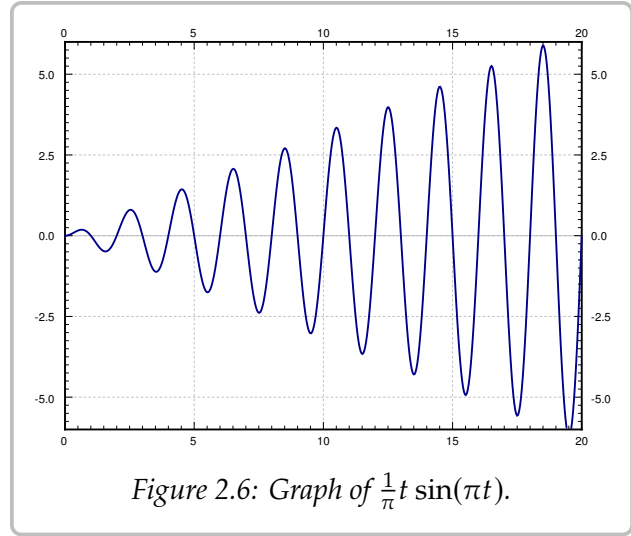


Figure 2.6: Graph of $\frac{1}{\pi}t \sin(\pi t)$.

2.6.2 Damped forced motion and practical resonance

In real life, things are not as simple as they were above. There is, of course, some damping. Our equation becomes

$$mx'' + cx' + kx = F_0 \cos(\omega t), \quad (2.8)$$

for some $c > 0$. We solved the homogeneous problem before. We let

$$p = \frac{c}{2m}, \quad \omega_0 = \sqrt{\frac{k}{m}}.$$

We replace equation (2.8) with

$$x'' + 2px' + \omega_0^2 x = \frac{F_0}{m} \cos(\omega t).$$

The roots of the characteristic equation of the associated homogeneous problem are $r_1, r_2 = -p \pm \sqrt{p^2 - \omega_0^2}$. The form of the general solution of the associated homogeneous equation depends on the sign of $p^2 - \omega_0^2$, or equivalently on the sign of $c^2 - 4km$, as before:

$$x_c = \begin{cases} C_1 e^{r_1 t} + C_2 e^{r_2 t} & \text{if } c^2 > 4km, \\ C_1 e^{-pt} + C_2 t e^{-pt} & \text{if } c^2 = 4km, \\ e^{-pt} (C_1 \cos(\omega_1 t) + C_2 \sin(\omega_1 t)) & \text{if } c^2 < 4km, \end{cases}$$

where $\omega_1 = \sqrt{\omega_0^2 - p^2}$. In any case, we see that $x_c(t) \rightarrow 0$ as $t \rightarrow \infty$.

*K. Billah and R. Scanlan, *Resonance, Tacoma Narrows Bridge Failure, and Undergraduate Physics Textbooks*, American Journal of Physics, 59(2), 1991, 118–124, <http://www.ketchum.org/billah/Billah-Scanlan.pdf>

Let us find a particular solution. There can be no conflicts when trying to solve for the undetermined coefficients by trying $x_p = A \cos(\omega t) + B \sin(\omega t)$. We plug in and solve for A and B . We get (the tedious details are left to reader)

$$((\omega_0^2 - \omega^2)B - 2\omega pA) \sin(\omega t) + ((\omega_0^2 - \omega^2)A + 2\omega pB) \cos(\omega t) = \frac{F_0}{m} \cos(\omega t).$$

We solve for A and B as usual:

$$A = \frac{(\omega_0^2 - \omega^2)F_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2}, \quad B = \frac{2\omega pF_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2}.$$

Thus our particular solution is

$$x_p = \frac{(\omega_0^2 - \omega^2)F_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2} \cos(\omega t) + \frac{2\omega pF_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2} \sin(\omega t).$$

The amplitude $C = \sqrt{A^2 + B^2}$ of this solution is

$$C = \frac{F_0}{m\sqrt{(2\omega p)^2 + (\omega_0^2 - \omega^2)^2}}.$$

Let us write the solution in the alternative form $x_p = C \cos(\omega t - \gamma)$ for amplitude C and phase shift γ where (if $\omega \neq \omega_0$)

$$\tan \gamma = \frac{B}{A} = \frac{2\omega p}{\omega_0^2 - \omega^2}.$$

Hence,

$$x_p = \frac{F_0}{m\sqrt{(2\omega p)^2 + (\omega_0^2 - \omega^2)^2}} \cos(\omega t - \gamma).$$

If $\omega = \omega_0$, then $A = 0$, $B = C = \frac{F_0}{2m\omega p}$, and $\gamma = \pi/2$.

For reasons we will explain in a moment, we call x_c the *transient solution* and denote it by x_{tr} . We call x_p the *steady periodic solution* and denote it by x_{sp} . The general solution is

$$x = x_c + x_p = x_{tr} + x_{sp}.$$

The transient solution $x_{tr} = x_c$ goes to zero as $t \rightarrow \infty$, as all the terms involve an exponential with a negative exponent. So for large t , the effect of x_{tr} is negligible and we see essentially only x_{sp} . Hence the name *transient*. Notice that x_{sp} involves no arbitrary constants, and the initial conditions only affect x_{tr} . Thus, the effect of the initial conditions is negligible after some period of time. We might as well focus on the steady periodic solution and ignore the transient solution. See [Figure 2.7](#) on the next page for a graph given several different initial conditions.

How fast x_{tr} goes to zero depends on p (and hence c). The bigger p is (the bigger c is), the “faster” x_{tr} becomes negligible. Conversely, the smaller the damping, the longer the “transient region.” This is consistent with the observation that when $c = 0$, the initial conditions affect the behavior for all time (i.e. an infinite “transient region”).

Let us describe what we mean by resonance when damping is present. There were no conflicts solving with undetermined coefficients, so no term in x_{sp} goes to infinity. Instead, we consider the amplitude C of the steady periodic solution x_{sp} , that is, the furthest the mass goes either way from the rest position. We plot C as a function of ω (fix all other parameters) and we find its maximum. The ω that gives this maximum is the *practical resonance frequency*. The maximal amplitude $C(\omega)$ is the *practical resonance amplitude*. Thus when damping is present we talk of *practical resonance* rather than pure resonance. Figure 2.8 shows a sample plot for three different values of c . Notice that the practical resonance amplitude grows as damping gets smaller, and practical resonance can disappear altogether when damping is large.

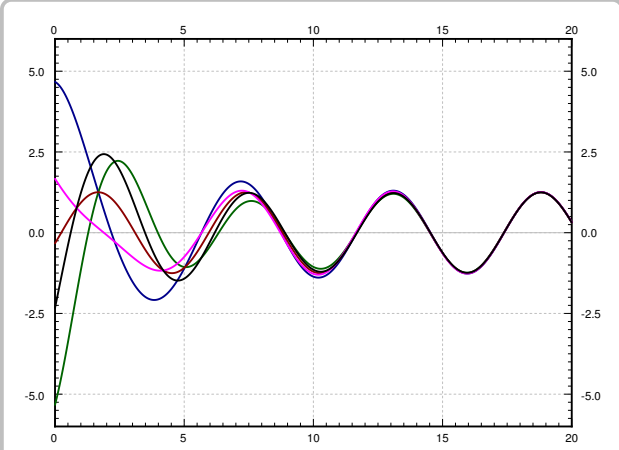


Figure 2.7: Solutions with different initial conditions for parameters $k = 1$, $m = 1$, $F_0 = 1$, $c = 0.7$, and $\omega = 1.1$.

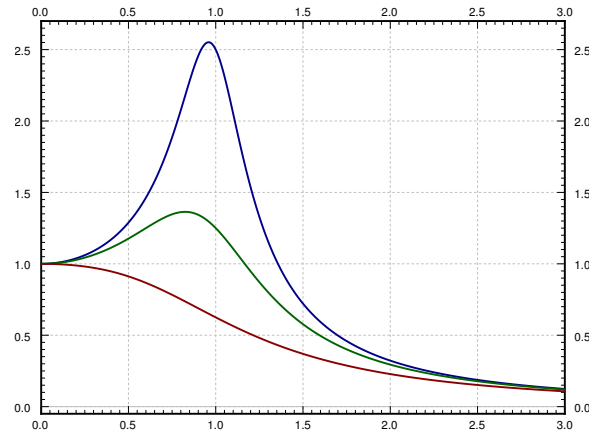


Figure 2.8: Graph of $C(\omega)$ showing practical resonance with parameters $k = 1$, $m = 1$, $F_0 = 1$. The top line is with $c = 0.4$, the middle line with $c = 0.8$, and the bottom line with $c = 1.6$.

To find the maximum of $C(\omega)$, we find the derivative $C'(\omega)$:

$$C'(\omega) = \frac{-2\omega(2p^2 + \omega^2 - \omega_0^2)F_0}{m((2\omega p)^2 + (\omega_0^2 - \omega^2)^2)^{3/2}}.$$

This is zero either when $\omega = 0$ or when $2p^2 + \omega^2 - \omega_0^2 = 0$. In other words, $C'(\omega) = 0$ when

$$\omega = \sqrt{\omega_0^2 - 2p^2} \quad \text{or} \quad \omega = 0.$$

If $\omega_0^2 - 2p^2$ is positive, then $\sqrt{\omega_0^2 - 2p^2}$ is the practical resonance frequency (the point where $C(\omega)$ is maximal). This conclusion follows by the first derivative test, for example, as then $C'(\omega) > 0$ for small ω in this case. If on the other hand $\omega_0^2 - 2p^2$ is not positive, then $C(\omega)$ achieves its maximum at $\omega = 0$, and there is no practical resonance since we assume $\omega > 0$ in our system. In this case, the amplitude gets larger as the forcing frequency gets smaller.

If practical resonance occurs, the frequency is smaller than ω_0 . As the damping c (and hence p) becomes smaller, the practical resonance frequency goes to ω_0 . So when damping is very small, ω_0 is a good estimate of the practical resonance frequency. This behavior agrees with the observation that when $c = 0$, then ω_0 is the resonance frequency.

Another interesting observation to make is that when $\omega \rightarrow \infty$, then $C \rightarrow 0$. This means that if the forcing frequency gets too high it does not manage to get the mass moving in the mass-spring system. This is quite reasonable intuitively. If we wiggle back and forth really fast while sitting on a swing, we will not get it moving at all, no matter how forceful. Fast vibrations just cancel each other out before the mass has any chance of responding by moving one way or the other.

The behavior is more complicated if the forcing function is not an exact cosine wave, but for example a square wave. A general periodic function will be the sum (superposition) of many cosine waves of different frequencies. The reader is encouraged to come back to this section once we have learned about the Fourier series.

2.6.3 Exercises

Exercise 2.6.1: Derive a formula for x_{sp} if the equation is $mx'' + cx' + kx = F_0 \sin(\omega t)$. Assume $c > 0$.

Exercise 2.6.2: Derive a formula for x_{sp} if the equation is $mx'' + cx' + kx = F_0 \cos(\omega t) + F_1 \cos(3\omega t)$. Assume $c > 0$.

Exercise 2.6.3: Take $mx'' + cx' + kx = F_0 \cos(\omega t)$. Fix $m > 0$, $k > 0$, and $F_0 > 0$. Consider the function $C(\omega)$. For what values of c (solve in terms of m , k , and F_0) will there be no practical resonance (that is, for what values of c is there no maximum of $C(\omega)$ for $\omega > 0$)?

Exercise 2.6.4: Take $mx'' + cx' + kx = F_0 \cos(\omega t)$. Fix $c > 0$, $k > 0$, and $F_0 > 0$. Consider the function $C(\omega)$. For what values of m (solve in terms of c , k , and F_0) will there be no practical resonance (that is, for what values of m is there no maximum of $C(\omega)$ for $\omega > 0$)?

Exercise 2.6.5: A water tower in an earthquake acts as a mass-spring system. Assume that the container on top is full and the water does not move around. The container then acts as the mass and the support acts as the spring, where the induced vibrations are horizontal. The container with water has a mass of $m = 10,000$ kg. It takes a force of 1000 newtons to displace the container 1 meter. For simplicity, assume no friction. When the earthquake hits, the water tower is at rest (it is not moving). The earthquake induces an external force $F(t) = mA\omega^2 \cos(\omega t)$ newtons.

- What is the natural frequency of the water tower?
- If ω is not the natural frequency, find a formula for the maximal amplitude of the resulting oscillations of the water container (the maximal deviation from the rest position). The motion will be a high-frequency wave modulated by a low-frequency wave, so simply find the constant in front of the sines.
- Suppose $A = 1$ and an earthquake with frequency 0.5 cycles per second comes. What is the amplitude of the oscillations? Suppose that if the water tower moves more than 1.5 meters from the rest position, the tower collapses. Will the tower collapse?

Exercise 2.6.101: A mass of 4 kg on a spring with $k = 4$ N/m and a damping constant $c = 1$ Ns/m. Suppose that $F_0 = 2$ N. Using the forcing function $F_0 \cos(\omega t)$, find the ω that causes practical resonance and find the practical resonance amplitude.

Exercise 2.6.102: Derive a formula for x_{sp} for $mx'' + cx' + kx = F_0 \cos(\omega t) + A$, where A is some constant. Assume $c > 0$.

Exercise 2.6.103: Suppose there is no damping in a mass and spring system with $m = 5$ kg, $k = 20$ N/m, and $F_0 = 5$ N. Suppose ω is chosen to be precisely the resonance frequency.

- Find ω .
- Find the amplitude of the oscillations at time $t = 100$ seconds, given the system is at rest at $t = 0$.

Chapter 3

Systems of ODEs

3.1 Introduction to systems of ODEs

Note: 1 to 1.5 lectures, §4.1 in [EP], §7.1 in [BD]

3.1.1 Systems

Often we do not have just one dependent variable and one equation. We will see that we may end up with a system of several equations and several dependent variables even if we start with a single equation.

Given several dependent variables, suppose $y_1, y_2, \dots, y_n$, we can have a differential equation involving all of them and their derivatives with respect to one independent variable x . For example, $y_1'' = f(y_1', y_2', y_1, y_2, x)$. Usually, when we have two dependent variables, we have two equations such as

$$\begin{aligned}y_1'' &= f_1(y_1', y_2', y_1, y_2, x), \\y_2'' &= f_2(y_1', y_2', y_1, y_2, x),\end{aligned}$$

for some known functions f_1 and f_2 . We call the above a *system of differential equations*. More precisely, the above is a *second-order system* of ODEs as second-order derivatives appear. An example of a *first-order system* is

$$\begin{aligned}x_1' &= g_1(x_1, x_2, x_3, t), \\x_2' &= g_2(x_1, x_2, x_3, t), \\x_3' &= g_3(x_1, x_2, x_3, t),\end{aligned}$$

where x_1, x_2, x_3 are the dependent variables, and t is the independent variable.

The terminology for systems is essentially the same as for single equations. For the system above, a *solution* is a set of three functions $x_1(t), x_2(t), x_3(t)$, such that

$$\begin{aligned}x_1'(t) &= g_1(x_1(t), x_2(t), x_3(t), t), \\x_2'(t) &= g_2(x_1(t), x_2(t), x_3(t), t), \\x_3'(t) &= g_3(x_1(t), x_2(t), x_3(t), t).\end{aligned}$$

We may also have an *initial condition*. As for single equations, we specify x_1 , x_2 , and x_3 for some fixed t . For example, $x_1(0) = a_1$, $x_2(0) = a_2$, $x_3(0) = a_3$, where a_1 , a_2 , and a_3 are some constants. For a second-order system, we must also specify the first derivatives at the initial point. If we find a solution with arbitrary constants in it, where solving for the constants gives a solution for any initial condition, we call this solution the *general solution*. Best to look at a simple example.

Example 3.1.1: Sometimes a system is easy to solve by solving for one variable and then for the second variable. Take the first-order system

$$\begin{aligned}y_1' &= y_1, \\y_2' &= y_1 - y_2,\end{aligned}$$

with y_1, y_2 as the dependent variables and x as the independent variable. Consider initial conditions $y_1(0) = 1$, $y_2(0) = 2$.

We note that $y_1 = C_1 e^x$ is the general solution of the first equation. We then plug this y_1 into the second equation and get the equation $y_2' = C_1 e^x - y_2$, which is a linear first-order equation that is easily solved for y_2 . By the integrating factor method, we get

$$e^x y_2 = \frac{C_1}{2} e^{2x} + C_2,$$

or $y_2 = \frac{C_1}{2} e^x + C_2 e^{-x}$. The general solution to the system is, therefore,

$$y_1 = C_1 e^x, \quad y_2 = \frac{C_1}{2} e^x + C_2 e^{-x}.$$

We solve for C_1 and C_2 given the initial conditions. We substitute $x = 0$ to find $1 = y_1(0) = C_1$ and $2 = y_2(0) = C_1/2 + C_2$, or in other words, $C_1 = 1$ and $C_2 = 3/2$. Thus the solution is $y_1 = e^x$, and $y_2 = (1/2)e^x + (3/2)e^{-x}$.

Generally, we will not be so lucky to be able to solve for each variable separately as in the example—we will need to solve for all variables at once. While we cannot always solve one variable at a time, we will try to salvage as much as possible from this technique. In a certain sense, to solve systems, we will still (try to) solve a bunch of single equations and put their solutions together. Let us not worry right now about how to solve systems yet.

We will mostly consider *linear systems*. [Example 3.1.1](#) is a *linear first-order system*. It is linear as none of the dependent variables or their derivatives appear in nonlinear functions or with powers higher than one (y_1, y_2, y_1', y_2' , constants, and functions of x can appear, but not $y_1 y_2$ or $(y_2')^2$ or y_1^3). Another, more complicated, example of a linear (second-order) system is

$$\begin{aligned}y_1'' &= e^x y_1' + x^2 y_1 + 5y_2 + \sin(x), \\y_2'' &= x y_1' - y_2' + 2y_1 + \cos(x).\end{aligned}$$

3.1.2 Applications

We consider some simple applications of systems and how to set up the equations.

Example 3.1.2: First, we consider salt and brine tanks, but this time water flows from one to the other and back. Imagine we have two tanks, each containing volume V liters of salt brine. The amount of salt in the first tank is x_1 grams, and the amount of salt in the second tank is x_2 grams. Let t denote time—the independent variable. The liquid in each tank is being constantly and perfectly mixed and flows or is pumped at a rate of r liters per second out of each tank into the other. See Figure 3.1.

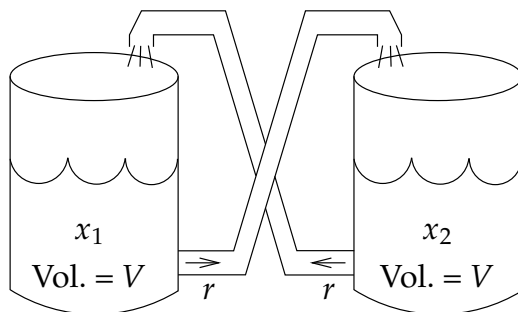


Figure 3.1: A closed system of two brine tanks.

The rate of change of x_1 , that is, x'_1 , is the rate of salt entering minus the rate leaving. The density of the salt in tank 2 is x_2/V , so the rate of salt entering tank 1 is x_2/V times r . The density of the salt in tank 1 is x_1/V , so the rate of salt leaving tank 1 is x_1/V times r . In other words,

$$x'_1 = \frac{x_2}{V}r - \frac{x_1}{V}r = \frac{r}{V}x_2 - \frac{r}{V}x_1 = \frac{r}{V}(x_2 - x_1).$$

Similarly, to find the rate x'_2 , the roles of x_1 and x_2 are reversed. All in all, the system of ODEs for this problem is

$$\begin{aligned} x'_1 &= \frac{r}{V}(x_2 - x_1), \\ x'_2 &= \frac{r}{V}(x_1 - x_2). \end{aligned}$$

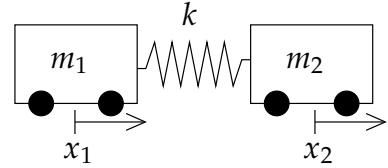
In this system, we cannot solve for x_1 or x_2 separately. We must solve for both x_1 and x_2 at once, which is intuitively clear—the amount of salt in one tank affects the amount in the other. We cannot know x_1 before we know x_2 , and vice versa.

We do not yet know how to find all the solutions, but intuitively we can at least find some solutions. Suppose we know that initially the tanks have the same amount of salt. That is, we have an initial condition such as $x_1(0) = x_2(0) = C$. Then clearly the amount of salt entering and leaving each tank is the same, so the amounts are not changing. In other words, $x_1 = C$ and $x_2 = C$ (the constant functions) is a solution: $x'_1 = x'_2 = 0$, and $\frac{r}{V}(x_2 - x_1) = \frac{r}{V}(x_1 - x_2) = 0$, so the equations are satisfied.

Let us think about the setup a little bit more without solving it. Suppose the initial conditions are $x_1(0) = A$ and $x_2(0) = B$, for two different constants A and B . Since no salt is coming in or out of this closed system, the total amount of salt is constant. That is, $x_1 + x_2$ is constant, and so it equals $A + B$. Intuitively, if $A > B$, more salt will flow out of tank one than into it. After a long time, we then expect the amount of salt in each tank to equalize. In other words, the solutions of both x_1 and x_2 should tend towards $\frac{A+B}{2}$ as t goes to ∞ . Once you know how to solve systems, you can check that this really is so.

Example 3.1.3: Let us look at a second-order example. We return to the mass and spring setup, but this time we consider two masses.

Consider one spring with constant k and two masses m_1 and m_2 . Think of the masses as carts on a straight track with no friction. Let x_1 be the displacement of the first cart and x_2 be the displacement of the second cart. That is, we put the two carts somewhere with no tension on the spring, we mark the position of the first and second cart, and we call those the zero positions. Then x_1 measures how far the first cart is from its zero position, and x_2 measures how far the second cart is from its zero position. The force exerted by the spring on the first cart is $k(x_2 - x_1)$ as $x_2 - x_1$ is how far the spring is stretched (or compressed) from the rest position. The force exerted on the second cart is the opposite, thus the same thing with a negative sign. Newton's second law states that force equals mass times acceleration, that is,



$$\begin{aligned} m_1 x_1'' &= k(x_2 - x_1), \\ m_2 x_2'' &= -k(x_2 - x_1). \end{aligned}$$

Again, we cannot solve for the x_1 or x_2 variables one at a time. Where the first cart goes depends on exactly where the second cart goes and vice versa.

3.1.3 Changing to first-order systems

To some degree, we need only be able to solve first-order systems. Consider an n^{th} -order differential equation

$$y^{(n)} = F(y^{(n-1)}, \dots, y', y, x).$$

We define new variables $u_1, u_2, \dots, u_n$ and write the system

$$\begin{aligned} u_1' &= u_2, \\ u_2' &= u_3, \\ &\vdots \\ u_{n-1}' &= u_n, \\ u_n' &= F(u_n, u_{n-1}, \dots, u_2, u_1, x). \end{aligned}$$

We solve this system for $u_1, u_2, \dots, u_n$. Once we have solved for the u , we can discard u_2 through u_n and let $y = u_1$. This y solves the original equation.

Example 3.1.4: Take $x''' = 2x'' + 8x' + x + t$. Letting $u_1 = x$, $u_2 = x'$, $u_3 = x''$, we find the system:

$$u_1' = u_2, \quad u_2' = u_3, \quad u_3' = 2u_3 + 8u_2 + u_1 + t.$$

Note why $x = u_1$ solves the original equation: The first two equations of the system give that $u_3' = u_2'' = u_1''' = x'''$. If we plug $u_3' = x'''$, $u_1 = x$, $u_2 = x'$, and $u_3 = x''$ into the third equation of the system, we recover the third-order equation we started with.

The same idea works for a system of higher-order differential equations. A system of k differential equations in k unknowns, all of order n , can be transformed into a first-order system of $n \times k$ equations and $n \times k$ unknowns.

Example 3.1.5: Consider the system from the example with carts, [Example 3.1.3](#):

$$m_1 x_1'' = k(x_2 - x_1), \quad m_2 x_2'' = -k(x_2 - x_1).$$

Let $u_1 = x_1$, $u_2 = x_1'$, $u_3 = x_2$, $u_4 = x_2'$. The second-order system becomes the first-order system

$$u_1' = u_2, \quad m_1 u_2' = k(u_3 - u_1), \quad u_3' = u_4, \quad m_2 u_4' = -k(u_3 - u_1).$$

Example 3.1.6: The idea works in reverse as well. Suppose we want to solve the system

$$x' = 2y - x, \quad y' = x,$$

for the initial conditions $x(0) = 1$, $y(0) = 0$. Let the independent variable be t .

If we differentiate the second equation, we get $y'' = x'$. We know what x' is in terms of x and y , and we know that $x = y'$. So,

$$y'' = x' = 2y - x = 2y - y'.$$

We now have the equation $y'' + y' - 2y = 0$. We know how to solve this equation and we find that $y = C_1 e^{-2t} + C_2 e^t$. Once we have y , we use the equation $y' = x$ to get x .

$$x = y' = -2C_1 e^{-2t} + C_2 e^t.$$

We solve for the initial conditions $1 = x(0) = -2C_1 + C_2$ and $0 = y(0) = C_1 + C_2$. Hence, $C_1 = -C_2$ and $1 = 3C_2$. So $C_1 = -1/3$ and $C_2 = 1/3$. Our solution is

$$x = \frac{2e^{-2t} + e^t}{3}, \quad y = \frac{-e^{-2t} + e^t}{3}.$$

Exercise 3.1.1: Plug in and check that this really is the solution.

It is useful to go back and forth between systems and higher-order equations for other reasons. For example, software for solving ODEs numerically (approximation) is generally for first-order systems. To use it, you take whatever ODE you want to solve and convert it to a first-order system. It is not very hard to adapt computer code for the Euler or Runge–Kutta method for first-order equations to handle first-order systems. We simply treat the dependent variable not as a number but as a vector. In many mathematical computer languages there is almost no distinction in syntax.

3.1.4 Autonomous systems and vector fields

A system where the equations do not depend on the independent variable is called an *autonomous system*. For example, the system $x' = 2y - x$, $y' = x$ is autonomous as the independent variable, say t , does not appear in the equations.

For autonomous systems, we can draw the so-called *direction field* or *vector field*, a plot similar to a slope field, but instead of giving a slope at each point, we give a direction (and a magnitude). The previous example, $x' = 2y - x$, $y' = x$, says that at the point (x, y) the direction in which we should travel to satisfy the equations should be the direction of the vector $(2y - x, x)$ with the speed equal to the magnitude of this vector. So we draw the vector $(2y - x, x)$ at the point (x, y) and we do this for many points on the xy -plane. For example, at the point $(1, 2)$, we draw the vector $(2(2) - 1, 1) = (3, 1)$, a vector pointing to the right and a little bit up, while at the point $(2, 1)$ we draw the vector $(2(1) - 2, 2) = (0, 2)$ a vector that points straight up. When drawing the vectors, we will scale down their size to fit many of them on the same direction field. If we drew the arrows at the actual size, the diagram would be a jumbled mess once we draw more than a couple of arrows. So we scale them all so that not even the longest one interferes with the others. We are mostly interested in their direction and relative size. See Figure 3.2 on the facing page. The diagrams we drew in § 1.6 for autonomous equations in one dimension are similar, but note how much more complicated things become when we allow just one extra dimension.

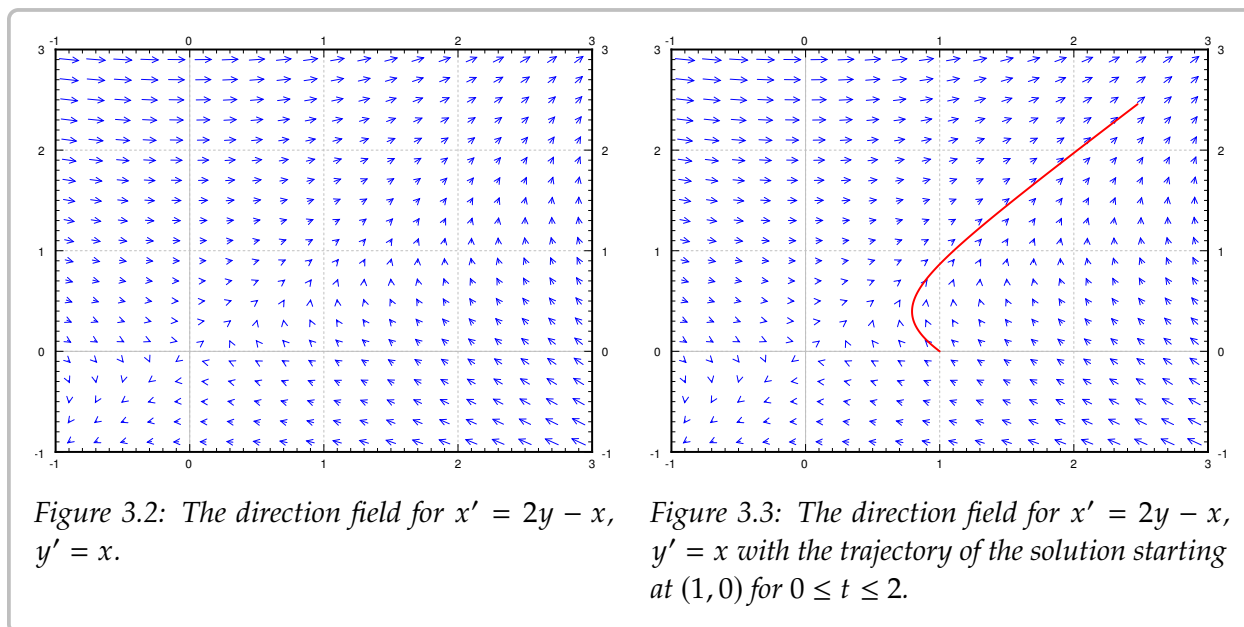
We can draw a path of the solution in the plane. Suppose the solution is given by $x = f(t)$, $y = g(t)$. We pick an interval of t (say $0 \leq t \leq 2$ for our example) and plot all the points $(f(t), g(t))$ for t in the selected range. The resulting picture is called the *phase portrait* (or phase plane portrait). The particular curve obtained is called the *trajectory* or *solution curve*. See an example plot in Figure 3.3 on the next page. In the figure the solution starts at $(1, 0)$ and travels along the vector field for a distance of 2 units of t . We solved this system precisely, so we compute $x(2)$ and $y(2)$ to find $x(2) \approx 2.475$ and $y(2) \approx 2.457$. This point corresponds to the top right end of the plotted solution curve in the figure.

We can draw phase portraits and trajectories in the xy -plane even if the system is not autonomous. In this case, however, we cannot draw the direction field, since the field changes as t changes. For each t we would get a different direction field.

3.1.5 Picard's theorem

Before going further, we mention that Picard's theorem on existence and uniqueness also holds for systems of ODEs. Let us restate this theorem in this setting. A general first-order system is of the form

$$\begin{aligned} x_1' &= F_1(x_1, x_2, \dots, x_n, t), \\ x_2' &= F_2(x_1, x_2, \dots, x_n, t), \\ &\vdots \\ x_n' &= F_n(x_1, x_2, \dots, x_n, t). \end{aligned} \tag{3.1}$$



Theorem 3.1.1 (Picard's theorem on existence and uniqueness for systems). *If for every $j = 1, 2, \dots, n$ and every $k = 1, 2, \dots, n$ each F_j is continuous and the derivative $\frac{\partial F_j}{\partial x_k}$ exists and is continuous near some $(x_1^0, x_2^0, \dots, x_n^0, t^0)$, then a solution to (3.1) subject to the initial condition $x_1(t^0) = x_1^0, x_2(t^0) = x_2^0, \dots, x_n(t^0) = x_n^0$ exists (at least for t in some small interval) and is unique.*

That is, a unique solution exists for any initial condition given that the system is reasonable (each F_j and its partial derivatives in the x variables are continuous). As for single equations, we may not have a solution for all time t , but it is guaranteed at least for some short period of time.

As we can change any n^{th} -order ODE into a first-order system, this theorem also provides the existence and uniqueness of solutions for higher-order equations.

3.1.6 Exercises

Exercise 3.1.2: Find the general solution of $x'_1 = x_2 - x_1 + t$, $x'_2 = x_2$.

Exercise 3.1.3: Find the general solution of $x'_1 = 3x_1 - x_2 + e^t$, $x'_2 = x_1$.

Exercise 3.1.4: Write $ay'' + by' + cy = f(x)$ as a first-order system of ODEs.

Exercise 3.1.5: Write $x'' + y^2y' - x^3 = \sin(t)$, $y'' + (x' + y')^2 - x = 0$ as a first-order system of ODEs.

Exercise 3.1.6: Suppose two masses on carts on frictionless surface are at displacements x_1 and x_2 as in [Example 3.1.3](#) on page 122. Suppose that a rocket applies force F in the positive direction on cart x_1 . Set up the system of equations.

Exercise 3.1.7: Suppose the tanks are as in [Example 3.1.2](#) on page 121, starting both at volume V , but now the rate of flow from tank 1 to tank 2 is r_1 , and rate of flow from tank 2 to tank one is r_2 . Notice that the volumes are now not constant. Set up the system of equations.

Exercise 3.1.101: Find the general solution to $y'_1 = 3y_1$, $y'_2 = y_1 + y_2$, $y'_3 = y_1 + y_3$.

Exercise 3.1.102: Solve $y' = 2x$, $x' = x + y$, $x(0) = 1$, $y(0) = 3$.

Exercise 3.1.103: Write $x''' = x + t$ as a first-order system.

Exercise 3.1.104: Write $y'_1 + y_1 + y_2 = t$, $y'_2 + y_1 - y_2 = t^2$ as a first-order system.

Exercise 3.1.105: Suppose two masses on carts on frictionless surface are at displacements x_1 and x_2 as in [Example 3.1.3](#) on page 122. Suppose initial displacement is $x_1(0) = x_2(0) = 0$, and initial velocity is $x'_1(0) = x'_2(0) = a$ for some number a . Use your intuition to solve the system, explain your reasoning.

Exercise 3.1.106: Suppose the tanks are as in [Example 3.1.2](#) on page 121 except that clean water flows in at a rate of s liters per second into tank 1, and brine flows out of tank 2 and into the sewer also at a rate of s liters per second. The rate of flow from tank 1 into tank 2 is still r , but the rate of flow from tank 2 back into tank 1 is $r - s$ (assume $r > s$).

- a) Draw the picture.
- b) Set up the system of equations.
- c) Intuitively, what happens as t goes to infinity, explain.

3.2 Matrices and linear systems

Note: 1.5 lectures, first part of §5.1 in [EP], §7.2 and §7.3 in [BD], see also [appendix A](#)

3.2.1 Matrices and vectors

Before we start talking about linear systems of ODEs, we need to talk about matrices, so let us review these briefly. A *matrix* is an $m \times n$ array of numbers (m rows and n columns). For example, we denote a 3×5 matrix as follows

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \end{bmatrix}.$$

The numbers a_{ij} are called *elements* or *entries*. We say a matrix is *square* if it has $m = n$, that is, it has the same number of rows and columns.

In the setting of matrices, by a *vector*, we usually mean a *column vector*, that is, an $m \times 1$ matrix. If we mean a *row vector*, we will explicitly say so (a row vector is a $1 \times n$ matrix). We usually denote matrices by upper case letters and vectors by lower case letters with an arrow such as $\vec{x}$ or $\vec{b}$. We write $\vec{0}$ for a vector of all zeros.

We define some operations on matrices. We want 1×1 matrices to really act like numbers, so our operations have to be compatible with this viewpoint. First, we can multiply a matrix by a *scalar* (a number). We simply multiply each entry in the matrix by the scalar. For example,

$$2 \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} = \begin{bmatrix} 2 & 4 & 6 \\ 8 & 10 & 12 \end{bmatrix}.$$

Matrix addition is also simple. We add matrices element by element. For example,

$$\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} + \begin{bmatrix} 1 & 1 & -1 \\ 0 & 2 & 4 \end{bmatrix} = \begin{bmatrix} 2 & 3 & 2 \\ 4 & 7 & 10 \end{bmatrix}.$$

If the sizes do not match, then addition is not defined. If we denote by 0 the matrix with all zero entries, by c, d scalars, and by A, B, C matrices, we have the following familiar rules:

$$\begin{aligned} A + 0 &= A = 0 + A, \\ A + B &= B + A, \\ (A + B) + C &= A + (B + C), \\ c(A + B) &= cA + cB, \\ (c + d)A &= cA + dA. \end{aligned}$$

Another useful operation for matrices is the so-called *transpose*. This operation just swaps rows and columns of a matrix. The transpose of A is denoted by A^T . Example:

$$\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix}^T = \begin{bmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix}$$

3.2.2 Matrix multiplication

Matrix multiplication is a bit more complicated. First we define the so-called *dot product* (or *inner product*) of two vectors. Usually this will be a row vector multiplied with a column vector of the same size. For the dot product, we multiply each pair of entries from the first and the second vector, and we sum these products. The result is a single number. For example,

$$\begin{bmatrix} a_1 & a_2 & a_3 \end{bmatrix} \cdot \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} = a_1b_1 + a_2b_2 + a_3b_3.$$

Similarly for larger (or smaller) vectors.

Armed with the dot product, we define the *product of matrices*. We denote by $\text{row}_i(A)$ the i^{th} row of A and by $\text{column}_j(A)$ the j^{th} column of A . For an $m \times n$ matrix A and an $n \times p$ matrix B , we can define the product AB . We let AB be an $m \times p$ matrix whose ij^{th} entry is the dot product

$$\text{row}_i(A) \cdot \text{column}_j(B).$$

Do note how the sizes match up: $m \times n$ multiplied by $n \times p$ is $m \times p$. Example:

$$\begin{aligned} \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} \begin{bmatrix} 1 & 0 & -1 \\ 1 & 1 & 1 \\ 1 & 0 & 0 \end{bmatrix} &= \\ &= \begin{bmatrix} 1 \cdot 1 + 2 \cdot 1 + 3 \cdot 1 & 1 \cdot 0 + 2 \cdot 1 + 3 \cdot 0 & 1 \cdot (-1) + 2 \cdot 1 + 3 \cdot 0 \\ 4 \cdot 1 + 5 \cdot 1 + 6 \cdot 1 & 4 \cdot 0 + 5 \cdot 1 + 6 \cdot 0 & 4 \cdot (-1) + 5 \cdot 1 + 6 \cdot 0 \end{bmatrix} = \begin{bmatrix} 6 & 2 & 1 \\ 15 & 5 & 1 \end{bmatrix} \end{aligned}$$

For multiplication, we want an analogue of a 1. This analogue is the so-called *identity matrix*. The identity matrix is a square matrix with 1s on the diagonal and zeros everywhere else. It is usually denoted by I . For each size, we have a different identity matrix. So sometimes we may denote the size as a subscript. For example, I_3 would be the 3×3 identity matrix

$$I = I_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

We have the following rules for matrix multiplication. Suppose that A, B, C are matrices of the correct sizes so that the following make sense. Let α denote a scalar (number). Then,

$$\begin{aligned} A(BC) &= (AB)C, \\ A(B + C) &= AB + AC, \\ (B + C)A &= BA + CA, \\ \alpha(AB) &= (\alpha A)B = A(\alpha B), \\ IA &= A = AI. \end{aligned}$$

A few warnings are in order.

- (i) $AB \neq BA$ in general (it may be true by fluke sometimes). That is, matrices do not commute. For example, take $A = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}$.
- (ii) $AB = AC$ does not necessarily imply $B = C$, even if A is not 0.
- (iii) $AB = 0$ does not necessarily mean that $A = 0$ or $B = 0$. Try, for example, $A = B = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$.

For the last two items to hold, we would need to “divide” by a matrix. This is where the *matrix inverse* comes in. Suppose that A and B are $n \times n$ matrices such that

$$AB = I = BA.$$

Then we call B the inverse of A and we denote B by A^{-1} . If the inverse of A exists, then we call A *invertible*. If A is not invertible, we sometimes say A is *singular*.

If A is invertible, then $AB = AC$ does imply that $B = C$ (in particular, the inverse of A is unique). We just multiply both sides by A^{-1} (on the left) to get $A^{-1}AB = A^{-1}AC$ or $IB = IC$ or $B = C$. It is also not hard to see that $(A^{-1})^{-1} = A$.

3.2.3 The determinant

For square matrices we define a useful quantity called the *determinant*. We define the determinant of a 1×1 matrix as the value of its only entry. For a 2×2 matrix, we define

$$\det \left(\begin{bmatrix} a & b \\ c & d \end{bmatrix} \right) \stackrel{\text{def}}{=} ad - bc.$$

Before trying to define the determinant for larger matrices, let us note the meaning of the determinant. Consider an $n \times n$ matrix as a mapping of the n -dimensional euclidean space $\mathbb{R}^n$ to itself, where $\vec{x}$ gets sent to $A\vec{x}$. In particular, a 2×2 matrix A is a mapping of the plane to itself. The determinant of A is the factor by which the area of objects changes. If we take the unit square (square of side 1) in the plane, then A takes the square to a parallelogram of area $|\det(A)|$. The sign of $\det(A)$ denotes changing of orientation (negative if the axes get flipped). For example, let

$$A = \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix}.$$

Then $\det(A) = 1 + 1 = 2$. Let us see where the (unit) square with vertices $(0, 0)$, $(1, 0)$, $(0, 1)$, and $(1, 1)$ gets sent. Clearly $(0, 0)$ gets sent to $(0, 0)$.

$$\begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ -1 \end{bmatrix}, \quad \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}, \quad \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 2 \\ 0 \end{bmatrix}.$$

The image of the square is another square with vertices $(0, 0)$, $(1, -1)$, $(1, 1)$, and $(2, 0)$. The image square has a side of length $\sqrt{2}$ and is therefore of area 2.

If you think back to high school geometry, you may have seen a formula for computing the area of a parallelogram with vertices $(0, 0)$, (a, c) , (b, d) and $(a + b, c + d)$. And it is precisely

$$\left| \det \begin{pmatrix} a & b \\ c & d \end{pmatrix} \right|.$$

The vertical lines above mean absolute value. The matrix $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$ carries the unit square to the given parallelogram.

Let us look at the determinant for larger matrices. We define A_{ij} as the matrix A with the i^{th} row and the j^{th} column deleted. To compute the determinant of a matrix, pick one row, say the i^{th} row and compute:

$$\det(A) = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det(A_{ij}).$$

For the first row, we get

$$\det(A) = a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + a_{13} \det(A_{13}) - \dots \begin{cases} +a_{1n} \det(A_{1n}) & \text{if } n \text{ is odd,} \\ -a_{1n} \det(A_{1n}) & \text{if } n \text{ even.} \end{cases}$$

We alternately add and subtract the determinants of the submatrices A_{ij} multiplied by a_{ij} for a fixed i and all j . For a 3×3 matrix, picking the first row, we get $\det(A) = a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + a_{13} \det(A_{13})$. For example,

$$\begin{aligned} \det \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} &= 1 \cdot \det \begin{pmatrix} 5 & 6 \\ 8 & 9 \end{pmatrix} - 2 \cdot \det \begin{pmatrix} 4 & 6 \\ 7 & 9 \end{pmatrix} + 3 \cdot \det \begin{pmatrix} 4 & 5 \\ 7 & 8 \end{pmatrix} \\ &= 1(5 \cdot 9 - 6 \cdot 8) - 2(4 \cdot 9 - 6 \cdot 7) + 3(4 \cdot 8 - 5 \cdot 7) = 0. \end{aligned}$$

The numbers $(-1)^{i+j} \det(A_{ij})$ are called *cofactors* of the matrix and this way of computing the determinant is called the *cofactor expansion*. No matter which row you pick, you always get the same number. It is also possible to compute the determinant by expanding along columns (picking a column instead of a row above). It is true that $\det(A) = \det(A^T)$.

A common notation for the determinant is a pair of vertical lines:

$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} = \det \begin{pmatrix} a & b \\ c & d \end{pmatrix}.$$

I personally find this notation confusing as vertical lines usually mean a positive quantity, while determinants can be negative. Also think about how to write the absolute value of a determinant. I will not use this notation in this book.

Think of the determinants telling you the scaling of a mapping. If B doubles the sizes of geometric objects and A triples them, then AB (which applies B to an object and then A) should make size go up by a factor of $6 = 3 \cdot 2$. This is true in general:

$$\det(AB) = \det(A) \det(B).$$

This property is one of the most useful, and it can be employed to actually compute determinants. A particularly interesting consequence is to note what it means for existence of inverses. Take A and B to be inverses of each other, that is, $AB = I$. Then

$$\det(A) \det(B) = \det(AB) = \det(I) = 1.$$

Neither $\det(A)$ nor $\det(B)$ can be zero. Conversely, if $\det(A)$ is not zero, then A is invertible. We state this observation as a theorem as it is very important in the context of this course.

Theorem 3.2.1. *An $n \times n$ matrix A is invertible if and only if $\det(A) \neq 0$.*

In fact, $\det(A^{-1}) \det(A) = 1$ says that $\det(A^{-1}) = \frac{1}{\det(A)}$. So we even know what the determinant of A^{-1} is before we know how to compute A^{-1} .

There is a simple formula for the inverse of a 2×2 matrix

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}.$$

Notice the determinant of the matrix $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$ in the denominator of the fraction. The formula only works if the determinant is nonzero, otherwise we are dividing by zero.

3.2.4 Solving linear systems

One application of matrices we will need is to solve systems of linear equations. This is best shown by example. Suppose that we have the following system of linear equations

$$\begin{aligned} 2x_1 + 2x_2 + 2x_3 &= 2, \\ x_1 + x_2 + 3x_3 &= 5, \\ x_1 + 4x_2 + x_3 &= 10. \end{aligned}$$

Without changing the solution, we could swap equations in this system, we could multiply any of the equations by a nonzero number, and we could add a multiple of one equation to another equation. It turns out these operations always suffice to find a solution.

It is easier to write the system as a matrix equation. The system above can be written as

$$\begin{bmatrix} 2 & 2 & 2 \\ 1 & 1 & 3 \\ 1 & 4 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 2 \\ 5 \\ 10 \end{bmatrix}.$$

To solve the system, we put the coefficient matrix (the matrix on the left-hand side of the equation) together with the vector on the right-hand side and get the so-called *augmented matrix*:

$$\left[\begin{array}{ccc|c} 2 & 2 & 2 & 2 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right].$$

We apply a sequence of the following three *elementary row operations*.

- (i) Swap two rows.
- (ii) Multiply a row by a nonzero number.
- (iii) Add a multiple of one row to another row.

We keep doing these operations until we get into a state where it is easy to read off the answer, or until we get into a contradiction indicating no solution, for example if we come up with an equation such as $0 = 1$.

Let us work through the example. First multiply the first row by $1/2$ to obtain

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right].$$

Now subtract the first row from the second and third row:

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 0 & 2 & 4 \\ 0 & 3 & 0 & 9 \end{array} \right]$$

Multiply the last row by $1/3$ and the second row by $1/2$:

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 0 & 1 & 2 \\ 0 & 1 & 0 & 3 \end{array} \right]$$

Swap rows 2 and 3:

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 2 \end{array} \right]$$

Subtract the last row from the first, then subtract the second row from the first:

$$\left[\begin{array}{ccc|c} 1 & 0 & 0 & -4 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 2 \end{array} \right]$$

Finally, we think about what equations this augmented matrix represents: $x_1 = -4$, $x_2 = 3$, and $x_3 = 2$. We try this solution in the original system and, voilà, it works!

Exercise 3.2.1: Check that the solution above really solves the given equations.

We write this equation in matrix notation as

$$A\vec{x} = \vec{b},$$

where A is the matrix $\begin{bmatrix} 2 & 2 & 2 \\ 1 & 1 & 3 \\ 1 & 4 & 1 \end{bmatrix}$ and $\vec{b}$ is the vector $\begin{bmatrix} 2 \\ 5 \\ 10 \end{bmatrix}$. The solution can also be computed via the inverse,

$$\vec{x} = A^{-1}A\vec{x} = A^{-1}\vec{b}.$$

It is possible that the solution is not unique, or that no solution exists. It is easy to tell if a solution does not exist. If during the row reduction you come up with a row where all the entries except the last one are zero (the last entry in a row corresponds to the right-hand side of the equation), then the system is *inconsistent* and has no solution. For example, for a system of 3 equations and 3 unknowns, if you find a row such as $[0 \ 0 \ 0 \mid 1]$ in the augmented matrix, you know the system is inconsistent. That row corresponds to $0 = 1$.

You generally try to use row operations until the following conditions are satisfied. The first (from the left) nonzero entry in each row is called the *leading entry*.

- (i) The leading entry in any row is strictly to the right of the leading entry of the row above.
- (ii) Any zero rows are below all the nonzero rows.
- (iii) All leading entries are 1.
- (iv) All the entries above and below a leading entry are zero.

Such a matrix is said to be in *reduced row echelon form*. The variables corresponding to columns with no leading entries are said to be *free variables*. Free variables mean that we can pick those variables to be anything we want and then solve for the rest of the unknowns.

Example 3.2.1: The following augmented matrix is in reduced row echelon form.

$$\left[\begin{array}{ccc|c} 1 & 2 & 0 & 3 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{array} \right]$$

Suppose the variables are x_1 , x_2 , and x_3 . Then x_2 is the free variable, $x_1 = 3 - 2x_2$, and $x_3 = 1$.

On the other hand, if during the row reduction process you come up with the matrix

$$\left[\begin{array}{ccc|c} 1 & 2 & 13 & 3 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 3 \end{array} \right],$$

there is no need to go further. The last row corresponds to the equation $0x_1 + 0x_2 + 0x_3 = 3$, which is preposterous. Hence, no solution exists.

3.2.5 Computing the inverse

If the matrix A is square and there exists a unique solution $\vec{x}$ to $A\vec{x} = \vec{b}$ for any $\vec{b}$ (there are no free variables), then A is invertible. Multiplying both sides by A^{-1} , you can see that $\vec{x} = A^{-1}\vec{b}$. So it is useful to compute the inverse if you want to solve the equation for many different right-hand sides $\vec{b}$.

We have a formula for the 2×2 inverse, but it is also not hard to compute inverses of larger matrices. While we will not have too much occasion to compute inverses for larger matrices than 2×2 by hand, let us touch on how to do it. Finding the inverse of A is actually just solving a bunch of linear equations. If we can solve $A\vec{x}_k = \vec{e}_k$ where $\vec{e}_k$ is the vector with all zeros except a 1 at the k^{th} position, then the inverse is the matrix with the columns $\vec{x}_k$ for $k = 1, 2, \dots, n$ (exercise: why?). Therefore, to find the inverse we write a larger $n \times 2n$ augmented matrix $[A \mid I]$, where I is the identity matrix. We then perform row reduction. The reduced row echelon form of $[A \mid I]$ will be of the form $[I \mid A^{-1}]$ if and only if A is invertible. We then just read off the inverse A^{-1} .

3.2.6 Exercises

Exercise 3.2.2: Solve $\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \vec{x} = \begin{bmatrix} 5 \\ 6 \end{bmatrix}$ by using matrix inverse.

Exercise 3.2.3: Compute determinant of $\begin{bmatrix} 9 & -2 & -6 \\ -8 & 3 & 6 \\ 10 & -2 & -6 \end{bmatrix}$.

Exercise 3.2.4: Compute determinant of $\begin{bmatrix} 1 & 2 & 3 & 1 \\ 4 & 0 & 5 & 0 \\ 6 & 0 & 7 & 0 \\ 8 & 0 & 10 & 1 \end{bmatrix}$. Hint: Expand along the proper row or column to make the calculations simpler.

Exercise 3.2.5: Compute inverse of $\begin{bmatrix} 1 & 2 & 3 \\ 1 & 1 & 1 \\ 0 & 1 & 0 \end{bmatrix}$.

Exercise 3.2.6: For which h is $\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & h \end{bmatrix}$ not invertible? Is there only one such h ? Are there several? Infinitely many?

Exercise 3.2.7: For which h is $\begin{bmatrix} h & 1 & 1 \\ 0 & h & 0 \\ 1 & 1 & h \end{bmatrix}$ not invertible? Find all such h .

Exercise 3.2.8: Solve $\begin{bmatrix} 9 & -2 & -6 \\ -8 & 3 & 6 \\ 10 & -2 & -6 \end{bmatrix} \vec{x} = \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}$.

Exercise 3.2.9: Solve $\begin{bmatrix} 5 & 3 & 7 \\ 8 & 4 & 4 \\ 6 & 3 & 3 \end{bmatrix} \vec{x} = \begin{bmatrix} 2 \\ 0 \\ 0 \end{bmatrix}$.

Exercise 3.2.10: Solve $\begin{bmatrix} 3 & 2 & 3 & 0 \\ 3 & 3 & 3 & 3 \\ 0 & 2 & 4 & 2 \\ 2 & 3 & 4 & 3 \end{bmatrix} \vec{x} = \begin{bmatrix} 2 \\ 0 \\ 4 \\ 1 \end{bmatrix}$.

Exercise 3.2.11: Find 3 nonzero 2×2 matrices A , B , and C such that $AB = AC$ but $B \neq C$.

Exercise 3.2.101: Compute determinant of $\begin{bmatrix} 1 & 1 & 1 \\ 2 & 3 & -5 \\ 1 & -1 & 0 \end{bmatrix}$

Exercise 3.2.102: Find t such that $\begin{bmatrix} 1 & t \\ -1 & 2 \end{bmatrix}$ is not invertible.

Exercise 3.2.103: Solve $\begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \vec{x} = \begin{bmatrix} 10 \\ 20 \end{bmatrix}$.

Exercise 3.2.104: Suppose a, b, c are nonzero numbers. Let $M = \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix}$, $N = \begin{bmatrix} a & 0 & 0 \\ 0 & b & 0 \\ 0 & 0 & c \end{bmatrix}$.

a) Compute M^{-1} .

b) Compute N^{-1} .

3.3 Linear systems of ODEs

Note: less than 1 lecture, second part of §5.1 in [EP], §7.4 in [BD]

A matrix- or vector-valued function is simply a matrix or a vector whose entries depend on some variable. If t is the independent variable, we write a *vector-valued function* $\vec{x}(t)$ as

$$\vec{x}(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{bmatrix}.$$

Similarly a *matrix-valued function* $A(t)$ is

$$A(t) = \begin{bmatrix} a_{11}(t) & a_{12}(t) & \cdots & a_{1n}(t) \\ a_{21}(t) & a_{22}(t) & \cdots & a_{2n}(t) \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1}(t) & a_{n2}(t) & \cdots & a_{nn}(t) \end{bmatrix}.$$

The derivative $A'(t)$ or $\frac{dA}{dt}$ is just the matrix-valued function whose ij^{th} entry is $a'_{ij}(t)$.

Rules of differentiation of matrix-valued functions are similar to rules for normal functions. Let $A(t)$ and $B(t)$ be matrix-valued functions. Let c be a scalar and C a constant matrix. Then

$$\begin{aligned} (A(t) + B(t))' &= A'(t) + B'(t), \\ (A(t)B(t))' &= A'(t)B(t) + A(t)B'(t), \\ (cA(t))' &= cA'(t), \\ (CA(t))' &= CA'(t), \\ (A(t)C)' &= A'(t)C. \end{aligned}$$

Note the order of the multiplication in the last two expressions.

A *first-order linear system of ODEs* is a system that can be written as the vector equation

$$\vec{x}'(t) = P(t)\vec{x}(t) + \vec{f}(t),$$

where $P(t)$ is a matrix-valued function, and $\vec{x}(t)$ and $\vec{f}(t)$ are vector-valued functions. We will often suppress the dependence on t and only write $\vec{x}' = P\vec{x} + \vec{f}$. A solution of the system is a vector-valued function $\vec{x}$ satisfying the vector equation.

For example, the equations

$$\begin{aligned} x_1' &= 2tx_1 + e^t x_2 + t^2, \\ x_2' &= \frac{x_1}{t} - x_2 + e^t, \end{aligned}$$

can be written as

$$\vec{x}' = \begin{bmatrix} 2t & e^t \\ 1/t & -1 \end{bmatrix} \vec{x} + \begin{bmatrix} t^2 \\ e^t \end{bmatrix}.$$

We will mostly concentrate on equations that are not just linear, but are in fact *constant-coefficient* equations. That is, the matrix P will be constant; it will not depend on t .

When $\vec{f} = \vec{0}$ (the zero vector), we say the system is *homogeneous*. For homogeneous linear systems, we have the principle of superposition, just like for single homogeneous linear equations.

Theorem 3.3.1 (Superposition). *Let $\vec{x}' = P\vec{x}$ be a linear homogeneous system of ODEs. Suppose that $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ are n solutions of the equation and $c_1, c_2, \dots, c_n$ are any constants, then*

$$\vec{x} = c_1\vec{x}_1 + c_2\vec{x}_2 + \dots + c_n\vec{x}_n, \quad (3.2)$$

is also a solution. Furthermore, if this is a system of n equations (P is $n \times n$), and $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ are linearly independent, then every solution $\vec{x}$ can be written as (3.2).

Linear independence for vector-valued functions is the same idea as for normal functions. The vector-valued functions $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ are linearly independent when

$$c_1\vec{x}_1 + c_2\vec{x}_2 + \dots + c_n\vec{x}_n = \vec{0}$$

has only the solution $c_1 = c_2 = \dots = c_n = 0$, where the equation must hold for all t .

Example 3.3.1: $\vec{x}_1 = \begin{bmatrix} t^2 \\ t \end{bmatrix}$, $\vec{x}_2 = \begin{bmatrix} 0 \\ 1+t \end{bmatrix}$, $\vec{x}_3 = \begin{bmatrix} -t^2 \\ 1 \end{bmatrix}$ are linearly dependent because $\vec{x}_1 + \vec{x}_3 = \vec{x}_2$, and this holds for all t . So $c_1 = 1$, $c_2 = -1$, and $c_3 = 1$ above will work.

On the other hand, if we change the example slightly to $\vec{x}_1 = \begin{bmatrix} t^2 \\ t \end{bmatrix}$, $\vec{x}_2 = \begin{bmatrix} 0 \\ t \end{bmatrix}$, $\vec{x}_3 = \begin{bmatrix} -t^2 \\ 1 \end{bmatrix}$, then the functions are linearly independent. First write $c_1\vec{x}_1 + c_2\vec{x}_2 + c_3\vec{x}_3 = \vec{0}$ and note that it has to hold for all t . We get that

$$c_1\vec{x}_1 + c_2\vec{x}_2 + c_3\vec{x}_3 = \begin{bmatrix} c_1t^2 - c_3t^2 \\ c_1t + c_2t + c_3 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}.$$

In other words $c_1t^2 - c_3t^2 = 0$ and $c_1t + c_2t + c_3 = 0$. If we set $t = 0$, then the second equation becomes $c_3 = 0$. But then the first equation becomes $c_1t^2 = 0$ for all t and so $c_1 = 0$. Thus the second equation is just $c_2t = 0$, which means $c_2 = 0$. So $c_1 = c_2 = c_3 = 0$ is the only solution and $\vec{x}_1$, $\vec{x}_2$, and $\vec{x}_3$ are linearly independent.

The linear combination $c_1\vec{x}_1 + c_2\vec{x}_2 + \dots + c_n\vec{x}_n$ could always be written as

$$X(t)\vec{c},$$

where $X(t)$ is the matrix with columns $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$, and $\vec{c}$ is the column vector with entries $c_1, c_2, \dots, c_n$. Assuming that $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ are linearly independent, the matrix-valued function $X(t)$ is called a *fundamental matrix*, or a *fundamental matrix solution*. Note that the fundamental matrix is not unique; if you take any constant invertible $n \times n$ matrix C , then $X(t)C$ is another fundamental matrix.

To solve nonhomogeneous first-order linear systems, we use the same technique as we applied to solve single linear nonhomogeneous equations.

Theorem 3.3.2. Let $\vec{x}' = P\vec{x} + \vec{f}$ be a linear system of ODEs. Suppose $\vec{x}_p$ is one particular solution. Then every solution can be written as

$$\vec{x} = \vec{x}_c + \vec{x}_p,$$

where $\vec{x}_c$ is a solution to the associated homogeneous equation ($\vec{x}' = P\vec{x}$).

The procedure for systems is the same as for single equations. We find a particular solution to the nonhomogeneous equation, then we find the general solution to the associated homogeneous equation, and finally we add the two together.

Alright, suppose you have found the general solution of $\vec{x}' = P\vec{x} + \vec{f}$. Next suppose you are given an initial condition of the form

$$\vec{x}(t_0) = \vec{b}$$

for some fixed t_0 and a constant vector $\vec{b}$. Let $X(t)$ be a fundamental matrix solution of the associated homogeneous equation (i.e. columns of $X(t)$ are solutions). The general solution can be written as

$$\vec{x}(t) = X(t)\vec{c} + \vec{x}_p(t).$$

We are seeking a vector $\vec{c}$ such that

$$\vec{b} = \vec{x}(t_0) = X(t_0)\vec{c} + \vec{x}_p(t_0).$$

In other words, we are solving for $\vec{c}$ the nonhomogeneous system of linear equations

$$X(t_0)\vec{c} = \vec{b} - \vec{x}_p(t_0).$$

Example 3.3.2: In [Example 3.1.1](#) on page 120, we solved the system

$$\begin{aligned} x_1' &= x_1, \\ x_2' &= x_1 - x_2, \end{aligned}$$

with initial conditions $x_1(0) = 1$, $x_2(0) = 2$. Let us consider this problem in the language of this section. The system is homogeneous, so $\vec{f}(t) = \vec{0}$. We write the system and the initial conditions as

$$\vec{x}' = \begin{bmatrix} 1 & 0 \\ 1 & -1 \end{bmatrix} \vec{x}, \quad \vec{x}(0) = \begin{bmatrix} 1 \\ 2 \end{bmatrix}.$$

Ignoring the initial condition for a moment, the general solution is $x_1 = c_1 e^t$ and $x_2 = \frac{c_1}{2} e^t + c_2 e^{-t}$. Letting $c_1 = 1$ and $c_2 = 0$, we obtain the solution $\vec{x}(t) = \begin{bmatrix} e^t \\ (1/2)e^t \end{bmatrix}$. Letting $c_1 = 0$ and $c_2 = 1$, we obtain $\vec{x}(t) = \begin{bmatrix} 0 \\ e^{-t} \end{bmatrix}$. These two solutions are linearly independent, as can be seen by setting $t = 0$, and noting that the resulting constant vectors are linearly independent. In matrix notation, a fundamental matrix solution is, therefore,

$$X(t) = \begin{bmatrix} e^t & 0 \\ \frac{1}{2}e^t & e^{-t} \end{bmatrix}.$$

To solve the initial value problem, we solve for $\vec{c}$ in the equation

$$X(0) \vec{c} = \vec{b},$$

or in other words,

$$\begin{bmatrix} 1 & 0 \\ \frac{1}{2} & 1 \end{bmatrix} \vec{c} = \begin{bmatrix} 1 \\ 2 \end{bmatrix}.$$

A single elementary row operation shows $\vec{c} = \begin{bmatrix} 1 \\ 3/2 \end{bmatrix}$. Our solution is

$$\vec{x}(t) = X(t) \vec{c} = \begin{bmatrix} e^t & 0 \\ \frac{1}{2}e^t & e^{-t} \end{bmatrix} \begin{bmatrix} 1 \\ \frac{3}{2} \end{bmatrix} = \begin{bmatrix} e^t \\ \frac{1}{2}e^t + \frac{3}{2}e^{-t} \end{bmatrix}.$$

This new solution agrees with our previous solution from § 3.1.

3.3.1 Exercises

Exercise 3.3.1: Write the system $x_1' = 2x_1 - 3tx_2 + \sin t$, $x_2' = e^t x_1 + 3x_2 + \cos t$ in the form $\vec{x}' = P(t)\vec{x} + \vec{f}(t)$.

Exercise 3.3.2:

- Verify that the system $\vec{x}' = \begin{bmatrix} 1 & 3 \\ 3 & 1 \end{bmatrix} \vec{x}$ has the two solutions $\begin{bmatrix} 1 \\ 1 \end{bmatrix} e^{4t}$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix} e^{-2t}$.
- Write down the general solution.
- Write down the general solution in the form $x_1 = ?$, $x_2 = ?$ (i.e. write down a formula for each element of the solution).

Exercise 3.3.3: Verify that $\begin{bmatrix} 1 \\ 1 \end{bmatrix} e^t$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix} e^t$ are linearly independent. Hint: Just plug in $t = 0$.

Exercise 3.3.4: Verify that $\begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} e^t$ and $\begin{bmatrix} 1 \\ -1 \\ 1 \end{bmatrix} e^t$ and $\begin{bmatrix} 1 \\ -1 \\ 1 \end{bmatrix} e^{2t}$ are linearly independent. Hint: You must be a bit more tricky than in the previous exercise.

Exercise 3.3.5: Verify that $\begin{bmatrix} t \\ t^2 \end{bmatrix}$ and $\begin{bmatrix} t^3 \\ t^4 \end{bmatrix}$ are linearly independent.

Exercise 3.3.6: Take the system $x_1' + x_2' = x_1$, $x_1' - x_2' = x_2$.

- Write it in the form $A\vec{x}' = B\vec{x}$ for matrices A and B .
- Compute A^{-1} and use that to write the system in the form $\vec{x}' = P\vec{x}$.

Exercise 3.3.101: Are $\begin{bmatrix} e^{2t} \\ e^t \end{bmatrix}$ and $\begin{bmatrix} e^t \\ e^{2t} \end{bmatrix}$ linearly independent? Justify.

Exercise 3.3.102: Are $\begin{bmatrix} \cosh(t) \\ 1 \end{bmatrix}$, $\begin{bmatrix} e^t \\ 1 \end{bmatrix}$, and $\begin{bmatrix} e^{-t} \\ 1 \end{bmatrix}$ linearly independent? Justify.

Exercise 3.3.103: Write $x' = 3x - y + e^t$, $y' = tx$ in matrix notation.

Exercise 3.3.104:

- Write $x_1' = 2tx_2$, $x_2' = 2tx_2$ in matrix notation.
- Solve and write the solution in matrix notation.

3.4 Eigenvalue method

Note: 2 lectures, §5.2 in [EP], part of §7.3, §7.5, and §7.6 in [BD]

In this section, we will learn how to solve linear homogeneous constant-coefficient systems of ODEs by the eigenvalue method. Suppose we have such a system

$$\vec{x}' = P\vec{x},$$

where P is a constant square matrix. We wish to adapt the method for the single constant-coefficient equation by trying the function $e^{\lambda t}$ for an unknown constant λ . However, $\vec{x}$ is a vector. So we try $\vec{x} = \vec{v}e^{\lambda t}$, where $\vec{v}$ is an unknown constant vector. We plug this $\vec{x}$ into the equation to get

$$\underbrace{\lambda \vec{v} e^{\lambda t}}_{\vec{x}'} = \underbrace{P \vec{v} e^{\lambda t}}_{P\vec{x}}.$$

We divide by $e^{\lambda t}$ and notice that we are looking for a scalar λ and a vector $\vec{v}$ that satisfy the equation

$$\lambda \vec{v} = P\vec{v}.$$

To solve this equation, we need a little bit more linear algebra, which we now review.

3.4.1 Eigenvalues and eigenvectors of a matrix

Let A be a constant square matrix. Suppose λ is a scalar and $\vec{v}$ is a nonzero vector such that

$$A\vec{v} = \lambda\vec{v}.$$

We call such a λ an *eigenvalue* of A , and we call $\vec{v}$ a corresponding *eigenvector*.

Example 3.4.1: The matrix $\begin{bmatrix} 2 & 1 \\ 0 & 1 \end{bmatrix}$ has an eigenvalue $\lambda = 2$ with a corresponding eigenvector $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$, because

$$\begin{bmatrix} 2 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 2 \\ 0 \end{bmatrix} = 2 \begin{bmatrix} 1 \\ 0 \end{bmatrix}.$$

Let us compute eigenvalues for any matrix. Rewrite the equation for an eigenvalue as $A\vec{v} - \lambda\vec{v} = \vec{0}$ and thus as

$$(A - \lambda I)\vec{v} = \vec{0}.$$

This equation has a nonzero solution $\vec{v}$ if and only if $A - \lambda I$ is not invertible. Were $A - \lambda I$ invertible, we could write $(A - \lambda I)^{-1}(A - \lambda I)\vec{v} = (A - \lambda I)^{-1}\vec{0}$, which implies $\vec{v} = \vec{0}$. Therefore, A has the eigenvalue λ if and only if λ solves the equation

$$\det(A - \lambda I) = 0.$$

Consequently, we can find an eigenvalue of A without finding a corresponding eigenvector at the same time. An eigenvector will have to be found later, once λ is known.

Example 3.4.2: Find all eigenvalues of $\begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$.

We write

$$\begin{aligned} \det\left(\begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix} - \lambda \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}\right) &= \det\left(\begin{bmatrix} 2-\lambda & 1 & 1 \\ 1 & 2-\lambda & 0 \\ 0 & 0 & 2-\lambda \end{bmatrix}\right) = \\ &= (2-\lambda)((2-\lambda)^2 - 1) = -(\lambda-1)(\lambda-2)(\lambda-3). \end{aligned}$$

Setting this to zero, we find that the eigenvalues are $\lambda = 1$, $\lambda = 2$, and $\lambda = 3$.

For an $n \times n$ matrix, the polynomial we get by computing $\det(A - \lambda I)$ is of degree n , and hence in general, we have n eigenvalues. Some may be repeated, some may be complex.

To find an eigenvector corresponding to an eigenvalue λ , we write

$$(A - \lambda I)\vec{v} = \vec{0},$$

and solve for a nontrivial (nonzero) vector $\vec{v}$. If λ is an eigenvalue, there will be at least one free variable, and so for each distinct eigenvalue λ , we can always find an eigenvector.

Example 3.4.3: Find an eigenvector of $\begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$ corresponding to the eigenvalue $\lambda = 3$.

We write

$$(A - \lambda I)\vec{v} = \left(\begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix} - 3 \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}\right) \begin{bmatrix} v_1 \\ v_2 \\ v_3 \end{bmatrix} = \begin{bmatrix} -1 & 1 & 1 \\ 1 & -1 & 0 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \\ v_3 \end{bmatrix} = \vec{0}.$$

To solve this system of linear equations, we write down the augmented matrix

$$\left[\begin{array}{ccc|c} -1 & 1 & 1 & 0 \\ 1 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \end{array} \right],$$

and we perform row operations (exercise: which ones?) until we get:

$$\left[\begin{array}{ccc|c} 1 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right].$$

The entries of $\vec{v}$ have to satisfy the equations $v_1 - v_2 = 0$, $v_3 = 0$, and v_2 is a free variable. We pick v_2 to be arbitrary (but nonzero), we let $v_1 = v_2$, and of course $v_3 = 0$. For example, if we pick $v_2 = 1$, then $\vec{v} = \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix}$. Let us verify that $\vec{v}$ really is an eigenvector corresponding to $\lambda = 3$:

$$\begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 3 \\ 3 \\ 0 \end{bmatrix} = 3 \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix}.$$

Yay! It worked.

Exercise 3.4.1 (easy): Are eigenvectors unique? Can you find a different eigenvector for $\lambda = 3$ in the example above? How are the two eigenvectors related?

Exercise 3.4.2: When the matrix is 2×2 you do not need to do row operations when computing an eigenvector, you can read it off from $A - \lambda I$ (if you have computed the eigenvalues correctly). Can you see why? Explain. Try it for the matrix $\begin{bmatrix} 2 & 1 \\ 1 & 2 \end{bmatrix}$.

3.4.2 The eigenvalue method with distinct real eigenvalues

OK, back to homogeneous constant-coefficient systems of ODEs. We have the system

$$\vec{x}' = P\vec{x}.$$

We find the eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$ of the matrix P , and corresponding eigenvectors $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$. The functions $\vec{v}_1 e^{\lambda_1 t}, \vec{v}_2 e^{\lambda_2 t}, \dots, \vec{v}_n e^{\lambda_n t}$ are solutions of the system of equations and hence $\vec{x} = c_1 \vec{v}_1 e^{\lambda_1 t} + c_2 \vec{v}_2 e^{\lambda_2 t} + \dots + c_n \vec{v}_n e^{\lambda_n t}$ is a solution.

Theorem 3.4.1. Take $\vec{x}' = P\vec{x}$. If P is an $n \times n$ constant matrix that has n distinct real eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$, then there exist n linearly independent corresponding eigenvectors $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$, and the general solution to $\vec{x}' = P\vec{x}$ can be written as

$$\vec{x} = c_1 \vec{v}_1 e^{\lambda_1 t} + c_2 \vec{v}_2 e^{\lambda_2 t} + \dots + c_n \vec{v}_n e^{\lambda_n t}.$$

The corresponding fundamental matrix solution is

$$X(t) = \begin{bmatrix} \vec{v}_1 e^{\lambda_1 t} & \vec{v}_2 e^{\lambda_2 t} & \dots & \vec{v}_n e^{\lambda_n t} \end{bmatrix}.$$

That is, $X(t)$ is the matrix whose j^{th} column is $\vec{v}_j e^{\lambda_j t}$.

Example 3.4.4: Consider the system

$$\vec{x}' = \begin{bmatrix} 2 & 1 & 1 \\ 1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix} \vec{x}.$$

Find the general solution.

Earlier, we found the eigenvalues are 1, 2, 3. We found the eigenvector $\begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix}$ for the eigenvalue 3. Similarly we find the eigenvector $\begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix}$ for the eigenvalue 1 and $\begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}$ for the eigenvalue 2 (exercise: check). The general solution is

$$\vec{x} = c_1 \begin{bmatrix} 1 \\ -1 \\ 0 \end{bmatrix} e^t + c_2 \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} e^{2t} + c_3 \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} e^{3t} = \begin{bmatrix} c_1 e^t + c_3 e^{3t} \\ -c_1 e^t + c_2 e^{2t} + c_3 e^{3t} \\ -c_2 e^{2t} \end{bmatrix}.$$

In terms of a fundamental matrix solution,

$$\vec{x} = X(t) \vec{c} = \begin{bmatrix} e^t & 0 & e^{3t} \\ -e^t & e^{2t} & e^{3t} \\ 0 & -e^{2t} & 0 \end{bmatrix} \begin{bmatrix} c_1 \\ c_2 \\ c_3 \end{bmatrix}.$$

Exercise 3.4.3: Check that this $\vec{x}$ really solves the system.

Note: If we write a single homogeneous linear constant-coefficient n^{th} -order equation as a first-order system (as we did in § 3.1), then the eigenvalue equation

$$\det(P - \lambda I) = 0$$

is essentially the same as the characteristic equation we got in § 2.2 and § 2.3.

3.4.3 Complex eigenvalues

A matrix may very well have complex eigenvalues even if all the entries are real. Take, for example,

$$\vec{x}' = \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \vec{x}.$$

Let us compute the eigenvalues of the matrix $P = \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix}$.

$$\det(P - \lambda I) = \det \left(\begin{bmatrix} 1-\lambda & 1 \\ -1 & 1-\lambda \end{bmatrix} \right) = (1-\lambda)^2 + 1 = \lambda^2 - 2\lambda + 2 = 0.$$

Thus $\lambda = 1 \pm i$. Corresponding eigenvectors are also complex. Start with $\lambda = 1 - i$.

$$(P - (1 - i)I)\vec{v} = \vec{0},$$

$$\begin{bmatrix} i & 1 \\ -1 & i \end{bmatrix} \vec{v} = \vec{0}.$$

The equations $iv_1 + v_2 = 0$ and $-v_1 + iv_2 = 0$ are multiples of each other. So we only need to consider one of them. After picking $v_2 = 1$, for example, we have an eigenvector $\vec{v} = \begin{bmatrix} i \\ 1 \end{bmatrix}$. In similar fashion, we find that $\begin{bmatrix} -i \\ 1 \end{bmatrix}$ is an eigenvector corresponding to the eigenvalue $1 + i$.

We could write the solution as

$$\vec{x} = c_1 \begin{bmatrix} i \\ 1 \end{bmatrix} e^{(1-i)t} + c_2 \begin{bmatrix} -i \\ 1 \end{bmatrix} e^{(1+i)t} = \begin{bmatrix} c_1 i e^{(1-i)t} - c_2 i e^{(1+i)t} \\ c_1 e^{(1-i)t} + c_2 e^{(1+i)t} \end{bmatrix}.$$

We would then need to look for complex values c_1 and c_2 to solve any initial conditions. It is perhaps not completely clear that we get a real solution. After solving for c_1 and c_2 , we could use Euler's formula and do the whole song and dance we did before, but we will not. We will apply the formula in a smarter way first to find independent real solutions.

We claim that we did not have to look for a second eigenvector (nor for the second eigenvalue). All complex eigenvalues come in pairs (because the matrix P is real).

First a small detour. The real part of a complex number z can be computed as $\frac{z + \bar{z}}{2}$, where the bar above z means $\overline{a + ib} = a - ib$. The bar of z is called the *complex conjugate* of z . If a is a real number, then $\bar{a} = a$. Similarly we bar whole vectors or matrices by taking

the complex conjugate of every entry. Suppose a matrix P is real. Then $\bar{P} = P$, and so $\overline{P\vec{x}} = \bar{P}\bar{\vec{x}} = P\bar{\vec{x}}$. Also the complex conjugate of 0 is still 0, therefore,

$$\vec{0} = \overline{\vec{0}} = \overline{(P - \lambda I)\vec{v}} = (P - \bar{\lambda}I)\bar{\vec{v}}.$$

In other words, if $\lambda = a + ib$ is an eigenvalue, then so is $\bar{\lambda} = a - ib$. And if $\vec{v}$ is an eigenvector corresponding to the eigenvalue λ , then $\bar{\vec{v}}$ is an eigenvector corresponding to the eigenvalue $\bar{\lambda}$.

Suppose $a + ib$ is a complex eigenvalue of P , and $\vec{v}$ is a corresponding eigenvector. Then

$$\vec{x}_1 = \vec{v}e^{(a+ib)t}$$

is a solution (complex-valued) of $\vec{x}' = P\vec{x}$. Euler's formula shows that $\overline{e^{a+ib}} = e^{a-ib}$, and so

$$\vec{x}_2 = \overline{\vec{x}_1} = \bar{\vec{v}}e^{(a-ib)t}$$

is also a solution. As $\vec{x}_1$ and $\vec{x}_2$ are solutions, the function

$$\vec{x}_3 = \operatorname{Re} \vec{x}_1 = \operatorname{Re} \vec{v}e^{(a+ib)t} = \frac{\vec{x}_1 + \overline{\vec{x}_1}}{2} = \frac{\vec{x}_1 + \vec{x}_2}{2} = \frac{1}{2}\vec{x}_1 + \frac{1}{2}\vec{x}_2$$

is also a solution. And $\vec{x}_3$ is real-valued! Similarly as $\operatorname{Im} z = \frac{z - \bar{z}}{2i}$ is the imaginary part, we find that

$$\vec{x}_4 = \operatorname{Im} \vec{x}_1 = \frac{\vec{x}_1 - \overline{\vec{x}_1}}{2i} = \frac{\vec{x}_1 - \vec{x}_2}{2i}.$$

is also a real-valued solution. It turns out that $\vec{x}_3$ and $\vec{x}_4$ are linearly independent. We will use Euler's formula to separate out the real and imaginary part.

Returning to our example problem,

$$\begin{aligned} \vec{x}_1 &= \begin{bmatrix} i \\ 1 \end{bmatrix} e^{(1-i)t} = \begin{bmatrix} i \\ 1 \end{bmatrix} (e^t \cos t - ie^t \sin t) \\ &= \begin{bmatrix} ie^t \cos t + e^t \sin t \\ e^t \cos t - ie^t \sin t \end{bmatrix} = \begin{bmatrix} e^t \sin t \\ e^t \cos t \end{bmatrix} + i \begin{bmatrix} e^t \cos t \\ -e^t \sin t \end{bmatrix}. \end{aligned}$$

The two real-valued linearly independent solutions we seek are

$$\operatorname{Re} \vec{x}_1 = \begin{bmatrix} e^t \sin t \\ e^t \cos t \end{bmatrix} \quad \text{and} \quad \operatorname{Im} \vec{x}_1 = \begin{bmatrix} e^t \cos t \\ -e^t \sin t \end{bmatrix}.$$

Exercise 3.4.4: Check that these really are solutions.

The general solution is

$$\vec{x} = c_1 \begin{bmatrix} e^t \sin t \\ e^t \cos t \end{bmatrix} + c_2 \begin{bmatrix} e^t \cos t \\ -e^t \sin t \end{bmatrix} = \begin{bmatrix} c_1 e^t \sin t + c_2 e^t \cos t \\ c_1 e^t \cos t - c_2 e^t \sin t \end{bmatrix}.$$

This solution is real-valued for real c_1 and c_2 . At this point, we would solve for any initial conditions we may have to find c_1 and c_2 .

We summarize the discussion as a theorem.

Theorem 3.4.2. *Let P be a real-valued constant matrix. If P has a complex eigenvalue $a + ib$ and a corresponding eigenvector $\vec{v}$, then P also has a complex eigenvalue $a - ib$ with a corresponding eigenvector $\bar{\vec{v}}$. Furthermore, $\vec{x}' = P\vec{x}$ has two linearly independent real-valued solutions*

$$\vec{x}_1 = \operatorname{Re} \vec{v} e^{(a+ib)t} \quad \text{and} \quad \vec{x}_2 = \operatorname{Im} \vec{v} e^{(a+ib)t}.$$

For each pair of complex eigenvalues $a + ib$ and $a - ib$, we get two real-valued linearly independent solutions. We then go on to the next eigenvalue, which is either a real eigenvalue or another complex eigenvalue pair. If we have n distinct eigenvalues (real or complex), then we end up with n linearly independent solutions. If we had only two equations ($n = 2$) as in the example above, then once we found two solutions we are finished, and our general solution is

$$\vec{x} = c_1 \vec{x}_1 + c_2 \vec{x}_2 = c_1 (\operatorname{Re} \vec{v} e^{(a+ib)t}) + c_2 (\operatorname{Im} \vec{v} e^{(a+ib)t}).$$

We can now find a real-valued general solution to any homogeneous system where the matrix has distinct eigenvalues. When we have repeated eigenvalues, matters get a bit more complicated and we will look at that situation in § 3.7.

3.4.4 Exercises

Exercise 3.4.5 (easy): *Let A be a 3×3 matrix with an eigenvalue of 3 and a corresponding eigenvector $\vec{v} = \begin{bmatrix} 1 \\ -1 \\ 3 \end{bmatrix}$. Find $A\vec{v}$.*

Exercise 3.4.6:

- Find the general solution of $x'_1 = 2x_1$, $x'_2 = 3x_2$ using the eigenvalue method (first write the system in the form $\vec{x}' = A\vec{x}$).*
- Solve the system by solving each equation separately and verify you get the same general solution.*

Exercise 3.4.7: *Find the general solution of $x'_1 = 3x_1 + x_2$, $x'_2 = 2x_1 + 4x_2$ using the eigenvalue method.*

Exercise 3.4.8: *Find the general solution of $x'_1 = x_1 - 2x_2$, $x'_2 = 2x_1 + x_2$ using the eigenvalue method. Do not use complex exponentials in your solution.*

Exercise 3.4.9:

- Compute eigenvalues and eigenvectors of $A = \begin{bmatrix} 9 & -2 & -6 \\ -8 & 3 & 6 \\ 10 & -2 & -6 \end{bmatrix}$.*
- Find the general solution of $\vec{x}' = A\vec{x}$.*

Exercise 3.4.10: *Compute eigenvalues and eigenvectors of $\begin{bmatrix} -2 & -1 & -1 \\ 3 & 2 & 1 \\ -3 & -1 & 0 \end{bmatrix}$.*

Exercise 3.4.11: Let a, b, c, d, e, f be numbers. Find the eigenvalues of $\begin{bmatrix} a & b & c \\ 0 & d & e \\ 0 & 0 & f \end{bmatrix}$.

Exercise 3.4.101:

a) Compute eigenvalues and eigenvectors of $A = \begin{bmatrix} 1 & 0 & 3 \\ -1 & 0 & 1 \\ 2 & 0 & 2 \end{bmatrix}$.

b) Solve the system $\vec{x}' = A\vec{x}$.

Exercise 3.4.102:

a) Compute eigenvalues and eigenvectors of $A = \begin{bmatrix} 1 & 1 \\ -1 & 0 \end{bmatrix}$.

b) Solve the system $\vec{x}' = A\vec{x}$.

Exercise 3.4.103: Solve $x_1' = x_2, x_2' = x_1$ using the eigenvalue method.

Exercise 3.4.104: Solve $x_1' = x_2, x_2' = -x_1$ using the eigenvalue method.

3.5 Two-dimensional systems and their vector fields

Note: 1 lecture, part of §6.2 in [EP], parts of §7.5 and §7.6 in [BD]

We take a moment to discuss constant-coefficient linear homogeneous systems in the plane. Much intuition can be obtained by studying this simple case. We use coordinates (x, y) for the plane as usual, and suppose $P = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ is a 2×2 matrix. Consider the system

$$\begin{bmatrix} x \\ y \end{bmatrix}' = P \begin{bmatrix} x \\ y \end{bmatrix} \quad \text{or} \quad \begin{bmatrix} x \\ y \end{bmatrix}' = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}. \quad (3.3)$$

The system is autonomous (compare this section to §1.6), and so we draw a vector field (see the end of §3.1). We will be able to visually understand this vectorfield and the solutions of the ODE in terms of the eigenvalues and eigenvectors of the matrix P . For this section, we assume that P has two distinct eigenvalues and two corresponding eigenvectors. We will also assume that P is nonsingular, that is, neither eigenvalue is zero.

Case 1. Suppose that the eigenvalues of P are real and positive. We find two corresponding eigenvectors and plot them in the plane. For example, take the matrix $\begin{bmatrix} 1 & 1 \\ 0 & 2 \end{bmatrix}$. The eigenvalues are 1 and 2 and corresponding eigenvectors are $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$. See Figure 3.4.

Let (x, y) be a point on the line determined by an eigenvector $\vec{v}$ for an eigenvalue λ . That is, $\begin{bmatrix} x \\ y \end{bmatrix} = \alpha \vec{v}$ for some scalar α . Then

$$\begin{bmatrix} x \\ y \end{bmatrix}' = P \begin{bmatrix} x \\ y \end{bmatrix} = P(\alpha \vec{v}) = \alpha(P\vec{v}) = \alpha\lambda\vec{v}.$$

The derivative is a multiple of $\vec{v}$ and hence points along the line determined by $\vec{v}$. As $\lambda > 0$, the derivative points in the direction of $\vec{v}$ when α is positive and in the opposite direction when α is negative. We draw the lines determined by the eigenvectors and arrows on the lines to indicate the directions. See Figure 3.5 on the following page.

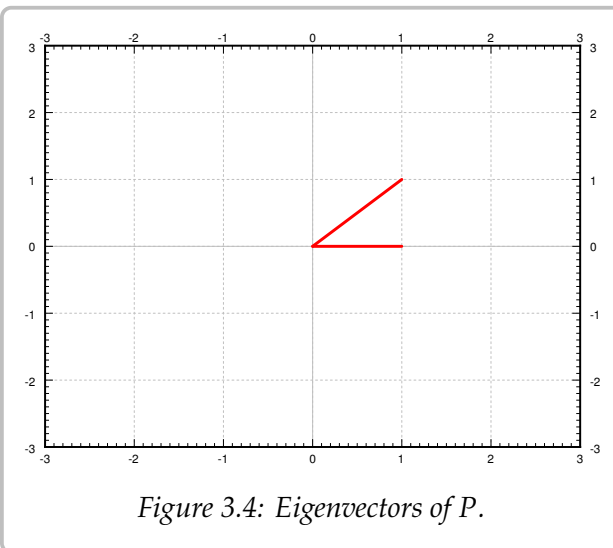


Figure 3.4: Eigenvectors of P .

We fill in the rest of the arrows for the vector field, and we draw a few solutions. See Figure 3.6 on the next page. The picture looks like a source with arrows coming out from the origin. Hence we call this type of picture a *source* or sometimes an *unstable node*.

Case 2. Suppose both eigenvalues are negative. For example, take the negation of the matrix from case 1, $\begin{bmatrix} -1 & -1 \\ 0 & -2 \end{bmatrix}$. The eigenvalues are -1 and -2 and corresponding eigenvectors are the same, $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$. The calculation and the picture are almost the same. The difference is that the eigenvalues are negative and arrows are reversed. We get the picture in Figure 3.7 on the following page. We call this type of picture a *sink* or a *stable node*.

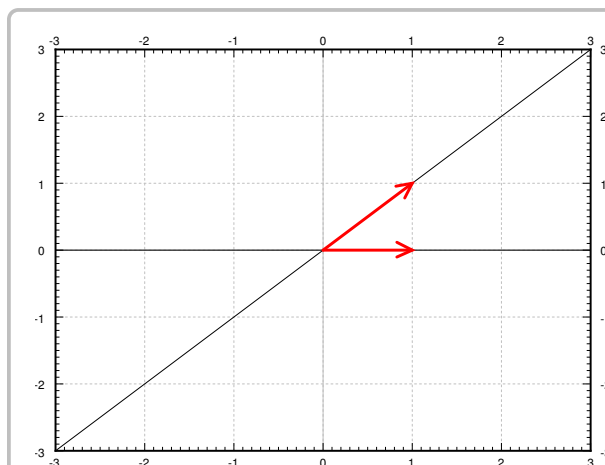
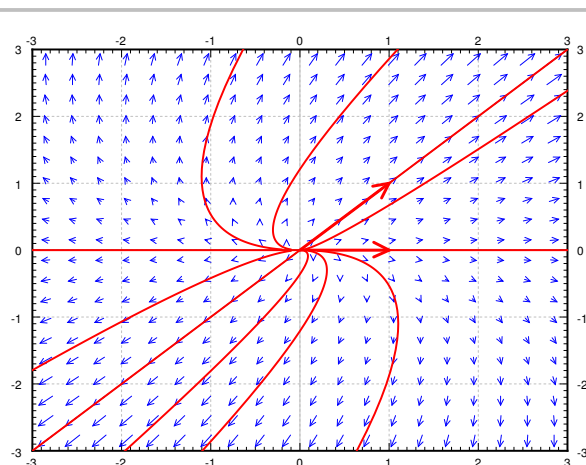
Figure 3.5: Eigenvectors of P with directions.

Figure 3.6: Example source vector field with eigenvectors and solutions.

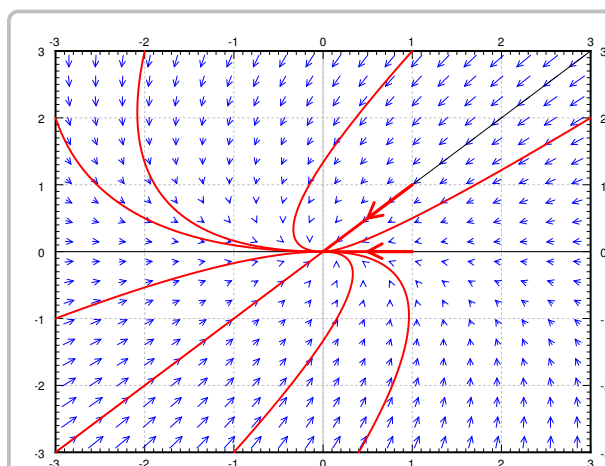


Figure 3.7: Example sink vector field with eigenvectors and solutions.

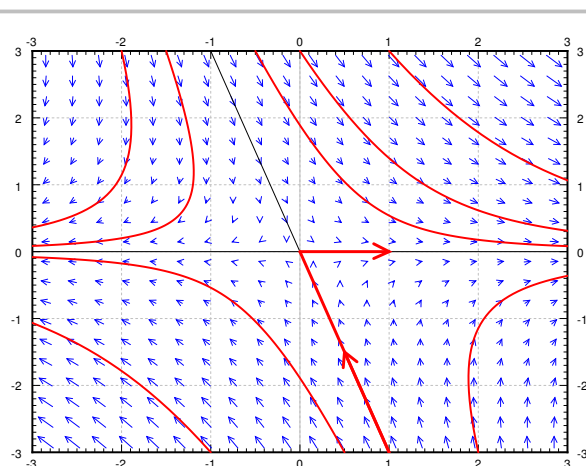


Figure 3.8: Example saddle vector field with eigenvectors and solutions.

Case 3. Suppose one eigenvalue is positive and one is negative. For example, consider $P = \begin{bmatrix} 1 & 1 \\ 0 & -2 \end{bmatrix}$. The eigenvalues are 1 and -2 and corresponding eigenvectors are $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -3 \end{bmatrix}$. We reverse the arrows on the line corresponding to the negative eigenvalue and we obtain the picture in Figure 3.8. We call this picture a *saddle point*.

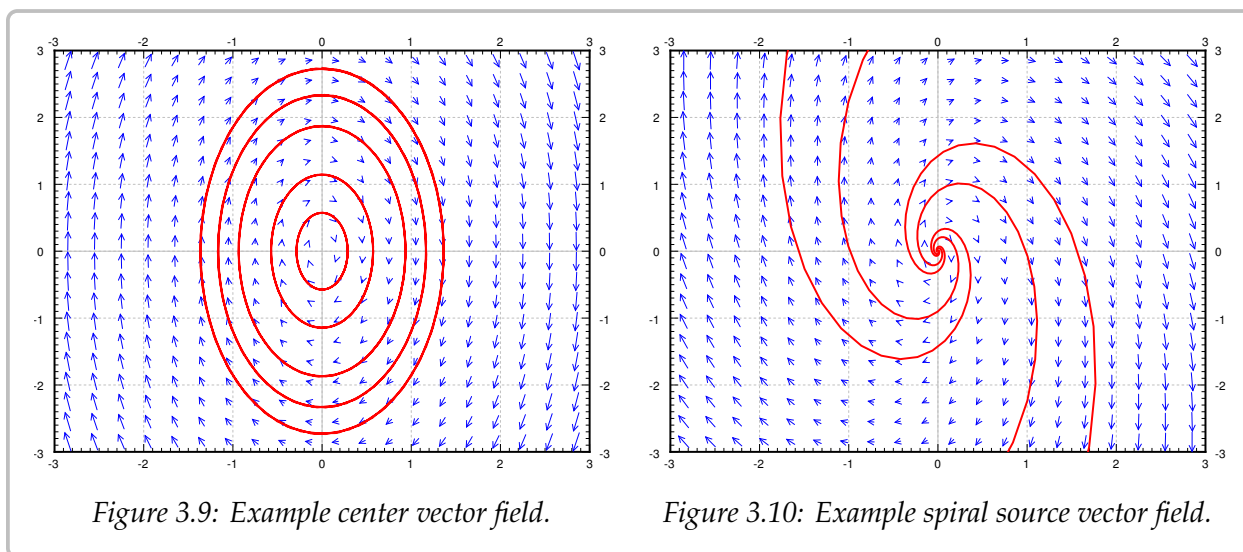
For the next three cases, we will assume the eigenvalues are complex. In this case the eigenvectors are also complex and we cannot just plot them in the plane.

Case 4. Suppose the eigenvalues are purely imaginary, that is, $\pm ib$. For example, let $P = \begin{bmatrix} 0 & 1 \\ -4 & 0 \end{bmatrix}$. The eigenvalues are $\pm 2i$ and corresponding eigenvectors are $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -2i \end{bmatrix}$. Consider the eigenvalue $2i$ and its eigenvector $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$. The real and imaginary parts of $\vec{v}e^{2it}$

are

$$\operatorname{Re} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{2it} = \begin{bmatrix} \cos(2t) \\ -2 \sin(2t) \end{bmatrix}, \quad \operatorname{Im} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{2it} = \begin{bmatrix} \sin(2t) \\ 2 \cos(2t) \end{bmatrix}.$$

We can take any linear combination of them to get other solutions, which one we take depends on the initial conditions. Now note that the real part is a parametric equation for an ellipse. Same with the imaginary part and in fact any linear combination of the two. This is what happens in general when the eigenvalues are purely imaginary. So when the eigenvalues are purely imaginary, we get *ellipses* for the solutions. This type of picture is sometimes called a *center*. See Figure 3.9.



Case 5. Now suppose the complex eigenvalues have a positive real part. That is, suppose the eigenvalues are $a \pm ib$ for some $a > 0$. For example, let $P = \begin{bmatrix} 1 & 1 \\ -4 & 1 \end{bmatrix}$. The eigenvalues turn out to be $1 \pm 2i$ and eigenvectors are $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -2i \end{bmatrix}$. We take $1 + 2i$ and its eigenvector $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$ and find the real and imaginary parts of $\vec{v}e^{(1+2i)t}$ are

$$\operatorname{Re} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{(1+2i)t} = e^t \begin{bmatrix} \cos(2t) \\ -2 \sin(2t) \end{bmatrix}, \quad \operatorname{Im} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{(1+2i)t} = e^t \begin{bmatrix} \sin(2t) \\ 2 \cos(2t) \end{bmatrix}.$$

Note the e^t in front of the solutions. The solutions grow in magnitude while spinning around the origin. Hence we get a *spiral source*. See Figure 3.10.

Case 6. Finally suppose the complex eigenvalues have a negative real part. That is, suppose the eigenvalues are $-a \pm ib$ for some $a > 0$. For example, let $P = \begin{bmatrix} -1 & 1 \\ 4 & -1 \end{bmatrix}$. The eigenvalues turn out to be $-1 \pm 2i$ and eigenvectors are $\begin{bmatrix} 1 \\ -2i \end{bmatrix}$ and $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$. We take $-1 - 2i$ and its eigenvector $\begin{bmatrix} 1 \\ 2i \end{bmatrix}$ and find the real and imaginary parts of $\vec{v}e^{(-1-2i)t}$ are

$$\operatorname{Re} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{(-1-2i)t} = e^{-t} \begin{bmatrix} \cos(2t) \\ 2 \sin(2t) \end{bmatrix}, \quad \operatorname{Im} \begin{bmatrix} 1 \\ 2i \end{bmatrix} e^{(-1-2i)t} = e^{-t} \begin{bmatrix} -\sin(2t) \\ 2 \cos(2t) \end{bmatrix}.$$

Note the e^{-t} in front of the solutions. The solutions shrink in magnitude while spinning around the origin. Hence we get a *spiral sink*. See Figure 3.11.

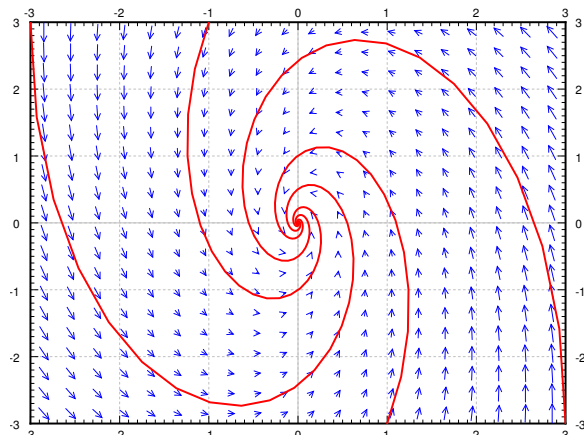


Figure 3.11: Example spiral sink vector field.

We summarize the behavior of linear homogeneous two-dimensional systems given by a nonsingular matrix in Table 3.1. Systems where one of the eigenvalues is zero (the matrix is singular) come up in practice from time to time, see Example 3.1.2 on page 121, and the pictures are somewhat different (simpler in a way). See the exercises. When the eigenvalues are real and repeated, the analysis is a little bit more difficult and we have not yet covered how to solve such systems, but the idea is roughly the same—a repeated positive eigenvalue means a source and a repeated negative eigenvalue means a sink.

Eigenvalues	Behavior
real and both positive	source / unstable node
real and both negative	sink / stable node
real and opposite signs	saddle
purely imaginary	center point / ellipses
complex with positive real part	spiral source
complex with negative real part	spiral sink

Table 3.1: Summary of behavior of linear homogeneous two-dimensional systems.

3.5.1 Exercises

Exercise 3.5.1: Take the equation $mx'' + cx' + kx = 0$, with $m > 0$, $c \geq 0$, $k > 0$ for the mass-spring system.

- Convert this to a system of first-order equations.
- Classify for what m, c, k do you get which behavior.
- Explain from physical intuition why you do not get all the different kinds of behavior here?

Exercise 3.5.2: What happens in the case when $P = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$? In this case the eigenvalue is repeated and there is only one independent eigenvector. Draw and describe the vector field.

Exercise 3.5.3: What happens in the case when $P = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$? Draw the vector field. Does this look like any of the pictures we have drawn?

Exercise 3.5.4: Which behaviors are possible if P is diagonal, that is, $P = \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix}$? You can assume that a and b are not zero.

Exercise 3.5.5: Take the system from [Example 3.1.2](#) on page 121, $x'_1 = \frac{r}{V}(x_2 - x_1)$, $x'_2 = \frac{r}{V}(x_1 - x_2)$. As we said, one of the eigenvalues is zero. What is the other eigenvalue, how does the picture look like, and what happens when t goes to infinity.

Exercise 3.5.101: Describe the behavior of the following systems without solving:

- $x' = x + y$, $y' = x - y$.
- $x'_1 = x_1 + x_2$, $x'_2 = 2x_2$.
- $x'_1 = -2x_2$, $x'_2 = 2x_1$.
- $x' = x + 3y$, $y' = -2x - 4y$.
- $x' = x - 4y$, $y' = -4x + y$.

Exercise 3.5.102: Suppose that $\vec{x}' = P\vec{x}$ where P is a 2 by 2 matrix with eigenvalues $2 \pm i$. Describe the behavior.

Exercise 3.5.103: Take $\begin{bmatrix} x \\ y \end{bmatrix}' = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}$. Draw the vector field and describe the behavior. Is it one of the behaviors that we have seen before?

3.6 Second-order systems and applications

Note: more than 2 lectures, §5.4 in [EP], not in [BD]

3.6.1 Undamped mass-spring systems

While we did say that we can usually only study first-order systems, it is sometimes more convenient to analyze the system in the way it arises naturally. For example, consider 3 masses connected by springs between two walls. We could pick any higher number, and the math would be essentially the same, but for simplicity, we pick 3 right now. We also assume no friction, that is, the system is undamped. The masses are m_1 , m_2 , and m_3 and the spring constants are k_1 , k_2 , k_3 , and k_4 . Let x_1 be the displacement from rest position of the first mass, and x_2 and x_3 the displacement of the second and third mass. We make, as usual, positive values go right (as x_1 grows, the first mass is moving right). See Figure 3.12.

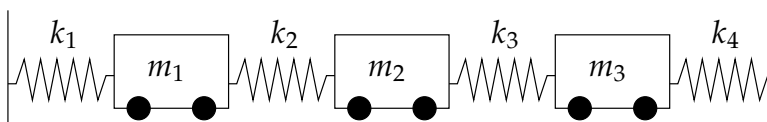


Figure 3.12: System of masses and springs.

This simple system turns up in unexpected places. For example, our world really consists of many small particles of matter interacting together. When we try the system above with many more masses, we obtain a good approximation to how an elastic material behaves. By somehow taking a limit of the number of masses going to infinity, we obtain the continuous one-dimensional wave equation (that we study in § 4.7). But we digress.

Let us set up the equations for the three mass system. By Hooke's law, the force acting on the mass equals the spring compression times the spring constant. By Newton's second law, force is mass times acceleration. So if we sum the forces acting on each mass, put the right sign in front of each term, depending on the direction in which it is acting, and set this equal to mass times the acceleration, we end up with the desired system of equations.

$$\begin{aligned} m_1 x_1'' &= -k_1 x_1 + k_2(x_2 - x_1) &= -(k_1 + k_2)x_1 + k_2 x_2, \\ m_2 x_2'' &= -k_2(x_2 - x_1) + k_3(x_3 - x_2) &= k_2 x_1 - (k_2 + k_3)x_2 + k_3 x_3, \\ m_3 x_3'' &= -k_3(x_3 - x_2) - k_4 x_3 &= k_3 x_2 - (k_3 + k_4)x_3. \end{aligned}$$

We define the matrices

$$M = \begin{bmatrix} m_1 & 0 & 0 \\ 0 & m_2 & 0 \\ 0 & 0 & m_3 \end{bmatrix} \quad \text{and} \quad K = \begin{bmatrix} -(k_1 + k_2) & k_2 & 0 \\ k_2 & -(k_2 + k_3) & k_3 \\ 0 & k_3 & -(k_3 + k_4) \end{bmatrix}.$$

We write the equation simply as

$$M\vec{x}'' = K\vec{x}.$$

At this point we could introduce 3 new variables and write out a system of 6 first-order equations. We claim this simple setup is easier to handle as a second-order system. We call $\vec{x}$ the *displacement vector*, M the *mass matrix*, and K the *stiffness matrix*.

Exercise 3.6.1: Repeat this setup for 4 masses (find the matrices M and K). Do it for 5 masses. Can you find a prescription to do it for n masses?

As with a single equation, we want to “divide by M .” This means computing the inverse of M . The masses are all nonzero and M is a diagonal matrix, so computing the inverse is easy:

$$M^{-1} = \begin{bmatrix} \frac{1}{m_1} & 0 & 0 \\ 0 & \frac{1}{m_2} & 0 \\ 0 & 0 & \frac{1}{m_3} \end{bmatrix}.$$

This fact follows readily by how we multiply diagonal matrices. As an exercise, you should verify that $MM^{-1} = M^{-1}M = I$.

Let $A = M^{-1}K$. We look at the system $\vec{x}'' = M^{-1}K\vec{x}$, or

$$\vec{x}'' = A\vec{x}.$$

Many real world systems can be modeled by this equation. For simplicity, we will only talk about the given masses-and-springs problem. We try a solution of the form

$$\vec{x} = \vec{v}e^{\alpha t}.$$

We compute that for this guess, $\vec{x}'' = \alpha^2\vec{v}e^{\alpha t}$. We plug our guess into the equation and get

$$\alpha^2\vec{v}e^{\alpha t} = A\vec{v}e^{\alpha t}.$$

We divide by $e^{\alpha t}$ to arrive at $\alpha^2\vec{v} = A\vec{v}$. Hence if α^2 is an eigenvalue of A and $\vec{v}$ is a corresponding eigenvector, we have found a solution.

In our example, and in other common applications, A has only real negative eigenvalues (and possibly a zero eigenvalue). So we study only this case. When an eigenvalue λ is negative, it means that $\alpha^2 = \lambda$ is negative. Hence there is some real number ω such that $-\omega^2 = \lambda$. Then $\alpha = \pm i\omega$. For $\alpha = i\omega$, the solution we guessed is $\vec{x} = \vec{v}e^{i\omega t}$ or

$$\vec{x} = \vec{v} (\cos(\omega t) + i \sin(\omega t)).$$

By taking the real and imaginary parts (note that $\vec{v}$ is real), we find that $\vec{v} \cos(\omega t)$ and $\vec{v} \sin(\omega t)$ are linearly independent solutions.

If an eigenvalue is zero, it turns out that both $\vec{v}$ and $\vec{v}t$ are solutions, where $\vec{v}$ is an eigenvector corresponding to the eigenvalue 0.

Exercise 3.6.2: Show that if A has a zero eigenvalue and $\vec{v}$ is a corresponding eigenvector, then $\vec{x} = \vec{v}(a + bt)$ is a solution of $\vec{x}'' = A\vec{x}$ for arbitrary constants a and b .

Theorem 3.6.1. Let A be a real $n \times n$ matrix with n distinct real negative (or zero) eigenvalues we denote by $-\omega_1^2 > -\omega_2^2 > \dots > -\omega_n^2$, and corresponding eigenvectors by $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$. If A is invertible (that is, if $\omega_1 > 0$), then

$$\vec{x}(t) = \sum_{i=1}^n \vec{v}_i (a_i \cos(\omega_i t) + b_i \sin(\omega_i t))$$

is the general solution of

$$\vec{x}'' = A\vec{x},$$

for some arbitrary constants a_i and b_i . If A has a zero eigenvalue, that is, $\omega_1 = 0$, and all other eigenvalues are distinct and negative, then the general solution can be written as

$$\vec{x}(t) = \vec{v}_1(a_1 + b_1 t) + \sum_{i=2}^n \vec{v}_i (a_i \cos(\omega_i t) + b_i \sin(\omega_i t)).$$

We use this solution and the setup from the introduction of this section even when some of the masses and springs are missing. For example, when there are only 2 masses and only 2 springs, simply take only the equations for the two masses and set all the spring constants for the springs that are missing to zero.

3.6.2 Examples

Example 3.6.1: Consider the setup in [Figure 3.13](#), with $m_1 = 2 \text{ kg}$, $m_2 = 1 \text{ kg}$, $k_1 = 4 \text{ N/m}$, and $k_2 = 2 \text{ N/m}$.

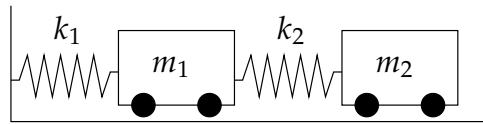


Figure 3.13: System of masses and springs.

The equations we write down are

$$\begin{bmatrix} 2 & 0 \\ 0 & 1 \end{bmatrix} \vec{x}'' = \begin{bmatrix} -(4+2) & 2 \\ 2 & -2 \end{bmatrix} \vec{x},$$

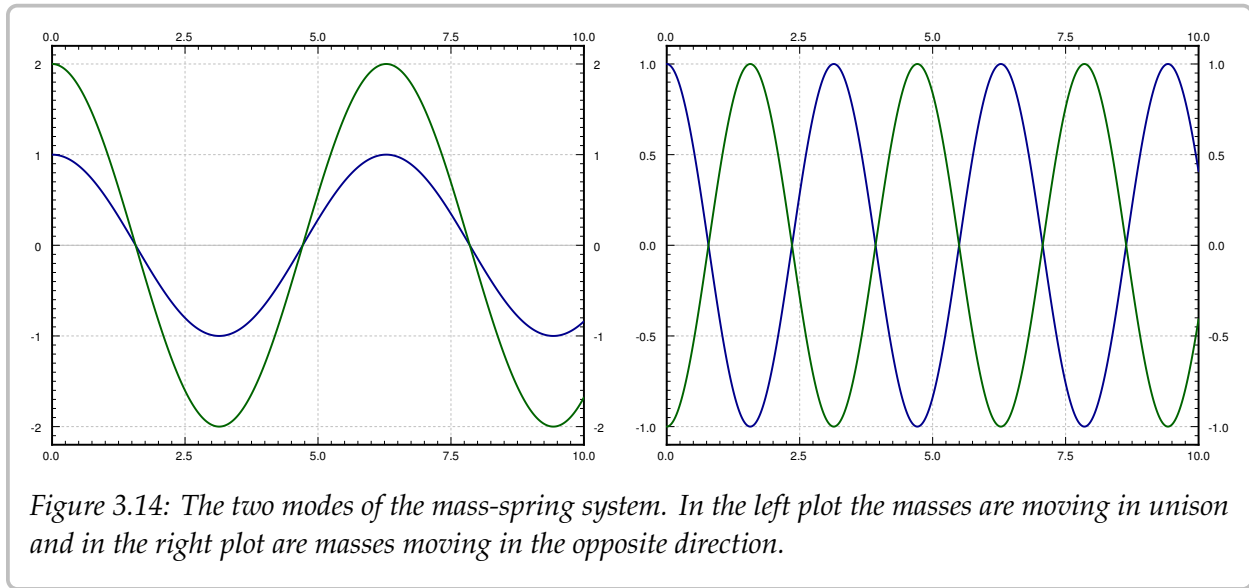
or

$$\vec{x}'' = \begin{bmatrix} -3 & 1 \\ 2 & -2 \end{bmatrix} \vec{x}.$$

We find the eigenvalues of A to be $\lambda = -1, -4$ (exercise). We find corresponding eigenvectors to be $\begin{bmatrix} 1 \\ 2 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$ respectively (exercise). We check the theorem and note that $\omega_1 = 1$ and $\omega_2 = 2$. Hence the general solution is

$$\vec{x} = \begin{bmatrix} 1 \\ 2 \end{bmatrix} (a_1 \cos(t) + b_1 \sin(t)) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (a_2 \cos(2t) + b_2 \sin(2t)).$$

The two terms in the solution represent the two so-called *natural* or *normal modes of oscillation*. And the two (angular) frequencies are the *natural frequencies*. The first natural frequency is 1, and the second is 2. The two modes are plotted in [Figure 3.14](#).



Let us write the solution as

$$\vec{x} = \begin{bmatrix} 1 \\ 2 \end{bmatrix} c_1 \cos(t - \alpha_1) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} c_2 \cos(2t - \alpha_2).$$

The first term,

$$\begin{bmatrix} 1 \\ 2 \end{bmatrix} c_1 \cos(t - \alpha_1) = \begin{bmatrix} c_1 \cos(t - \alpha_1) \\ 2c_1 \cos(t - \alpha_1) \end{bmatrix},$$

corresponds to the mode where the masses move synchronously in the same direction.

The second term,

$$\begin{bmatrix} 1 \\ -1 \end{bmatrix} c_2 \cos(2t - \alpha_2) = \begin{bmatrix} c_2 \cos(2t - \alpha_2) \\ -c_2 \cos(2t - \alpha_2) \end{bmatrix},$$

corresponds to the mode where the masses move synchronously but in opposite directions.

The general solution is a combination of the two modes. That is, the initial conditions determine the amplitude and phase shift of each mode. As an example, suppose we have initial conditions

$$\vec{x}(0) = \begin{bmatrix} 1 \\ -1 \end{bmatrix}, \quad \vec{x}'(0) = \begin{bmatrix} 0 \\ 6 \end{bmatrix}.$$

We use the a_j, b_j constants to solve for initial conditions. First

$$\begin{bmatrix} 1 \\ -1 \end{bmatrix} = \vec{x}(0) = \begin{bmatrix} 1 \\ 2 \end{bmatrix} a_1 + \begin{bmatrix} 1 \\ -1 \end{bmatrix} a_2 = \begin{bmatrix} a_1 + a_2 \\ 2a_1 - a_2 \end{bmatrix}.$$

We solve (exercise) to find $a_1 = 0, a_2 = 1$. To find the b_1 and b_2 , we differentiate first:

$$\vec{x}' = \begin{bmatrix} 1 \\ 2 \end{bmatrix} (-a_1 \sin(t) + b_1 \cos(t)) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (-2a_2 \sin(2t) + 2b_2 \cos(2t)).$$

Now we solve:

$$\begin{bmatrix} 0 \\ 6 \end{bmatrix} = \vec{x}'(0) = \begin{bmatrix} 1 \\ 2 \end{bmatrix} b_1 + \begin{bmatrix} 1 \\ -1 \end{bmatrix} 2b_2 = \begin{bmatrix} b_1 + 2b_2 \\ 2b_1 - 2b_2 \end{bmatrix}.$$

Again solve (exercise) to find $b_1 = 2, b_2 = -1$. So our solution is

$$\vec{x} = \begin{bmatrix} 1 \\ 2 \end{bmatrix} 2 \sin(t) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (\cos(2t) - \sin(2t)) = \begin{bmatrix} 2 \sin(t) + \cos(2t) - \sin(2t) \\ 4 \sin(t) - \cos(2t) + \sin(2t) \end{bmatrix}.$$

The graphs of the two displacements, x_1 and x_2 of the two carts is in [Figure 3.15](#).

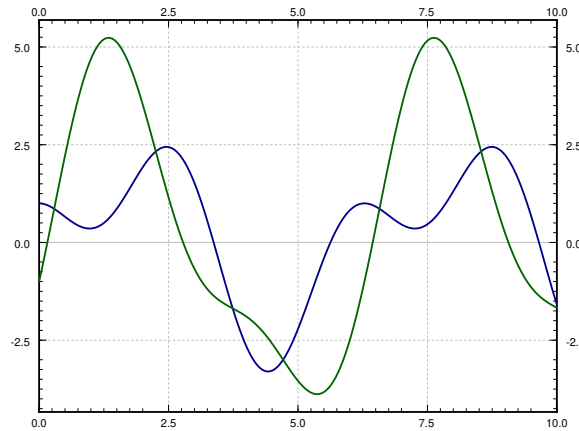
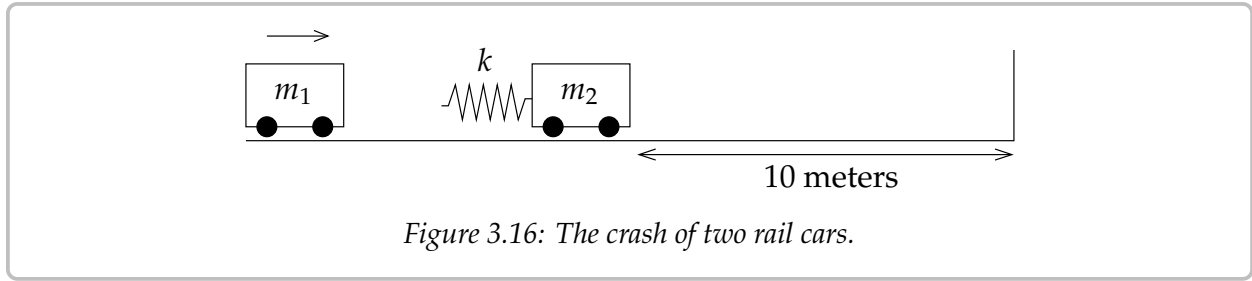


Figure 3.15: Superposition of the two modes given the initial conditions.

Example 3.6.2: We have two toy rail cars. Car 1 of mass 2 kg is traveling at 3 m/s towards the second rail car of mass 1 kg. There is a bumper on the second rail car that engages at the moment the cars hit (it connects to two cars) and does not let go. The bumper acts like a spring of spring constant $k = 2 \text{ N/m}$. The second car is 10 meters from a wall. See [Figure 3.16](#) on the next page.

We want to ask several questions. At what time after the cars link does impact with the wall happen? What is the speed of car 2 when it hits the wall?

OK, let us first set the system up. Let $t = 0$ be the time when the two cars link up. Let x_1 be the displacement of the first car from the position at $t = 0$, and let x_2 be the displacement



of the second car from its original location. Then the time when $x_2(t) = 10$ is exactly the time when impact with the wall occurs. For this t , $x'_2(t)$ is the speed at impact. This system acts just like the system of the previous example but without k_1 . Hence the equation is

$$\begin{bmatrix} 2 & 0 \\ 0 & 1 \end{bmatrix} \vec{x}'' = \begin{bmatrix} -2 & 2 \\ 2 & -2 \end{bmatrix} \vec{x},$$

or

$$\vec{x}'' = \begin{bmatrix} -1 & 1 \\ 2 & -2 \end{bmatrix} \vec{x}.$$

We compute the eigenvalues of A . It is not hard to see that the eigenvalues are 0 and -3 (exercise). Furthermore, eigenvectors are $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -2 \end{bmatrix}$ respectively (exercise). Then $\omega_1 = 0$, $\omega_2 = \sqrt{3}$, and by the second part of the theorem the general solution is

$$\begin{aligned} \vec{x} &= \begin{bmatrix} 1 \\ 1 \end{bmatrix} (a_1 + b_1 t) + \begin{bmatrix} 1 \\ -2 \end{bmatrix} \left(a_2 \cos(\sqrt{3} t) + b_2 \sin(\sqrt{3} t) \right) \\ &= \begin{bmatrix} a_1 + b_1 t + a_2 \cos(\sqrt{3} t) + b_2 \sin(\sqrt{3} t) \\ a_1 + b_1 t - 2a_2 \cos(\sqrt{3} t) - 2b_2 \sin(\sqrt{3} t) \end{bmatrix}. \end{aligned}$$

We now apply the initial conditions. First the cars start at position 0 so $x_1(0) = 0$ and $x_2(0) = 0$. The first car is traveling at 3 m/s, so $x'_1(0) = 3$ and the second car starts at rest, so $x'_2(0) = 0$. The first conditions says

$$\vec{0} = \vec{x}(0) = \begin{bmatrix} a_1 + a_2 \\ a_1 - 2a_2 \end{bmatrix}.$$

It is not hard to see that $a_1 = a_2 = 0$. We set $a_1 = 0$ and $a_2 = 0$ in $\vec{x}(t)$ and differentiate to get

$$\vec{x}'(t) = \begin{bmatrix} b_1 + \sqrt{3} b_2 \cos(\sqrt{3} t) \\ b_1 - 2\sqrt{3} b_2 \cos(\sqrt{3} t) \end{bmatrix}.$$

So

$$\begin{bmatrix} 3 \\ 0 \end{bmatrix} = \vec{x}'(0) = \begin{bmatrix} b_1 + \sqrt{3} b_2 \\ b_1 - 2\sqrt{3} b_2 \end{bmatrix}.$$

Solving these two equations, we find $b_1 = 2$ and $b_2 = \frac{1}{\sqrt{3}}$. Hence the position of our cars is (until the impact with the wall)

$$\vec{x} = \begin{bmatrix} 2t + \frac{1}{\sqrt{3}} \sin(\sqrt{3}t) \\ 2t - \frac{2}{\sqrt{3}} \sin(\sqrt{3}t) \end{bmatrix}.$$

Note how the presence of the zero eigenvalue resulted in a term containing t . This means that the cars will be traveling in the positive direction as time grows, which is what we expect.

What we are really interested in is the second expression, the one for x_2 . We have $x_2(t) = 2t - \frac{2}{\sqrt{3}} \sin(\sqrt{3}t)$. See Figure 3.17 for the plot of x_2 versus time.

Just from the graph we can see that the time of impact will be a little more than 5 seconds from time zero. For this we have to solve the equation $10 = x_2(t) = 2t - \frac{2}{\sqrt{3}} \sin(\sqrt{3}t)$. Using a computer (or even a graphing calculator) we find that $t_{\text{impact}} \approx 5.22$ seconds.

The speed of the second car is $x_2' = 2 - 2\cos(\sqrt{3}t)$. At the time of impact (5.22 seconds from $t = 0$) we get $x_2'(t_{\text{impact}}) \approx 3.85$. The maximum speed is the maximum of $2 - 2\cos(\sqrt{3}t)$, which is 4. We are traveling at almost the maximum speed when we hit the wall.

Suppose that Bob is a tiny person sitting on car 2. Bob has a Martini in his hand and would like not to spill it. Let us suppose Bob would not spill his Martini when the first car links up with car 2, but if car 2 hits the wall at any speed greater than zero, Bob will spill his drink. Suppose Bob can move car 2 a few meters towards or away from the wall (he cannot go all the way to the wall, nor can he get out of the way of the first car). Is there a “safe” distance for him to be at? A distance such that the impact with the wall is at zero speed?

The answer is yes. On Figure 3.17, note the “plateau” between $t = 3$ and $t = 4$. There is a point where the speed is zero. To find it we solve $x_2'(t) = 0$. This is when $\cos(\sqrt{3}t) = 1$ or in other words when $t = \frac{2\pi}{\sqrt{3}}, \frac{4\pi}{\sqrt{3}}, \dots$ and so on. We plug in the first value to obtain $x_2\left(\frac{2\pi}{\sqrt{3}}\right) = \frac{4\pi}{\sqrt{3}} \approx 7.26$. So a “safe” distance is about 7 and a quarter meters from the wall.

Alternatively Bob could move away from the wall towards the incoming car 2, where another safe distance is $x_2\left(\frac{4\pi}{\sqrt{3}}\right) = \frac{8\pi}{\sqrt{3}} \approx 14.51$ and so on. We can use all the different t such that $x_2'(t) = 0$. Of course $t = 0$ is also a solution, corresponding to $x_2 = 0$, but that means standing right at the wall.

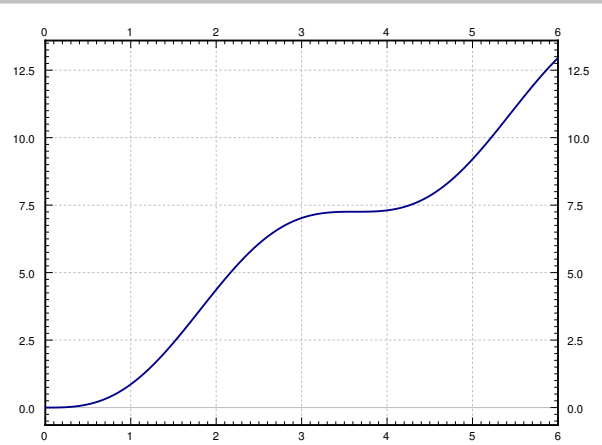


Figure 3.17: Position of the second car in time (ignoring the wall).

3.6.3 Forced oscillations

Finally we move to forced oscillations. Suppose that now our system is

$$\vec{x}'' = A\vec{x} + \vec{F} \cos(\omega t). \quad (3.4)$$

That is, we are adding periodic forcing to the system in the direction of the vector $\vec{F}$.

As before, this system just requires us to find one particular solution $\vec{x}_p$, add it to the general solution of the associated homogeneous system $\vec{x}_c$, and we will have the general solution to (3.4). Let us suppose that ω is not one of the natural frequencies of $\vec{x}'' = A\vec{x}$, then we can guess

$$\vec{x}_p = \vec{c} \cos(\omega t),$$

where $\vec{c}$ is an unknown constant vector. Note that we do not need to use sine since there are only second derivatives. We solve for $\vec{c}$ to find $\vec{x}_p$. This is really just the method of *undetermined coefficients* for systems. Let us differentiate $\vec{x}_p$ twice to get

$$\vec{x}_p'' = -\omega^2 \vec{c} \cos(\omega t).$$

Plug $\vec{x}_p$ and $\vec{x}_p''$ into equation (3.4):

$$\underbrace{\vec{x}_p''}_{-\omega^2 \vec{c} \cos(\omega t)} = \underbrace{A\vec{x}_p}_{A\vec{c} \cos(\omega t)} + \vec{F} \cos(\omega t).$$

We cancel out the cosine and rearrange the equation to obtain

$$(A + \omega^2 I)\vec{c} = -\vec{F}.$$

So

$$\vec{c} = (A + \omega^2 I)^{-1}(-\vec{F}).$$

Of course this is possible only if $(A + \omega^2 I) = (A - (-\omega^2)I)$ is invertible. That matrix is invertible if and only if $-\omega^2$ is not an eigenvalue of A . That is true if and only if ω is not a natural frequency of the system.

We simplified things a little bit. If we wish to have the forcing term to be in the units of force, say Newtons, then we must write

$$M\vec{x}'' = K\vec{x} + \vec{G} \cos(\omega t).$$

If we then write things in terms of $A = M^{-1}K$, we have

$$\vec{x}'' = M^{-1}K\vec{x} + M^{-1}\vec{G} \cos(\omega t) \quad \text{or} \quad \vec{x}'' = A\vec{x} + \vec{F} \cos(\omega t),$$

where $\vec{F} = M^{-1}\vec{G}$.

Example 3.6.3: Let us take the example in [Figure 3.13](#) on page 154 with the same parameters as before: $m_1 = 2$, $m_2 = 1$, $k_1 = 4$, and $k_2 = 2$. Now suppose that there is a force $2 \cos(3t)$ acting on the second cart.

The equation is

$$\begin{bmatrix} 2 & 0 \\ 0 & 1 \end{bmatrix} \vec{x}'' = \begin{bmatrix} -(4+2) & 2 \\ 2 & -2 \end{bmatrix} \vec{x} + \begin{bmatrix} 0 \\ 2 \end{bmatrix} \cos(3t) \quad \text{or} \quad \vec{x}'' = \begin{bmatrix} -3 & 1 \\ 2 & -2 \end{bmatrix} \vec{x} + \begin{bmatrix} 0 \\ 2 \end{bmatrix} \cos(3t).$$

We solved the associated homogeneous equation before and found the complementary solution to be

$$\vec{x}_c = \begin{bmatrix} 1 \\ 2 \end{bmatrix} (a_1 \cos(t) + b_1 \sin(t)) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (a_2 \cos(2t) + b_2 \sin(2t)).$$

The natural frequencies are 1 and 2. As 3 is not a natural frequency, we try $\vec{c} \cos(3t)$. We invert $(A + 3^2 I)$:

$$\left(\begin{bmatrix} -3 & 1 \\ 2 & -2 \end{bmatrix} + 3^2 I \right)^{-1} = \begin{bmatrix} 6 & 1 \\ 2 & 7 \end{bmatrix}^{-1} = \begin{bmatrix} \frac{7}{40} & \frac{-1}{40} \\ \frac{-1}{20} & \frac{3}{20} \end{bmatrix}.$$

Hence,

$$\vec{c} = (A + \omega^2 I)^{-1}(-\vec{F}) = \begin{bmatrix} \frac{7}{40} & \frac{-1}{40} \\ \frac{-1}{20} & \frac{3}{20} \end{bmatrix} \begin{bmatrix} 0 \\ -2 \end{bmatrix} = \begin{bmatrix} \frac{1}{20} \\ \frac{-3}{10} \end{bmatrix}.$$

Combining with the general solution of the associated homogeneous problem, we get that the general solution to $\vec{x}'' = A\vec{x} + \vec{F} \cos(\omega t)$ is

$$\vec{x} = \vec{x}_c + \vec{x}_p = \begin{bmatrix} 1 \\ 2 \end{bmatrix} (a_1 \cos(t) + b_1 \sin(t)) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (a_2 \cos(2t) + b_2 \sin(2t)) + \begin{bmatrix} \frac{1}{20} \\ \frac{-3}{10} \end{bmatrix} \cos(3t).$$

We would then solve for the constants a_1, a_2, b_1 , and b_2 using any given initial conditions.

Note that given force $\vec{f}$, we write the equation as $M\vec{x}'' = K\vec{x} + \vec{f}$ to get the units right. Then we write $\vec{x}'' = M^{-1}K\vec{x} + M^{-1}\vec{f}$. The term $\vec{g} = M^{-1}\vec{f}$ in $\vec{x}'' = A\vec{x} + \vec{g}$ is in units of force per unit mass.

If ω is a natural frequency of the system, *resonance* may occur, because we will have to try a particular solution of the form

$$\vec{x}_p = \vec{c} t \sin(\omega t) + \vec{d} \cos(\omega t).$$

That is assuming that the eigenvalues of the coefficient matrix are distinct. Next, note that the amplitude of this solution grows without bound as t grows.

3.6.4 Exercises

Exercise 3.6.3: Find a particular solution to

$$\vec{x}'' = \begin{bmatrix} -3 & 1 \\ 2 & -2 \end{bmatrix} \vec{x} + \begin{bmatrix} 0 \\ 2 \end{bmatrix} \cos(2t).$$

Exercise 3.6.4 (challenging): Consider the example in [Figure 3.13](#) on page 154 with the same parameters as before: $m_1 = 2$, $k_1 = 4$, and $k_2 = 2$, except for m_2 , which is unknown. Suppose that there is a force $\cos(5t)$ acting on the first mass. Find an m_2 such that there exists a particular solution where the first mass does not move.

Note: This idea is called **dynamic damping**. In practice there will be a small amount of damping and so any transient solution will disappear and after long enough time, the first mass will always come to a stop.

Exercise 3.6.5: Let us take the [Example 3.6.2](#) on page 156, but that at time of impact, car 2 is moving to the left with speed 3 m/s.

- Find the behavior of the system after linkup.
- Will the second car hit the wall, or will it be moving away from the wall as time goes on?
- At what speed would the first car have to be traveling for the system to essentially stay in place after linkup?

Exercise 3.6.6: Let us take the example in [Figure 3.13](#) on page 154 with parameters $m_1 = m_2 = 1$, $k_1 = k_2 = 1$. Does there exist a set of initial conditions for which the first cart moves but the second cart does not? If so, find those conditions. If not, argue why not.

Exercise 3.6.101: Find the general solution to $\begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \vec{x}'' = \begin{bmatrix} -3 & 0 & 0 \\ 2 & -4 & 0 \\ 0 & 6 & -3 \end{bmatrix} \vec{x} + \begin{bmatrix} \cos(2t) \\ 0 \\ 0 \end{bmatrix}$.

Exercise 3.6.102: Suppose there are three carts of equal mass m and connected by two springs of constant k (and no connections to walls). Set up the system and find its general solution.

Exercise 3.6.103: Suppose a cart of mass 2 kg is attached by a spring of constant $k = 1$ to a cart of mass 3 kg, which is attached to the wall by a spring also of constant $k = 1$. Suppose that the initial position of the first cart is 1 meter in the positive direction from the rest position, and the second cart starts at the rest position. The carts are not moving and are let go. Find the position of the second cart as a function of time.

3.7 Repeated eigenvalues

Note: 1 or 1.5 lectures, §5.5 in [EP], §7.8 in [BD]

It may happen that a matrix A has some “repeated” eigenvalues. That is, the characteristic equation $\det(A - \lambda I) = 0$ may have repeated roots. This is actually unlikely to happen for a random matrix. If we take a small perturbation of A (we change the entries of A slightly), we get a matrix with distinct eigenvalues. As any system we want to solve in practice is an approximation to reality anyway, it is not absolutely indispensable to know how to solve these corner cases. On the other hand, these cases do come up in applications from time to time. Furthermore, if we have distinct but very close eigenvalues, the behavior is similar to that of repeated eigenvalues, and so understanding that case will give us insight into what is going on.

3.7.1 Geometric multiplicity

Take the diagonal matrix

$$A = \begin{bmatrix} 3 & 0 \\ 0 & 3 \end{bmatrix}.$$

A has an eigenvalue 3 of multiplicity 2. We call the multiplicity of the eigenvalue in the characteristic equation the *algebraic multiplicity*. In this case, there also exist 2 linearly independent eigenvectors, $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$ corresponding to the eigenvalue 3. This means that the so-called *geometric multiplicity* of this eigenvalue is also 2.

In all the theorems where we required a matrix to have n distinct eigenvalues, we only really needed to have n linearly independent eigenvectors. For example, $\vec{x}' = A\vec{x}$ has the general solution

$$\vec{x} = c_1 \begin{bmatrix} 1 \\ 0 \end{bmatrix} e^{3t} + c_2 \begin{bmatrix} 0 \\ 1 \end{bmatrix} e^{3t}.$$

We restate the theorem about real eigenvalues. In the theorem, we will repeat eigenvalues according to (algebraic) multiplicity. So for the matrix A above, we would say that it has eigenvalues 3 and 3.

Theorem 3.7.1. Suppose the $n \times n$ matrix A has n real eigenvalues (not necessarily distinct), $\lambda_1, \lambda_2, \dots, \lambda_n$, and there are n linearly independent corresponding eigenvectors $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$. Then the general solution to $\vec{x}' = A\vec{x}$ can be written as

$$\vec{x} = c_1 \vec{v}_1 e^{\lambda_1 t} + c_2 \vec{v}_2 e^{\lambda_2 t} + \dots + c_n \vec{v}_n e^{\lambda_n t}.$$

The *geometric multiplicity* of an eigenvalue is equal to the number of corresponding linearly independent eigenvectors. The geometric multiplicity is always less than or equal to the algebraic multiplicity. The theorem handles the case when these two multiplicities are equal for all eigenvalues. If for an eigenvalue the geometric multiplicity is equal to the algebraic multiplicity, then we say the eigenvalue is *complete*.

In other words, the hypothesis of the theorem could be stated as saying that if all the eigenvalues of A are complete, then there are n linearly independent eigenvectors and thus we have the given general solution.

If the geometric multiplicity of an eigenvalue is 2 or greater, then the set of linearly independent eigenvectors is not unique up to multiples as it was before. For example, for the diagonal matrix $A = \begin{bmatrix} 3 & 0 \\ 0 & 3 \end{bmatrix}$ we could also pick eigenvectors $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$, or in fact any pair of two linearly independent vectors. The number of linearly independent eigenvectors corresponding to λ is the number of free variables we obtain when solving $A\vec{v} = \lambda\vec{v}$. We pick specific values for those free variables to obtain eigenvectors. If you pick different values, you may get different eigenvectors.

3.7.2 Defective eigenvalues

If an $n \times n$ matrix has less than n linearly independent eigenvectors, it is said to be *deficient*. Then there is at least one eigenvalue with its geometric multiplicity that is less than its algebraic multiplicity. We call this eigenvalue *defective* and the difference between the two multiplicities we call the *defect*. In other words, the defect is the number of eigenvectors we are “missing.”

Example 3.7.1: The matrix

$$\begin{bmatrix} 3 & 1 \\ 0 & 3 \end{bmatrix}$$

has an eigenvalue 3 of algebraic multiplicity 2. As $A - \lambda I = \begin{bmatrix} 3 & 1 \\ 0 & 3 \end{bmatrix} - 3 \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$, to compute the eigenvectors, we must solve

$$\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \vec{0}.$$

So $v_2 = 0$. All eigenvectors are of the form $\begin{bmatrix} v_1 \\ 0 \end{bmatrix}$. Any two such vectors are linearly dependent, and hence the geometric multiplicity of the eigenvalue is 1. Therefore, the defect is 1, and we can no longer apply the eigenvalue method directly to a system of ODEs with such a coefficient matrix.

To solve such an ODE, we need a new idea. Roughly, the key observation we will use is that if λ is an eigenvalue of A of algebraic multiplicity m , then it is possible to find certain m linearly independent vectors solving $(A - \lambda I)^k \vec{v} = \vec{0}$ for various powers k . We will call these *generalized eigenvectors*.

We continue with the example equation $\vec{x}' = A\vec{x}$ where $A = \begin{bmatrix} 3 & 1 \\ 0 & 3 \end{bmatrix}$. We found an eigenvalue $\lambda = 3$ of (algebraic) multiplicity 2 and defect 1. We found one eigenvector $\vec{v} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$. We have one solution

$$\vec{x}_1 = \vec{v}e^{3t} = \begin{bmatrix} 1 \\ 0 \end{bmatrix} e^{3t}.$$

We are now stuck, we get no other solutions from standard eigenvectors. But we need two linearly independent solutions to find the general solution of the equation.

In the spirit of repeated roots of the characteristic equation for a single equation, we try another solution of the form

$$\vec{x}_2 = (\vec{v}_2 + \vec{v}_1 t) e^{3t}.$$

We differentiate to get

$$\vec{x}_2' = \vec{v}_1 e^{3t} + 3(\vec{v}_2 + \vec{v}_1 t) e^{3t} = (3\vec{v}_2 + \vec{v}_1) e^{3t} + 3\vec{v}_1 t e^{3t}.$$

As we are assuming that $\vec{x}_2$ is a solution, $\vec{x}_2'$ must equal $A\vec{x}_2$. So we compute $A\vec{x}_2$:

$$A\vec{x}_2 = A(\vec{v}_2 + \vec{v}_1 t) e^{3t} = A\vec{v}_2 e^{3t} + A\vec{v}_1 t e^{3t}.$$

By looking at the coefficients of e^{3t} and $t e^{3t}$, we see $3\vec{v}_2 + \vec{v}_1 = A\vec{v}_2$ and $3\vec{v}_1 = A\vec{v}_1$. This means that

$$(A - 3I)\vec{v}_2 = \vec{v}_1, \quad \text{and} \quad (A - 3I)\vec{v}_1 = \vec{0}.$$

Therefore, $\vec{x}_2$ is a solution if these two equations are satisfied. The second equation is satisfied if $\vec{v}_1$ is an eigenvector, and we found the eigenvector above, so let $\vec{v}_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$. So, if we can find a $\vec{v}_2$ that solves $(A - 3I)\vec{v}_2 = \vec{v}_1$, then we are done. This is just a bunch of linear equations to solve and we are by now very good at that. Let us solve $(A - 3I)\vec{v}_2 = \vec{v}_1$. Write

$$\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} a \\ b \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}.$$

By inspection, we see that letting $a = 0$ (a could be anything in fact) and $b = 1$ does the job. Hence we can take $\vec{v}_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$. Our general solution to $\vec{x}' = A\vec{x}$ is

$$\vec{x} = c_1 \begin{bmatrix} 1 \\ 0 \end{bmatrix} e^{3t} + c_2 \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \end{bmatrix} t \right) e^{3t} = \begin{bmatrix} c_1 e^{3t} + c_2 t e^{3t} \\ c_2 e^{3t} \end{bmatrix}.$$

Let us check that we really do have the solution. First $x_1' = c_1 3e^{3t} + c_2 e^{3t} + 3c_2 t e^{3t} = 3x_1 + x_2$. Good. Now $x_2' = 3c_2 e^{3t} = 3x_2$. Good.

In the example, if we plug $(A - 3I)\vec{v}_2 = \vec{v}_1$ into $(A - 3I)\vec{v}_1 = \vec{0}$ we find

$$(A - 3I)(A - 3I)\vec{v}_2 = \vec{0}, \quad \text{or} \quad (A - 3I)^2 \vec{v}_2 = \vec{0}.$$

If $(A - 3I)\vec{w} \neq \vec{0}$ and $(A - 3I)^2 \vec{w} = \vec{0}$, then $(A - 3I)\vec{w}$ is an eigenvector, a multiple of $\vec{v}_1$. In this 2×2 case, $(A - 3I)^2$ is just the zero matrix (exercise). So any vector $\vec{w}$ solves $(A - 3I)^2 \vec{w} = \vec{0}$ and we just need a $\vec{w}$ such that $(A - 3I)\vec{w} \neq \vec{0}$. Then we could use $\vec{w}$ for $\vec{v}_2$ and $(A - 3I)\vec{w}$ for $\vec{v}_1$.

Note this example system $\vec{x}' = A\vec{x}$ has a simpler solution since A is a so-called *upper triangular matrix*, that is, every entry below the diagonal is zero. In particular, the equation for x_2 does not depend on x_1 . Mind you, not every defective matrix is triangular.

Exercise 3.7.1: Solve $\vec{x}' = \begin{bmatrix} 3 & 1 \\ 0 & 3 \end{bmatrix} \vec{x}$ by first solving for x_2 and then for x_1 independently. Check that you got the same solution as we did above.

Let us describe the general algorithm. Suppose that λ is an eigenvalue of multiplicity 2, defect 1. First find an eigenvector $\vec{v}_1$ of λ . That is, $\vec{v}_1$ solves $(A - \lambda I)\vec{v}_1 = \vec{0}$. Then, find a vector $\vec{v}_2$ such that

$$(A - \lambda I)\vec{v}_2 = \vec{v}_1.$$

This gives us two linearly independent solutions

$$\begin{aligned}\vec{x}_1 &= \vec{v}_1 e^{\lambda t}, \\ \vec{x}_2 &= (\vec{v}_2 + \vec{v}_1 t) e^{\lambda t}.\end{aligned}$$

Example 3.7.2: Consider the system

$$\vec{x}' = \begin{bmatrix} 2 & -5 & 0 \\ 0 & 2 & 0 \\ -1 & 4 & 1 \end{bmatrix} \vec{x}.$$

Compute the eigenvalues,

$$0 = \det(A - \lambda I) = \det \left(\begin{bmatrix} 2 - \lambda & -5 & 0 \\ 0 & 2 - \lambda & 0 \\ -1 & 4 & 1 - \lambda \end{bmatrix} \right) = (2 - \lambda)^2(1 - \lambda).$$

The eigenvalues are 1 and 2, where 2 has multiplicity 2. We leave it to the reader to find that $\begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}$ is an eigenvector for the eigenvalue $\lambda = 1$.

We focus on $\lambda = 2$. We compute eigenvectors:

$$\vec{0} = (A - 2I)\vec{v} = \begin{bmatrix} 0 & -5 & 0 \\ 0 & 0 & 0 \\ -1 & 4 & -1 \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \\ v_3 \end{bmatrix}.$$

The first equation says that $v_2 = 0$, so the last equation is $-v_1 - v_3 = 0$. Let v_3 be the free variable to find that $v_1 = -v_3$. Perhaps let $v_3 = -1$ to find an eigenvector $\begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix}$. Problem is that setting v_3 to anything else just gets multiples of this vector and so we have a defect of 1. Let $\vec{v}_1$ be the eigenvector and we look for a generalized eigenvector $\vec{v}_2$:

$$(A - 2I)\vec{v}_2 = \vec{v}_1,$$

or

$$\begin{bmatrix} 0 & -5 & 0 \\ 0 & 0 & 0 \\ -1 & 4 & -1 \end{bmatrix} \begin{bmatrix} a \\ b \\ c \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix},$$

where we used a, b, c as components of $\vec{v}_2$ for simplicity. The first equation says $-5b = 1$ so $b = -1/5$. The second equation says nothing. The last equation is $-a + 4b - c = -1$, or

$a + 4/5 + c = 1$, or $a + c = 1/5$. We let c be the free variable, and we choose $c = 0$. We find $\vec{v}_2 = \begin{bmatrix} 1/5 \\ -1/5 \\ 0 \end{bmatrix}$.

The general solution is therefore,

$$\vec{x} = c_1 \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} e^t + c_2 \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} e^{2t} + c_3 \left(\begin{bmatrix} 1/5 \\ -1/5 \\ 0 \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} t \right) e^{2t}.$$

This machinery can also be generalized to higher multiplicities and higher defects. We will not go over this method in detail, but we sketch the idea. Suppose that A has an eigenvalue λ of multiplicity m . We find vectors such that

$$(A - \lambda I)^k \vec{v} = \vec{0}, \quad \text{but} \quad (A - \lambda I)^{k-1} \vec{v} \neq \vec{0}.$$

Such vectors are called *generalized eigenvectors* (then $\vec{v}_1 = (A - \lambda I)^{k-1} \vec{v}$ is an eigenvector). For the eigenvector $\vec{v}_1$ there is a chain of generalized eigenvectors $\vec{v}_2$ through $\vec{v}_k$ such that:

$$\begin{aligned} (A - \lambda I) \vec{v}_1 &= \vec{0}, \\ (A - \lambda I) \vec{v}_2 &= \vec{v}_1, \\ &\vdots \\ (A - \lambda I) \vec{v}_k &= \vec{v}_{k-1}. \end{aligned}$$

Really once you find the $\vec{v}_k$ such that $(A - \lambda I)^k \vec{v}_k = \vec{0}$ but $(A - \lambda I)^{k-1} \vec{v}_k \neq \vec{0}$, you find the entire chain since you can compute the rest, $\vec{v}_{k-1} = (A - \lambda I) \vec{v}_k$, $\vec{v}_{k-2} = (A - \lambda I) \vec{v}_{k-1}$, etc. We form the linearly independent solutions

$$\begin{aligned} \vec{x}_1 &= \vec{v}_1 e^{\lambda t}, \\ \vec{x}_2 &= (\vec{v}_2 + \vec{v}_1 t) e^{\lambda t}, \\ &\vdots \\ \vec{x}_k &= \left(\vec{v}_k + \vec{v}_{k-1} t + \vec{v}_{k-2} \frac{t^2}{2} + \cdots + \vec{v}_2 \frac{t^{k-2}}{(k-2)!} + \vec{v}_1 \frac{t^{k-1}}{(k-1)!} \right) e^{\lambda t}. \end{aligned}$$

Recall that $k! = 1 \cdot 2 \cdot 3 \cdots (k-1) \cdot k$ is the factorial. If you have an eigenvalue of geometric multiplicity ℓ , you will have to find ℓ such chains (some of them might be short: just the single eigenvector equation). We go until we form m linearly independent solutions where m is the algebraic multiplicity. We do not quite know which specific eigenvectors go with which chain, so start by finding $\vec{v}_k$ first for the longest possible chain and go from there.

For example, if λ is an eigenvalue of A of algebraic multiplicity 3 and defect 2, then solve

$$(A - \lambda I) \vec{v}_1 = \vec{0}, \quad (A - \lambda I) \vec{v}_2 = \vec{v}_1, \quad (A - \lambda I) \vec{v}_3 = \vec{v}_2.$$

That is, find $\vec{v}_3$ such that $(A - \lambda I)^3 \vec{v}_3 = \vec{0}$, but $(A - \lambda I)^2 \vec{v}_3 \neq \vec{0}$. Then you are done as $\vec{v}_2 = (A - \lambda I)\vec{v}_3$ and $\vec{v}_1 = (A - \lambda I)\vec{v}_2$. The 3 linearly independent solutions are

$$\vec{x}_1 = \vec{v}_1 e^{\lambda t}, \quad \vec{x}_2 = (\vec{v}_2 + \vec{v}_1 t) e^{\lambda t}, \quad \vec{x}_3 = \left(\vec{v}_3 + \vec{v}_2 t + \vec{v}_1 \frac{t^2}{2} \right) e^{\lambda t}.$$

If, on the other hand, A has an eigenvalue λ of algebraic multiplicity 3 and defect 1, then solve

$$(A - \lambda I)\vec{v}_1 = \vec{0}, \quad (A - \lambda I)\vec{v}_2 = \vec{0}, \quad (A - \lambda I)\vec{v}_3 = \vec{v}_2.$$

Here $\vec{v}_1$ and $\vec{v}_2$ are actual honest eigenvectors, and $\vec{v}_3$ is a generalized eigenvector. So there are two chains. To solve, first find a $\vec{v}_3$ such that $(A - \lambda I)^2 \vec{v}_3 = \vec{0}$, but $(A - \lambda I)\vec{v}_3 \neq \vec{0}$. Then $\vec{v}_2 = (A - \lambda I)\vec{v}_3$ is going to be an eigenvector. Then solve for an eigenvector $\vec{v}_1$ that is linearly independent from $\vec{v}_2$. You get 3 linearly independent solutions

$$\vec{x}_1 = \vec{v}_1 e^{\lambda t}, \quad \vec{x}_2 = \vec{v}_2 e^{\lambda t}, \quad \vec{x}_3 = (\vec{v}_3 + \vec{v}_2 t) e^{\lambda t}.$$

3.7.3 Exercises

Exercise 3.7.2: Let $A = \begin{bmatrix} 5 & -3 \\ 3 & -1 \end{bmatrix}$. Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.3: Let $A = \begin{bmatrix} 5 & -4 & 4 \\ 0 & 3 & 0 \\ -2 & 4 & -1 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.4: Let $A = \begin{bmatrix} 2 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$ in two different ways and verify you get the same answer.

Exercise 3.7.5: Let $A = \begin{bmatrix} 0 & 1 & 2 \\ -1 & -2 & -2 \\ -4 & 4 & 7 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.6: Let $A = \begin{bmatrix} 0 & 4 & -2 \\ -1 & -4 & 1 \\ 0 & 0 & -2 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.7: Let $A = \begin{bmatrix} 2 & 1 & -1 \\ -1 & 0 & 2 \\ -1 & -2 & 4 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.8: Suppose that A is a 2×2 matrix with a repeated eigenvalue λ . Suppose that there are two linearly independent eigenvectors. Show that $A = \lambda I$.

Exercise 3.7.101: Let $A = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.102: Let $A = \begin{bmatrix} 1 & 3 & 3 \\ 1 & 1 & 0 \\ -1 & 1 & 2 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.103: Let $A = \begin{bmatrix} 2 & 0 & 0 \\ -1 & -1 & 9 \\ 0 & -1 & 5 \end{bmatrix}$.

- What are the eigenvalues?
- What is/are the defect(s) of the eigenvalue(s)?
- Find the general solution of $\vec{x}' = A\vec{x}$.

Exercise 3.7.104: Let $A = \begin{bmatrix} a & a \\ b & c \end{bmatrix}$, where a , b , and c are unknowns. Suppose that 5 is a doubled eigenvalue of defect 1, and suppose that $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ is a corresponding eigenvector. Find A and show that there is only one such matrix A .

3.8 Matrix exponentials

Note: 2 lectures, may need to refer to § 7.1, §5.6 in [EP], §7.7 in [BD]

3.8.1 Definition

There is another way of finding a fundamental matrix solution of a system. Consider the constant-coefficient equation

$$\vec{x}' = P\vec{x}.$$

If this would be just one equation (when P is a number or a 1×1 matrix), then the solution would be

$$\vec{x} = e^{Pt}.$$

That does not (yet) make sense if P is a larger matrix, but essentially the same computation that led to the above works for matrices when we define e^{Pt} properly. First we write the Taylor series for e^{at} for some number a :

$$e^{at} = 1 + at + \frac{(at)^2}{2} + \frac{(at)^3}{6} + \frac{(at)^4}{24} + \cdots = \sum_{k=0}^{\infty} \frac{(at)^k}{k!}.$$

Recall $k! = 1 \cdot 2 \cdot 3 \cdots k$ is the factorial, and $0! = 1$. We differentiate this series term by term

$$\frac{d}{dt}(e^{at}) = 0 + a + a^2t + \frac{a^3t^2}{2} + \frac{a^4t^3}{6} + \cdots = a \left(1 + at + \frac{(at)^2}{2} + \frac{(at)^3}{6} + \cdots \right) = ae^{at}.$$

Maybe we can try the same trick with matrices. For an $n \times n$ matrix A , we define the *matrix exponential* as

$$e^A \stackrel{\text{def}}{=} I + A + \frac{1}{2}A^2 + \frac{1}{6}A^3 + \cdots + \frac{1}{k!}A^k + \cdots$$

We will not worry about convergence—the series really does always converge. We usually write Pt as tP by convention when P is a matrix. With this small change and by the exact same calculation as above, we have that

$$\frac{d}{dt}(e^{tP}) = Pe^{tP}.$$

Now P and hence e^{tP} is an $n \times n$ matrix. What we are looking for is a vector. In the 1×1 case we would at this point multiply by an arbitrary constant to get the general solution. In the matrix case we multiply by a column vector $\vec{c}$.

Theorem 3.8.1. Let P be an $n \times n$ matrix. Then the general solution to $\vec{x}' = P\vec{x}$ is

$$\vec{x} = e^{tP}\vec{c},$$

where $\vec{c}$ is an arbitrary constant vector. In fact, $\vec{x}(0) = \vec{c}$.

Check:

$$\frac{d}{dt}\vec{x} = \frac{d}{dt}\left(e^{tP}\vec{c}\right) = P e^{tP}\vec{c} = P\vec{x}.$$

Hence e^{tP} is a fundamental matrix solution of the homogeneous system. If we can compute the matrix exponential, we have another method of solving constant-coefficient homogeneous systems. It also makes it easy to solve for initial conditions. To solve $\vec{x}' = P\vec{x}$, $\vec{x}(0) = \vec{b}$, we take the solution

$$\vec{x} = e^{tP}\vec{b}.$$

This equation follows because $e^{0P} = I$, so $\vec{x}(0) = e^{0P}\vec{b} = \vec{b}$.

We mention a drawback of matrix exponentials. In general $e^{A+B} \neq e^A e^B$. The trouble is that matrices do not commute, that is, in general $AB \neq BA$. If you would try to rewrite $e^A e^B$ as e^{A+B} using the Taylor series, you will see why the lack of commutativity becomes a problem. It is true that if $AB = BA$, that is, if A and B commute, then $e^{A+B} = e^A e^B$. We will find this fact useful. We restate this as a theorem to make a point.

Theorem 3.8.2. *If $AB = BA$, then $e^{A+B} = e^A e^B$. Otherwise, $e^{A+B} \neq e^A e^B$ in general.*

3.8.2 Simple cases

In some instances it may work to just plug into the series definition. Suppose the matrix is diagonal. For example, $D = \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix}$. Then

$$D^k = \begin{bmatrix} a^k & 0 \\ 0 & b^k \end{bmatrix},$$

and

$$\begin{aligned} e^D &= I + D + \frac{1}{2}D^2 + \frac{1}{6}D^3 + \cdots \\ &= \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} + \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix} + \frac{1}{2} \begin{bmatrix} a^2 & 0 \\ 0 & b^2 \end{bmatrix} + \frac{1}{6} \begin{bmatrix} a^3 & 0 \\ 0 & b^3 \end{bmatrix} + \cdots = \begin{bmatrix} e^a & 0 \\ 0 & e^b \end{bmatrix}. \end{aligned}$$

In particular,

$$e^I = \begin{bmatrix} e & 0 \\ 0 & e \end{bmatrix} \quad \text{and} \quad e^{aI} = \begin{bmatrix} e^a & 0 \\ 0 & e^a \end{bmatrix}.$$

This makes exponentials of certain other matrices easy to compute. For example, the matrix $A = \begin{bmatrix} 5 & 4 \\ -1 & 1 \end{bmatrix}$ can be written as $3I + B$ where $B = \begin{bmatrix} 2 & 4 \\ -1 & -2 \end{bmatrix}$. Notice that $B^2 = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}$. So $B^k = 0$ for all $k \geq 2$. Therefore, $e^B = I + B$. Suppose we actually want to compute e^{tA} . The matrices $3tI$ and tB commute (exercise: check this) and $e^{tB} = I + tB$, since $(tB)^2 = t^2 B^2 = 0$. We write

$$\begin{aligned} e^{tA} &= e^{3tI+tB} = e^{3tI} e^{tB} = \begin{bmatrix} e^{3t} & 0 \\ 0 & e^{3t} \end{bmatrix} (I + tB) = \\ &= \begin{bmatrix} e^{3t} & 0 \\ 0 & e^{3t} \end{bmatrix} \begin{bmatrix} 1+2t & 4t \\ -t & 1-2t \end{bmatrix} = \begin{bmatrix} (1+2t)e^{3t} & 4te^{3t} \\ -te^{3t} & (1-2t)e^{3t} \end{bmatrix}. \end{aligned}$$

We found a fundamental matrix solution for the system $\vec{x}' = A\vec{x}$. Note that this matrix has a repeated eigenvalue with a defect; there is only one eigenvector for the eigenvalue 3. So we found a perhaps easier way to handle this case. In fact, if a matrix A is 2×2 and has an eigenvalue λ of multiplicity 2, then either $A = \lambda I$, or $A = \lambda I + B$ where $B^2 = 0$. This is a good exercise.

Exercise 3.8.1: Suppose that A is 2×2 and λ is the only eigenvalue. Show that $(A - \lambda I)^2 = 0$, and therefore that we can write $A = \lambda I + B$, where $B^2 = 0$ (and possibly $B = 0$). Hint: First write down what does it mean for the eigenvalue to be of multiplicity 2. You will get an equation for the entries. Now compute the square of B .

Matrices B such that $B^k = 0$ for some k are called *nilpotent*. Computation of the matrix exponential for nilpotent matrices is easy by just writing down the first k terms of the Taylor series.

3.8.3 General matrices

In general, the exponential is not as easy to compute as above. We usually cannot write a matrix as a sum of commuting matrices where the exponential is simple for each one. But fear not, it is still not too difficult provided we can find enough eigenvectors. First we need the following interesting result about matrix exponentials. For two square matrices A and B , with B invertible, we have

$$e^{BAB^{-1}} = Be^AB^{-1}.$$

This can be seen by writing down the Taylor series. First

$$(BAB^{-1})^2 = BAB^{-1}BAB^{-1} = BAIAB^{-1} = BA^2B^{-1}.$$

And by the same reasoning $(BAB^{-1})^k = BA^kB^{-1}$. Now write the Taylor series for $e^{BAB^{-1}}$:

$$\begin{aligned} e^{BAB^{-1}} &= I + BAB^{-1} + \frac{1}{2}(BAB^{-1})^2 + \frac{1}{6}(BAB^{-1})^3 + \dots \\ &= BB^{-1} + BAB^{-1} + \frac{1}{2}BA^2B^{-1} + \frac{1}{6}BA^3B^{-1} + \dots \\ &= B\left(I + A + \frac{1}{2}A^2 + \frac{1}{6}A^3 + \dots\right)B^{-1} \\ &= Be^AB^{-1}. \end{aligned}$$

Given a square matrix A , we can usually write $A = EDE^{-1}$, where D is diagonal and E invertible. This procedure is called *diagonalization*. If we can do that, the computation of the exponential becomes easy as e^D is just taking the exponential of the entries on the diagonal. Adding t into the mix, we can then compute the exponential

$$e^{tA} = Ee^{tD}E^{-1}.$$

To diagonalize A we need n linearly independent eigenvectors of A . Otherwise, this method of computing the exponential does not work and we need to be trickier, but we will not get into such details. Let E be the matrix with the eigenvectors as columns. Let $\lambda_1, \lambda_2, \dots, \lambda_n$ be the eigenvalues and let $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ be the eigenvectors, then $E = [\vec{v}_1 \ \vec{v}_2 \ \cdots \ \vec{v}_n]$. Make a diagonal matrix D with the eigenvalues on the diagonal:

$$D = \begin{bmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{bmatrix}.$$

We compute

$$\begin{aligned} AE &= A[\vec{v}_1 \ \vec{v}_2 \ \cdots \ \vec{v}_n] \\ &= [A\vec{v}_1 \ A\vec{v}_2 \ \cdots \ A\vec{v}_n] \\ &= [\lambda_1\vec{v}_1 \ \lambda_2\vec{v}_2 \ \cdots \ \lambda_n\vec{v}_n] \\ &= [\vec{v}_1 \ \vec{v}_2 \ \cdots \ \vec{v}_n]D \\ &= ED. \end{aligned}$$

The columns of E are linearly independent as these are linearly independent eigenvectors of A . Hence E is invertible. Since $AE = ED$, we multiply on the right by E^{-1} and we get

$$A = EDE^{-1}.$$

This means that $e^A = Ee^DE^{-1}$. Multiplying the matrix by t , we obtain

$$e^{tA} = Ee^{tD}E^{-1} = E \begin{bmatrix} e^{\lambda_1 t} & 0 & \cdots & 0 \\ 0 & e^{\lambda_2 t} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & e^{\lambda_n t} \end{bmatrix} E^{-1}. \quad (3.5)$$

The formula (3.5), therefore, gives the formula for computing a fundamental matrix solution e^{tA} for the system $\vec{x}' = A\vec{x}$, in the case where we have n linearly independent eigenvectors.

This computation still works when the eigenvalues and eigenvectors are complex, though then you have to compute with complex numbers. It is clear from the definition that if A is real, then e^{tA} is real. So you will only need complex numbers in the computation and not for the result. You may need to apply [Euler's formula](#) to simplify the result. If simplified properly, the final matrix will not have any complex numbers in it.

Example 3.8.1: Compute a fundamental matrix solution using the matrix exponential for the system

$$\begin{bmatrix} x \\ y \end{bmatrix}' = \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

Then compute the particular solution for the initial conditions $x(0) = 4$ and $y(0) = 2$.

Let A be the coefficient matrix $\begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix}$. We first compute (exercise) that the eigenvalues are 3 and -1 and corresponding eigenvectors are $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$. Hence the diagonalization of A is

$$\underbrace{\begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix}}_A = \underbrace{\begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}}_E \underbrace{\begin{bmatrix} 3 & 0 \\ 0 & -1 \end{bmatrix}}_D \underbrace{\begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}^{-1}}_{E^{-1}}.$$

We write

$$\begin{aligned} e^{tA} &= E e^{tD} E^{-1} = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} e^{3t} & 0 \\ 0 & e^{-t} \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}^{-1} \\ &= \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} e^{3t} & 0 \\ 0 & e^{-t} \end{bmatrix} \frac{-1}{2} \begin{bmatrix} -1 & -1 \\ -1 & 1 \end{bmatrix} \\ &= \frac{-1}{2} \begin{bmatrix} e^{3t} & e^{-t} \\ e^{3t} & -e^{-t} \end{bmatrix} \begin{bmatrix} -1 & -1 \\ -1 & 1 \end{bmatrix} \\ &= \frac{-1}{2} \begin{bmatrix} -e^{3t} - e^{-t} & -e^{3t} + e^{-t} \\ -e^{3t} + e^{-t} & -e^{3t} - e^{-t} \end{bmatrix} = \begin{bmatrix} \frac{e^{3t} + e^{-t}}{2} & \frac{e^{3t} - e^{-t}}{2} \\ \frac{e^{3t} - e^{-t}}{2} & \frac{e^{3t} + e^{-t}}{2} \end{bmatrix}. \end{aligned}$$

The initial conditions are $x(0) = 4$ and $y(0) = 2$. Hence, by the property that $e^{0A} = I$, we find that the particular solution we are looking for is $e^{tA} \vec{b}$ where $\vec{b}$ is $\begin{bmatrix} 4 \\ 2 \end{bmatrix}$. The particular solution we are looking for is

$$\begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} \frac{e^{3t} + e^{-t}}{2} & \frac{e^{3t} - e^{-t}}{2} \\ \frac{e^{3t} - e^{-t}}{2} & \frac{e^{3t} + e^{-t}}{2} \end{bmatrix} \begin{bmatrix} 4 \\ 2 \end{bmatrix} = \begin{bmatrix} 2e^{3t} + 2e^{-t} + e^{3t} - e^{-t} \\ 2e^{3t} - 2e^{-t} + e^{3t} + e^{-t} \end{bmatrix} = \begin{bmatrix} 3e^{3t} + e^{-t} \\ 3e^{3t} - e^{-t} \end{bmatrix}.$$

3.8.4 Fundamental matrix solutions

If you can compute a fundamental matrix solution in a different way, you can use this to find the matrix exponential e^{tA} . A fundamental matrix solution of a system of ODEs is not unique. The exponential is the fundamental matrix solution with the property that for $t = 0$ we get the identity matrix. So we must find the right fundamental matrix solution. Let X be any fundamental matrix solution to $\vec{x}' = A\vec{x}$. We claim

$$e^{tA} = X(t) [X(0)]^{-1}.$$

If we plug $t = 0$ into $X(t) [X(0)]^{-1}$, we get the identity as needed. We can multiply a fundamental matrix solution on the right by any constant invertible matrix and we still get a fundamental matrix solution. All we are doing is changing what are the arbitrary constants in the general solution $\vec{x}(t) = X(t) \vec{c}$.

3.8.5 Approximations

If you think about it, the computation of any fundamental matrix solution X using the eigenvalue method is just as difficult as the computation of e^{tA} . So perhaps we did not gain much by this new tool. However, the Taylor series expansion actually gives us a way to approximate solutions, which the eigenvalue method did not.

The simplest thing we can do is to just compute the series up to a certain number of terms. There are better ways to approximate the exponential*. In many cases, however, few terms of the Taylor series give a reasonable approximation for the exponential and may suffice for the application. For example, let us compute the first 4 terms of the series for the matrix $A = \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix}$.

$$\begin{aligned} e^{tA} &\approx I + tA + \frac{t^2}{2}A^2 + \frac{t^3}{6}A^3 = I + t \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix} + t^2 \begin{bmatrix} \frac{5}{2} & 2 \\ 2 & \frac{5}{2} \end{bmatrix} + t^3 \begin{bmatrix} \frac{13}{6} & \frac{7}{3} \\ \frac{7}{3} & \frac{13}{6} \end{bmatrix} = \\ &= \begin{bmatrix} 1 + t + \frac{5}{2}t^2 + \frac{13}{6}t^3 & 2t + 2t^2 + \frac{7}{3}t^3 \\ 2t + 2t^2 + \frac{7}{3}t^3 & 1 + t + \frac{5}{2}t^2 + \frac{13}{6}t^3 \end{bmatrix}. \end{aligned}$$

Just like the scalar version of the Taylor series approximation, the approximation will be better for small t and worse for larger t . For larger t , we will have to compute more terms. Let us see how we stack up against the real solution with $t = 0.1$. The approximate solution is approximately (rounded to 8 decimal places)

$$e^{0.1A} \approx I + 0.1A + \frac{0.1^2}{2}A^2 + \frac{0.1^3}{6}A^3 = \begin{bmatrix} 1.12716667 & 0.22233333 \\ 0.22233333 & 1.12716667 \end{bmatrix}.$$

Plugging $t = 0.1$ into the real solution (rounded to 8 decimal places), we get

$$e^{0.1A} = \begin{bmatrix} 1.12734811 & 0.22251069 \\ 0.22251069 & 1.12734811 \end{bmatrix}.$$

Not bad at all! Although if we take the same approximation for $t = 1$, we get

$$I + A + \frac{1}{2}A^2 + \frac{1}{6}A^3 = \begin{bmatrix} 6.66666667 & 6.33333333 \\ 6.33333333 & 6.66666667 \end{bmatrix},$$

while the real value is (again rounded to 8 decimal places)

$$e^A = \begin{bmatrix} 10.22670818 & 9.85882874 \\ 9.85882874 & 10.22670818 \end{bmatrix}.$$

So the approximation is not very good once we get up to $t = 1$. To get a good approximation at $t = 1$ (say up to 2 decimal places), we would need to go up to the 11th power (exercise).

*C. Moler and C.F. Van Loan, *Nineteen Dubious Ways to Compute the Exponential of a Matrix, Twenty-Five Years Later*, SIAM Review 45 (1), 2003, 3–49

3.8.6 Exercises

Exercise 3.8.2: Using the matrix exponential, find a fundamental matrix solution for the system $x' = 3x + y$, $y' = x + 3y$.

Exercise 3.8.3: Find e^{tA} for the matrix $A = \begin{bmatrix} 2 & 3 \\ 0 & 2 \end{bmatrix}$.

Exercise 3.8.4: Find a fundamental matrix solution for the system $x'_1 = 7x_1 + 4x_2 + 12x_3$, $x'_2 = x_1 + 2x_2 + x_3$, $x'_3 = -3x_1 - 2x_2 - 5x_3$. Then find the solution that satisfies $\vec{x}(0) = \begin{bmatrix} 0 \\ 1 \\ -2 \end{bmatrix}$.

Exercise 3.8.5: Compute the matrix exponential e^A for $A = \begin{bmatrix} 1 & 2 \\ 0 & 1 \end{bmatrix}$.

Exercise 3.8.6 (challenging): Suppose $AB = BA$. Show that under this assumption, $e^{A+B} = e^A e^B$.

Exercise 3.8.7: Use [Exercise 3.8.6](#) to show that $(e^A)^{-1} = e^{-A}$. In particular, e^A is invertible even if A is not.

Exercise 3.8.8: Let A be a 2×2 matrix with eigenvalues -1 , 1 , and corresponding eigenvectors $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$, $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$.

- Find matrix A with these properties.
- Find a fundamental matrix solution to $\vec{x}' = A\vec{x}$.
- Solve the system in with initial conditions $\vec{x}(0) = \begin{bmatrix} 2 \\ 3 \end{bmatrix}$.

Exercise 3.8.9: Suppose that A is an $n \times n$ matrix with a repeated eigenvalue λ of multiplicity n with n linearly independent eigenvectors. Show that the matrix is diagonal, in fact, $A = \lambda I$. Hint: Use diagonalization and the fact that the identity matrix commutes with every other matrix.

Exercise 3.8.10: Let $A = \begin{bmatrix} -1 & -1 \\ 1 & -3 \end{bmatrix}$.

- Find e^{tA} .
- Solve $\vec{x}' = A\vec{x}$, $\vec{x}(0) = \begin{bmatrix} 1 \\ -2 \end{bmatrix}$.

Exercise 3.8.11: Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$. Approximate e^{tA} by expanding the power series up to the third order.

Exercise 3.8.12: For any positive integer n , find a formula (or a recipe) for A^n for the following matrices:

- $\begin{bmatrix} 3 & 0 \\ 0 & 9 \end{bmatrix}$
- $\begin{bmatrix} 5 & 2 \\ 4 & 7 \end{bmatrix}$
- $\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$
- $\begin{bmatrix} 2 & 1 \\ 0 & 2 \end{bmatrix}$

Exercise 3.8.101: Compute e^{tA} where $A = \begin{bmatrix} 1 & -2 \\ -2 & 1 \end{bmatrix}$.

Exercise 3.8.102: Compute e^{tA} where $A = \begin{bmatrix} 1 & -3 & 2 \\ -2 & 1 & 2 \\ -1 & -3 & 4 \end{bmatrix}$.

Exercise 3.8.103:

a) Compute e^{tA} where $A = \begin{bmatrix} 3 & -1 \\ 1 & 1 \end{bmatrix}$.

b) Solve $\vec{x}' = A\vec{x}$ for $\vec{x}(0) = \begin{bmatrix} 1 \\ 2 \end{bmatrix}$.

Exercise 3.8.104: Compute the first 3 terms (up to the second degree) of the Taylor expansion of e^{tA} where $A = \begin{bmatrix} 2 & 3 \\ 2 & 2 \end{bmatrix}$. Write it as a single matrix. Then use it to approximate $e^{0.1A}$.

Exercise 3.8.105: For any positive integer n , find a formula (or a recipe) for A^n for the following matrices:

a) $\begin{bmatrix} 7 & 4 \\ -5 & -2 \end{bmatrix}$

b) $\begin{bmatrix} -3 & 4 \\ -6 & 7 \end{bmatrix}$

c) $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$

3.9 Nonhomogeneous systems

Note: 3 lectures (may have to skip a little), somewhat different from §5.7 in [EP], §7.9 in [BD]

3.9.1 First-order constant-coefficient

Integrating factor

We first focus on the nonhomogeneous first-order equation

$$\vec{x}'(t) = A\vec{x}(t) + \vec{f}(t),$$

where A is a constant matrix. The first method we look at is the *integrating factor method*. For simplicity, we rewrite the equation as

$$\vec{x}'(t) + P\vec{x}(t) = \vec{f}(t),$$

where $P = -A$. We multiply both sides of the equation by e^{tP} (being mindful that we are dealing with matrices that may not commute) to obtain

$$e^{tP}\vec{x}'(t) + e^{tP}P\vec{x}(t) = e^{tP}\vec{f}(t).$$

We notice that $Pe^{tP} = e^{tP}P$. This fact follows by writing down the series definition of e^{tP} :

$$\begin{aligned} Pe^{tP} &= P \left(I + tP + \frac{1}{2}(tP)^2 + \dots \right) = P + tP^2 + \frac{1}{2}t^2P^3 + \dots = \\ &= \left(I + tP + \frac{1}{2}(tP)^2 + \dots \right) P = e^{tP}P. \end{aligned}$$

So $\frac{d}{dt}(e^{tP}) = Pe^{tP} = e^{tP}P$. The product rule says

$$\frac{d}{dt}(e^{tP}\vec{x}(t)) = e^{tP}\vec{x}'(t) + e^{tP}P\vec{x}(t),$$

and so

$$\frac{d}{dt}(e^{tP}\vec{x}(t)) = e^{tP}\vec{f}(t).$$

We can now integrate. That is, we integrate each component of the vector separately

$$e^{tP}\vec{x}(t) = \int e^{tP}\vec{f}(t) dt + \vec{c}.$$

Recall from [Exercise 3.8.7](#) that $(e^{tP})^{-1} = e^{-tP}$. Therefore, we obtain

$$\vec{x}(t) = e^{-tP} \int e^{tP}\vec{f}(t) dt + e^{-tP}\vec{c}.$$

Perhaps it is better understood as a definite integral. In this case it will be easy to also solve for the initial conditions. Consider the equation with initial conditions

$$\vec{x}'(t) + P\vec{x}(t) = \vec{f}(t), \quad \vec{x}(0) = \vec{b}.$$

The solution can then be written as

$$\boxed{\vec{x}(t) = e^{-tP} \int_0^t e^{sP} \vec{f}(s) ds + e^{-tP} \vec{b}.} \quad (3.6)$$

Again, the integration means that each component of the vector $e^{sP} \vec{f}(s)$ is integrated separately. It is not hard to see that (3.6) really does satisfy the initial condition $\vec{x}(0) = \vec{b}$.

$$\vec{x}(0) = e^{-0P} \int_0^0 e^{sP} \vec{f}(s) ds + e^{-0P} \vec{b} = I\vec{b} = \vec{b}.$$

Example 3.9.1: Suppose that we have the system

$$\begin{aligned} x_1' + 5x_1 - 3x_2 &= e^t, \\ x_2' + 3x_1 - x_2 &= 0, \end{aligned}$$

with initial conditions $x_1(0) = 1, x_2(0) = 0$.

Let us write the system as

$$\vec{x}' + \begin{bmatrix} 5 & -3 \\ 3 & -1 \end{bmatrix} \vec{x} = \begin{bmatrix} e^t \\ 0 \end{bmatrix}, \quad \vec{x}(0) = \begin{bmatrix} 1 \\ 0 \end{bmatrix}.$$

The matrix $P = \begin{bmatrix} 5 & -3 \\ 3 & -1 \end{bmatrix}$ has a doubled eigenvalue 2 with defect 1, and we leave it as an exercise to double-check we computed e^{tP} correctly. Once we have e^{tP} , we find e^{-tP} , simply by negating t .

$$e^{tP} = \begin{bmatrix} (1+3t)e^{2t} & -3te^{2t} \\ 3te^{2t} & (1-3t)e^{2t} \end{bmatrix}, \quad e^{-tP} = \begin{bmatrix} (1-3t)e^{-2t} & 3te^{-2t} \\ -3te^{-2t} & (1+3t)e^{-2t} \end{bmatrix}.$$

Instead of computing the whole formula at once, let us do it in stages. First

$$\begin{aligned} \int_0^t e^{sP} \vec{f}(s) ds &= \int_0^t \begin{bmatrix} (1+3s)e^{2s} & -3se^{2s} \\ 3se^{2s} & (1-3s)e^{2s} \end{bmatrix} \begin{bmatrix} e^s \\ 0 \end{bmatrix} ds \\ &= \int_0^t \begin{bmatrix} (1+3s)e^{3s} \\ 3se^{3s} \end{bmatrix} ds \\ &= \begin{bmatrix} \int_0^t (1+3s)e^{3s} ds \\ \int_0^t 3se^{3s} ds \end{bmatrix} \\ &= \begin{bmatrix} te^{3t} \\ \frac{(3t-1)e^{3t}+1}{3} \end{bmatrix} \quad (\text{used integration by parts}). \end{aligned}$$

Then

$$\begin{aligned}
 \vec{x}(t) &= e^{-tP} \int_0^t e^{sP} \vec{f}(s) ds + e^{-tP} \vec{b} \\
 &= \begin{bmatrix} (1-3t)e^{-2t} & 3te^{-2t} \\ -3te^{-2t} & (1+3t)e^{-2t} \end{bmatrix} \begin{bmatrix} te^{3t} \\ \frac{(3t-1)e^{3t}+1}{3} \end{bmatrix} + \begin{bmatrix} (1-3t)e^{-2t} & 3te^{-2t} \\ -3te^{-2t} & (1+3t)e^{-2t} \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} \\
 &= \begin{bmatrix} te^{-2t} \\ -\frac{e^t}{3} + \left(\frac{1}{3} + t\right)e^{-2t} \end{bmatrix} + \begin{bmatrix} (1-3t)e^{-2t} \\ -3te^{-2t} \end{bmatrix} \\
 &= \begin{bmatrix} (1-2t)e^{-2t} \\ -\frac{e^t}{3} + \left(\frac{1}{3} - 2t\right)e^{-2t} \end{bmatrix}.
 \end{aligned}$$

Phew!

Let us check that this really works.

$$x'_1 + 5x_1 - 3x_2 = (4te^{-2t} - 4e^{-2t}) + 5(1-2t)e^{-2t} + e^t - (1-6t)e^{-2t} = e^t.$$

Similarly (exercise) $x'_2 + 3x_1 - x_2 = 0$. The initial conditions are also satisfied (exercise).

For systems, the integrating factor method only works if P does not depend on t , that is, P is constant. The problem is that in general

$$\frac{d}{dt} \left[e^{\int P(t) dt} \right] \neq P(t) e^{\int P(t) dt},$$

because matrix multiplication is not commutative.

Eigenvector decomposition

For the next method, note that eigenvectors of a matrix give the directions in which the matrix acts like a scalar. If we solve the system along these directions, the computations are simpler as we treat the matrix as a scalar. We then put those solutions together to get the general solution for the system.

Take the equation

$$\vec{x}'(t) = A\vec{x}(t) + \vec{f}(t). \quad (3.7)$$

Assume A has n linearly independent eigenvectors $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$. Write

$$\vec{x}(t) = \vec{v}_1 \xi_1(t) + \vec{v}_2 \xi_2(t) + \dots + \vec{v}_n \xi_n(t). \quad (3.8)$$

That is, we wish to write our solution as a linear combination of eigenvectors of A . If we solve for the scalar functions ξ_1 through ξ_n , we have our solution $\vec{x}$. Let us decompose $\vec{f}$ in terms of the eigenvectors as well. We wish to write

$$\vec{f}(t) = \vec{v}_1 g_1(t) + \vec{v}_2 g_2(t) + \dots + \vec{v}_n g_n(t). \quad (3.9)$$

That is, we wish to find g_1 through g_n that satisfy (3.9). Since all the eigenvectors are independent, the matrix $E = [\vec{v}_1 \quad \vec{v}_2 \quad \dots \quad \vec{v}_n]$ is invertible. Write the equation (3.9) as

$\vec{f} = E\vec{g}$, where the components of $\vec{g}$ are the functions g_1 through g_n . Then $\vec{g} = E^{-1}\vec{f}$. Hence it is always possible to find $\vec{g}$ when there are n linearly independent eigenvectors.

We plug (3.8) into (3.7), and note that $A\vec{v}_k = \lambda_k\vec{v}_k$:

$$\begin{aligned} \overbrace{\vec{v}_1\xi'_1 + \vec{v}_2\xi'_2 + \cdots + \vec{v}_n\xi'_n}^{\vec{x}'} &= \overbrace{A(\vec{v}_1\xi_1 + \vec{v}_2\xi_2 + \cdots + \vec{v}_n\xi_n)}^{A\vec{x}} + \overbrace{\vec{v}_1g_1 + \vec{v}_2g_2 + \cdots + \vec{v}_ng_n}^{\vec{f}} \\ &= A\vec{v}_1\xi_1 + A\vec{v}_2\xi_2 + \cdots + A\vec{v}_n\xi_n + \vec{v}_1g_1 + \vec{v}_2g_2 + \cdots + \vec{v}_ng_n \\ &= \vec{v}_1\lambda_1\xi_1 + \vec{v}_2\lambda_2\xi_2 + \cdots + \vec{v}_n\lambda_n\xi_n + \vec{v}_1g_1 + \vec{v}_2g_2 + \cdots + \vec{v}_ng_n \\ &= \vec{v}_1(\lambda_1\xi_1 + g_1) + \vec{v}_2(\lambda_2\xi_2 + g_2) + \cdots + \vec{v}_n(\lambda_n\xi_n + g_n). \end{aligned}$$

If we identify the coefficients of the vectors $\vec{v}_1$ through $\vec{v}_n$, we get the equations

$$\begin{aligned} \xi'_1 &= \lambda_1\xi_1 + g_1, \\ \xi'_2 &= \lambda_2\xi_2 + g_2, \\ &\vdots \\ \xi'_n &= \lambda_n\xi_n + g_n. \end{aligned}$$

Each one of these equations is independent of the others. They are all linear first-order equations and can easily be solved by the standard integrating factor method for single equations. That is, for the k^{th} equation we write

$$\xi'_k(t) - \lambda_k\xi_k(t) = g_k(t).$$

We use the integrating factor $e^{-\lambda_k t}$ to find that

$$\frac{d}{dt} \left[\xi_k(t) e^{-\lambda_k t} \right] = e^{-\lambda_k t} g_k(t).$$

We integrate and solve for ξ_k to get

$$\xi_k(t) = e^{\lambda_k t} \int e^{-\lambda_k t} g_k(t) dt + C_k e^{\lambda_k t}.$$

If we are looking for just any particular solution, we can set C_k to be zero. If we leave these constants in, we get the general solution. Write $\vec{x}(t) = \vec{v}_1\xi_1(t) + \vec{v}_2\xi_2(t) + \cdots + \vec{v}_n\xi_n(t)$, and we are done.

As always, it is perhaps better to write these integrals as definite integrals. Suppose that we have an initial condition $\vec{x}(0) = \vec{b}$. Take $\vec{a} = E^{-1}\vec{b}$ to find $\vec{b} = \vec{v}_1a_1 + \vec{v}_2a_2 + \cdots + \vec{v}_na_n$, just like before. Then if we write

$$\xi_k(t) = e^{\lambda_k t} \int_0^t e^{-\lambda_k s} g_k(s) ds + a_k e^{\lambda_k t},$$

we get the particular solution $\vec{x}(t) = \vec{v}_1\xi_1(t) + \vec{v}_2\xi_2(t) + \cdots + \vec{v}_n\xi_n(t)$ satisfying $\vec{x}(0) = \vec{b}$, because $\xi_k(0) = a_k$.

We remark that the technique we just outlined is the eigenvalue method applied to nonhomogeneous systems. If a system is homogeneous, that is, if $\vec{f} = \vec{0}$, then the equations we get are $\xi'_k = \lambda_k \xi_k$, and so $\xi_k = C_k e^{\lambda_k t}$ are the solutions and that is precisely what we got in § 3.4.

Example 3.9.2: Let $A = \begin{bmatrix} 1 & 3 \\ 3 & 1 \end{bmatrix}$. Solve $\vec{x}' = A\vec{x} + \vec{f}$ where $\vec{f}(t) = \begin{bmatrix} 2e^t \\ 2t \end{bmatrix}$ for $\vec{x}(0) = \begin{bmatrix} 3/16 \\ -5/16 \end{bmatrix}$.

The eigenvalues of A are -2 and 4 and corresponding eigenvectors are $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ respectively. This calculation is left as an exercise. We write down the matrix E of the eigenvectors and compute its inverse (using the inverse formula for 2×2 matrices)

$$E = \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix}, \quad E^{-1} = \frac{1}{2} \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix}.$$

We are looking for a solution of the form $\vec{x} = \begin{bmatrix} 1 \\ -1 \end{bmatrix} \xi_1 + \begin{bmatrix} 1 \\ 1 \end{bmatrix} \xi_2$. We first need to write $\vec{f}$ in terms of the eigenvectors. That is we wish to write $\vec{f} = \begin{bmatrix} 2e^t \\ 2t \end{bmatrix} = \begin{bmatrix} 1 \\ -1 \end{bmatrix} g_1 + \begin{bmatrix} 1 \\ 1 \end{bmatrix} g_2$. Thus

$$\begin{bmatrix} g_1 \\ g_2 \end{bmatrix} = E^{-1} \begin{bmatrix} 2e^t \\ 2t \end{bmatrix} = \frac{1}{2} \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} 2e^t \\ 2t \end{bmatrix} = \begin{bmatrix} e^t - t \\ e^t + t \end{bmatrix}.$$

So $g_1 = e^t - t$ and $g_2 = e^t + t$.

We further need to write $\vec{x}(0)$ in terms of the eigenvectors. That is, we wish to write $\vec{x}(0) = \begin{bmatrix} 3/16 \\ -5/16 \end{bmatrix} = \begin{bmatrix} 1 \\ -1 \end{bmatrix} a_1 + \begin{bmatrix} 1 \\ 1 \end{bmatrix} a_2$. Hence

$$\begin{bmatrix} a_1 \\ a_2 \end{bmatrix} = E^{-1} \begin{bmatrix} 3/16 \\ -5/16 \end{bmatrix} = \begin{bmatrix} 1/4 \\ -1/16 \end{bmatrix}.$$

So $a_1 = 1/4$ and $a_2 = -1/16$. We plug our $\vec{x}$ into the equation and get

$$\begin{aligned} \overbrace{\begin{bmatrix} 1 \\ -1 \end{bmatrix} \xi'_1 + \begin{bmatrix} 1 \\ 1 \end{bmatrix} \xi'_2}^{\vec{x}'} &= \overbrace{A \begin{bmatrix} 1 \\ -1 \end{bmatrix} \xi_1 + A \begin{bmatrix} 1 \\ 1 \end{bmatrix} \xi_2}^{A\vec{x}} + \overbrace{\begin{bmatrix} 1 \\ -1 \end{bmatrix} g_1 + \begin{bmatrix} 1 \\ 1 \end{bmatrix} g_2}^{\vec{f}} \\ &= \begin{bmatrix} 1 \\ -1 \end{bmatrix} (-2\xi_1) + \begin{bmatrix} 1 \\ 1 \end{bmatrix} 4\xi_2 + \begin{bmatrix} 1 \\ -1 \end{bmatrix} (e^t - t) + \begin{bmatrix} 1 \\ 1 \end{bmatrix} (e^t + t). \end{aligned}$$

We get the two equations

$$\begin{aligned} \xi'_1 &= -2\xi_1 + e^t - t, & \text{where } \xi_1(0) &= a_1 = \frac{1}{4}, \\ \xi'_2 &= 4\xi_2 + e^t + t, & \text{where } \xi_2(0) &= a_2 = \frac{-1}{16}. \end{aligned}$$

We solve with the integrating factor method. Computation of the integral is left as an exercise to the student. You will need integration by parts.

$$\xi_1 = e^{-2t} \int e^{2t} (e^t - t) dt + C_1 e^{-2t} = \frac{e^t}{3} - \frac{t}{2} + \frac{1}{4} + C_1 e^{-2t},$$

where C_1 is the constant of integration. As $\xi_1(0) = 1/4$, then $1/4 = 1/3 + 1/4 + C_1$ and $C_1 = -1/3$. Similarly,

$$\xi_2 = e^{4t} \int e^{-4t} (e^t + t) dt + C_2 e^{4t} = -\frac{e^t}{3} - \frac{t}{4} - \frac{1}{16} + C_2 e^{4t}.$$

As $\xi_2(0) = -1/16$, we have $-1/16 = -1/3 - 1/16 + C_2$ and hence $C_2 = 1/3$. The solution is

$$\vec{x}(t) = \begin{bmatrix} 1 \\ -1 \end{bmatrix} \underbrace{\left(\frac{e^t - e^{-2t}}{3} + \frac{1 - 2t}{4} \right)}_{\xi_1} + \begin{bmatrix} 1 \\ 1 \end{bmatrix} \underbrace{\left(\frac{e^{4t} - e^t}{3} - \frac{4t + 1}{16} \right)}_{\xi_2} = \begin{bmatrix} \frac{e^{4t} - e^{-2t}}{3} + \frac{3 - 12t}{16} \\ \frac{e^{-2t} + e^{4t} - 2e^t}{3} + \frac{4t - 5}{16} \end{bmatrix}.$$

That is, $x_1 = \frac{e^{4t} - e^{-2t}}{3} + \frac{3 - 12t}{16}$ and $x_2 = \frac{e^{-2t} + e^{4t} - 2e^t}{3} + \frac{4t - 5}{16}$.

Exercise 3.9.1: Check that x_1 and x_2 solve the problem. Check both that they satisfy the differential equation and that they satisfy the initial conditions.

Undetermined coefficients

The method of undetermined coefficients also works for systems. The only difference is that we use unknown vectors rather than just numbers. Same caveats apply to undetermined coefficients for systems as for single equations. This method does not always work. Furthermore, if the right-hand side is complicated, we have to solve for lots of variables. Each element of an unknown vector is an unknown number. In a system of 3 equations with say 4 unknown vectors (this would not be uncommon), we already have 12 unknown numbers to solve for. The method can turn into a lot of tedious work if done by hand. As the method is essentially the same as for single equations, let us just do an example.

Example 3.9.3: Let $A = \begin{bmatrix} -1 & 0 \\ -2 & 1 \end{bmatrix}$. Find a particular solution of $\vec{x}' = A\vec{x} + \vec{f}$ where $\vec{f}(t) = \begin{bmatrix} e^t \\ t \end{bmatrix}$.

One can solve this system in an easier way (can you see how?), but for the purposes of the example, we use the eigenvalue method and undetermined coefficients. The eigenvalues of A are -1 and 1 . Corresponding eigenvectors are $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$ respectively. Hence, our complementary solution is

$$\vec{x}_c = \alpha_1 \begin{bmatrix} 1 \\ 1 \end{bmatrix} e^{-t} + \alpha_2 \begin{bmatrix} 0 \\ 1 \end{bmatrix} e^t,$$

for some arbitrary constants α_1 and α_2 .

We would next want to guess a particular solution of the form

$$\vec{x} = \vec{a}e^t + \vec{b}t + \vec{c}.$$

However, something of the form $\vec{a}e^t$ appears in the complementary solution. Because we do not yet know if the vector $\vec{a}$ is a multiple of $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$, we do not know if a conflict arises. It is possible that there is no conflict, but to be safe we also try $\vec{b}te^t$. Here we find the crux of the difference between a single equation and systems. We try *both* terms $\vec{a}e^t$ and $\vec{b}te^t$ in the solution, not just the term $\vec{b}te^t$. Therefore, we try

$$\vec{x} = \vec{a}e^t + \vec{b}te^t + \vec{c}t + \vec{d}.$$

We have 8 unknowns as $\vec{a} = \begin{bmatrix} a_1 \\ a_2 \end{bmatrix}$, $\vec{b} = \begin{bmatrix} b_1 \\ b_2 \end{bmatrix}$, $\vec{c} = \begin{bmatrix} c_1 \\ c_2 \end{bmatrix}$, and $\vec{d} = \begin{bmatrix} d_1 \\ d_2 \end{bmatrix}$. We plug $\vec{x}$ into the equation. First let us compute $\vec{x}'$.

$$\vec{x}' = (\vec{a} + \vec{b})e^t + \vec{b}te^t + \vec{c} = \begin{bmatrix} a_1 + b_1 \\ a_2 + b_2 \end{bmatrix}e^t + \begin{bmatrix} b_1 \\ b_2 \end{bmatrix}te^t + \begin{bmatrix} c_1 \\ c_2 \end{bmatrix}.$$

Now $\vec{x}'$ must equal $A\vec{x} + \vec{f}$, which is

$$\begin{aligned} A\vec{x} + \vec{f} &= A\vec{a}e^t + A\vec{b}te^t + A\vec{c}t + A\vec{d} + \vec{f} \\ &= \begin{bmatrix} -a_1 \\ -2a_1 + a_2 \end{bmatrix}e^t + \begin{bmatrix} -b_1 \\ -2b_1 + b_2 \end{bmatrix}te^t + \begin{bmatrix} -c_1 \\ -2c_1 + c_2 \end{bmatrix}t + \begin{bmatrix} -d_1 \\ -2d_1 + d_2 \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \end{bmatrix}e^t + \begin{bmatrix} 0 \\ 1 \end{bmatrix}t \\ &= \begin{bmatrix} -a_1 + 1 \\ -2a_1 + a_2 \end{bmatrix}e^t + \begin{bmatrix} -b_1 \\ -2b_1 + b_2 \end{bmatrix}te^t + \begin{bmatrix} -c_1 \\ -2c_1 + c_2 + 1 \end{bmatrix}t + \begin{bmatrix} -d_1 \\ -2d_1 + d_2 \end{bmatrix}. \end{aligned}$$

We identify the coefficients of e^t , te^t , t and any constant vectors in $\vec{x}'$ and in $A\vec{x} + \vec{f}$ to find the equations:

$$\begin{aligned} a_1 + b_1 &= -a_1 + 1, & 0 &= -c_1, \\ a_2 + b_2 &= -2a_1 + a_2, & 0 &= -2c_1 + c_2 + 1, \\ b_1 &= -b_1, & c_1 &= -d_1, \\ b_2 &= -2b_1 + b_2, & c_2 &= -2d_1 + d_2. \end{aligned}$$

We could write the 8×9 augmented matrix and do row reduction, but it is easier to solve these in an ad hoc manner. Immediately, we see $b_1 = 0$, $c_1 = 0$, $d_1 = 0$. Plugging these back in, we get $c_2 = -1$ and $d_2 = -1$. The remaining equations that tell us something are

$$\begin{aligned} a_1 &= -a_1 + 1, \\ a_2 + b_2 &= -2a_1 + a_2. \end{aligned}$$

So $a_1 = 1/2$ and $b_2 = -1$. Finally, a_2 can be arbitrary and still satisfy the equations. We are looking for just a single solution, so presumably the simplest one is when $a_2 = 0$. Therefore,

$$\vec{x} = \vec{a}e^t + \vec{b}te^t + \vec{c}t + \vec{d} = \begin{bmatrix} 1/2 \\ 0 \end{bmatrix}e^t + \begin{bmatrix} 0 \\ -1 \end{bmatrix}te^t + \begin{bmatrix} 0 \\ -1 \end{bmatrix}t + \begin{bmatrix} 0 \\ -1 \end{bmatrix} = \begin{bmatrix} \frac{1}{2}e^t \\ -te^t - t - 1 \end{bmatrix}.$$

That is, $x_1 = \frac{1}{2}e^t$, $x_2 = -te^t - t - 1$. We add this particular solution to the complementary solution to get the general solution of the problem:

$$x_1 = \frac{1}{2}e^t + \alpha_1 e^{-t} \quad \text{and} \quad x_2 = -te^t - t - 1 + \alpha_1 e^{-t} + \alpha_2 e^t.$$

Notice that both $\vec{a}e^t$ and $\vec{b}te^t$ really were needed.

Exercise 3.9.2: Check that x_1 and x_2 solve the problem. Then set $a_2 = 1$ and check this solution as well. What is the difference between the two solutions (one with $a_2 = 0$ and one with $a_2 = 1$)?

As you can see, other than the handling of conflicts, undetermined coefficients works exactly the same as it did for single equations. However, the computations can get out of hand pretty quickly for systems. The equation we considered was pretty simple.

3.9.2 First-order variable-coefficient

Variation of parameters

Just as for a single equation, there is the method of variation of parameters. For constant-coefficient systems, it is essentially the same thing as the integrating factor method we discussed earlier. However, this method works for any linear system, even if it is not constant-coefficient, provided we somehow solve the associated homogeneous problem.

Consider the equation

$$\vec{x}' = A(t)\vec{x} + \vec{f}(t). \quad (3.10)$$

Suppose we somehow solved the associated homogeneous equation $\vec{x}' = A(t)\vec{x}$ and found a fundamental matrix solution $X(t)$. The general solution to the associated homogeneous equation is $X(t)\vec{c}$ for a constant vector $\vec{c}$. As in variation of parameters for a single equation, we try the solution to the nonhomogeneous equation of the form

$$\vec{x}_p = X(t)\vec{u}(t),$$

where $\vec{u}(t)$ is a vector-valued function instead of a constant. We substitute $\vec{x}_p$ into (3.10):

$$\underbrace{X'(t)\vec{u}(t) + X(t)\vec{u}'(t)}_{\vec{x}_p'(t)} = \underbrace{A(t)X(t)\vec{u}(t)}_{A(t)\vec{x}_p(t)} + \vec{f}(t).$$

As $X(t)$ is a fundamental matrix solution to the homogeneous problem, that is, $X'(t) = A(t)X(t)$, we find

$$\cancel{X'(t)\vec{u}(t)} + X(t)\vec{u}'(t) = \cancel{X'(t)\vec{u}(t)} + \vec{f}(t).$$

Hence, $X(t)\vec{u}'(t) = \vec{f}(t)$. We compute $[X(t)]^{-1}$, and then $\vec{u}'(t) = [X(t)]^{-1}\vec{f}(t)$. We integrate to obtain $\vec{u}$, and we have the particular solution $\vec{x}_p = X(t)\vec{u}(t)$. Hence, we have the formula

$$\vec{x}_p = X(t) \int [X(t)]^{-1} \vec{f}(t) dt.$$

If A is a constant matrix and $X(t) = e^{tA}$, then $[X(t)]^{-1} = e^{-tA}$. We obtain a solution $\vec{x}_p = e^{tA} \int e^{-tA} \vec{f}(t) dt$, which is precisely what we got using the integrating factor method.

Example 3.9.4: Find a particular solution to

$$\vec{x}' = \frac{1}{t^2 + 1} \begin{bmatrix} t & -1 \\ 1 & t \end{bmatrix} \vec{x} + \begin{bmatrix} t \\ 1 \end{bmatrix} (t^2 + 1). \quad (3.11)$$

Here, $A = \frac{1}{t^2+1} \begin{bmatrix} t & -1 \\ 1 & t \end{bmatrix}$ is most definitely not constant. Perhaps by a lucky guess, we find that $X = \begin{bmatrix} 1 & -t \\ t & 1 \end{bmatrix}$ solves $X'(t) = A(t)X(t)$. Once we know the complementary solution, we can find a solution to (3.11). First, we compute

$$[X(t)]^{-1} = \frac{1}{t^2 + 1} \begin{bmatrix} 1 & t \\ -t & 1 \end{bmatrix}.$$

Next, we know a particular solution to (3.11) is

$$\begin{aligned} \vec{x}_p &= X(t) \int [X(t)]^{-1} \vec{f}(t) dt \\ &= \begin{bmatrix} 1 & -t \\ t & 1 \end{bmatrix} \int \frac{1}{t^2 + 1} \begin{bmatrix} 1 & t \\ -t & 1 \end{bmatrix} \begin{bmatrix} t \\ 1 \end{bmatrix} (t^2 + 1) dt \\ &= \begin{bmatrix} 1 & -t \\ t & 1 \end{bmatrix} \int \begin{bmatrix} 2t \\ -t^2 + 1 \end{bmatrix} dt \\ &= \begin{bmatrix} 1 & -t \\ t & 1 \end{bmatrix} \begin{bmatrix} t^2 \\ -\frac{1}{3}t^3 + t \end{bmatrix} \\ &= \begin{bmatrix} \frac{1}{3}t^4 \\ \frac{2}{3}t^3 + t \end{bmatrix}. \end{aligned}$$

Adding the complementary solution, we find the general solution to (3.11):

$$\vec{x} = \begin{bmatrix} 1 & -t \\ t & 1 \end{bmatrix} \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} + \begin{bmatrix} \frac{1}{3}t^4 \\ \frac{2}{3}t^3 + t \end{bmatrix} = \begin{bmatrix} c_1 - c_2t + \frac{1}{3}t^4 \\ c_2 + (c_1 + 1)t + \frac{2}{3}t^3 \end{bmatrix}.$$

Exercise 3.9.3: Check that $x_1 = \frac{1}{3}t^4$ and $x_2 = \frac{2}{3}t^3 + t$ really solve (3.11).

In the variation of parameters, as in the integrating factor method, we can obtain the general solution by adding in constants of integration. Doing so would add $X(t)\vec{c}$ for a vector $\vec{c}$ of arbitrary constants. But that is precisely the complementary solution.

3.9.3 Second-order constant-coefficients

Undetermined coefficients

We have already seen a simple example of the method of undetermined coefficients for second-order systems in § 3.6. This method is essentially the same as undetermined

coefficients for first-order systems. There are some simplifications that we can make, as we did in § 3.6. Consider the equation

$$\vec{x}'' = A\vec{x} + \vec{F}(t),$$

where A is a constant matrix. If $\vec{F}(t)$ is of the form $\vec{F}_0 \cos(\omega t)$, then as two derivatives of cosine is again cosine, we do not need to introduce sines, and we try a solution of the form

$$\vec{x}_p = \vec{c} \cos(\omega t).$$

If $\vec{F}$ is a sum of cosines, recall the superposition principle. If $\vec{F}(t) = \vec{F}_0 \cos(\omega_0 t) + \vec{F}_1 \cos(\omega_1 t)$, then we would try $\vec{a} \cos(\omega_0 t)$ for the problem $\vec{x}'' = A\vec{x} + \vec{F}_0 \cos(\omega_0 t)$, and we would try $\vec{b} \cos(\omega_1 t)$ for the problem $\vec{x}'' = A\vec{x} + \vec{F}_1 \cos(\omega_1 t)$. Then we sum the solutions.

If there is duplication with the complementary solution, or the equation is of the form $\vec{x}'' = A\vec{x}' + B\vec{x} + \vec{F}(t)$, then we need to do the same thing as we do for first-order systems.

You can never go wrong by putting in more terms than needed into your guess. Those extra coefficients will turn out to be zero. But it is useful to save some time and effort.

Eigenvector decomposition

If we have the system

$$\vec{x}'' = A\vec{x} + \vec{f}(t),$$

we can do *eigenvector decomposition*, just like for first-order systems.

Let $\lambda_1, \lambda_2, \dots, \lambda_n$ be the eigenvalues and $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ be eigenvectors. Again form the matrix $E = [\vec{v}_1 \ \vec{v}_2 \ \cdots \ \vec{v}_n]$. Write

$$\vec{x}(t) = \vec{v}_1 \xi_1(t) + \vec{v}_2 \xi_2(t) + \cdots + \vec{v}_n \xi_n(t).$$

Decompose $\vec{f}$ in terms of the eigenvectors

$$\vec{f}(t) = \vec{v}_1 g_1(t) + \vec{v}_2 g_2(t) + \cdots + \vec{v}_n g_n(t),$$

where, again, $\vec{g} = E^{-1} \vec{f}$.

We plug in and as before we obtain

$$\begin{aligned} \overbrace{\vec{v}_1 \xi_1'' + \vec{v}_2 \xi_2'' + \cdots + \vec{v}_n \xi_n''}^{\vec{x}''} &= \overbrace{A(\vec{v}_1 \xi_1 + \vec{v}_2 \xi_2 + \cdots + \vec{v}_n \xi_n)}^{A\vec{x}} + \overbrace{\vec{v}_1 g_1 + \vec{v}_2 g_2 + \cdots + \vec{v}_n g_n}^{\vec{f}} \\ &= A\vec{v}_1 \xi_1 + A\vec{v}_2 \xi_2 + \cdots + A\vec{v}_n \xi_n + \vec{v}_1 g_1 + \vec{v}_2 g_2 + \cdots + \vec{v}_n g_n \\ &= \vec{v}_1 \lambda_1 \xi_1 + \vec{v}_2 \lambda_2 \xi_2 + \cdots + \vec{v}_n \lambda_n \xi_n + \vec{v}_1 g_1 + \vec{v}_2 g_2 + \cdots + \vec{v}_n g_n \\ &= \vec{v}_1 (\lambda_1 \xi_1 + g_1) + \vec{v}_2 (\lambda_2 \xi_2 + g_2) + \cdots + \vec{v}_n (\lambda_n \xi_n + g_n). \end{aligned}$$

We identify the coefficients of the eigenvectors to get the equations

$$\begin{aligned}\xi_1'' &= \lambda_1 \xi_1 + g_1, \\ \xi_2'' &= \lambda_2 \xi_2 + g_2, \\ &\vdots \\ \xi_n'' &= \lambda_n \xi_n + g_n.\end{aligned}$$

Each one of these equations is independent of the others. We solve each equation using the methods of [chapter 2](#). We write $\vec{x}(t) = \vec{v}_1 \xi_1(t) + \vec{v}_2 \xi_2(t) + \cdots + \vec{v}_n \xi_n(t)$ to find a particular solution. If we find the general solutions for ξ_1 through ξ_n , then $\vec{x}(t) = \vec{v}_1 \xi_1(t) + \vec{v}_2 \xi_2(t) + \cdots + \vec{v}_n \xi_n(t)$ is the general solution as well.

Example 3.9.5: Let us do the example from [§ 3.6](#) using this method. The equation is

$$\vec{x}'' = \begin{bmatrix} -3 & 1 \\ 2 & -2 \end{bmatrix} \vec{x} + \begin{bmatrix} 0 \\ 2 \end{bmatrix} \cos(3t).$$

The eigenvalues are -1 and -4 , with eigenvectors $\begin{bmatrix} 1 \\ 2 \end{bmatrix}$ and $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$. Therefore $E = \begin{bmatrix} 1 & 1 \\ 2 & -1 \end{bmatrix}$ and $E^{-1} = \frac{1}{3} \begin{bmatrix} 1 & 1 \\ 2 & -1 \end{bmatrix}$. Therefore,

$$\begin{bmatrix} g_1 \\ g_2 \end{bmatrix} = E^{-1} \vec{f}(t) = \frac{1}{3} \begin{bmatrix} 1 & 1 \\ 2 & -1 \end{bmatrix} \begin{bmatrix} 0 \\ 2 \cos(3t) \end{bmatrix} = \begin{bmatrix} \frac{2}{3} \cos(3t) \\ -\frac{2}{3} \cos(3t) \end{bmatrix}.$$

So after the whole song and dance of plugging in, the equations we get are

$$\xi_1'' = -\xi_1 + \frac{2}{3} \cos(3t), \quad \xi_2'' = -4 \xi_2 - \frac{2}{3} \cos(3t).$$

For each equation we use the method of undetermined coefficients. We try $C_1 \cos(3t)$ for the first equation and $C_2 \cos(3t)$ for the second equation. We plug in to get

$$\begin{aligned}-9C_1 \cos(3t) &= -C_1 \cos(3t) + \frac{2}{3} \cos(3t), \\ -9C_2 \cos(3t) &= -4C_2 \cos(3t) - \frac{2}{3} \cos(3t).\end{aligned}$$

We solve each of these equations separately. We get $-9C_1 = -C_1 + 2/3$ and $-9C_2 = -4C_2 - 2/3$. And hence $C_1 = -1/12$ and $C_2 = 2/15$. So our particular solution is

$$\vec{x} = \begin{bmatrix} 1 \\ 2 \end{bmatrix} \left(\frac{-1}{12} \cos(3t) \right) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} \left(\frac{2}{15} \cos(3t) \right) = \begin{bmatrix} 1/20 \\ -3/10 \end{bmatrix} \cos(3t).$$

This solution matches what we got previously in [§ 3.6](#).

3.9.4 Exercises

Exercise 3.9.4: Find a particular solution to $x' = x + 2y + 2t$, $y' = 3x + 2y - 4$

- a) using integrating factor method,
- b) using eigenvector decomposition,
- c) using undetermined coefficients.

Exercise 3.9.5: Find the general solution to $x' = 4x + y - 1$, $y' = x + 4y - e^t$

- a) using integrating factor method,
- b) using eigenvector decomposition,
- c) using undetermined coefficients.

Exercise 3.9.6: Find the general solution to $x_1'' = -6x_1 + 3x_2 + \cos(t)$, $x_2'' = 2x_1 - 7x_2 + 3\cos(t)$

- a) using eigenvector decomposition,
- b) using undetermined coefficients.

Exercise 3.9.7: Find the general solution to $x_1'' = -6x_1 + 3x_2 + \cos(2t)$, $x_2'' = 2x_1 - 7x_2 + 3\cos(2t)$

- a) using eigenvector decomposition,
- b) using undetermined coefficients.

Exercise 3.9.8: Take the equation $\vec{x}' = \begin{bmatrix} \frac{1}{t} & -1 \\ 1 & \frac{1}{t} \end{bmatrix} \vec{x} + \begin{bmatrix} t^2 \\ -t \end{bmatrix}$.

- a) Check that $\vec{x}_c = c_1 \begin{bmatrix} t \sin t \\ -t \cos t \end{bmatrix} + c_2 \begin{bmatrix} t \cos t \\ t \sin t \end{bmatrix}$ is the complementary solution.
- b) Use variation of parameters to find a particular solution.

Exercise 3.9.101: Find a particular solution to $x' = 5x + 4y + t$, $y' = x + 8y - t$

- a) using integrating factor method,
- b) using eigenvector decomposition,
- c) using undetermined coefficients.

Exercise 3.9.102: Find a particular solution to $x' = y + e^t$, $y' = x + e^t$

- a) using integrating factor method,
- b) using eigenvector decomposition,
- c) using undetermined coefficients.

Exercise 3.9.103: Solve $x_1' = x_2 + t$, $x_2' = x_1 + t$ with initial conditions $x_1(0) = 1$, $x_2(0) = 2$ using eigenvector decomposition.

Exercise 3.9.104: Solve $x_1'' = -3x_1 + x_2 + t$, $x_2'' = 9x_1 + 5x_2 + \cos(t)$ with initial conditions $x_1(0) = 0$, $x_2(0) = 0$, $x_1'(0) = 0$, $x_2'(0) = 0$ using eigenvector decomposition.

Chapter 4

Fourier series and PDEs

4.1 Boundary value problems

Note: 2 lectures, similar to §3.8 in [EP], §10.1 and §11.1 in [BD]

4.1.1 Boundary value problems

Before we tackle the Fourier series, we study the so-called *boundary value problems* (or *endpoint problems*). Consider

$$x'' + \lambda x = 0, \quad x(a) = 0, \quad x(b) = 0,$$

for some constant λ , where $x(t)$ is defined for t in the interval $[a, b]$. Previously we specified the value of the solution and its derivative at a single point. Now we specify the value of the solution at two different points. As $x = 0$ is a solution, existence of solutions is not a problem. Uniqueness of solutions is another issue. The general solution to $x'' + \lambda x = 0$ has two arbitrary constants[†]. It is, therefore, natural (but wrong) to believe that requiring two conditions guarantees a unique solution.

Example 4.1.1: Take $\lambda = 1$, $a = 0$, $b = \pi$. That is,

$$x'' + x = 0, \quad x(0) = 0, \quad x(\pi) = 0.$$

Then $x = \sin t$ is another solution (besides $x = 0$) satisfying both boundary conditions. There are more. Write down the general solution of the differential equation, which is $x = A \cos t + B \sin t$. The condition $x(0) = 0$ forces $A = 0$. Letting $x(\pi) = 0$ does not give us any more information, as $x = B \sin t$ already satisfies both boundary conditions. Hence, there are infinitely many solutions of the form $x = B \sin t$, where B is an arbitrary constant.

Example 4.1.2: On the other hand, consider $\lambda = 2$. That is,

$$x'' + 2x = 0, \quad x(0) = 0, \quad x(\pi) = 0.$$

[†]See [subsection 0.2.4](#) on page 13 or [Example 2.2.1](#) on page 85 and [Example 2.2.3](#) on page 88.

Then the general solution is $x = A \cos(\sqrt{2}t) + B \sin(\sqrt{2}t)$. Letting $x(0) = 0$ still forces $A = 0$. We apply the second condition to find $0 = x(\pi) = B \sin(\sqrt{2}\pi)$. As $\sin(\sqrt{2}\pi) \neq 0$, we obtain $B = 0$. Therefore, $x = 0$ is the unique solution to this problem.

What is going on? We will be interested in finding which constants λ allow a nonzero solution, and we will be interested in finding those solutions. This problem is an analogue of finding eigenvalues and eigenvectors of matrices.

4.1.2 Eigenvalue problems

For basic Fourier series theory, we will need the following three eigenvalue problems:

$$x'' + \lambda x = 0, \quad x(a) = 0, \quad x(b) = 0, \quad (4.1)$$

$$x'' + \lambda x = 0, \quad x'(a) = 0, \quad x'(b) = 0, \quad (4.2)$$

and

$$x'' + \lambda x = 0, \quad x(a) = x(b), \quad x'(a) = x'(b). \quad (4.3)$$

A number λ is called an *eigenvalue* of (4.1) (resp. (4.2) or (4.3)) if and only if there exists a nonzero (not identically zero) solution to (4.1) (resp. (4.2) or (4.3)) given that specific λ . A nonzero solution is called a corresponding *eigenfunction*. We will consider more general equations and boundary conditions, but we will postpone this until [chapter 5](#).

Note the similarity to eigenvalues and eigenvectors of matrices. The similarity is not just coincidental. If we think of the equations as differential operators, then we are doing the same exact thing. Think of a function $x(t)$ as a vector with infinitely many components (one for each t). Let $L = -\frac{d^2}{dt^2}$ be the linear operator. Then the eigenvalue/eigenfunction pair should be λ and nonzero x such that $Lx = \lambda x$. In other words, we are looking for nonzero functions x satisfying certain endpoint conditions that solve $(L - \lambda)x = 0$. A lot of the formalism from linear algebra still applies here, though we will not pursue this line of reasoning too far.

Example 4.1.3: Find the eigenvalues and eigenfunctions of

$$x'' + \lambda x = 0, \quad x(0) = 0, \quad x(\pi) = 0.$$

We have to handle the cases $\lambda > 0$, $\lambda = 0$, and $\lambda < 0$ separately. First, suppose that $\lambda > 0$. Then the general solution to $x'' + \lambda x = 0$ is

$$x = A \cos(\sqrt{\lambda}t) + B \sin(\sqrt{\lambda}t).$$

The condition $x(0) = 0$ implies immediately $A = 0$. Next

$$0 = x(\pi) = B \sin(\sqrt{\lambda}\pi).$$

If B is zero, then x is not a nonzero solution. So to get a nonzero solution, we must have that $\sin(\sqrt{\lambda}\pi) = 0$. Hence, $\sqrt{\lambda}\pi$ must be an integer multiple of π . In other words,

$\sqrt{\lambda} = k$ for a positive integer k . Hence, the positive eigenvalues are k^2 for all integers $k \geq 1$. Corresponding eigenfunctions can be taken as $x = \sin(kt)$. Just like for eigenvectors, constant multiples of an eigenfunction are also eigenfunctions, so we only need to pick one.

Now, suppose that $\lambda = 0$. In this case, the equation is $x'' = 0$, and its general solution is $x = At + B$. The condition $x(0) = 0$ implies that $B = 0$, and $x(\pi) = 0$ implies that $A = 0$. This means that $\lambda = 0$ is *not* an eigenvalue.

Finally, suppose that $\lambda < 0$. In this case, we have the general solution*

$$x = A \cosh(\sqrt{-\lambda} t) + B \sinh(\sqrt{-\lambda} t).$$

Letting $x(0) = 0$ implies that $A = 0$ (recall $\cosh 0 = 1$ and $\sinh 0 = 0$). So our solution must be $x = B \sinh(\sqrt{-\lambda} t)$ and satisfy $x(\pi) = 0$. This is only possible if B is zero. Why? Because $\sinh \xi$ is only zero when $\xi = 0$. You should plot $\sinh$ to see this fact. We can also see this from the definition of $\sinh$. We get $0 = \sinh \xi = \frac{e^\xi - e^{-\xi}}{2}$. Hence $e^\xi = e^{-\xi}$, which implies $\xi = -\xi$ and that is only true if $\xi = 0$. So there are no negative eigenvalues.

In summary, the eigenvalues and corresponding eigenfunctions are

$$\lambda_k = k^2 \quad \text{with an eigenfunction} \quad x_k = \sin(kt) \quad \text{for all integers } k \geq 1.$$

Example 4.1.4: Compute the eigenvalues and eigenfunctions of

$$x'' + \lambda x = 0, \quad x'(0) = 0, \quad x'(\pi) = 0.$$

Again we have to handle the cases $\lambda > 0$, $\lambda = 0$, $\lambda < 0$ separately. First suppose that $\lambda > 0$. The general solution to $x'' + \lambda x = 0$ is $x = A \cos(\sqrt{\lambda} t) + B \sin(\sqrt{\lambda} t)$. So

$$x' = -A\sqrt{\lambda} \sin(\sqrt{\lambda} t) + B\sqrt{\lambda} \cos(\sqrt{\lambda} t).$$

The condition $x'(0) = 0$ implies immediately $B = 0$. Next

$$0 = x'(\pi) = -A\sqrt{\lambda} \sin(\sqrt{\lambda} \pi).$$

Again A cannot be zero if λ is to be an eigenvalue, and $\sin(\sqrt{\lambda} \pi)$ is only zero if $\sqrt{\lambda} = k$ for a positive integer k . Hence, the positive eigenvalues are again k^2 for all integers $k \geq 1$. And the corresponding eigenfunctions can be taken as $x = \cos(kt)$.

Now suppose that $\lambda = 0$. In this case, the equation is $x'' = 0$ and the general solution is $x = At + B$ so $x' = A$. The condition $x'(0) = 0$ implies that $A = 0$. The condition $x'(\pi) = 0$ also implies $A = 0$. Hence B could be anything (let us take it to be 1). So $\lambda = 0$ is an eigenvalue and $x = 1$ is a corresponding eigenfunction.

Finally, let $\lambda < 0$. In this case, the general solution is $x = A \cosh(\sqrt{-\lambda} t) + B \sinh(\sqrt{-\lambda} t)$ and

$$x' = A\sqrt{-\lambda} \sinh(\sqrt{-\lambda} t) + B\sqrt{-\lambda} \cosh(\sqrt{-\lambda} t).$$

*Recall that $\cosh s = \frac{1}{2}(e^s + e^{-s})$ and $\sinh s = \frac{1}{2}(e^s - e^{-s})$. As an exercise, try the computation with the general solution written as $x = Ae^{\sqrt{-\lambda} t} + Be^{-\sqrt{-\lambda} t}$ (for different A and B of course).

We have already seen (with roles of A and B switched) that for this expression to be zero at $t = 0$ and $t = \pi$, we must have $A = B = 0$. Hence, there are no negative eigenvalues.

In summary, the eigenvalues and corresponding eigenfunctions are

$$\lambda_k = k^2 \quad \text{with an eigenfunction} \quad x_k = \cos(kt) \quad \text{for all integers } k \geq 1,$$

and there is another eigenvalue

$$\lambda_0 = 0 \quad \text{with an eigenfunction} \quad x_0 = 1.$$

The following problem is the one that leads to the general Fourier series.

Example 4.1.5: Compute the eigenvalues and eigenfunctions of

$$x'' + \lambda x = 0, \quad x(-\pi) = x(\pi), \quad x'(-\pi) = x'(\pi).$$

We have not specified the values or the derivatives at the endpoints, but rather that they are the same at the beginning and at the end of the interval.

We skip $\lambda < 0$. The computations are the same as before, and again we find that there are no negative eigenvalues.

For $\lambda = 0$, the general solution is $x = At + B$. The condition $x(-\pi) = x(\pi)$ implies that $A = 0$ ($A\pi + B = -A\pi + B$ implies $A = 0$). The second condition $x'(-\pi) = x'(\pi)$ says nothing about B and hence $\lambda = 0$ is an eigenvalue with a corresponding eigenfunction $x = 1$.

For $\lambda > 0$ we get that $x = A \cos(\sqrt{\lambda} t) + B \sin(\sqrt{\lambda} t)$. Now

$$\underbrace{A \cos(-\sqrt{\lambda} \pi) + B \sin(-\sqrt{\lambda} \pi)}_{x(-\pi)} = \underbrace{A \cos(\sqrt{\lambda} \pi) + B \sin(\sqrt{\lambda} \pi)}_{x(\pi)}.$$

We remember that $\cos(-\theta) = \cos(\theta)$ and $\sin(-\theta) = -\sin(\theta)$. Therefore,

$$A \cos(\sqrt{\lambda} \pi) - B \sin(\sqrt{\lambda} \pi) = A \cos(\sqrt{\lambda} \pi) + B \sin(\sqrt{\lambda} \pi).$$

Hence either $B = 0$ or $\sin(\sqrt{\lambda} \pi) = 0$. Similarly (exercise), if we differentiate x and plug in the second condition, we find that $A = 0$ or $\sin(\sqrt{\lambda} \pi) = 0$. Therefore, unless we want A and B to both be zero (which we do not), we must have $\sin(\sqrt{\lambda} \pi) = 0$. Hence, $\sqrt{\lambda}$ is an integer and the eigenvalues are yet again $\lambda = k^2$ for an integer $k \geq 1$. In this case, however, $x = A \cos(kt) + B \sin(kt)$ is an eigenfunction for any A and any B . So we have two linearly independent eigenfunctions $\sin(kt)$ and $\cos(kt)$. Remember that for a matrix, we can also have two eigenvectors corresponding to a single eigenvalue if the eigenvalue is repeated.

In summary, the eigenvalues and corresponding eigenfunctions are

$$\begin{aligned} \lambda_k &= k^2 & \text{with eigenfunctions} & \quad \cos(kt) \quad \text{and} \quad \sin(kt) & \quad \text{for all integers } k \geq 1, \\ \lambda_0 &= 0 & \text{with an eigenfunction} & \quad x_0 = 1. \end{aligned}$$

4.1.3 Orthogonality of eigenfunctions

Something that will be very useful in the next section is the *orthogonality* property of the eigenfunctions. This is an analogue of the following fact about eigenvectors of a matrix. A matrix is called *symmetric* if $A = A^T$ (it is equal to its transpose). *Eigenvectors for two distinct eigenvalues of a symmetric matrix are orthogonal.* The differential operators we are dealing with act much like a symmetric matrix. We, therefore, get the following theorem.

Theorem 4.1.1. *Suppose that $x_1(t)$ and $x_2(t)$ are two eigenfunctions of the problem (4.1), (4.2), or (4.3) for two different eigenvalues λ_1 and λ_2 . Then they are orthogonal in the sense that*

$$\int_a^b x_1(t)x_2(t) dt = 0.$$

The terminology comes from the fact that the integral is a type of inner product. We will expand on this in the next section. The theorem has a very short, elegant, and illuminating proof so we give it here. First, we have the following two equations.

$$x_1'' + \lambda_1 x_1 = 0 \quad \text{and} \quad x_2'' + \lambda_2 x_2 = 0.$$

Multiply the first by x_2 and the second by x_1 and subtract to get

$$(\lambda_1 - \lambda_2)x_1x_2 = x_2''x_1 - x_2x_1''.$$

Integrate both sides of the equation:

$$\begin{aligned} (\lambda_1 - \lambda_2) \int_a^b x_1(t)x_2(t) dt &= \int_a^b (x_2''(t)x_1(t) - x_2(t)x_1''(t)) dt \\ &= \int_a^b \frac{d}{dt} (x_2'(t)x_1(t) - x_2(t)x_1'(t)) dt \\ &= \left[x_2'(t)x_1(t) - x_2(t)x_1'(t) \right]_{t=a}^b = 0. \end{aligned}$$

The last equality holds because of the boundary conditions. For example, if we consider (4.1) we have $x_1(a) = x_1(b) = x_2(a) = x_2(b) = 0$ and so $x_2'x_1 - x_2x_1'$ is zero at both a and b . As $\lambda_1 \neq \lambda_2$, the theorem follows.

Exercise 4.1.1 (easy): *Finish the proof of the theorem (check the last equality in the proof) for the cases (4.2) and (4.3).*

The function $\sin(nt)$ is an eigenfunction for the problem $x'' + \lambda x = 0$, $x(0) = 0$, $x(\pi) = 0$. Hence for positive integers n and m , we have the integrals

$$\int_0^\pi \sin(mt) \sin(nt) dt = 0, \quad \text{when } m \neq n.$$

Similarly, still assuming that m and n are positive integers,

$$\int_0^\pi \cos(mt) \cos(nt) dt = 0, \quad \text{when } m \neq n, \quad \text{and} \quad \int_0^\pi \cos(nt) dt = 0.$$

Finally, we also get

$$\int_{-\pi}^\pi \sin(mt) \sin(nt) dt = 0, \quad \text{when } m \neq n, \quad \text{and} \quad \int_{-\pi}^\pi \sin(nt) dt = 0,$$

$$\int_{-\pi}^\pi \cos(mt) \cos(nt) dt = 0, \quad \text{when } m \neq n, \quad \text{and} \quad \int_{-\pi}^\pi \cos(nt) dt = 0,$$

and

$$\int_{-\pi}^\pi \cos(mt) \sin(nt) dt = 0 \quad (\text{even if } m = n).$$

4.1.4 Fredholm alternative

We now touch on a very useful theorem in the theory of differential equations. The theorem holds in a more general setting than we are going to state it, but for our purposes the following statement is sufficient. We will give a slightly more general version in [chapter 5](#).

Theorem 4.1.2 (Fredholm alternative*). *Exactly one of the following statements holds. Either*

$$x'' + \lambda x = 0, \quad x(a) = 0, \quad x(b) = 0 \tag{4.4}$$

has a nonzero solution, or

$$x'' + \lambda x = f(t), \quad x(a) = 0, \quad x(b) = 0 \tag{4.5}$$

has a unique solution for every function f continuous on $[a, b]$.

The theorem is also true for the other types of boundary conditions we considered. The theorem means that if λ is not an eigenvalue, the nonhomogeneous equation (4.5) has a unique solution for every right-hand side. On the other hand, if λ is an eigenvalue, then (4.5) need not have a solution for every f , and furthermore, even if it happens to have a solution, the solution is not unique.

We also want to reinforce the idea here that linear differential operators have much in common with matrices. So it is no surprise that there is a finite-dimensional version of Fredholm alternative for matrices as well. Let A be an $n \times n$ matrix. The Fredholm alternative then states that either $(A - \lambda I)\vec{x} = \vec{0}$ has a nontrivial solution, or $(A - \lambda I)\vec{x} = \vec{b}$ has a unique solution for every $\vec{b}$.

A lot of intuition from linear algebra can be applied to linear differential operators, but one must be careful of course. For example, one difference we have already seen is that in general a differential operator will have infinitely many eigenvalues, while a matrix has only finitely many.

*Named after the Swedish mathematician [Erik Ivar Fredholm](#) (1866–1927).

4.1.5 Application

Let us consider a physical application of an endpoint problem. Suppose we have a tightly stretched quickly spinning elastic string or rope of uniform linear density ρ , for example in kg/m. Let us put this problem into the xy -plane and both x and y are in meters. The x -axis represents the position on the string. The string rotates at angular velocity ω , in radians/s. Imagine that the whole xy -plane rotates at angular velocity ω . This way, the string stays in this xy -plane and y measures its deflection from the equilibrium position, $y = 0$, on the x -axis. Hence the graph of y gives the shape of the string. We consider an ideal string with no volume, just a mathematical curve. We suppose the tension on the string is a constant T in Newtons. Assuming that the deflection is small, we can use Newton's second law (let us skip the derivation) to get the equation

$$Ty'' + \rho\omega^2 y = 0.$$

To check the units, notice that the units of y'' are m/m^2 , as the derivative is in terms of x .

Let L be the length of the string (in meters) and the string is fixed at both endpoints. Hence, $y(0) = 0$ and $y(L) = 0$. See [Figure 4.1](#).

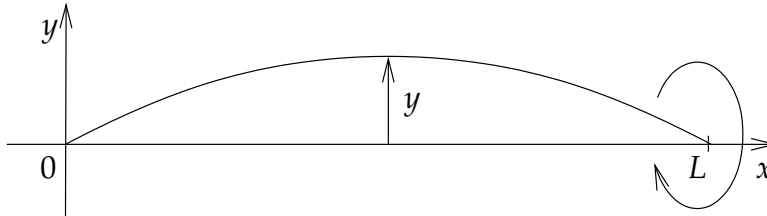


Figure 4.1: Whirling string.

We rewrite the equation as $y'' + \frac{\rho\omega^2}{T}y = 0$. The setup is similar to [Example 4.1.3](#) on page 190, except for the interval length being L instead of π . We are looking for eigenvalues of $y'' + \lambda y = 0, y(0) = 0, y(L) = 0$ where $\lambda = \frac{\rho\omega^2}{T}$. As before, there are no nonpositive eigenvalues. With $\lambda > 0$, the general solution to the equation is $y = A \cos(\sqrt{\lambda}x) + B \sin(\sqrt{\lambda}x)$. The condition $y(0) = 0$ implies that $A = 0$ as before. The condition $y(L) = 0$ implies that $\sin(\sqrt{\lambda}L) = 0$ and hence $\sqrt{\lambda}L = k\pi$ for some integer $k > 0$, so

$$\frac{\rho\omega^2}{T} = \lambda = \frac{k^2\pi^2}{L^2}.$$

What does this say about the shape of the string? It says that for all parameters ρ, ω, T not satisfying the equation above, the string is in the equilibrium position, $y = 0$. When $\frac{\rho\omega^2}{T} = \frac{k^2\pi^2}{L^2}$, then the string will “pop out” some distance B . We cannot compute B with the information we have.

Let us assume that ρ and T are fixed and we are changing ω . For most values of ω , the string is in the equilibrium state. When the angular velocity ω hits a value $\omega = \frac{k\pi\sqrt{T}}{L\sqrt{\rho}}$, then the string pops out and has the shape of a sine wave crossing the x -axis $k - 1$ times between the endpoints. For example, at $k = 1$, the string does not cross the x -axis and the shape looks like in Figure 4.1 on the previous page. On the other hand, when $k = 3$ the string crosses the x -axis 2 times, see Figure 4.2. When ω changes again, the string returns to the equilibrium position. The higher the angular velocity, the more times it crosses the x -axis when it is popped out.

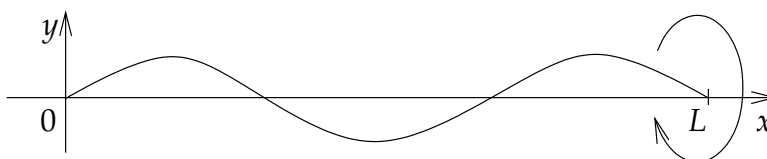


Figure 4.2: Whirling string at the third eigenvalue ($k = 3$).

For another example, if you have a spinning jump rope (then $k = 1$ as it is completely “popped out”) and you pull on the ends to increase the tension, then the velocity also increases for the rope to stay “popped out”.

4.1.6 Exercises

Hint for the following exercises: Note that if $\lambda > 0$, then $\cos(\sqrt{\lambda}(t - a))$ and $\sin(\sqrt{\lambda}(t - a))$ are also solutions of the homogeneous equation.

Exercise 4.1.2: Compute all eigenvalues and eigenfunctions of $x'' + \lambda x = 0$, $x(a) = 0$, $x(b) = 0$ (assume $a < b$).

Exercise 4.1.3: Compute all eigenvalues and eigenfunctions of $x'' + \lambda x = 0$, $x'(a) = 0$, $x'(b) = 0$ (assume $a < b$).

Exercise 4.1.4: Compute all eigenvalues and eigenfunctions of $x'' + \lambda x = 0$, $x'(a) = 0$, $x(b) = 0$ (assume $a < b$).

Exercise 4.1.5: Compute all eigenvalues and eigenfunctions of $x'' + \lambda x = 0$, $x(a) = x(b)$, $x'(a) = x'(b)$ (assume $a < b$).

Exercise 4.1.6: We skipped the case of $\lambda < 0$ for the boundary value problem $x'' + \lambda x = 0$, $x(-\pi) = x(\pi)$, $x'(-\pi) = x'(\pi)$. Finish the calculation and show that there are no negative eigenvalues.

Exercise 4.1.101: Consider a spinning string of length 2 and linear density 0.1 and tension 3. Find the smallest angular velocity when the string pops out.

Exercise 4.1.102: Suppose $x'' + \lambda x = 0$ and $x(0) = 1$, $x(1) = 1$. Find all λ for which there is more than one solution. Also find the corresponding solutions (only for the eigenvalues).

Exercise 4.1.103: Suppose $x'' + x = 0$ and $x(0) = 0$, $x'(\pi) = 1$. Find all the solution(s) if any exist.

Exercise 4.1.104: Consider $x' + \lambda x = 0$ and $x(0) = 0$, $x(1) = 0$. Why does it not have any eigenvalues? Why does every first-order equation with two endpoint conditions such as the one above have no eigenvalues?

Exercise 4.1.105 (challenging): Suppose $x''' + \lambda x = 0$ and $x(0) = 0$, $x'(0) = 0$, $x(1) = 0$. Suppose that $\lambda > 0$. Find an equation that all such eigenvalues must satisfy. Hint: Note that $-\sqrt[3]{\lambda}$ is a root of $r^3 + \lambda = 0$.

4.2 The trigonometric series

Note: 2 lectures, §9.1 in [EP], §10.2 in [BD]

4.2.1 Periodic functions and motivation

As motivation for studying Fourier series, consider the problem

$$x'' + \omega_0^2 x = f(t), \quad (4.6)$$

for some periodic function $f(t)$. In §2.6, we found the general solution to

$$x'' + \omega_0^2 x = F_0 \cos(\omega t). \quad (4.7)$$

One way to solve (4.6) is to decompose $f(t)$ as a sum of cosines (and sines) and then solve many problems of the form (4.7). We then use the principle of superposition to sum up all these solutions to get a solution to (4.6).

Before we proceed, let us talk a little bit more in detail about periodic functions. A function is said to be *periodic* with period P if $f(t) = f(t + P)$ for all t . For brevity, we say $f(t)$ is P -periodic. Note that a P -periodic function is also $2P$ -periodic, $3P$ -periodic and so on. For example, $\cos(t)$ and $\sin(t)$ are 2π -periodic. So are $\cos(kt)$ and $\sin(kt)$ for all integers k . The constant functions are an extreme example. They are periodic for any period (exercise).

Normally, we start with a function $f(t)$ defined on some interval $[-L, L]$, and we want to *extend* $f(t)$ *periodically* to make it a $2L$ -periodic function. We do this extension by defining a new function $F(t)$ such that for t in $[-L, L]$, $F(t) = f(t)$. For t in $[L, 3L]$, we define $F(t) = f(t - 2L)$, for t in $[-3L, -L]$, $F(t) = f(t + 2L)$, and so on. To make that work, we needed $f(-L) = f(L)$. We could have also started with f defined only on the half-open interval $(-L, L]$ and then define $f(-L) = f(L)$.

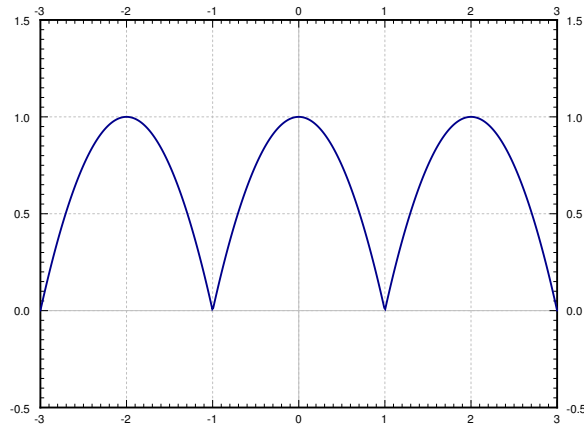
Example 4.2.1: Define $f(t) = 1 - t^2$ on $[-1, 1]$. Extend $f(t)$ periodically to a 2-periodic function. For $1 \leq t \leq 3$, we get $f(t) = 1 - (t - 2)^2$. For $-3 \leq t \leq -1$, we get $f(t) = 1 - (t + 2)^2$. For $3 \leq t \leq 5$, we get $f(t) = 1 - (t - 4)^2$. And so on. See Figure 4.3 on the next page.

You should be careful to distinguish between $f(t)$ and its extension. A common mistake is to assume that a formula for $f(t)$ holds for its extension. It can be confusing when the formula for $f(t)$ is periodic, but with perhaps a different period.

Exercise 4.2.1: Define $f(t) = \cos t$ on $[-\pi/2, \pi/2]$. Take the π -periodic extension and sketch its graph. How does it compare to the graph of $\cos t$?

4.2.2 Inner product and eigenvector decomposition

Suppose A is a *symmetric matrix*, that is, $A^T = A$. As we remarked before, eigenvectors of A are then orthogonal. Here the word *orthogonal* means that if $\vec{v}$ and $\vec{w}$ are two eigenvectors

Figure 4.3: Periodic extension of the function $1 - t^2$.

of A for distinct eigenvalues, then $\langle \vec{v}, \vec{w} \rangle = 0$. In this case, the inner product $\langle \vec{v}, \vec{w} \rangle$ is the *dot product*, which can be computed as $\vec{v}^T \vec{w}$.

To decompose a vector $\vec{v}$ in terms of mutually orthogonal vectors $\vec{w}_1$ and $\vec{w}_2$, we write

$$\vec{v} = a_1 \vec{w}_1 + a_2 \vec{w}_2.$$

Let us find the formula for a_1 and a_2 . We compute,

$$\langle \vec{v}, \vec{w}_1 \rangle = \langle a_1 \vec{w}_1 + a_2 \vec{w}_2, \vec{w}_1 \rangle = a_1 \langle \vec{w}_1, \vec{w}_1 \rangle + a_2 \underbrace{\langle \vec{w}_2, \vec{w}_1 \rangle}_{=0} = a_1 \langle \vec{w}_1, \vec{w}_1 \rangle.$$

Therefore,

$$a_1 = \frac{\langle \vec{v}, \vec{w}_1 \rangle}{\langle \vec{w}_1, \vec{w}_1 \rangle}.$$

Similarly,

$$a_2 = \frac{\langle \vec{v}, \vec{w}_2 \rangle}{\langle \vec{w}_2, \vec{w}_2 \rangle}.$$

You probably remember this formula from vector calculus.

Example 4.2.2: Write $\vec{v} = \begin{bmatrix} 2 \\ 3 \end{bmatrix}$ as a linear combination of $\vec{w}_1 = \begin{bmatrix} 1 \\ -1 \end{bmatrix}$ and $\vec{w}_2 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$.

Note that $\vec{w}_1$ and $\vec{w}_2$ are orthogonal as $\langle \vec{w}_1, \vec{w}_2 \rangle = 1(1) + (-1)1 = 0$. Then

$$\begin{aligned} a_1 &= \frac{\langle \vec{v}, \vec{w}_1 \rangle}{\langle \vec{w}_1, \vec{w}_1 \rangle} = \frac{2(1) + 3(-1)}{1(1) + (-1)(-1)} = \frac{-1}{2}, \\ a_2 &= \frac{\langle \vec{v}, \vec{w}_2 \rangle}{\langle \vec{w}_2, \vec{w}_2 \rangle} = \frac{2 + 3}{1 + 1} = \frac{5}{2}. \end{aligned}$$

Hence,

$$\begin{bmatrix} 2 \\ 3 \end{bmatrix} = \frac{-1}{2} \begin{bmatrix} 1 \\ -1 \end{bmatrix} + \frac{5}{2} \begin{bmatrix} 1 \\ 1 \end{bmatrix}.$$

4.2.3 The trigonometric series

Instead of decomposing a vector in terms of eigenvectors of a matrix, we decompose a function in terms of eigenfunctions of a certain eigenvalue problem. The eigenvalue problem we use for the Fourier series is

$$x'' + \lambda x = 0, \quad x(-\pi) = x(\pi), \quad x'(-\pi) = x'(\pi).$$

We computed that eigenfunctions are $1, \cos(kt), \sin(kt)$. That is, we want to find a representation of a 2π -periodic function $f(t)$ as

$$\begin{aligned} f(t) &= \frac{a_0}{2} + \sum_{n=1}^{\infty} [a_n \cos(nt) + b_n \sin(nt)] \\ &= \frac{a_0}{2} + a_1 \cos(t) + b_1 \sin(t) + a_2 \cos(2t) + b_2 \sin(2t) + \cdots \end{aligned}$$

This series is called the *Fourier series** or the *trigonometric series* for $f(t)$. The term $a_n \cos(nt) + b_n \sin(nt)$ is sometimes called the n^{th} harmonic. We write the coefficient of the eigenfunction 1 as $\frac{a_0}{2}$ for convenience. We could also think of $1 = \cos(0t)$, so that we only need to look at $\cos(kt)$ and $\sin(kt)$.

As for matrices, we want to find a *projection* of $f(t)$ onto the subspaces given by the eigenfunctions. So we want to define an *inner product of functions*. For example, to find a_n , we want to compute $\langle f(t), \cos(nt) \rangle$. We define the inner product as

$$\langle f(t), g(t) \rangle \stackrel{\text{def}}{=} \int_{-\pi}^{\pi} f(t) g(t) dt.$$

With this inner product, we saw in the previous section that the eigenfunctions $\cos(kt)$ (including the constant eigenfunction), and $\sin(kt)$ are *orthogonal*, that is,

$$\begin{aligned} \langle \cos(mt), \cos(nt) \rangle &= 0 && \text{for } m \neq n, \\ \langle \sin(mt), \sin(nt) \rangle &= 0 && \text{for } m \neq n, \\ \langle \sin(mt), \cos(nt) \rangle &= 0 && \text{for all } m \text{ and } n. \end{aligned}$$

For $n = 1, 2, 3, \dots$, we have

$$\begin{aligned} \langle \cos(nt), \cos(nt) \rangle &= \int_{-\pi}^{\pi} \cos(nt) \cos(nt) dt = \pi, \\ \langle \sin(nt), \sin(nt) \rangle &= \int_{-\pi}^{\pi} \sin(nt) \sin(nt) dt = \pi, \end{aligned}$$

by elementary calculus. For the constant, we get

$$\langle 1, 1 \rangle = \int_{-\pi}^{\pi} 1 \cdot 1 dt = 2\pi.$$

*Named after the French mathematician [Jean Baptiste Joseph Fourier](#) (1768–1830).

The coefficients are given by

$$\begin{aligned} a_n &= \frac{\langle f(t), \cos(nt) \rangle}{\langle \cos(nt), \cos(nt) \rangle} = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \cos(nt) dt, \\ b_n &= \frac{\langle f(t), \sin(nt) \rangle}{\langle \sin(nt), \sin(nt) \rangle} = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \sin(nt) dt. \end{aligned}$$

Compare these expressions with the finite-dimensional example. For a_0 , we get a similar formula

$$a_0 = 2 \frac{\langle f(t), 1 \rangle}{\langle 1, 1 \rangle} = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) dt.$$

Let us check the formulas via the orthogonality properties. Suppose for a moment that

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nt) + b_n \sin(nt).$$

Then for $m \geq 1$, we have

$$\begin{aligned} \langle f(t), \cos(mt) \rangle &= \left\langle \frac{a_0}{2} + \sum_{n=1}^{\infty} [a_n \cos(nt) + b_n \sin(nt)], \cos(mt) \right\rangle \\ &= \frac{a_0}{2} \langle 1, \cos(mt) \rangle + \sum_{n=1}^{\infty} [a_n \langle \cos(nt), \cos(mt) \rangle + b_n \langle \sin(nt), \cos(mt) \rangle] \\ &= a_m \langle \cos(mt), \cos(mt) \rangle. \end{aligned}$$

Hence, $a_m = \frac{\langle f(t), \cos(mt) \rangle}{\langle \cos(mt), \cos(mt) \rangle}$.

Exercise 4.2.2: Carry out the calculation for a_0 and b_m .

Example 4.2.3: Take the function

$$f(t) = t$$

for t in $(-\pi, \pi]$. Extend $f(t)$ periodically and write it as a Fourier series. This function is called the *sawtooth*, and it finds many applications, for example, in electronic music.

The plot of the extended periodic function is given in [Figure 4.4](#) on the following page. Let us compute the coefficients. We start with a_0 ,

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} t dt = 0.$$

We will often use the result from calculus that says that the integral of an odd function over a symmetric interval is zero. Recall that an *odd function* is a function $\varphi(t)$ such that $\varphi(-t) = -\varphi(t)$. For example, the functions t , $\sin t$, or (importantly for us) $t \cos(nt)$ are all odd functions. Thus

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} t \cos(nt) dt = 0.$$

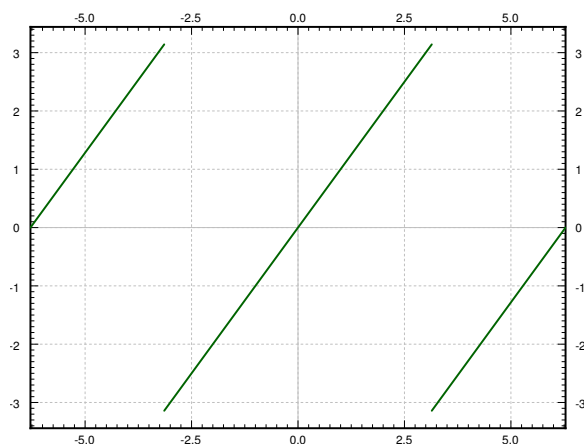


Figure 4.4: The graph of the sawtooth function.

We move to b_n . Another useful fact from calculus is that the integral of an even function over a symmetric interval is twice the integral of the same function over half the interval. An *even function* is a function $\varphi(t)$ such that $\varphi(-t) = \varphi(t)$. For example, $t \sin(nt)$ is even.

$$\begin{aligned}
 b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} t \sin(nt) dt \\
 &= \frac{2}{\pi} \int_0^{\pi} t \sin(nt) dt \\
 &= \frac{2}{\pi} \left(\left[\frac{-t \cos(nt)}{n} \right]_{t=0}^{\pi} + \frac{1}{n} \int_0^{\pi} \cos(nt) dt \right) \\
 &= \frac{2}{\pi} \left(\frac{-\pi \cos(n\pi)}{n} + 0 \right) \\
 &= \frac{-2 \cos(n\pi)}{n} = \frac{2(-1)^{n+1}}{n}.
 \end{aligned}$$

We have used that

$$\cos(n\pi) = (-1)^n = \begin{cases} 1 & \text{if } n \text{ even,} \\ -1 & \text{if } n \text{ odd.} \end{cases}$$

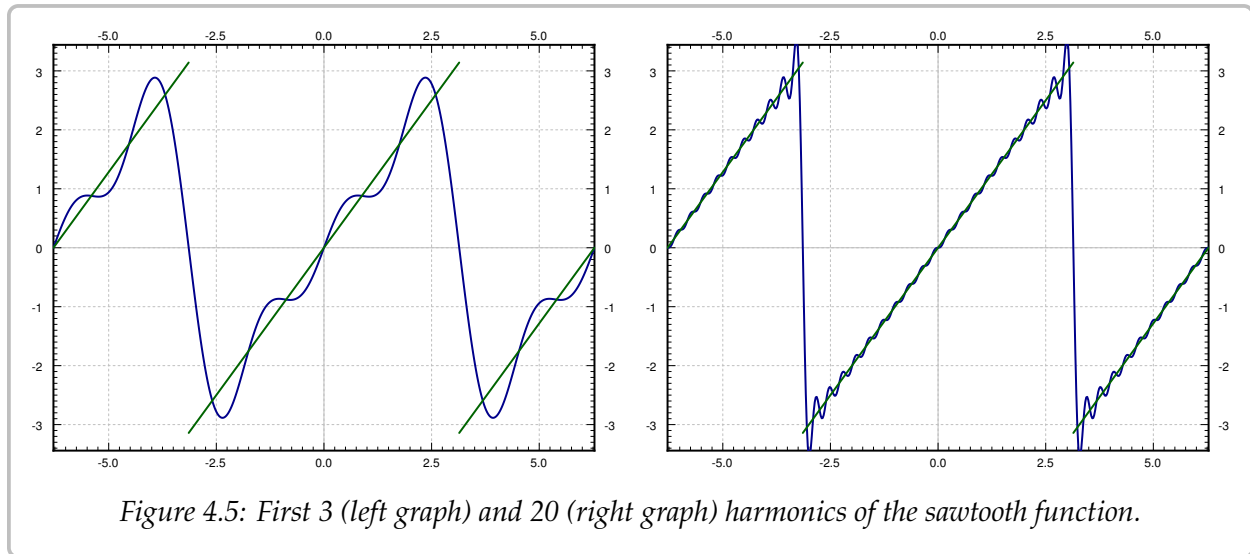
The series, therefore, is

$$\sum_{n=1}^{\infty} \frac{2(-1)^{n+1}}{n} \sin(nt).$$

More explicitly, the first 3 harmonics of the series for $f(t)$ are

$$2 \sin(t) - \sin(2t) + \frac{2}{3} \sin(3t) + \dots$$

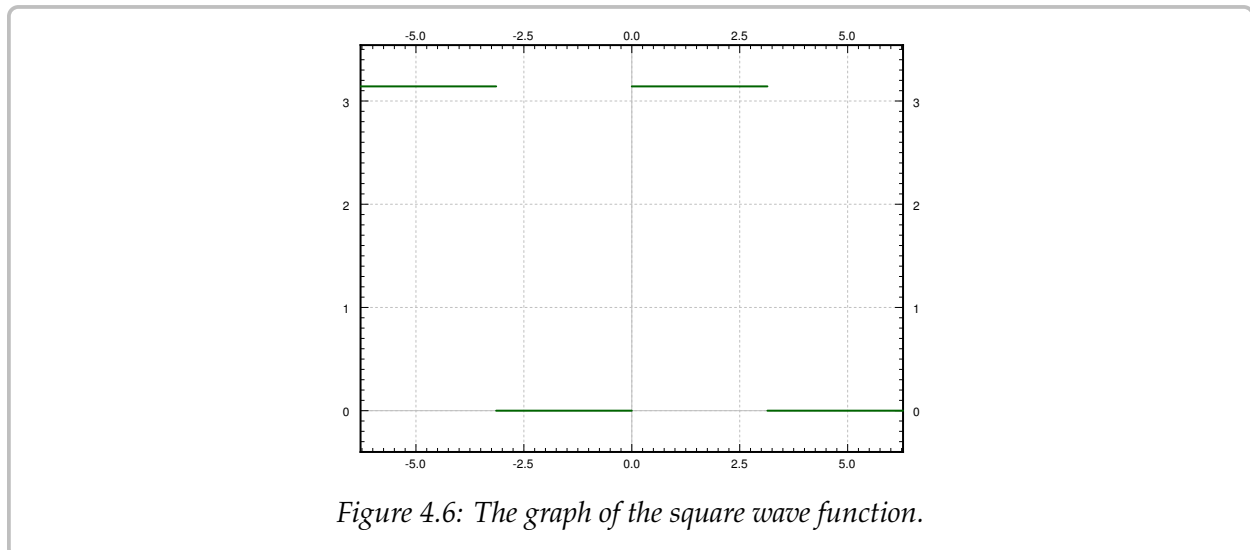
The plot of these first three terms of the series, along with a plot of the first 20 terms is given in [Figure 4.5](#) on the next page.



Example 4.2.4: Take the function

$$f(t) = \begin{cases} 0 & \text{if } -\pi < t \leq 0, \\ \pi & \text{if } 0 < t \leq \pi. \end{cases}$$

Extend $f(t)$ periodically and write it as a Fourier series. This function or its variants appear often in applications and the function is called the *square wave*. It is a signal generated by simply periodically flipping a switch on or off.



The plot of the extended periodic function is given in [Figure 4.6](#). Now we compute the coefficients. We start with a_0

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) dt = \frac{1}{\pi} \int_0^{\pi} \pi dt = \pi.$$

Next,

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \cos(nt) dt = \frac{1}{\pi} \int_0^{\pi} \pi \cos(nt) dt = 0.$$

And finally,

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \sin(nt) dt \\ &= \frac{1}{\pi} \int_0^{\pi} \pi \sin(nt) dt \\ &= \left[\frac{-\cos(nt)}{n} \right]_{t=0}^{\pi} \\ &= \frac{1 - \cos(\pi n)}{n} = \frac{1 - (-1)^n}{n} = \begin{cases} \frac{2}{n} & \text{if } n \text{ is odd,} \\ 0 & \text{if } n \text{ is even.} \end{cases} \end{aligned}$$

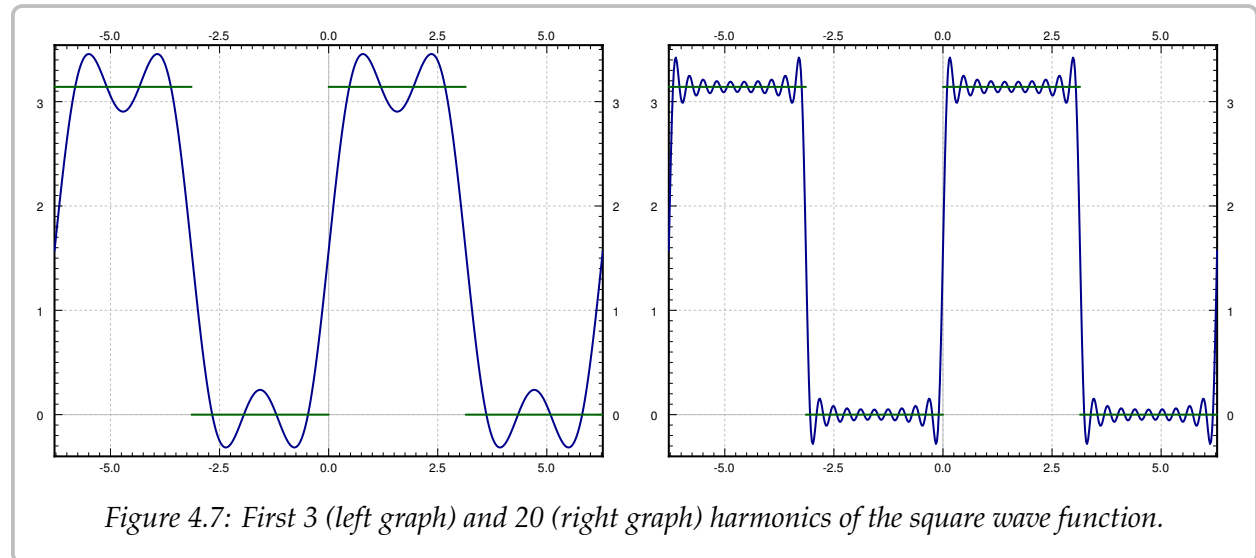
The Fourier series is

$$\frac{\pi}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{2}{n} \sin(nt) = \frac{\pi}{2} + \sum_{k=1}^{\infty} \frac{2}{2k-1} \sin((2k-1)t).$$

The first 3 harmonics of the series for $f(t)$ are

$$\frac{\pi}{2} + 2 \sin(t) + \frac{2}{3} \sin(3t) + \dots$$

The plot of these first three and also of the first 20 terms of the series is given in [Figure 4.7](#).



We have so far skirted the issue of convergence. For example, if $f(t)$ is the square wave function, the equation

$$f(t) = \frac{\pi}{2} + \sum_{k=1}^{\infty} \frac{2}{2k-1} \sin((2k-1)t).$$

is only an equality for such t where $f(t)$ is continuous. We do not get an equality for $t = -\pi, 0, \pi$ and all the other discontinuities of $f(t)$. It is not hard to see that when t is an integer multiple of π (which gives all the discontinuities), then $\sin((2k-1)t) = 0$ and so

$$\frac{\pi}{2} + \sum_{k=1}^{\infty} \frac{2}{2k-1} \sin((2k-1)t) = \frac{\pi}{2}.$$

We redefine $f(t)$ on $[-\pi, \pi]$ as

$$f(t) = \begin{cases} 0 & \text{if } -\pi < t < 0, \\ \pi & \text{if } 0 < t < \pi, \\ \pi/2 & \text{if } t = -\pi, t = 0, \text{ or } t = \pi, \end{cases}$$

and extend periodically. The series equals this new extended $f(t)$ everywhere, including the discontinuities. We will generally not worry about changing the function values at several (finitely many) points.

We will say more about convergence in the next section. Roughly speaking, if $f(t)$ is a nice enough function then the series converges to $f(t)$ wherever $f(t)$ is continuous. Let us, however, briefly mention an effect of the discontinuity. We zoom in near the discontinuity in the square wave and we plot the first 100 harmonics. See [Figure 4.8](#) on the following page. While the series is a very good approximation away from the discontinuities, the error (the overshoot) near the discontinuity at $t = \pi$ does not seem to be getting any smaller as we take more and more harmonics. This behavior is known as the *Gibbs phenomenon*. The region where the error is large does get smaller, however, the more terms in the series we take.

We can think of a periodic function as a “signal” that is a superposition of many signals of pure frequency. For example, we could think of the square wave as a tone of certain base frequency coming through the speaker of your computer. This base frequency is called the *fundamental frequency* and that is the “musical note” that you identify this sound as. But the square wave will also be a superposition of many different pure tones of frequencies that are multiples of the fundamental frequency. In music, the higher frequencies are called the *overtones*. The set of all the frequencies that appear is called the *spectrum* of the signal. On the other hand, if the signal were a simple sine wave instead of the square wave, then it would be only the pure tone of fundamental frequency (no overtones). The simplest way to make sound using a computer is the square wave, and the sound is very different from a pure tone. If you ever played video games from the 1980s or so, then you heard what square waves sound like.

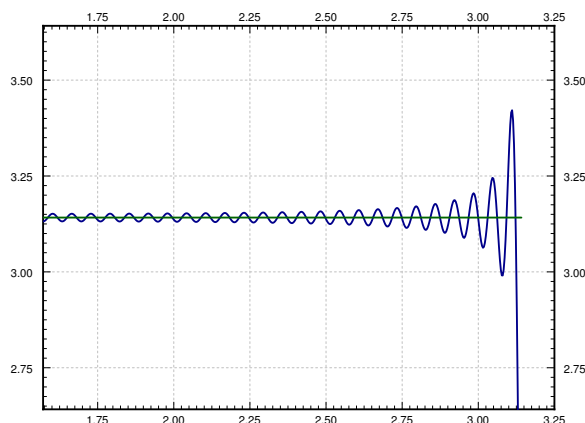


Figure 4.8: Gibbs phenomenon in action.

4.2.4 Exercises

Exercise 4.2.3: Suppose $f(t)$ is defined on $[-\pi, \pi]$ as $\sin(5t) + \cos(3t)$. Extend periodically and compute the Fourier series of $f(t)$.

Exercise 4.2.4: Suppose $f(t)$ is defined on $[-\pi, \pi]$ as $|t|$. Extend periodically and compute the Fourier series of $f(t)$.

Exercise 4.2.5: Suppose $f(t)$ is defined on $[-\pi, \pi]$ as $|t|^3$. Extend periodically and compute the Fourier series of $f(t)$.

Exercise 4.2.6: Suppose $f(t)$ is defined on $(-\pi, \pi]$ as

$$f(t) = \begin{cases} -1 & \text{if } -\pi < t \leq 0, \\ 1 & \text{if } 0 < t \leq \pi. \end{cases}$$

Extend periodically and compute the Fourier series of $f(t)$.

Exercise 4.2.7: Suppose $f(t)$ is defined on $(-\pi, \pi]$ as t^3 . Extend periodically and compute the Fourier series of $f(t)$.

Exercise 4.2.8: Suppose $f(t)$ is defined on $[-\pi, \pi]$ as t^2 . Extend periodically and compute the Fourier series of $f(t)$.

There is another form of the Fourier series using complex exponentials e^{nt} for $n = \dots, -2, -1, 0, 1, 2, \dots$ instead of $\cos(nt)$ and $\sin(nt)$ for positive n . This form may be easier to work with sometimes. It is certainly more compact to write, and there is only one formula for the coefficients. On the downside, the coefficients are complex numbers.

Exercise 4.2.9: Let

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos(nt) + b_n \sin(nt) \right].$$

Use Euler's formula $e^{i\theta} = \cos(\theta) + i \sin(\theta)$ to show that there exist complex numbers c_m such that

$$f(t) = \sum_{m=-\infty}^{\infty} c_m e^{imt}.$$

Note that the sum now ranges over all the integers including negative ones. Do not worry about convergence in this calculation. Hint: It may be better to start from the complex exponential form and write the series as

$$c_0 + \sum_{m=1}^{\infty} \left[c_m e^{imt} + c_{-m} e^{-imt} \right].$$

Exercise 4.2.101: Suppose $f(t)$ is defined on $[-\pi, \pi]$ as $f(t) = \sin(t)$. Extend periodically and compute the Fourier series.

Exercise 4.2.102: Suppose $f(t)$ is defined on $(-\pi, \pi]$ as $f(t) = \sin(\pi t)$. Extend periodically and compute the Fourier series.

Exercise 4.2.103: Suppose $f(t)$ is defined on $(-\pi, \pi]$ as $f(t) = \sin^2(t)$. Extend periodically and compute the Fourier series.

Exercise 4.2.104: Suppose $f(t)$ is defined on $(-\pi, \pi]$ as $f(t) = t^4$. Extend periodically and compute the Fourier series.

4.3 More on the Fourier series

Note: 2 lectures, §9.2–§9.3 in [EP], §10.3 in [BD]

4.3.1 $2L$ -periodic functions

We computed the Fourier series for a 2π -periodic function, but what about functions of different periods? Well, fear not, the computation is a simple case of change of variables. We just rescale the independent axis. Consider a $2L$ -periodic function $f(t)$. The L is called the *half period*. Let $s = \frac{\pi}{L}t$. Then the function

$$g(s) = f\left(\frac{L}{\pi}s\right)$$

is 2π -periodic and we know what to do with it. We must also rescale all our sines and cosines. In the series, we use $\frac{\pi}{L}t$ as the variable. That is, we want to write

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right].$$

If we change variables to s , we see that

$$g(s) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos(ns) + b_n \sin(ns) \right].$$

We compute a_n and b_n as before. After we write down the integrals, we change variables from s back to t , noting also that $ds = \frac{\pi}{L} dt$.

$$\begin{aligned} a_0 &= \frac{1}{\pi} \int_{-\pi}^{\pi} g(s) ds = \frac{1}{L} \int_{-L}^L f(t) dt, \\ a_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} g(s) \cos(ns) ds = \frac{1}{L} \int_{-L}^L f(t) \cos\left(\frac{n\pi}{L}t\right) dt, \\ b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} g(s) \sin(ns) ds = \frac{1}{L} \int_{-L}^L f(t) \sin\left(\frac{n\pi}{L}t\right) dt. \end{aligned}$$

The two most common half periods that show up in examples are π and 1 because of the simplicity of the formulas. We should stress that we have done no new mathematics—we have only changed variables. If you understand the Fourier series for 2π -periodic functions, you understand it for $2L$ -periodic functions. You can think of it as just using different units for time. All that we are doing is moving some constants around, but all the mathematics is the same.

Example 4.3.1: Let

$$f(t) = |t| \quad \text{for } -1 < t \leq 1,$$

extended periodically. The plot of the periodic extension is given in Figure 4.9. Compute the Fourier series of $f(t)$.

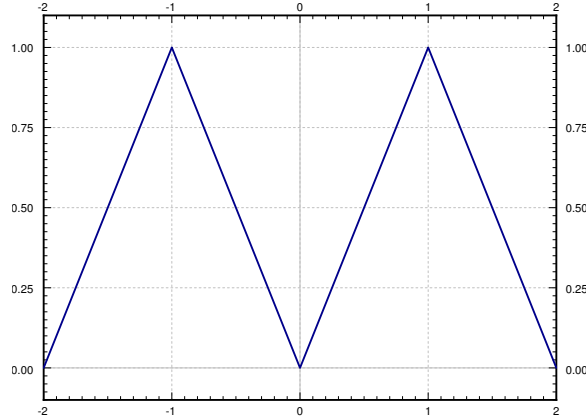


Figure 4.9: Periodic extension of the function $f(t)$.

First, we recognize that f is 2-periodic and so $L = 1$. We want to write $f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\pi t) + b_n \sin(n\pi t)$. We start with a_n for $n \geq 1$. We note that $|t| \cos(n\pi t)$ is even and for $0 \leq t \leq 1$, $f(t) = |t| = t$. Hence,

$$\begin{aligned} a_n &= \int_{-1}^1 f(t) \cos(n\pi t) dt \\ &= 2 \int_0^1 t \cos(n\pi t) dt \\ &= 2 \left[\frac{t}{n\pi} \sin(n\pi t) \right]_{t=0}^1 - 2 \int_0^1 \frac{1}{n\pi} \sin(n\pi t) dt \\ &= 0 + \frac{2}{n^2\pi^2} \left[\cos(n\pi t) \right]_{t=0}^1 = \frac{2((-1)^n - 1)}{n^2\pi^2} = \begin{cases} 0 & \text{if } n \text{ is even,} \\ \frac{-4}{n^2\pi^2} & \text{if } n \text{ is odd.} \end{cases} \end{aligned}$$

Next we find a_0 :

$$a_0 = \int_{-1}^1 |t| dt = 1.$$

You should be able to find this integral by thinking about the integral as the area under the graph without doing any computation at all. Finally, we find b_n . Notice that $|t| \sin(n\pi t)$ is odd, and so

$$b_n = \int_{-1}^1 f(t) \sin(n\pi t) dt = 0.$$

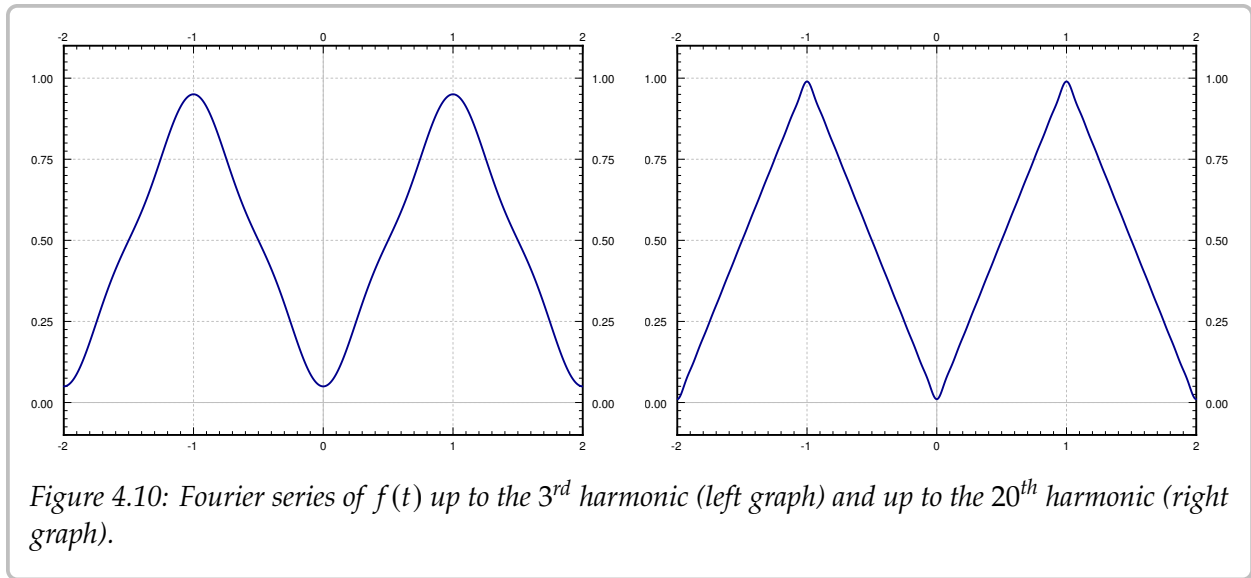
Hence, the series is

$$\frac{1}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{n^2\pi^2} \cos(n\pi t).$$

The first few terms of the series up to the 3rd harmonic are

$$\frac{1}{2} - \frac{4}{\pi^2} \cos(\pi t) - \frac{4}{9\pi^2} \cos(3\pi t) - \dots$$

The plot of these few terms and also a plot up to the 20th harmonic is given in [Figure 4.10](#). You should notice how close the graph is to the real function. You should also notice that there is no “Gibbs phenomenon” present as there are no discontinuities.



4.3.2 Convergence

We will need the one-sided limits of functions. We will use the following notation

$$f(c-) = \lim_{t \uparrow c} f(t), \quad \text{and} \quad f(c+) = \lim_{t \downarrow c} f(t).$$

If you are unfamiliar with this notation, $\lim_{t \uparrow c} f(t)$ means we are taking a limit of $f(t)$ as t approaches c from below (i.e. $t < c$) and $\lim_{t \downarrow c} f(t)$ means we are taking a limit of $f(t)$ as t approaches c from above (i.e. $t > c$). For example, for the square wave function

$$f(t) = \begin{cases} 0 & \text{if } -\pi < t \leq 0, \\ \pi & \text{if } 0 < t \leq \pi, \end{cases} \quad (4.8)$$

we have $f(0-) = 0$ and $f(0+) = \pi$.

Let $f(t)$ be a function defined on an interval $[a, b]$. Suppose that we find finitely many points $a = t_0, t_1, t_2, \dots, t_k = b$ in the interval, such that $f(t)$ is continuous on the intervals $(t_0, t_1), (t_1, t_2), \dots, (t_{k-1}, t_k)$. Also suppose that all the one-sided limits exist, that is, all of $f(t_0+), f(t_1-), f(t_1+), f(t_2-), f(t_2+), \dots, f(t_k-)$ exist and are finite. Then we say $f(t)$ is *piecewise continuous*.

If moreover, $f(t)$ is differentiable at all but finitely many points, and $f'(t)$ is piecewise continuous, then $f(t)$ is said to be *piecewise smooth*.

Example 4.3.2: The square wave function, (4.8) extended periodically, is piecewise smooth on $[-\pi, \pi]$ or any other finite interval, so we just say that $f(t)$ is piecewise smooth without mentioning an interval.

Example 4.3.3: The function $f(t) = |t|$ is piecewise smooth.

Example 4.3.4: The function $f(t) = \frac{1}{t}$ is not piecewise smooth on $[-1, 1]$ (or any other interval containing zero). In fact, it is not even piecewise continuous.

Example 4.3.5: The function $f(t) = \sqrt[3]{t}$ is not piecewise smooth on $[-1, 1]$ (or any other interval containing zero). The function $f(t)$ is continuous, but its derivative $f'(t) = \frac{1}{3t^{2/3}}$ is unbounded near zero and hence not piecewise continuous.

Piecewise smooth functions have an easy answer on the convergence of the Fourier series.

Theorem 4.3.1. Suppose $f(t)$ is a $2L$ -periodic piecewise smooth function. Let

$$\frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right]$$

be the Fourier series for $f(t)$. Then the series converges for all t . If $f(t)$ is continuous at t , then

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right].$$

Otherwise,

$$\frac{f(t-) + f(t+)}{2} = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right].$$

If we happen to have that $f(t) = \frac{f(t-) + f(t+)}{2}$ at all the discontinuities, the Fourier series converges to $f(t)$ everywhere. We can always just redefine $f(t)$ by changing the value at each discontinuity appropriately. Then we can write an equals sign between $f(t)$ and the series without any worry. We mentioned this fact briefly at the end of the last section.

The theorem does not say how fast the series converges. Think back to the discussion of the Gibbs phenomenon in the last section. The closer you get to the discontinuity, the more terms you need to take to get an accurate approximation to the function.

4.3.3 Differentiation and integration of Fourier series

Not only does Fourier series converge nicely, but it is easy to differentiate and integrate the series. We can do this just by differentiating or integrating term by term.

Theorem 4.3.2. *Suppose*

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right]$$

is a piecewise smooth continuous function and the derivative $f'(t)$ is piecewise smooth. Then the derivative can be obtained by differentiating term by term,

$$f'(t) = \sum_{n=1}^{\infty} \left[\frac{-a_n n \pi}{L} \sin\left(\frac{n\pi}{L}t\right) + \frac{b_n n \pi}{L} \cos\left(\frac{n\pi}{L}t\right) \right].$$

It is important that the function is continuous. It can have corners, but no jumps. Otherwise, the differentiated series will fail to converge. For an exercise, take the series obtained for the square wave and try to differentiate the series. Similarly to differentiation, integration of Fourier series is also done term by term.

Theorem 4.3.3. *Suppose*

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right]$$

is a piecewise smooth function. Then the antiderivative is obtained by antidifferentiating term by term and so

$$F(t) = \frac{a_0 t}{2} + C + \sum_{n=1}^{\infty} \left[\frac{a_n L}{n\pi} \sin\left(\frac{n\pi}{L}t\right) + \frac{-b_n L}{n\pi} \cos\left(\frac{n\pi}{L}t\right) \right],$$

where $F'(t) = f(t)$ and C is an arbitrary constant.

Note that the series for $F(t)$ is no longer a Fourier series as it contains the $\frac{a_0 t}{2}$ term. The antiderivative of a periodic function need no longer be periodic and so we should not expect a Fourier series. Unless, of course, $a_0 = 0$.

4.3.4 Rates of convergence and smoothness

We consider an example of a periodic function differentiable once but not twice.

Example 4.3.6: Take the function

$$f(t) = \begin{cases} (t+1)t & \text{if } -1 < t \leq 0, \\ (1-t)t & \text{if } 0 < t \leq 1, \end{cases}$$

and extend to a 2-periodic function. The derivative $f'(t)$ exists for all t , but $f''(t)$ does not exist if t is an integer. In particular, $f'(t) = 2t + 1$ if $-1 \leq t \leq 0$ and $f'(t) = 1 - 2t$ if $0 \leq t \leq 1$, so $f'(t)$ has a “corner” at 0. See [Figure 4.11](#) on the next page for the plot of $f(t)$ and $f'(t)$.

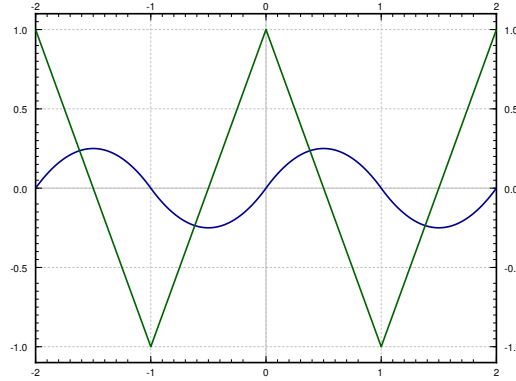


Figure 4.11: Smooth 2-periodic function (the smooth line) with its nonsmooth derivative (the jagged line).

Exercise 4.3.1: Compute $f''(0+)$ and $f''(0-)$.

Let us compute the Fourier-series coefficients. The actual computation involves several integrations by parts and is left to the reader.

$$\begin{aligned}
 a_0 &= \int_{-1}^1 f(t) dt = \int_{-1}^0 (t+1)t dt + \int_0^1 (1-t)t dt = 0, \\
 a_n &= \int_{-1}^1 f(t) \cos(n\pi t) dt = \int_{-1}^0 (t+1)t \cos(n\pi t) dt + \int_0^1 (1-t)t \cos(n\pi t) dt = 0, \\
 b_n &= \int_{-1}^1 f(t) \sin(n\pi t) dt = \int_{-1}^0 (t+1)t \sin(n\pi t) dt + \int_0^1 (1-t)t \sin(n\pi t) dt \\
 &= \frac{4(1 - (-1)^n)}{\pi^3 n^3} = \begin{cases} \frac{8}{\pi^3 n^3} & \text{if } n \text{ is odd,} \\ 0 & \text{if } n \text{ is even.} \end{cases}
 \end{aligned}$$

That is, the series is

$$\sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{8}{\pi^3 n^3} \sin(n\pi t).$$

This series converges very fast. If we plot up to the third harmonic, that is,

$$\frac{8}{\pi^3} \sin(\pi t) + \frac{8}{27\pi^3} \sin(3\pi t),$$

the plot is almost indistinguishable from the plot of $f(t)$ in Figure 4.11. In fact, the coefficient $\frac{8}{27\pi^3}$ is already just 0.0096 (approximately). The reason for this behavior is the n^3 term in the denominator. The coefficients b_n go to zero as fast as $1/n^3$ goes to zero.

For functions constructed piecewise from polynomials as above, it is generally true that if you have one derivative but not two derivatives, the Fourier coefficients will go to zero

approximately like $1/n^3$. If you have only a continuous function that is not differentiable, then the Fourier coefficients will go to zero as $1/n^2$. If you have discontinuities, then the Fourier coefficients will go to zero approximately as $1/n$. For more general functions the story is somewhat more complicated but the same idea holds: The more derivatives you have, the faster the coefficients go to zero. Similar reasoning works in reverse. If the coefficients go to zero like $1/n^2$ (or faster), you always obtain a continuous function. If they go to zero like $1/n^3$ (or faster), you obtain an everywhere differentiable function.

To justify this behavior, take for example the function defined by the Fourier series

$$f(t) = \sum_{n=1}^{\infty} \frac{1}{n^3} \sin(nt).$$

When we differentiate term by term, we notice

$$f'(t) = \sum_{n=1}^{\infty} \frac{1}{n^2} \cos(nt).$$

Therefore, the coefficients now go down like $1/n^2$, which means that we have a continuous function. The derivative of $f'(t)$ is defined at most points, but there are points where $f'(t)$ is not differentiable. It has corners, but no jumps. If we differentiate again (where we can), we find that the function $f''(t)$ now fails to be continuous (has jumps)

$$f''(t) = \sum_{n=1}^{\infty} \frac{-1}{n} \sin(nt).$$

This function is similar to the sawtooth. If we tried to differentiate the series again, we would obtain

$$\sum_{n=1}^{\infty} -\cos(nt),$$

which does not converge!

Exercise 4.3.2: Use a computer to plot the series we obtained for $f(t)$, $f'(t)$ and $f''(t)$. That is, plot say the first 5 harmonics of the functions. At what points does $f''(t)$ have the discontinuities?

4.3.5 Exercises

Exercise 4.3.3: Let

$$f(t) = \begin{cases} 0 & \text{if } -1 < t \leq 0, \\ t & \text{if } 0 < t \leq 1, \end{cases}$$

extended periodically.

- Compute the Fourier series for $f(t)$.
- Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.4: Let

$$f(t) = \begin{cases} -t & \text{if } -1 < t \leq 0, \\ t^2 & \text{if } 0 < t \leq 1, \end{cases}$$

extended periodically.

- Compute the Fourier series for $f(t)$.
- Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.5: Let

$$f(t) = \begin{cases} \frac{-t}{10} & \text{if } -10 < t \leq 0, \\ \frac{t}{10} & \text{if } 0 < t \leq 10, \end{cases}$$

extended periodically (period is 20).

- Compute the Fourier series for $f(t)$.
- Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.6: Let $f(t) = \sum_{n=1}^{\infty} \frac{1}{n^3} \cos(nt)$. Is $f(t)$ continuous and differentiable everywhere? Find the derivative (if it exists everywhere) or justify why $f(t)$ is not differentiable everywhere.

Exercise 4.3.7: Let $f(t) = \sum_{n=1}^{\infty} \frac{(-1)^n}{n} \sin(nt)$. Is $f(t)$ differentiable everywhere? Find the derivative (if it exists everywhere) or justify why $f(t)$ is not differentiable everywhere.

Exercise 4.3.8: Let

$$f(t) = \begin{cases} 0 & \text{if } -2 < t \leq 0, \\ t & \text{if } 0 < t \leq 1, \\ -t + 2 & \text{if } 1 < t \leq 2, \end{cases}$$

extended periodically.

- Compute the Fourier series for $f(t)$.
- Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.9: Let

$$f(t) = e^t \quad \text{for } -1 < t \leq 1$$

extended periodically.

- Compute the Fourier series for $f(t)$.
- Write out the series explicitly up to the 3rd harmonic.
- What does the series converge to at $t = 1$.

Exercise 4.3.10: Let

$$f(t) = t^2 \quad \text{for } -1 < t \leq 1$$

extended periodically.

a) Compute the Fourier series for $f(t)$.

b) By plugging in $t = 0$, evaluate $\sum_{n=1}^{\infty} \frac{(-1)^n}{n^2} = -1 + \frac{1}{4} - \frac{1}{9} + \cdots$.

c) Now evaluate $\sum_{n=1}^{\infty} \frac{1}{n^2} = 1 + \frac{1}{4} + \frac{1}{9} + \cdots$.

Exercise 4.3.11: Let

$$f(t) = \begin{cases} 0 & \text{if } -3 < t \leq 0, \\ t & \text{if } 0 < t \leq 3, \end{cases}$$

extended periodically. Suppose $F(t)$ is the function given by the Fourier series of f . Without computing the Fourier series evaluate

a) $F(2)$

b) $F(-2)$

c) $F(4)$

d) $F(-4)$

e) $F(3)$

f) $F(-9)$

Exercise 4.3.101: Let

$$f(t) = t^2 \quad \text{for } -2 < t \leq 2$$

extended periodically.

a) Compute the Fourier series for $f(t)$.

b) Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.102: Let

$$f(t) = t \quad \text{for } -\lambda < t \leq \lambda \quad (\text{for some } \lambda > 0)$$

extended periodically.

a) Compute the Fourier series for $f(t)$.

b) Write out the series explicitly up to the 3rd harmonic.

Exercise 4.3.103: Let

$$f(t) = \frac{1}{2} + \sum_{n=1}^{\infty} \frac{1}{n(n^2 + 1)} \sin(n\pi t).$$

Compute $f'(t)$.

Exercise 4.3.104: Let

$$f(t) = \frac{1}{2} + \sum_{n=1}^{\infty} \frac{1}{n^3} \cos(nt).$$

- a) Find the antiderivative.
- b) Is the antiderivative periodic?

Exercise 4.3.105: Let

$$f(t) = t/2 \quad \text{for } -\pi < t < \pi$$

extended periodically.

- a) Compute the Fourier series for $f(t)$.
- b) Plug in $t = \pi/2$ to find a series representation for $\pi/4$.
- c) Using the first 4 terms of the result from part b) approximate $\pi/4$.

Exercise 4.3.106: Let

$$f(t) = \begin{cases} 0 & \text{if } -2 < t \leq 0, \\ 2 & \text{if } 0 < t \leq 2, \end{cases}$$

extended periodically. Suppose $F(t)$ is the function given by the Fourier series of f . Without computing the Fourier series evaluate

- | | | |
|------------|------------|------------|
| a) $F(0)$ | b) $F(-1)$ | c) $F(1)$ |
| d) $F(-2)$ | e) $F(4)$ | f) $F(-9)$ |

4.4 Sine and cosine series

Note: 2 lectures, §9.3 in [EP], §10.4 in [BD]

4.4.1 Odd and even periodic functions

You may have noticed by now that an odd function has no cosine terms in the Fourier series and an even function has no sine terms in the Fourier series. This observation is not a coincidence. Let us look at even and odd periodic functions in more detail.

Recall that a function $f(t)$ is *odd* if $f(-t) = -f(t)$. A function $f(t)$ is *even* if $f(-t) = f(t)$. For example, $\cos(nt)$ is even and $\sin(nt)$ is odd. Similarly, the function t^k is even if k is even and odd if k is odd.

Exercise 4.4.1: Take two functions $f(t)$ and $g(t)$ and define their product $h(t) = f(t)g(t)$.

- a) Suppose both $f(t)$ and $g(t)$ are odd. Is $h(t)$ odd or even?
- b) Suppose one is even and one is odd. Is $h(t)$ odd or even?
- c) Suppose both are even. Is $h(t)$ odd or even?

If $f(t)$ and $g(t)$ are both odd, then $f(t) + g(t)$ is odd. Similarly for even functions. On the other hand, if $f(t)$ is odd and $g(t)$ even, then we cannot say anything about the sum $f(t) + g(t)$. In fact, the Fourier series of any function is a sum of an odd function (the sine terms) and an even function (the cosine terms).

In this section, we consider odd and even periodic functions. We have previously defined the $2L$ -periodic extension of a function defined on the interval $[-L, L]$. Sometimes we are only interested in the function on the interval $[0, L]$, and it would be convenient to have an odd (resp. even) function. If the function is odd (resp. even), all the cosine (resp. sine) terms disappear. What we will do is take the odd (resp. even) extension of the function to $[-L, L]$ and then extend periodically to a $2L$ -periodic function.

Take a function $f(t)$ defined on $[0, L]$. On $(-L, L]$ define the functions

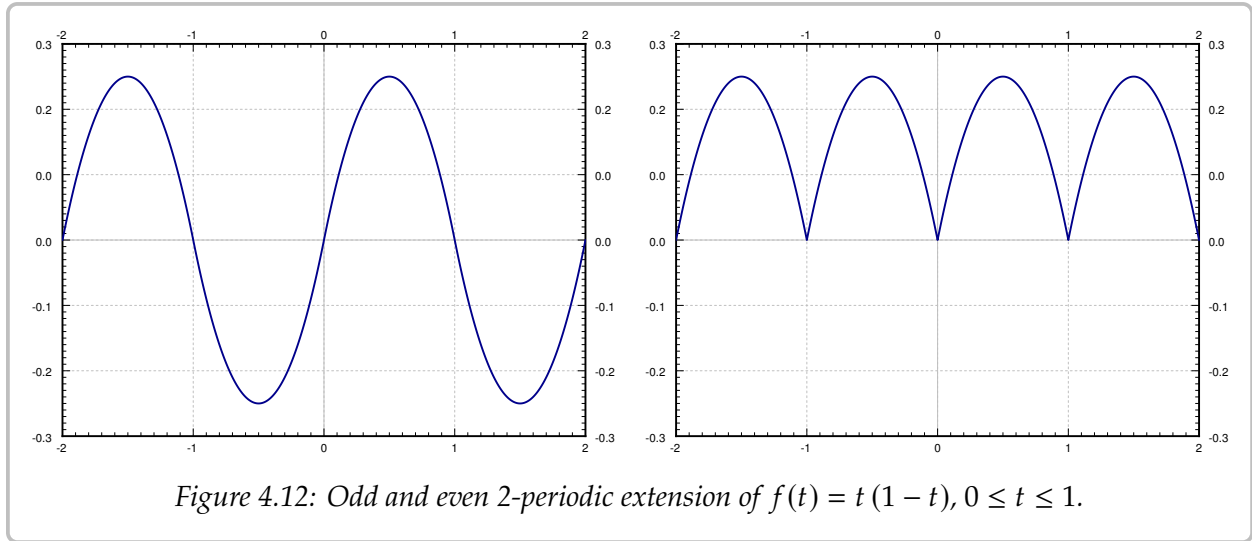
$$F_{\text{odd}}(t) \stackrel{\text{def}}{=} \begin{cases} f(t) & \text{if } 0 \leq t \leq L, \\ -f(-t) & \text{if } -L < t < 0, \end{cases}$$

$$F_{\text{even}}(t) \stackrel{\text{def}}{=} \begin{cases} f(t) & \text{if } 0 \leq t \leq L, \\ f(-t) & \text{if } -L < t < 0. \end{cases}$$

Extend $F_{\text{odd}}(t)$ and $F_{\text{even}}(t)$ to be $2L$ -periodic. Then $F_{\text{odd}}(t)$ is called the *odd periodic extension* of $f(t)$, and $F_{\text{even}}(t)$ is called the *even periodic extension* of $f(t)$. For the odd extension, we generally assume that $f(0) = f(L) = 0$.

Exercise 4.4.2: Check that $F_{\text{odd}}(t)$ is odd and $F_{\text{even}}(t)$ is even. For F_{odd} , assume $f(0) = f(L) = 0$.

Example 4.4.1: Take the function $f(t) = t(1 - t)$ defined on $[0, 1]$. Figure 4.12 on the next page shows the plots of the odd and even periodic extensions of $f(t)$.



4.4.2 Sine and cosine series

Let us write the Fourier series for an odd $2L$ -periodic function $f(t)$. First, we compute the coefficients a_n (including $n = 0$) and get

$$a_n = \frac{1}{L} \int_{-L}^L f(t) \cos\left(\frac{n\pi}{L}t\right) dt = 0.$$

That is, the Fourier series of an odd function has no cosine terms. The integral is zero as $f(t) \cos\left(\frac{n\pi}{L}t\right)$ is an odd function (product of an odd and an even function is odd) and the integral of an odd function over a symmetric interval is zero. The function $f(t) \sin\left(\frac{n\pi}{L}t\right)$ is even as it is the product of two odd functions. The integral of an even function over a symmetric interval $[-L, L]$ is twice the integral of the function over the interval $[0, L]$. So

$$b_n = \frac{1}{L} \int_{-L}^L f(t) \sin\left(\frac{n\pi}{L}t\right) dt = \frac{2}{L} \int_0^L f(t) \sin\left(\frac{n\pi}{L}t\right) dt.$$

The Fourier series of an odd $f(t)$ is then

$$\sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}t\right).$$

Now suppose $f(t)$ is an even $2L$ -periodic function. For the same exact reasons as above, we find that $b_n = 0$ and

$$a_n = \frac{2}{L} \int_0^L f(t) \cos\left(\frac{n\pi}{L}t\right) dt.$$

The formula still works for $n = 0$, in which case it becomes

$$a_0 = \frac{2}{L} \int_0^L f(t) dt.$$

The Fourier series of an even $f(t)$ is then

$$\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos\left(\frac{n\pi}{L}t\right).$$

An interesting consequence is that the coefficients of the Fourier series of an odd (or even) function can be computed by just integrating over the half-interval $[0, L]$. Therefore, we can compute the Fourier series of the odd (or even) extension of a function by computing certain integrals over the interval where the original function is defined.

Theorem 4.4.1. *Let $f(t)$ be a piecewise smooth function defined on $[0, L]$. Then the odd periodic extension of $f(t)$ has the Fourier series*

$$F_{\text{odd}}(t) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}t\right),$$

where

$$b_n = \frac{2}{L} \int_0^L f(t) \sin\left(\frac{n\pi}{L}t\right) dt.$$

The even periodic extension of $f(t)$ has the Fourier series

$$F_{\text{even}}(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos\left(\frac{n\pi}{L}t\right),$$

where

$$a_n = \frac{2}{L} \int_0^L f(t) \cos\left(\frac{n\pi}{L}t\right) dt.$$

We call the series $\sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}t\right)$ the *sine series* of $f(t)$ and we call the series $\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos\left(\frac{n\pi}{L}t\right)$ the *cosine series* of $f(t)$. We often do not actually care what happens outside of $[0, L]$. We simply pick whichever series fits our problem better.

It is not necessary to start with the full Fourier series to obtain the sine and cosine series. The sine series is really the eigenfunction expansion of $f(t)$ using eigenfunctions of the eigenvalue problem $x'' + \lambda x = 0$, $x(0) = 0$, $x(L) = 0$. The cosine series is the eigenfunction expansion of $f(t)$ using eigenfunctions of the eigenvalue problem $x'' + \lambda x = 0$, $x'(0) = 0$, $x'(L) = 0$. We would, therefore, get the same formulas by defining the inner product

$$\langle f(t), g(t) \rangle = \int_0^L f(t)g(t) dt,$$

and following the procedure of § 4.2. This point of view is useful, as we commonly use a specific series that arose because our underlying question led to a certain eigenvalue

problem. If the eigenvalue problem is not one of the three we covered so far, you can still do an eigenfunction expansion, generalizing the results of this chapter. We will deal with such a generalization in [chapter 5](#).

Example 4.4.2: Find the Fourier series of the even periodic extension of the function $f(t) = t^2$ for $0 \leq t \leq \pi$.

We want to write

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nt),$$

where

$$a_0 = \frac{2}{\pi} \int_0^{\pi} t^2 dt = \frac{2\pi^2}{3},$$

and

$$\begin{aligned} a_n &= \frac{2}{\pi} \int_0^{\pi} t^2 \cos(nt) dt = \frac{2}{\pi} \left[t^2 \frac{1}{n} \sin(nt) \right]_0^{\pi} - \frac{4}{n\pi} \int_0^{\pi} t \sin(nt) dt \\ &= \frac{4}{n^2\pi} \left[t \cos(nt) \right]_0^{\pi} - \frac{4}{n^2\pi} \int_0^{\pi} \cos(nt) dt = \frac{4(-1)^n}{n^2}. \end{aligned}$$

Note that we “detected” the continuity of the extension since the coefficients decay as $\frac{1}{n^2}$. That is, the even periodic extension of t^2 has no jump discontinuities. It does have corners, since the derivative, which is an odd function and a sine series, has jumps; it has a Fourier series whose coefficients decay only as $\frac{1}{n}$.

Explicitly, the first few terms of the series for $f(t)$ are

$$\frac{\pi^2}{3} - 4 \cos(t) + \cos(2t) - \frac{4}{9} \cos(3t) + \cdots$$

Exercise 4.4.3:

- Compute the derivative of the even periodic extension of $f(t)$ above and verify it has jump discontinuities. Use the actual definition of $f(t)$, not its cosine series!
- Why is it that the derivative of the even periodic extension of $f(t)$ is the odd periodic extension of $f'(t)$?

4.4.3 Application

Fourier series ties in to the boundary value problems that we studied earlier. Consider the boundary value problem

$$x''(t) + \lambda x(t) = f(t), \quad 0 < t < L,$$

with the *Dirichlet boundary conditions* $x(0) = 0$, $x(L) = 0$. The Fredholm alternative ([Theorem 4.1.2](#) on page 194) says that as long as λ is not an eigenvalue of the underlying

homogeneous problem, there exists a unique solution. Eigenfunctions of this eigenvalue problem are the functions $\sin\left(\frac{n\pi}{L}t\right)$. To find the solution, we first find the Fourier sine series for $f(t)$. We write $x(t)$ also as a sine series, but with unknown coefficients. We substitute the series for $x(t)$ into the equation and solve for the unknown coefficients. If we have the *Neumann boundary conditions* $x'(0) = 0$, $x'(L) = 0$, we do the same procedure using the cosine series. Let us see how this method works on examples.

Example 4.4.3: Take the boundary value problem

$$x''(t) + 2x(t) = f(t), \quad 0 < t < 1,$$

where $f(t) = t$ on $0 < t < 1$, and satisfying the Dirichlet boundary conditions $x(0) = 0$, $x(1) = 0$. We write $f(t)$ as a sine series

$$f(t) = \sum_{n=1}^{\infty} c_n \sin(n\pi t).$$

Compute

$$c_n = 2 \int_0^1 t \sin(n\pi t) dt = \frac{2(-1)^{n+1}}{n\pi}.$$

We write $x(t)$ as

$$x(t) = \sum_{n=1}^{\infty} b_n \sin(n\pi t).$$

We plug in to obtain

$$\begin{aligned} x''(t) + 2x(t) &= \underbrace{\sum_{n=1}^{\infty} -b_n n^2 \pi^2 \sin(n\pi t)}_{x''} + 2 \underbrace{\sum_{n=1}^{\infty} b_n \sin(n\pi t)}_x \\ &= \sum_{n=1}^{\infty} b_n (2 - n^2 \pi^2) \sin(n\pi t) \\ &= f(t) = \sum_{n=1}^{\infty} \frac{2(-1)^{n+1}}{n\pi} \sin(n\pi t). \end{aligned}$$

Therefore,

$$b_n (2 - n^2 \pi^2) = \frac{2(-1)^{n+1}}{n\pi}$$

or

$$b_n = \frac{2(-1)^{n+1}}{n\pi(2 - n^2 \pi^2)}.$$

That $2 - n^2\pi^2$ is not zero for any n , and that we can solve for b_n , is precisely because 2 is not an eigenvalue of the problem. We have thus obtained a Fourier series for the solution

$$x(t) = \sum_{n=1}^{\infty} \frac{2(-1)^{n+1}}{n\pi(2 - n^2\pi^2)} \sin(n\pi t).$$

See Figure 4.13 for a graph of the solution. Notice that because the eigenfunctions satisfy the boundary conditions, and x is written in terms of the eigenfunctions, then x satisfies the boundary conditions.

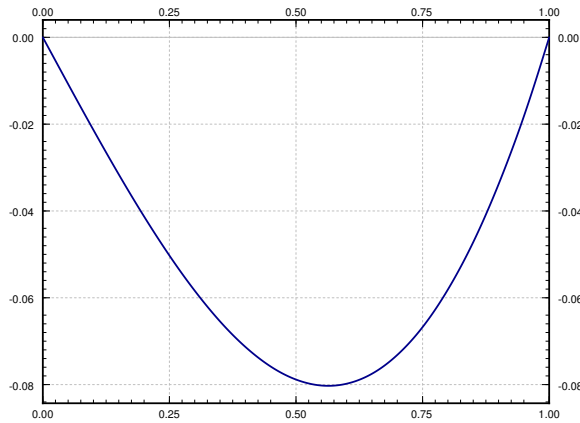


Figure 4.13: Plot of the solution of $x'' + 2x = t$, $x(0) = 0$, $x(1) = 0$.

Example 4.4.4: We handle the Neumann conditions with cosine series. Take the boundary value problem

$$x''(t) + 2x(t) = f(t), \quad 0 < t < 1,$$

where again $f(t) = t$ on $0 < t < 1$, but now satisfying the Neumann boundary conditions $x'(0) = 0$, $x'(1) = 0$. We write $f(t)$ as a cosine series

$$f(t) = \frac{c_0}{2} + \sum_{n=1}^{\infty} c_n \cos(n\pi t),$$

where

$$c_0 = 2 \int_0^1 t \, dt = 1,$$

and

$$c_n = 2 \int_0^1 t \cos(n\pi t) \, dt = \frac{2((-1)^n - 1)}{\pi^2 n^2} = \begin{cases} \frac{-4}{\pi^2 n^2} & \text{if } n \text{ odd,} \\ 0 & \text{if } n \text{ even.} \end{cases}$$

We write $x(t)$ as a cosine series

$$x(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\pi t).$$

We plug in to obtain

$$\begin{aligned}
 x''(t) + 2x(t) &= \left(\sum_{n=1}^{\infty} -a_n n^2 \pi^2 \cos(n\pi t) \right) + a_0 + 2 \left(\sum_{n=1}^{\infty} a_n \cos(n\pi t) \right) \\
 &= a_0 + \sum_{n=1}^{\infty} a_n (2 - n^2 \pi^2) \cos(n\pi t) \\
 &= f(t) = \frac{1}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{\pi^2 n^2} \cos(n\pi t).
 \end{aligned}$$

Therefore, $a_0 = \frac{1}{2}$, $a_n = 0$ for n even ($n \geq 2$), and for n odd, we have

$$a_n (2 - n^2 \pi^2) = \frac{-4}{\pi^2 n^2},$$

or

$$a_n = \frac{-4}{n^2 \pi^2 (2 - n^2 \pi^2)}.$$

The Fourier series for the solution $x(t)$ is

$$x(t) = \frac{1}{4} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{n^2 \pi^2 (2 - n^2 \pi^2)} \cos(n\pi t).$$

4.4.4 Exercises

Exercise 4.4.4: Take $f(t) = (t - 1)^2$ defined on $0 \leq t \leq 1$.

- a) Sketch the plot of the even periodic extension of f .
- b) Sketch the plot of the odd periodic extension of f .

Exercise 4.4.5: Find the Fourier series of both the odd and even periodic extension of the function $f(t) = (t - 1)^2$ for $0 \leq t \leq 1$. Can you tell which extension is continuous from the Fourier-series coefficients?

Exercise 4.4.6: Find the Fourier series of both the odd and even periodic extension of the function $f(t) = t$ for $0 \leq t \leq \pi$.

Exercise 4.4.7: Find the Fourier series of the even periodic extension of the function $f(t) = \sin t$ for $0 \leq t \leq \pi$.

Exercise 4.4.8: Consider

$$x''(t) + 4x(t) = f(t),$$

where $f(t) = 1$ on $0 < t < 1$.

- a) Solve for the Dirichlet conditions $x(0) = 0, x(1) = 0$.
- b) Solve for the Neumann conditions $x'(0) = 0, x'(1) = 0$.

Exercise 4.4.9: Consider

$$x''(t) + 9x(t) = f(t),$$

for $f(t) = \sin(2\pi t)$ on $0 < t < 1$.

- a) Solve for the Dirichlet conditions $x(0) = 0, x(1) = 0$.
- b) Solve for the Neumann conditions $x'(0) = 0, x'(1) = 0$.

Exercise 4.4.10: Consider

$$x''(t) + 3x(t) = f(t), \quad x(0) = 0, \quad x(1) = 0,$$

where $f(t) = \sum_{n=1}^{\infty} b_n \sin(n\pi t)$. Write the solution $x(t)$ as a Fourier series, where the coefficients are given in terms of b_n .

Exercise 4.4.11: Let $f(t) = t^2(2-t)$ for $0 \leq t \leq 2$. Let $F(t)$ be the odd periodic extension. Compute $F(1), F(2), F(3), F(-1), F(9/2), F(101), F(103)$. Note: Do **not** compute the sine series.

Exercise 4.4.101: Let $f(t) = t/3$ on $0 \leq t < 3$.

- a) Find the Fourier series of the even periodic extension.
- b) Find the Fourier series of the odd periodic extension.

Exercise 4.4.102: Let $f(t) = \cos(2t)$ on $0 \leq t < \pi$.

- a) Find the Fourier series of the even periodic extension.
- b) Find the Fourier series of the odd periodic extension.

Exercise 4.4.103: Let $f(t)$ be defined on $0 \leq t < 1$. Consider the average of the two extensions $g(t) = \frac{F_{\text{odd}}(t) + F_{\text{even}}(t)}{2}$.

- a) What is $g(t)$ if $0 \leq t < 1$ (Justify!)
- b) What is $g(t)$ if $-1 < t < 0$ (Justify!)

Exercise 4.4.104: Let $f(t) = \sum_{n=1}^{\infty} \frac{1}{n^2} \sin(nt)$. Solve $x'' - x = f(t)$ for the Dirichlet conditions $x(0) = 0$ and $x(\pi) = 0$.

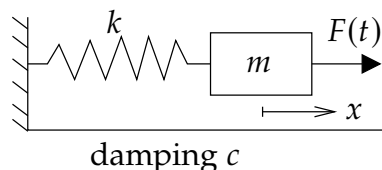
Exercise 4.4.105 (challenging): Let $f(t) = t + \sum_{n=1}^{\infty} \frac{1}{2^n} \sin(nt)$. Solve $x'' + \pi x = f(t)$ for the Dirichlet conditions $x(0) = 0$ and $x(\pi) = 1$. Hint: Note that $\frac{t}{\pi}$ satisfies the given Dirichlet conditions.

4.5 Applications of Fourier series

Note: 2 lectures, §9.4 in [EP], not in [BD]

4.5.1 Periodically forced oscillation

We return to forced oscillations. Consider a mass-spring system as before, where we have a mass m on a spring with spring constant k , with damping c , and a force $F(t)$ applied to the mass. Suppose the forcing function $F(t)$ is $2L$ -periodic for some $L > 0$. We saw this problem in chapter 2 with $F(t) = F_0 \cos(\omega t)$. The equation that governs this particular setup is



$$mx''(t) + cx'(t) + kx(t) = F(t). \quad (4.9)$$

The general solution of (4.9) consists of the complementary solution x_c , which solves the associated homogeneous equation $mx'' + cx' + kx = 0$, and a particular solution of (4.9) we call x_p . For $c > 0$, the complementary solution x_c will decay as time goes by. Therefore, we are mostly interested in a particular solution x_p that does not decay and is periodic with the same period as $F(t)$. We call this particular solution the *steady periodic solution* and we write it as x_{sp} as before. What is new in this section is that we consider an arbitrary forcing function $F(t)$ instead of a simple cosine.

For simplicity, suppose $c = 0$. The problem with $c > 0$ is very similar. The equation

$$mx'' + kx = 0$$

has the general solution

$$x(t) = A \cos(\omega_0 t) + B \sin(\omega_0 t),$$

where $\omega_0 = \sqrt{\frac{k}{m}}$. Any solution to $mx''(t) + kx(t) = F(t)$ is of the form $A \cos(\omega_0 t) + B \sin(\omega_0 t) + x_{sp}$. The steady periodic solution x_{sp} has the same period as $F(t)$.

In the spirit of the last section and the idea of undetermined coefficients, we first write

$$F(t) = \frac{c_0}{2} + \sum_{n=1}^{\infty} \left[c_n \cos\left(\frac{n\pi}{L}t\right) + d_n \sin\left(\frac{n\pi}{L}t\right) \right].$$

Then we write a proposed steady periodic solution x as

$$x(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right],$$

where a_n and b_n are unknowns. We plug x into the differential equation and solve for a_n and b_n in terms of c_n and d_n . This process is perhaps best understood by example.

Example 4.5.1: Suppose that $k = 2$ and $m = 1$. The units are again the mks units (meters-kilograms-seconds). There is a jetpack strapped to the mass, which fires with a force of 1 newton for 1 second and then is off for 1 second, and so on. We want to find the steady periodic solution.

The equation is, therefore,

$$x'' + 2x = F(t),$$

where $F(t)$ is the step function

$$F(t) = \begin{cases} 0 & \text{if } -1 < t < 0, \\ 1 & \text{if } 0 < t < 1, \end{cases}$$

extended periodically. We write

$$F(t) = \frac{c_0}{2} + \sum_{n=1}^{\infty} [c_n \cos(n\pi t) + d_n \sin(n\pi t)].$$

We compute

$$\begin{aligned} c_n &= \int_{-1}^1 F(t) \cos(n\pi t) dt = \int_0^1 \cos(n\pi t) dt = 0 \quad \text{for } n \geq 1, \\ c_0 &= \int_{-1}^1 F(t) dt = \int_0^1 dt = 1, \\ d_n &= \int_{-1}^1 F(t) \sin(n\pi t) dt \\ &= \int_0^1 \sin(n\pi t) dt \\ &= \left[\frac{-\cos(n\pi t)}{n\pi} \right]_{t=0}^1 \\ &= \frac{1 - (-1)^n}{\pi n} = \begin{cases} \frac{2}{\pi n} & \text{if } n \text{ odd,} \\ 0 & \text{if } n \text{ even.} \end{cases} \end{aligned}$$

So

$$F(t) = \frac{1}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{2}{\pi n} \sin(n\pi t).$$

We want to try

$$x(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} [a_n \cos(n\pi t) + b_n \sin(n\pi t)].$$

Once we plug x into the differential equation $x'' + 2x = F(t)$, it is clear that $a_n = 0$ for $n \geq 1$ as there are no corresponding terms in the series for $F(t)$. Similarly, $b_n = 0$ for even n .

Hence we try

$$x(t) = \frac{a_0}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} b_n \sin(n\pi t).$$

We plug into the differential equation and obtain

$$\begin{aligned} x'' + 2x &= \left(\sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} -b_n n^2 \pi^2 \sin(n\pi t) \right) + a_0 + 2 \left(\sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} b_n \sin(n\pi t) \right) \\ &= a_0 + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} b_n (2 - n^2 \pi^2) \sin(n\pi t) \\ &= F(t) = \frac{1}{2} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{2}{\pi n} \sin(n\pi t). \end{aligned}$$

So $a_0 = \frac{1}{2}$, $b_n = 0$ for even n , and for odd n , we get

$$b_n = \frac{2}{\pi n (2 - n^2 \pi^2)}.$$

The steady periodic solution has the Fourier series

$$x_{sp}(t) = \frac{1}{4} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{2}{\pi n (2 - n^2 \pi^2)} \sin(n\pi t).$$

We know this is the steady periodic solution as it contains no terms of the complementary solution and it is periodic with the same period as $F(t)$ itself. See [Figure 4.14](#) on the facing page for the plot of this solution.

4.5.2 Resonance

Just as when the forcing function was a simple cosine, we may encounter resonance. We assume $c = 0$ and so we discuss only pure resonance. Let $F(t)$ be $2L$ -periodic and consider

$$mx''(t) + kx(t) = F(t).$$

When we expand $F(t)$ and find that some of its terms coincide with the complementary solution to $mx'' + kx = 0$, we cannot use those terms in the guess. Just like before, they disappear when we plug them into the left-hand side, and we get a contradictory equation (such as $0 = 1$). That is, suppose

$$x_c = A \cos(\omega_0 t) + B \sin(\omega_0 t),$$

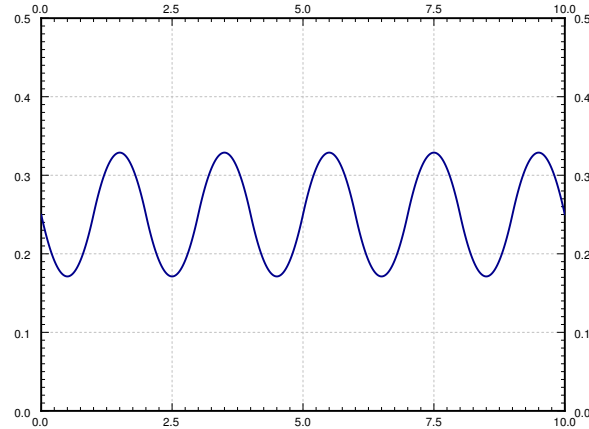


Figure 4.14: Plot of the steady periodic solution x_{sp} of Example 4.5.1.

where $\omega_0 = \frac{N\pi}{L}$ for some positive integer N . We have to modify our guess and try

$$x(t) = \frac{a_0}{2} + t \left(a_N \cos\left(\frac{N\pi}{L}t\right) + b_N \sin\left(\frac{N\pi}{L}t\right) \right) + \sum_{\substack{n=1 \\ n \neq N}}^{\infty} \left[a_n \cos\left(\frac{n\pi}{L}t\right) + b_n \sin\left(\frac{n\pi}{L}t\right) \right].$$

In other words, we multiply the offending term by t . From then on, we proceed as before.

The solution is not a Fourier series (it is not even periodic) since it contains these terms multiplied by t . The terms $t \left(a_N \cos\left(\frac{N\pi}{L}t\right) + b_N \sin\left(\frac{N\pi}{L}t\right) \right)$ eventually dominate and lead to wild oscillations. As before, this behavior is called *pure resonance* or just *resonance*.

Note that we may hit the resonance frequency with infinitely many terms (overtones) of the forcing function F . That is, suppose we use the same “shape” for F , and we change the base frequency (we change the L). Then different terms from the Fourier series of F may interfere with the complementary solution, that is, $\frac{n\pi}{L} = \omega_0$ for some n , and possibly cause resonance. Theoretically, infinitely many base frequencies could cause resonance. However, we should note that since everything is an approximation to any real life application, only the first few terms of F , and hence only a few such frequencies, would matter.

Example 4.5.2: We want to solve the equation

$$2x'' + 18\pi^2 x = F(t), \quad (4.10)$$

where

$$F(t) = \begin{cases} -1 & \text{if } -1 < t < 0, \\ 1 & \text{if } 0 < t < 1, \end{cases}$$

extended periodically. We note that

$$F(t) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{4}{\pi n} \sin(n\pi t).$$

Exercise 4.5.1: Compute the Fourier series of F to verify the equation above.

As $\sqrt{\frac{k}{m}} = \sqrt{\frac{18\pi^2}{2}} = 3\pi$, the solution to (4.10) is

$$x(t) = c_1 \cos(3\pi t) + c_2 \sin(3\pi t) + x_p(t)$$

for some particular solution x_p .

If we just try an x_p given as a Fourier series with $\sin(n\pi t)$ as usual, the complementary equation, $2x'' + 18\pi^2 x = 0$, eats our 3rd harmonic. That is, the term with $\sin(3\pi t)$ is already in our complementary solution. Therefore, we pull that term out and multiply it by t . We also add a cosine term to get everything right. That is, we try

$$x_p(t) = a_3 t \cos(3\pi t) + b_3 t \sin(3\pi t) + \sum_{\substack{n=1 \\ n \text{ odd} \\ n \neq 3}}^{\infty} b_n \sin(n\pi t).$$

We compute the second derivative.

$$\begin{aligned} x_p''(t) = & -6a_3\pi \sin(3\pi t) - 9\pi^2 a_3 t \cos(3\pi t) + 6b_3\pi \cos(3\pi t) - 9\pi^2 b_3 t \sin(3\pi t) \\ & + \sum_{\substack{n=1 \\ n \text{ odd} \\ n \neq 3}}^{\infty} (-n^2\pi^2 b_n) \sin(n\pi t). \end{aligned}$$

We now plug into the left-hand side of the differential equation.

$$\begin{aligned} 2x_p'' + 18\pi^2 x_p = & -12a_3\pi \sin(3\pi t) - 18\pi^2 a_3 t \cos(3\pi t) + 12b_3\pi \cos(3\pi t) - 18\pi^2 b_3 t \sin(3\pi t) \\ & + 18\pi^2 a_3 t \cos(3\pi t) + 18\pi^2 b_3 t \sin(3\pi t) \\ & + \sum_{\substack{n=1 \\ n \text{ odd} \\ n \neq 3}}^{\infty} (-2n^2\pi^2 b_n + 18\pi^2 b_n) \sin(n\pi t). \end{aligned}$$

We simplify,

$$2x_p'' + 18\pi^2 x_p = -12a_3\pi \sin(3\pi t) + 12b_3\pi \cos(3\pi t) + \sum_{\substack{n=1 \\ n \text{ odd} \\ n \neq 3}}^{\infty} (-2n^2\pi^2 b_n + 18\pi^2 b_n) \sin(n\pi t).$$

This series must equal the series for $F(t)$. We equate the coefficients and solve for a_3 and b_n :

$$\begin{aligned} a_3 &= \frac{4/(3\pi)}{-12\pi} = \frac{-1}{9\pi^2}, \\ b_3 &= 0, \\ b_n &= \frac{4}{n\pi(18\pi^2 - 2n^2\pi^2)} = \frac{2}{\pi^3 n(9 - n^2)} \quad \text{for } n \text{ odd and } n \neq 3. \end{aligned}$$

That is,

$$x_p(t) = \frac{-1}{9\pi^2} t \cos(3\pi t) + \sum_{\substack{n=1 \\ n \text{ odd} \\ n \neq 3}}^{\infty} \frac{2}{\pi^3 n(9 - n^2)} \sin(n\pi t).$$

When $c > 0$, you do not have to worry about pure resonance. There are never any conflicts, and you do not need to multiply any terms by t . There is a corresponding concept of practical resonance, and it is very similar to the ideas we already explored in [chapter 2](#). Basically, what happens in practical resonance is that one of the coefficients in the series for x_{sp} can get very big. Let us not go into details here.

4.5.3 Exercises

Exercise 4.5.2: Let $F(t) = \frac{1}{2} + \sum_{n=1}^{\infty} \frac{1}{n^2} \cos(n\pi t)$. Find the steady periodic solution to $x'' + 2x = F(t)$. Express your solution as a Fourier series.

Exercise 4.5.3: Let $F(t) = \sum_{n=1}^{\infty} \frac{1}{n^3} \sin(n\pi t)$. Find the steady periodic solution to $x'' + x' + x = F(t)$. Express your solution as a Fourier series.

Exercise 4.5.4: Let $F(t) = \sum_{n=1}^{\infty} \frac{1}{n^2} \cos(n\pi t)$. Find the steady periodic solution to $x'' + 4x = F(t)$. Express your solution as a Fourier series.

Exercise 4.5.5: Let $F(t) = t$ for $-1 < t < 1$ and extended periodically. Find the steady periodic solution to $x'' + x = F(t)$. Express your solution as a series.

Exercise 4.5.6: Let $F(t) = t$ for $-1 < t < 1$ and extended periodically. Find the steady periodic solution to $x'' + \pi^2 x = F(t)$. Express your solution as a series.

Exercise 4.5.101: Let $F(t) = \sin(2\pi t) + 0.1 \cos(10\pi t)$. Find the steady periodic solution to $x'' + \sqrt{2}x = F(t)$. Express your solution as a Fourier series.

Exercise 4.5.102: Let $F(t) = \sum_{n=1}^{\infty} e^{-n} \cos(2nt)$. Find the steady periodic solution to $x'' + 3x = F(t)$. Express your solution as a Fourier series.

Exercise 4.5.103: Let $F(t) = |t|$ for $-1 \leq t \leq 1$ extended periodically. Find the steady periodic solution to $x'' + \sqrt{3}x = F(t)$. Express your solution as a series.

Exercise 4.5.104: Let $F(t) = |t|$ for $-1 \leq t \leq 1$ extended periodically. Find the steady periodic solution to $x'' + \pi^2 x = F(t)$. Express your solution as a series.

4.6 PDEs, separation of variables, and the heat equation

Note: 2 lectures, §9.5 in [EP], §10.5 in [BD]

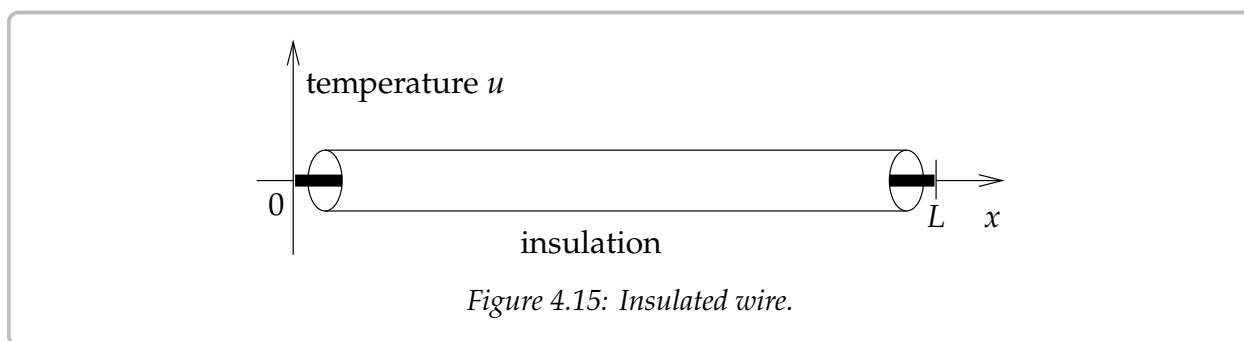
Recall that a *partial differential equation* or *PDE* is an equation containing the partial derivatives with respect to *several* independent variables. Solving PDEs will be our main application of Fourier series.

A PDE is said to be *linear* if the dependent variable and its derivatives appear at most to the first power and in no functions. We will only talk about linear PDEs. Together with a PDE, we usually specify some *boundary conditions*, where the value of the solution or its derivatives is given along the boundary of a region, and/or some *initial conditions* where the value of the solution or its derivatives is given for some initial time. Sometimes such conditions are mixed together and we will refer to them simply as *side conditions*.

We will study three specific partial differential equations, each one representing a general class of equations. First, we will study the *heat equation*, which is an example of a *parabolic PDE*. Next, we will study the *wave equation*, which is an example of a *hyperbolic PDE*. Finally, we will study the *Laplace equation*, which is an example of an *elliptic PDE*. Each of our examples will illustrate behavior that is typical for the whole class.

4.6.1 Heat on an insulated wire

We start with the heat equation. Consider a wire (or a thin metal rod) of length L insulated along its length except at the endpoints. Let x denote the position along the wire and let t denote time. See Figure 4.15.



Let $u(x, t)$ denote the temperature at point x at time t . The equation governing this setup is the so-called *one-dimensional heat equation*:

$$\frac{\partial u}{\partial t} = k \frac{\partial^2 u}{\partial x^2},$$

where $k > 0$ is a constant (the *thermal diffusivity* of the material). That is, the change in heat with respect to time at some point is proportional to the second derivative of the heat in

the x direction—along the wire. This makes sense; if at a fixed t the graph of the heat distribution has a maximum (the graph is concave down and the second x derivative is negative), then heat should flow away from the maximum and so the t derivative should also be negative. Similarly at a minimum, heat wants to flow in.

We generally use a more convenient notation for partial derivatives. We write u_t instead of $\frac{\partial u}{\partial t}$, and we write u_{xx} instead of $\frac{\partial^2 u}{\partial x^2}$. With this notation the heat equation becomes

$$u_t = ku_{xx}.$$

The region in which we will solve the heat equation is given by

$$0 < x < L \quad \text{and} \quad t > 0.$$

We must also have some side conditions on the boundaries of that region. We assume that the ends of the wire are either exposed and touching some body of constant heat, or the ends are insulated. If the ends of the wire are kept at temperature 0, then the conditions are

$$u(0, t) = 0 \quad \text{and} \quad u(L, t) = 0 \quad \text{for } t > 0.$$

If, on the other hand, the ends are insulated, the conditions are

$$u_x(0, t) = 0 \quad \text{and} \quad u_x(L, t) = 0 \quad \text{for } t > 0.$$

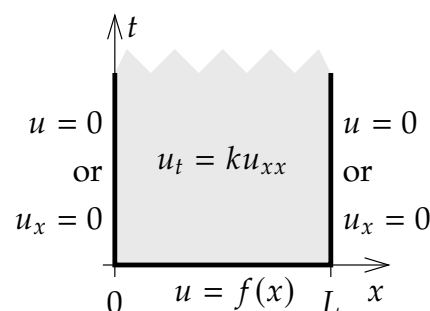
Let us see why that is so. If u_x is positive at some point x_0 , then at a particular time, u is smaller to the left of x_0 and higher to the right of x_0 . Heat is flowing from high temperature to low temperature, that is, to the left. On the other hand, if u_x is negative, then heat is again flowing from high temperature to low temperature, that is, to the right. So when u_x is zero, we are at a point where heat is not flowing in either direction. In other words, $u_x(0, t) = 0$ means no heat is flowing in or out of the wire at the point $x = 0$.

We have two conditions along the x -axis as there are two derivatives in the x direction. These side conditions are said to be *homogeneous* (i.e. u or a derivative of u is set to zero).

We also need an initial condition—the temperature distribution at time $t = 0$. That is,

$$u(x, 0) = f(x) \quad \text{for } 0 < x < L,$$

for some known function $f(x)$. This initial condition is not a homogeneous side condition.



4.6.2 Separation of variables

The heat equation is linear as u and its derivatives do not appear to any powers or in any functions, and it is homogeneous—there is no term independent of u . Thus the principle of superposition still applies for the heat equation (without side conditions): If u_1 and u_2 are solutions and c_1, c_2 are constants, then $u = c_1u_1 + c_2u_2$ is also a solution.

Exercise 4.6.1: Verify the principle of superposition for the heat equation.

Superposition preserves some side conditions. If u_1 and u_2 are solutions that satisfy $u(0, t) = 0$ and $u(L, t) = 0$, and c_1, c_2 are constants, then $u = c_1u_1 + c_2u_2$ is still a solution that satisfies $u(0, t) = 0$ and $u(L, t) = 0$. Similarly for the side conditions $u_x(0, t) = 0$ and $u_x(L, t) = 0$. In general, superposition preserves all homogeneous side conditions.

The method of *separation of variables* is to try to find solutions that are products of functions of one variable. For the heat equation, we try to find solutions of the form

$$u(x, t) = X(x)T(t).$$

That the desired particular solution we are looking for is of this form is too much to hope for. What is perfectly reasonable to ask, however, is to find enough “building-block” solutions of the form $u(x, t) = X(x)T(t)$ using this procedure so that the desired solution to the PDE is somehow constructed from these building blocks by the use of superposition.

Let us try to solve the heat equation problem

$$u_t = ku_{xx}, \quad \text{with } u(0, t) = 0, \quad u(L, t) = 0, \quad \text{and } u(x, 0) = f(x).$$

We guess $u(x, t) = X(x)T(t)$. We will try to make this guess satisfy the differential equation, $u_t = ku_{xx}$, and the homogeneous side conditions, $u(0, t) = 0$ and $u(L, t) = 0$. Then, as superposition preserves the differential equation and the homogeneous side conditions, we will try to build up a solution from these building blocks to solve the nonhomogeneous initial condition $u(x, 0) = f(x)$.

First, we plug $u(x, t) = X(x)T(t)$ into the heat equation to obtain

$$X(x)T'(t) = kX''(x)T(t).$$

We rewrite as

$$\frac{T'(t)}{kT(t)} = \frac{X''(x)}{X(x)}.$$

This equation must hold for all x and all t . But the left-hand side does not depend on x and the right-hand side does not depend on t . Hence, each side must be a constant. Let us call this constant $-\lambda$ (the minus sign is for convenience later). We obtain the two equations

$$\frac{T'(t)}{kT(t)} = -\lambda = \frac{X''(x)}{X(x)}.$$

In other words,

$$\begin{aligned} X''(x) + \lambda X(x) &= 0, \\ T'(t) + \lambda kT(t) &= 0. \end{aligned}$$

The boundary condition $u(0, t) = 0$ implies $X(0)T(t) = 0$. We are looking for a nontrivial solution, and so we can assume that $T(t)$ is not identically zero. Hence $X(0) = 0$. Similarly, $u(L, t) = 0$ implies $X(L) = 0$. We are looking for nontrivial solutions X of the eigenvalue

problem $X'' + \lambda X = 0$, $X(0) = 0$, $X(L) = 0$. We have previously found that the only eigenvalues are $\lambda_n = \frac{n^2\pi^2}{L^2}$, for integers $n \geq 1$, where eigenfunctions are $\sin\left(\frac{n\pi}{L}x\right)$. Hence, let us pick the solutions

$$X_n(x) = \sin\left(\frac{n\pi}{L}x\right).$$

The corresponding T_n must satisfy the equation

$$T_n'(t) + \frac{n^2\pi^2}{L^2}kT_n(t) = 0.$$

This is one of our [fundamental equations](#), and the solution is an exponential:

$$T_n(t) = e^{\frac{-n^2\pi^2}{L^2}kt},$$

where we picked the particular solution where conveniently $T_n(0) = 1$. Our building-block solutions are

$$u_n(x, t) = X_n(x)T_n(t) = \sin\left(\frac{n\pi}{L}x\right)e^{\frac{-n^2\pi^2}{L^2}kt}.$$

We note that $u_n(x, 0) = \sin\left(\frac{n\pi}{L}x\right)$. We write $f(x)$ as the sine series

$$f(x) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}x\right).$$

That is, we find the Fourier series of the odd periodic extension of $f(x)$. We used the sine series as it corresponds to the eigenvalue problem for $X(x)$ above. Finally, we use superposition to write the solution as

$$u(x, t) = \sum_{n=1}^{\infty} b_n u_n(x, t) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}x\right) e^{\frac{-n^2\pi^2}{L^2}kt}.$$

Why does this solution work? First note that it is a solution to the heat equation by superposition. It satisfies $u(0, t) = 0$ and $u(L, t) = 0$, because $x = 0$ or $x = L$ makes all the sines vanish. Finally, plugging in $t = 0$, we notice that $T_n(0) = 1$, and so

$$u(x, 0) = \sum_{n=1}^{\infty} b_n u_n(x, 0) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}x\right) = f(x).$$

Example 4.6.1: Consider an insulated wire of length 1 whose ends are embedded in ice (temperature 0). Let $k = 0.003$ and let the initial heat distribution be $u(x, 0) = 50x(1 - x)$. See [Figure 4.16](#) on the following page. Suppose we want to find the temperature function $u(x, t)$. In particular, suppose we want to find when (at what time t) the maximum temperature in the wire drops to one half of the initial maximum of 12.5.

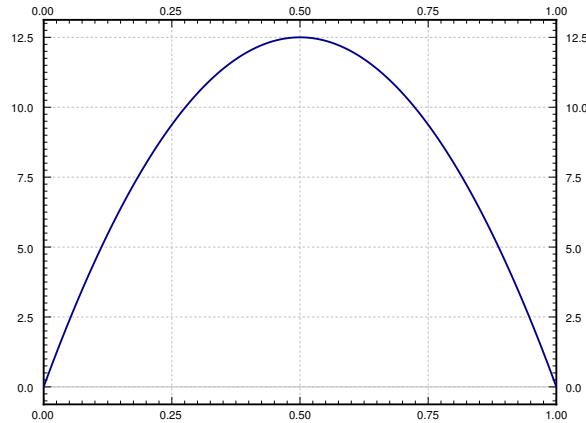


Figure 4.16: Initial distribution of temperature in the wire.

We are solving the following PDE problem:

$$\begin{aligned} u_t &= 0.003 u_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ u(0, t) &= u(1, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 50 x (1 - x) && \text{for } 0 < x < 1. \end{aligned}$$

Write $f(x) = 50 x (1 - x)$ for $0 < x < 1$ as a sine series $f(x) = \sum_{n=1}^{\infty} b_n \sin(n\pi x)$, where

$$b_n = 2 \int_0^1 50 x (1 - x) \sin(n\pi x) dx = \frac{200}{\pi^3 n^3} - \frac{200(-1)^n}{\pi^3 n^3} = \begin{cases} 0 & \text{if } n \text{ even,} \\ \frac{400}{\pi^3 n^3} & \text{if } n \text{ odd.} \end{cases}$$

We plug these coefficients into the series for $u(x, t)$ to obtain the solution

$$u(x, t) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{400}{\pi^3 n^3} \sin(n\pi x) e^{-n^2 \pi^2 0.003 t}.$$

We plot the solution in [Figure 4.17](#) on the next page for $0 \leq t \leq 100$.

Finally, we answer the question about the maximum temperature. It is relatively easy to see that the maximum temperature at any fixed time is always at $x = 0.5$, in the middle of the wire. The plot of $u(x, t)$ confirms this intuition. If we plug in $x = 0.5$, we get

$$u(0.5, t) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{400}{\pi^3 n^3} \sin(n\pi 0.5) e^{-n^2 \pi^2 0.003 t}.$$

For $n = 3$ and higher (remember n is only odd), the terms of the series are insignificant compared to the first term. The first term in the series is already a very good approximation of the function. Hence

$$u(0.5, t) \approx \frac{400}{\pi^3} e^{-\pi^2 0.003 t}.$$

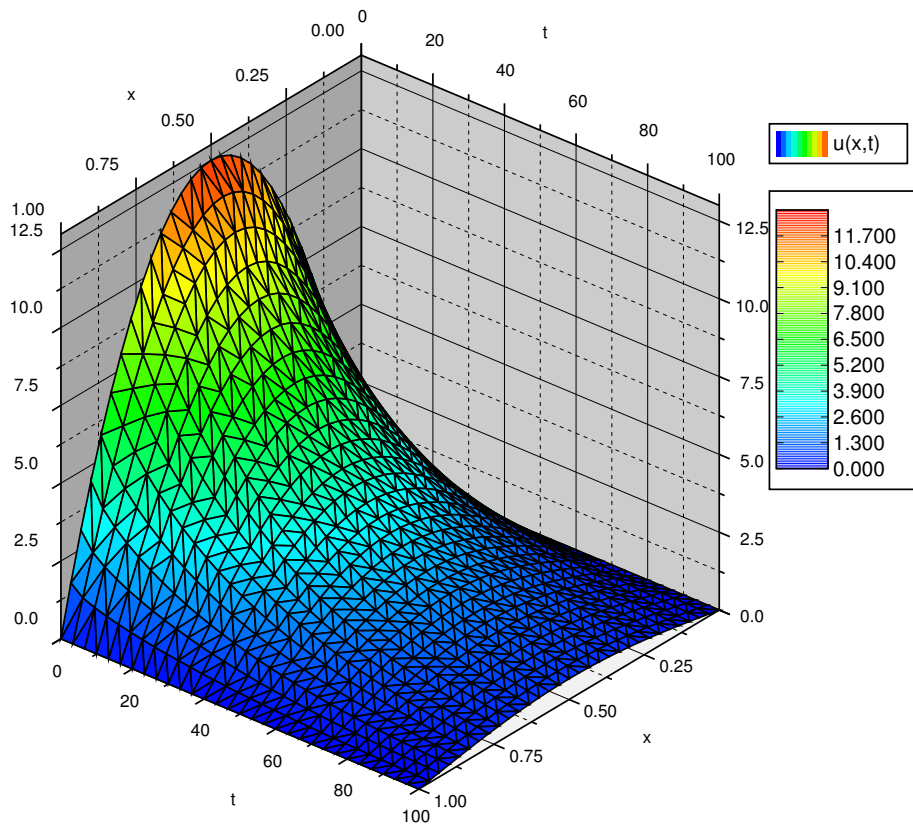


Figure 4.17: Plot of the temperature $u(x, t)$ of the wire at position x at time t for $0 \leq t \leq 100$. Notice the side conditions $u(0, t) = u(1, t) = 0$ and how the exponential makes the temperature decay with time.

The approximation gets better and better as t gets larger, because the other terms decay much faster. We plot the function $u(0.5, t)$, the temperature at the midpoint of the wire at time t , and its approximation by the first term in Figure 4.18 on the following page.

After $t = 5$ or so, it would be hard to tell the difference between the first term of the series for $u(x, t)$ and the real solution $u(x, t)$. This behavior is a general feature of solving the heat equation. If you are interested in behavior for large enough t , only the first one or two terms may be necessary.

We get back to the question of when the maximum temperature is one half of the initial maximum temperature. That is, when the temperature at the midpoint is $12.5/2 = 6.25$. The graph suggests that the approximation by the first term will be close enough. We solve

$$6.25 = \frac{400}{\pi^3} e^{-\pi^2 0.003 t}.$$

That is,

$$t = \frac{\ln \frac{6.25 \pi^3}{400}}{-\pi^2 0.003} \approx 24.5.$$

So the maximum temperature drops to half at about $t = 24.5$.

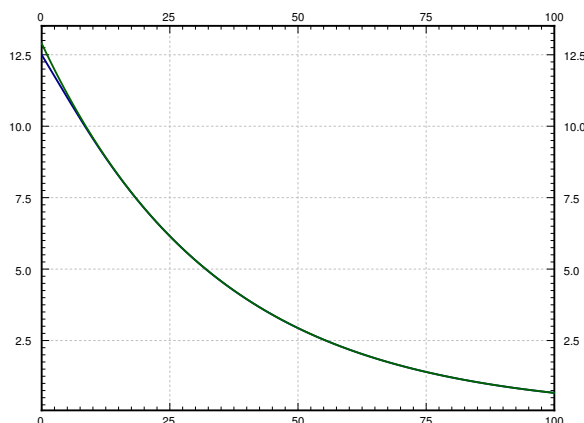


Figure 4.18: Temperature at the midpoint of the wire (the bottom curve), and the approximation of this temperature by using only the first term in the series (top curve).

We mention an interesting behavior of the solution to the heat equation. The heat equation “smooths” out the function $f(x)$ as t grows. For a fixed t , the solution is a Fourier series with coefficients $b_n e^{-\frac{n^2 \pi^2}{L^2} kt}$. If $t > 0$, then these coefficients go to zero faster than any $\frac{1}{n^p}$ for any power p . In other words, the Fourier series has infinitely many derivatives everywhere. Thus even if the function $f(x)$ has jumps and corners, then for a fixed $t > 0$, the solution $u(x, t)$ as a function of x is as smooth as we want it to be.

Example 4.6.2: When the initial condition is already a sine series, then there is no need to compute anything, you just need to plug in. Consider

$$u_t = 0.3 u_{xx}, \quad u(0, t) = u(1, t) = 0, \quad u(x, 0) = 0.1 \sin(\pi x) + \sin(2\pi x).$$

The solution is then

$$u(x, t) = 0.1 \sin(\pi x) e^{-0.3\pi^2 t} + \sin(2\pi x) e^{-1.2\pi^2 t}.$$

4.6.3 Insulated ends

Now suppose the ends of the wire are insulated. In this case, we are solving the problem

$$u_t = k u_{xx} \quad \text{with} \quad u_x(0, t) = 0, \quad u_x(L, t) = 0, \quad \text{and} \quad u(x, 0) = f(x).$$

Yet again we try a solution of the form $u(x, t) = X(x)T(t)$. By the same procedure as before, we plug into the heat equation and arrive at the following two equations:

$$\begin{aligned} X''(x) + \lambda X(x) &= 0, \\ T'(t) + \lambda k T(t) &= 0. \end{aligned}$$

At this point, the story changes slightly. The boundary condition $u_x(0, t) = 0$ implies $X'(0)T(t) = 0$. Hence $X'(0) = 0$. Similarly, $u_x(L, t) = 0$ implies $X'(L) = 0$. We want nontrivial solutions X of the eigenvalue problem $X'' + \lambda X = 0$, $X'(0) = 0$, $X'(L) = 0$. We previously found that the only eigenvalues are $\lambda_n = \frac{n^2\pi^2}{L^2}$, for integers $n \geq 0$, where eigenfunctions are $\cos\left(\frac{n\pi}{L}x\right)$ (we include the constant eigenfunction). We pick the solutions

$$X_n(x) = \cos\left(\frac{n\pi}{L}x\right) \quad \text{and} \quad X_0(x) = 1.$$

The corresponding T_n must satisfy the equation

$$T_n'(t) + \frac{n^2\pi^2}{L^2}kT_n(t) = 0.$$

For $n \geq 1$, as before,

$$T_n(t) = e^{\frac{-n^2\pi^2}{L^2}kt}.$$

For $n = 0$, we have $T_0'(t) = 0$ and hence $T_0(t) = 1$. Our building-block solutions are

$$u_n(x, t) = X_n(x)T_n(t) = \cos\left(\frac{n\pi}{L}x\right) e^{\frac{-n^2\pi^2}{L^2}kt}$$

and

$$u_0(x, t) = 1.$$

We note that $u_n(x, 0) = \cos\left(\frac{n\pi}{L}x\right)$. We write f using the cosine series

$$f(x) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos\left(\frac{n\pi}{L}x\right).$$

That is, we find the Fourier series of the even periodic extension of $f(x)$.

We use superposition to write the solution as

$$u(x, t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n u_n(x, t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos\left(\frac{n\pi}{L}x\right) e^{\frac{-n^2\pi^2}{L^2}kt}.$$

Example 4.6.3: Try the same equation as before, but with insulated ends. We are solving the following PDE problem

$$\begin{aligned} u_t &= 0.003 u_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ u_x(0, t) &= u_x(1, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 50x(1-x) && \text{for } 0 < x < 1. \end{aligned}$$

For this problem, we must find the cosine series of $u(x, 0)$. For $0 < x < 1$, we have

$$50x(1-x) = \frac{25}{3} + \sum_{\substack{n=2 \\ n \text{ even}}}^{\infty} \left(\frac{-200}{\pi^2 n^2}\right) \cos(n\pi x).$$

The calculation is left to the reader. Hence, the solution to the PDE problem, plotted in Figure 4.19, is given by the series

$$u(x, t) = \frac{25}{3} + \sum_{\substack{n=2 \\ n \text{ even}}}^{\infty} \left(\frac{-200}{\pi^2 n^2} \right) \cos(n\pi x) e^{-n^2 \pi^2 0.003 t}.$$

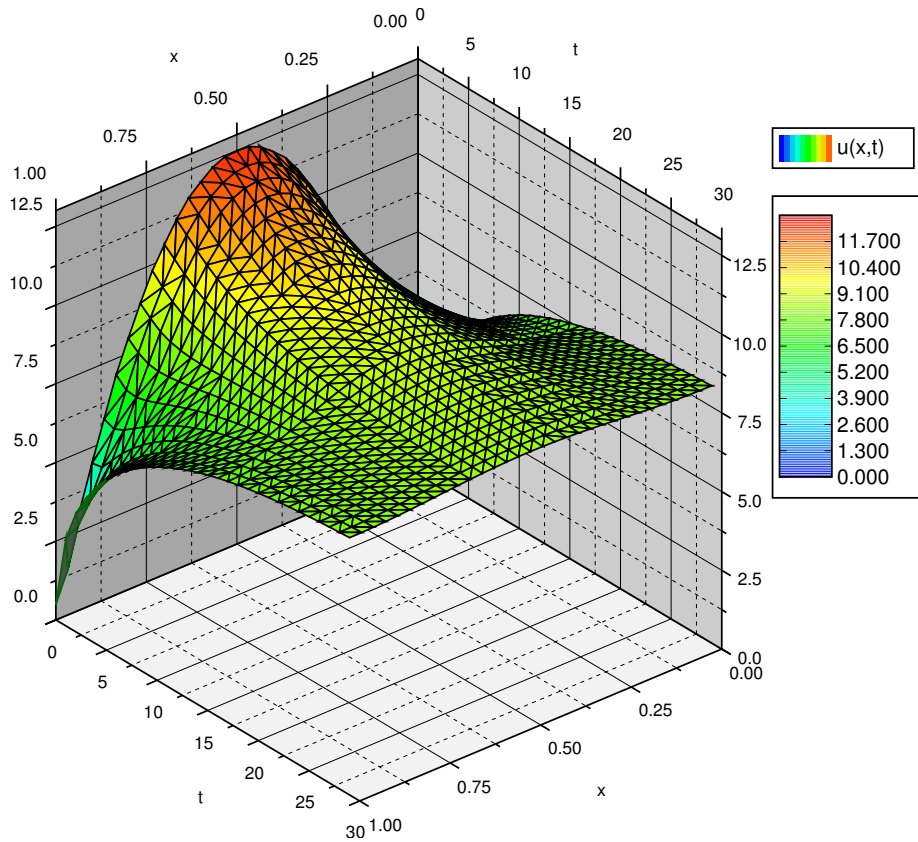


Figure 4.19: Plot of the temperature of the insulated wire at position x at time t .

Note in the graph that as time goes on, the temperature evens out across the wire. Eventually, all the terms except the constant die out, and you will be left with a uniform temperature of $\frac{25}{3} \approx 8.33$ along the entire length of the wire.

Let us expand on the last point. The constant term in the series is

$$\frac{a_0}{2} = \frac{1}{L} \int_0^L f(x) dx.$$

In other words, $\frac{a_0}{2}$ is the average value of $f(x)$, that is, the average of the initial temperature. As the wire is insulated everywhere, no heat can get out, no heat can get in. So the temperature tries to distribute evenly over time, and the average temperature must always be the same, in particular, it is always $\frac{a_0}{2}$. As time goes to infinity, the temperature goes to the constant $\frac{a_0}{2}$ everywhere.

4.6.4 Exercises

Exercise 4.6.2: Consider a wire of length 2, with $k = 0.001$ and an initial temperature distribution $u(x, 0) = 50x$. Both ends are embedded in ice (temperature 0). Find the solution as a series.

Exercise 4.6.3: Find a series solution of

$$\begin{aligned} u_t &= u_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ u(0, t) &= u(1, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 100 && \text{for } 0 < x < 1. \end{aligned}$$

Exercise 4.6.4: Find a series solution of

$$\begin{aligned} u_t &= u_{xx} && \text{for } 0 < x < \pi \text{ and } t > 0, \\ u_x(0, t) &= u_x(\pi, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 3 \cos(x) + \cos(3x) && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.6.5: Find a series solution of

$$\begin{aligned} u_t &= \frac{1}{3} u_{xx} && \text{for } 0 < x < \pi \text{ and } t > 0, \\ u_x(0, t) &= u_x(\pi, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= \frac{10x}{\pi} && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.6.6: Find a series solution of

$$\begin{aligned} u_t &= u_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ u(0, t) &= 0, \quad u(1, t) = 100 && \text{for } t > 0, \\ u(x, 0) &= \sin(\pi x) && \text{for } 0 < x < 1. \end{aligned}$$

Hint: Use the fact that $u(x, t) = 100x$ is a solution satisfying $u_t = u_{xx}$, $u(0, t) = 0$, $u(1, t) = 100$. Then use superposition.

Exercise 4.6.7: Find the steady-state temperature solution as a function of x alone, by letting $t \rightarrow \infty$ in the solution from exercises 4.6.5 and 4.6.6. Verify that it satisfies the equation $u_{xx} = 0$.

Exercise 4.6.8: Use separation of variables to find a nontrivial solution to $u_{xx} + u_{yy} = 0$, where $u(x, 0) = 0$ and $u(0, y) = 0$. Hint: Try $u(x, y) = X(x)Y(y)$.

Exercise 4.6.9 (challenging): Suppose that one end of the wire is insulated (say at $x = 0$) and the other end is kept at zero temperature. That is, find a series solution of

$$\begin{aligned} u_t &= ku_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\ u_x(0, t) &= u(L, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= f(x) && \text{for } 0 < x < L. \end{aligned}$$

Express any coefficients in the series by integrals of $f(x)$.

Exercise 4.6.10 (challenging): Suppose that the wire is circular and insulated, so there are no ends. You can think of this as simply connecting the two ends and making sure the solution matches up at the ends. That is, find a series solution of

$$\begin{aligned} u_t &= ku_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\ u(0, t) &= u(L, t), \quad u_x(0, t) = u_x(L, t) && \text{for } t > 0, \\ u(x, 0) &= f(x) && \text{for } 0 < x < L. \end{aligned}$$

Express any coefficients in the series by integrals of $f(x)$.

Exercise 4.6.11: Consider a wire insulated on both ends, $L = 1$, $k = 1$, and $u(x, 0) = \cos^2(\pi x)$.

- Find the solution $u(x, t)$. Hint: a trig identity.
- Find the average temperature.
- Initially the temperature variation is 1 (maximum minus the minimum). Find the time when the variation is $1/2$.

Exercise 4.6.101: Find a series solution of

$$\begin{aligned} u_t &= 3u_{xx} && \text{for } 0 < x < \pi \text{ and } t > 0, \\ u(0, t) &= u(\pi, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 5 \sin(x) + 2 \sin(5x) && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.6.102: Find a series solution of

$$\begin{aligned} u_t &= 0.1u_{xx} && \text{for } 0 < x < \pi \text{ and } t > 0, \\ u_x(0, t) &= u_x(\pi, t) = 0 && \text{for } t > 0, \\ u(x, 0) &= 1 + 2 \cos(x) && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.6.103: Use separation of variables to find a nontrivial solution to $u_{xt} = u_{xx}$.

Exercise 4.6.104: Use a variation on separation of variables to find a nontrivial solution to $u_x + u_t = u$. Hint: Try $u(x, t) = X(x) + T(t)$.

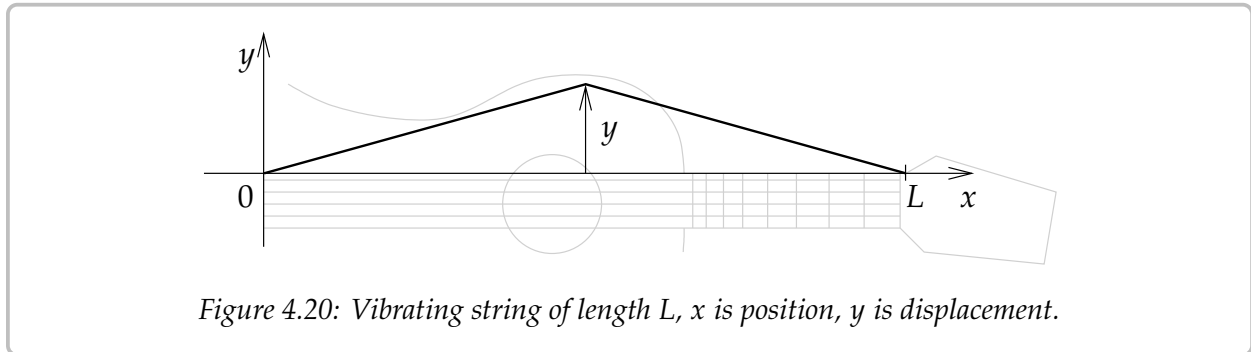
Exercise 4.6.105: Suppose that the temperature on the wire is fixed at 0 at the ends, $L = 1$, $k = 1$, and $u(x, 0) = 100 \sin(2\pi x)$.

- What is the temperature at $x = 1/2$ at any time?
- What is the maximum and the minimum temperature on the wire at $t = 0$?
- At what time is the maximum temperature on the wire exactly one half of the initial maximum at $t = 0$?

4.7 One-dimensional wave equation

Note: 1 lecture, §9.6 in [EP], §10.7 in [BD]

Imagine a tensioned guitar string of length L that can vibrate. We will only consider vibrations in one direction. Let x denote the position along the string, let t denote time, and let $y(x, t)$ denote the displacement of the string from the rest position. See Figure 4.20.



The equation that governs this setup is the so-called *one-dimensional wave equation*:

$$y_{tt} = a^2 y_{xx},$$

for some constant $a > 0$. The intuition is similar to the heat equation, replacing velocity with acceleration: The acceleration at a specific point is proportional to the second derivative of the shape of the string. In other words, when the string is concave down then y_{xx} is negative and the string wants to accelerate downwards, so y_{tt} should be negative. And vice versa. The wave equation is an example of a hyperbolic PDE.

We will again solve for y in the region $0 < x < L$ and $t > 0$. Assume that the ends of the string are fixed in place as on the guitar:

$$y(0, t) = 0 \quad \text{and} \quad y(L, t) = 0 \quad \text{for } t > 0.$$

We have two conditions along the x -axis as there are two derivatives in the x direction. There are also two derivatives along the t direction and hence we need two further conditions here. We need to know the initial position and the initial velocity of the string. That is, for some known functions $f(x)$ and $g(x)$, we impose

$$y(x, 0) = f(x) \quad \text{and} \quad y_t(x, 0) = g(x) \quad \text{for } 0 < x < L.$$

The equation is linear and homogeneous, so superposition works just as it did for the heat equation. Superposition also preserves the homogeneous side conditions $y(0, t) = 0$ and $y(L, t) = 0$. Again, we will use separation of variables to find enough building-block solutions to get the particular solution also solving the nonhomogeneous initial conditions. There is one change. We will solve two separate problems and add their solutions.

The two problems we will solve are

$$\begin{aligned}
 w_{tt} &= a^2 w_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\
 w(0, t) &= w(L, t) = 0 && \text{for } t > 0, \\
 w(x, 0) &= 0 && \text{for } 0 < x < L, \\
 w_t(x, 0) &= g(x) && \text{for } 0 < x < L,
 \end{aligned} \tag{4.11}$$

and

$$\begin{aligned}
 z_{tt} &= a^2 z_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\
 z(0, t) &= z(L, t) = 0 && \text{for } t > 0, \\
 z(x, 0) &= f(x) && \text{for } 0 < x < L, \\
 z_t(x, 0) &= 0 && \text{for } 0 < x < L.
 \end{aligned} \tag{4.12}$$

The principle of superposition implies that $y = w + z$ solves the wave equation and the homogeneous side conditions. Furthermore, $y(x, 0) = w(x, 0) + z(x, 0) = f(x)$ and $y_t(x, 0) = w_t(x, 0) + z_t(x, 0) = g(x)$. Hence, y is a solution to

$$\begin{aligned}
 y_{tt} &= a^2 y_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\
 y(0, t) &= y(L, t) = 0 && \text{for } t > 0, \\
 y(x, 0) &= f(x) && \text{for } 0 < x < L, \\
 y_t(x, 0) &= g(x) && \text{for } 0 < x < L.
 \end{aligned} \tag{4.13}$$

The reason for all this complexity is that superposition only works for homogeneous conditions such as $y(0, t) = y(L, t) = 0$, $y(x, 0) = 0$, or $y_t(x, 0) = 0$. Therefore, we can use separation of variables to find many building-block solutions solving all the homogeneous conditions. We can then use them to construct a solution satisfying the remaining nonhomogeneous condition.

Let us start with (4.11). We try a solution of the form $w(x, t) = X(x)T(t)$ again. We plug into the wave equation to obtain

$$X(x)T''(t) = a^2 X''(x)T(t).$$

Rewriting, we get

$$\frac{T''(t)}{a^2 T(t)} = \frac{X''(x)}{X(x)}.$$

Again, the left-hand side depends only on t and the right-hand side depends only on x . So both sides equal a constant, which we denote by $-\lambda$:

$$\frac{T''(t)}{a^2 T(t)} = -\lambda = \frac{X''(x)}{X(x)}.$$

We solve to get two ordinary differential equations

$$\begin{aligned}
 X''(x) + \lambda X(x) &= 0, \\
 T''(t) + \lambda a^2 T(t) &= 0.
 \end{aligned}$$

The condition $0 = w(0, t) = X(0)T(t)$ implies $X(0) = 0$ and $w(L, t) = 0$ implies that $X(L) = 0$. Therefore, the only nontrivial solutions for the first equation are when $\lambda = \lambda_n = \frac{n^2\pi^2}{L^2}$ and they are

$$X_n(x) = \sin\left(\frac{n\pi}{L}x\right).$$

The general solution for T for this particular λ_n is

$$T_n(t) = A \cos\left(\frac{n\pi a}{L}t\right) + B \sin\left(\frac{n\pi a}{L}t\right).$$

We also have the condition that $w(x, 0) = 0$ or $X(x)T(0) = 0$. This implies that $T(0) = 0$, which in turn forces $A = 0$. It is convenient to pick $B = \frac{L}{n\pi a}$ (you will see why in a moment) and hence

$$T_n(t) = \frac{L}{n\pi a} \sin\left(\frac{n\pi a}{L}t\right).$$

Our building-block solutions are

$$w_n(x, t) = \frac{L}{n\pi a} \sin\left(\frac{n\pi}{L}x\right) \sin\left(\frac{n\pi a}{L}t\right).$$

We differentiate in t :

$$\frac{\partial w_n}{\partial t}(x, t) = \sin\left(\frac{n\pi}{L}x\right) \cos\left(\frac{n\pi a}{L}t\right).$$

Hence,

$$\frac{\partial w_n}{\partial t}(x, 0) = \sin\left(\frac{n\pi}{L}x\right).$$

We expand $g(x)$ in terms of these sines as

$$g(x) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}x\right).$$

Using superposition we write the solution to (4.11) as a series

$$w(x, t) = \sum_{n=1}^{\infty} b_n w_n(x, t) = \sum_{n=1}^{\infty} b_n \frac{L}{n\pi a} \sin\left(\frac{n\pi}{L}x\right) \sin\left(\frac{n\pi a}{L}t\right).$$

Exercise 4.7.1: Check that $w(x, 0) = 0$ and $w_t(x, 0) = g(x)$.

We solve (4.12) similarly. We again try $z(x, t) = X(x)T(t)$. The procedure works exactly the same at first. We obtain

$$X''(x) + \lambda X(x) = 0,$$

$$T''(t) + \lambda a^2 T(t) = 0,$$

and the conditions $X(0) = 0, X(L) = 0$. Again, $\lambda = \lambda_n = \frac{n^2\pi^2}{L^2}$ and

$$X_n(x) = \sin\left(\frac{n\pi}{L}x\right).$$

This time, the condition on T is $T'(0) = 0$. Thus we get that $B = 0$, and we take

$$T_n(t) = \cos\left(\frac{n\pi a}{L}t\right).$$

Our building-block solution is

$$z_n(x, t) = \sin\left(\frac{n\pi}{L}x\right) \cos\left(\frac{n\pi a}{L}t\right).$$

As $z_n(x, 0) = \sin\left(\frac{n\pi}{L}x\right)$, we expand $f(x)$ in terms of these sines as

$$f(x) = \sum_{n=1}^{\infty} c_n \sin\left(\frac{n\pi}{L}x\right).$$

We write down the solution to (4.12) as a series

$$z(x, t) = \sum_{n=1}^{\infty} c_n z_n(x, t) = \sum_{n=1}^{\infty} c_n \sin\left(\frac{n\pi}{L}x\right) \cos\left(\frac{n\pi a}{L}t\right).$$

Exercise 4.7.2: Fill in the details in the derivation of the solution of (4.12). Check that the solution satisfies all the side conditions.

Putting these two solutions together, let us state the result as a theorem.

Theorem 4.7.1. Take the problem

$$\begin{aligned} y_{tt} &= a^2 y_{xx} && \text{for } 0 < x < L \text{ and } t > 0, \\ y(0, t) &= y(L, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= f(x) && \text{for } 0 < x < L, \\ y_t(x, 0) &= g(x) && \text{for } 0 < x < L, \end{aligned} \tag{4.14}$$

where

$$f(x) = \sum_{n=1}^{\infty} c_n \sin\left(\frac{n\pi}{L}x\right) \quad \text{and} \quad g(x) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{L}x\right).$$

Then the solution $y(x, t)$ can be written as a sum of the solutions of (4.11) and (4.12):

$$\begin{aligned} y(x, t) &= \sum_{n=1}^{\infty} \left[b_n \frac{L}{n\pi a} \sin\left(\frac{n\pi}{L}x\right) \sin\left(\frac{n\pi a}{L}t\right) + c_n \sin\left(\frac{n\pi}{L}x\right) \cos\left(\frac{n\pi a}{L}t\right) \right] \\ &= \sum_{n=1}^{\infty} \sin\left(\frac{n\pi}{L}x\right) \left[b_n \frac{L}{n\pi a} \sin\left(\frac{n\pi a}{L}t\right) + c_n \cos\left(\frac{n\pi a}{L}t\right) \right]. \end{aligned}$$

Example 4.7.1: Consider a string of length 2 plucked in the middle. It has an initial shape given in Figure 4.21 on the next page. That is,

$$f(x) = \begin{cases} 0.1x & \text{if } 0 \leq x \leq 1, \\ 0.1(2-x) & \text{if } 1 < x \leq 2. \end{cases}$$

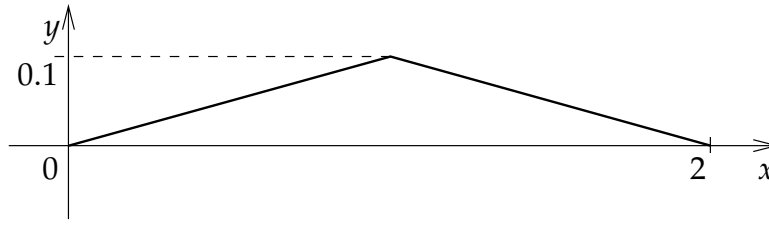


Figure 4.21: Initial shape of a plucked string from [Example 4.7.1](#).

Let the string start at rest ($g(x) = 0$), and let $a = 1$ for simplicity. In other words, we wish to solve the problem:

$$\begin{aligned} y_{tt} &= y_{xx} && \text{for } 0 < x < 2 \text{ and } t > 0, \\ y(0, t) &= y(2, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= f(x) && \text{for } 0 < x < 2, \\ y_t(x, 0) &= 0 && \text{for } 0 < x < 2. \end{aligned}$$

We leave the details of computing the sine series of $f(x)$ to the reader. The series is

$$f(x) = \sum_{n=1}^{\infty} \frac{0.8}{n^2\pi^2} \sin\left(\frac{n\pi}{2}\right) \sin\left(\frac{n\pi}{2}x\right).$$

Note that $\sin\left(\frac{n\pi}{2}\right)$ is the sequence $1, 0, -1, 0, 1, 0, -1, \dots$ for $n = 1, 2, 3, 4, \dots$. Therefore,

$$f(x) = \frac{0.8}{\pi^2} \sin\left(\frac{\pi}{2}x\right) - \frac{0.8}{9\pi^2} \sin\left(\frac{3\pi}{2}x\right) + \frac{0.8}{25\pi^2} \sin\left(\frac{5\pi}{2}x\right) - \dots$$

The solution $y(x, t)$ is given by

$$\begin{aligned} y(x, t) &= \sum_{n=1}^{\infty} \frac{0.8}{n^2\pi^2} \sin\left(\frac{n\pi}{2}\right) \sin\left(\frac{n\pi}{2}x\right) \cos\left(\frac{n\pi}{2}t\right) \\ &= \sum_{m=1}^{\infty} \frac{0.8(-1)^{m+1}}{(2m-1)^2\pi^2} \sin\left(\frac{(2m-1)\pi}{2}x\right) \cos\left(\frac{(2m-1)\pi}{2}t\right) \\ &= \frac{0.8}{\pi^2} \sin\left(\frac{\pi}{2}x\right) \cos\left(\frac{\pi}{2}t\right) - \frac{0.8}{9\pi^2} \sin\left(\frac{3\pi}{2}x\right) \cos\left(\frac{3\pi}{2}t\right) \\ &\quad + \frac{0.8}{25\pi^2} \sin\left(\frac{5\pi}{2}x\right) \cos\left(\frac{5\pi}{2}t\right) - \dots \end{aligned}$$

See [Figure 4.22](#) on the following page for a plot for $0 < t < 3$. Notice that unlike the heat equation, the solution does not become “smoother,” the “sharp edges” remain. We will see the reason for this behavior in the next section where we derive the solution to the wave equation in a different way.

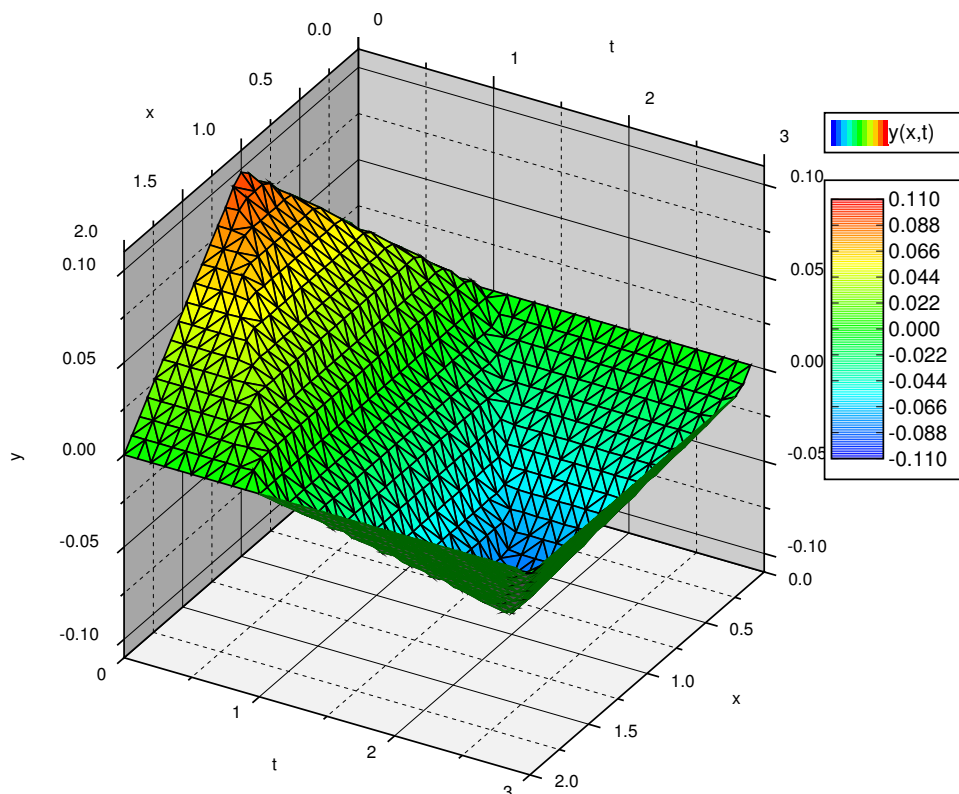


Figure 4.22: Shape of the plucked string for $0 < t < 3$.

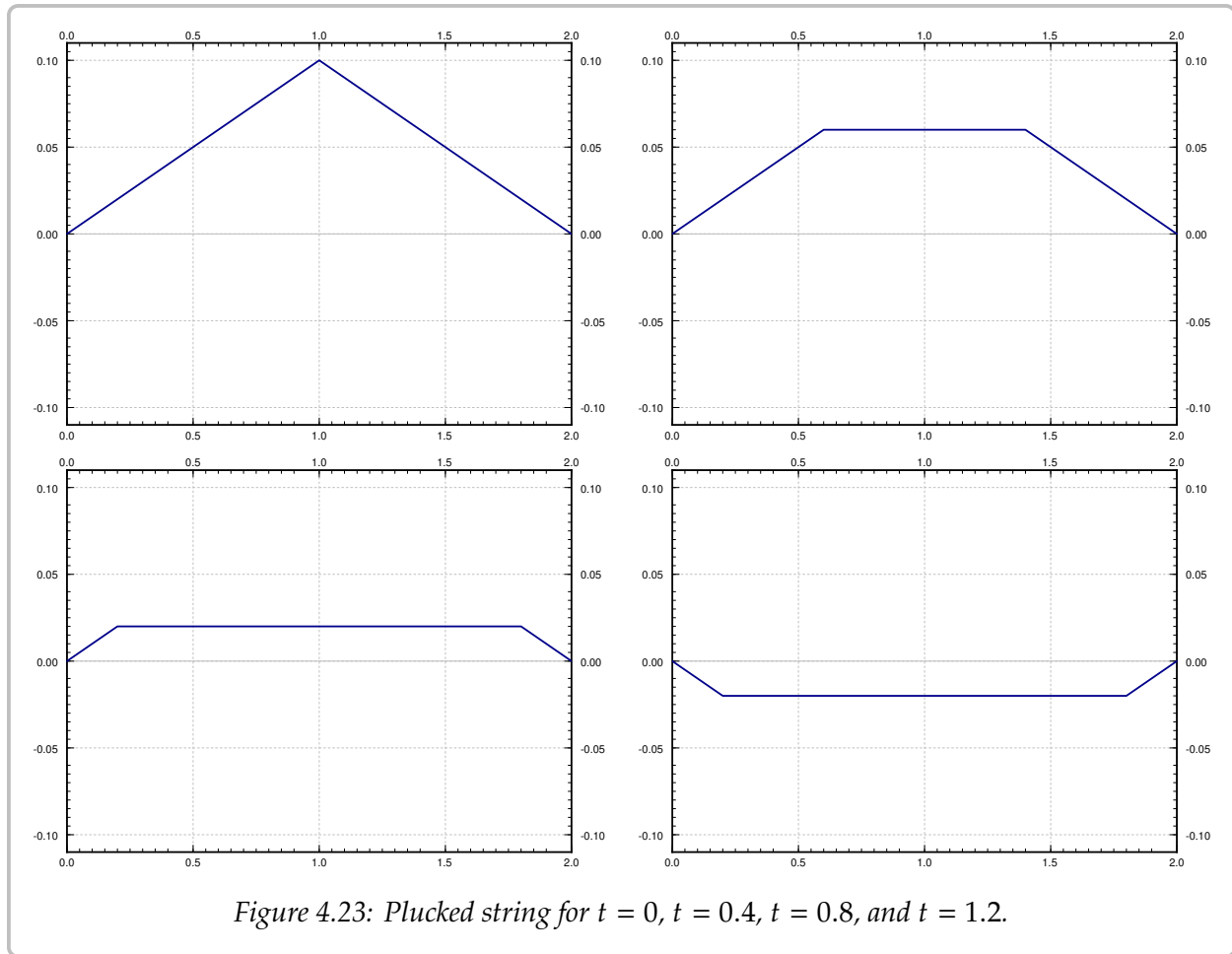
Make sure you understand what the plot, such as the one in the figure, is telling you. For each fixed t , you can think of the function $y(x, t)$ as just a function of x . This function gives you the shape of the string at time t . See Figure 4.23 on the next page for plots of y as a function of x at several different values of t . On these plots, you can see the sharp edges remaining much better.

One thing to take away from all this is how a guitar sounds. Notice that the (angular) frequencies that come up in the solution are $n \frac{\pi a}{L}$. That is, there is a certain base *fundamental frequency* $\frac{\pi a}{L}$, and then we also get all the multiples of this frequency, which in music are called the *overtones*. Which overtones appear and with what amplitude is what musicians call the *timbre* of the note. Mathematicians usually call this the *spectrum*. Because all the frequencies are multiples of one frequency (the fundamental), we get a nice pleasing sound.

The fundamental frequency $\frac{\pi a}{L}$ increases as we decrease length L . That is, if we place a finger on the fingerboard and then pluck a string we get a higher note. The constant a is given by

$$a = \sqrt{\frac{T}{\rho}},$$

where T is tension and ρ is the linear density of the string. Tightening the string (turning



the tuning peg on a guitar) increases a and hence produces a higher fundamental frequency (a higher note). On the other hand, using a heavier string reduces a and produces a lower fundamental frequency (a lower note). A bass guitar has longer thicker strings, while a ukulele has short strings made of lighter material.

Something rather interesting is the almost-symmetry between space and time. In its simplest form, we see this symmetry in the solutions

$$\sin\left(\frac{n\pi}{L}x\right)\sin\left(\frac{n\pi a}{L}t\right).$$

Except for the constant a , this solution looks the same if we flip time and space. In general, the solution for a fixed x is a Fourier series in t , for a fixed t it is a Fourier series in x , and the coefficients are related. If the shape $f(x)$ or the initial velocity $g(x)$ has lots of corners, then the sound wave will have lots of corners. That is because the Fourier coefficients of the initial shape decay to zero (as $n \rightarrow \infty$) at the same rate as the Fourier coefficients of the wave in time (for some fixed x). So if you use a sharp object to pluck the string, you get a sharper sound with lots of high-frequency components, while if you use your thumb, you get a softer sound without so many high overtones. Similarly, if you pluck close to

the bridge (close to one end of the string), you are getting a pluck that looks more like the sawtooth, and you get an even sharper sound.

In fact, if you look at the formula for the solution, you see that for any fixed x , we get an almost arbitrary Fourier series in t , everything except the constant term. In theory, you can obtain any timbre you want by plucking the string in just the right way. Of course, we are considering an ideal string of no stiffness and no air resistance. Those variables clearly impact the sound as well.

4.7.1 Exercises

Exercise 4.7.3: Solve

$$\begin{aligned} y_{tt} &= 9y_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ y(0, t) &= y(1, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= \sin(3\pi x) + \frac{1}{4}\sin(6\pi x) && \text{for } 0 < x < 1, \\ y_t(x, 0) &= 0 && \text{for } 0 < x < 1. \end{aligned}$$

Exercise 4.7.4: Solve

$$\begin{aligned} y_{tt} &= 4y_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ y(0, t) &= y(1, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= \sin(3\pi x) + \frac{1}{4}\sin(6\pi x) && \text{for } 0 < x < 1, \\ y_t(x, 0) &= \sin(9\pi x) && \text{for } 0 < x < 1. \end{aligned}$$

Exercise 4.7.5: Derive the solution for a general plucked string of length L and any constant a (in the equation $y_{tt} = a^2 y_{xx}$), where we raise the string some distance b at the midpoint and let go.

Exercise 4.7.6: Imagine that a stringed musical instrument fell on the floor. Suppose that the length of the string is 1 and $a = 1$. When the musical instrument hit the ground, the string was in rest position and hence $y(x, 0) = 0$. However, the string was moving at some velocity at impact ($t = 0$), say $y_t(x, 0) = -1$. Find the solution $y(x, t)$ for the shape of the string at time t .

Exercise 4.7.7 (challenging): Suppose that you have a vibrating string and that there is air resistance proportional to the velocity. That is, you have

$$\begin{aligned} y_{tt} &= a^2 y_{xx} - ky_t && \text{for } 0 < x < 1 \text{ and } t > 0, \\ y(0, t) &= y(1, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= f(x) && \text{for } 0 < x < 1, \\ y_t(x, 0) &= 0 && \text{for } 0 < x < 1. \end{aligned}$$

Suppose that $0 < k < 2\pi a$. Derive a series solution to the problem. Any coefficients in the series should be expressed as integrals of $f(x)$.

Exercise 4.7.8: Suppose you touch the guitar string exactly in the middle to ensure another condition $y(L/2, t) = 0$ for all time. Which multiples of the fundamental frequency $\frac{\pi a}{L}$ show up in the solution?

Exercise 4.7.101: Solve

$$\begin{aligned} y_{tt} &= y_{xx} && \text{for } 0 < x < \pi, t > 0, \\ y(0, t) &= y(\pi, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= \sin(x) && \text{for } 0 < x < \pi, \\ y_t(x, 0) &= \sin(x) && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.7.102: Solve

$$\begin{aligned} y_{tt} &= 25y_{xx} && \text{for } 0 < x < 2 \text{ and } t > 0, \\ y(0, t) &= y(2, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= 0 && \text{for } 0 < x < 2, \\ y_t(x, 0) &= \sin(\pi x) + 0.1 \sin(2\pi x) && \text{for } 0 < x < 2. \end{aligned}$$

Exercise 4.7.103: Solve

$$\begin{aligned} y_{tt} &= 2y_{xx} && \text{for } 0 < x < \pi \text{ and } t > 0, \\ y(0, t) &= y(\pi, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= x && \text{for } 0 < x < \pi, \\ y_t(x, 0) &= 0 && \text{for } 0 < x < \pi. \end{aligned}$$

Exercise 4.7.104: What happens when $a = 0$? Find a solution to $y_{tt} = 0$, $y(0, t) = y(\pi, t) = 0$, $y(x, 0) = \sin(2x)$, $y_t(x, 0) = \sin(x)$.

4.8 D'Alembert solution of the wave equation

Note: 1 lecture, different from §9.6 in [EP], part of §10.7 in [BD]

We have solved the wave equation by using Fourier series. But it is often more convenient to use the so-called *d'Alembert solution to the wave equation**. While this solution can be derived using Fourier series as well, it is really an awkward use of those concepts. It is easier and more instructive to derive this solution by making a correct change of variables to get an equation that can be solved by simple integration.

Suppose we wish to solve the wave equation

$$y_{tt} = a^2 y_{xx} \quad (4.15)$$

in the region $0 < x < L$ and $t > 0$ subject to the side conditions

$$\begin{aligned} y(0, t) = y(L, t) &= 0 && \text{for } t > 0, \\ y(x, 0) &= f(x) && \text{for } 0 < x < L, \\ y_t(x, 0) &= g(x) && \text{for } 0 < x < L. \end{aligned} \quad (4.16)$$

4.8.1 Change of variables

We will transform the equation into a simpler form where it can be solved by simple integration. We change variables to $\xi = x - at$, $\eta = x + at$. The chain rule says:

$$\begin{aligned} \frac{\partial}{\partial x} &= \frac{\partial \xi}{\partial x} \frac{\partial}{\partial \xi} + \frac{\partial \eta}{\partial x} \frac{\partial}{\partial \eta} = \frac{\partial}{\partial \xi} + \frac{\partial}{\partial \eta}, \\ \frac{\partial}{\partial t} &= \frac{\partial \xi}{\partial t} \frac{\partial}{\partial \xi} + \frac{\partial \eta}{\partial t} \frac{\partial}{\partial \eta} = -a \frac{\partial}{\partial \xi} + a \frac{\partial}{\partial \eta}. \end{aligned}$$

We compute

$$\begin{aligned} y_{xx} &= \frac{\partial^2 y}{\partial x^2} = \left(\frac{\partial}{\partial \xi} + \frac{\partial}{\partial \eta} \right) \left(\frac{\partial y}{\partial \xi} + \frac{\partial y}{\partial \eta} \right) = \frac{\partial^2 y}{\partial \xi^2} + 2 \frac{\partial^2 y}{\partial \xi \partial \eta} + \frac{\partial^2 y}{\partial \eta^2}, \\ y_{tt} &= \frac{\partial^2 y}{\partial t^2} = \left(-a \frac{\partial}{\partial \xi} + a \frac{\partial}{\partial \eta} \right) \left(-a \frac{\partial y}{\partial \xi} + a \frac{\partial y}{\partial \eta} \right) = a^2 \frac{\partial^2 y}{\partial \xi^2} - 2a^2 \frac{\partial^2 y}{\partial \xi \partial \eta} + a^2 \frac{\partial^2 y}{\partial \eta^2}. \end{aligned}$$

In the computations above, we used the fact from calculus that $\frac{\partial^2 y}{\partial \xi \partial \eta} = \frac{\partial^2 y}{\partial \eta \partial \xi}$. We plug what we got into the wave equation,

$$0 = a^2 y_{xx} - y_{tt} = 4a^2 \frac{\partial^2 y}{\partial \xi \partial \eta} = 4a^2 y_{\xi\eta}.$$

Therefore, the wave equation (4.15) transforms into $y_{\xi\eta} = 0$. It is easy to find the general solution to this new equation by integrating twice. Keeping ξ constant, we integrate with

*Named after the French mathematician [Jean le Rond d'Alembert](#) (1717–1783).

respect to η first* and note that the constant of integration depends on ξ ; for each ξ , we may get a different constant of integration. We get $y_\xi = C(\xi)$. Next, we integrate with respect to ξ and note that the constant of integration depends on η . Thus, $y = \int C(\xi) d\xi + B(\eta)$. The solution must, therefore, be of the following form for some functions $A(\xi)$ and $B(\eta)$:

$$y = A(\xi) + B(\eta) = A(x - at) + B(x + at).$$

The solution is a superposition of two functions (waves) traveling at speed a in opposite directions. The coordinates ξ and η are called the *characteristic coordinates*, and a similar technique can be applied to more complicated hyperbolic PDEs. In § 1.9 it is used to solve first-order linear PDEs. Basically, to solve the wave equation (or more general hyperbolic equations) we find certain characteristic curves along which the equation is really just an ODE, or a pair of ODEs. In this case, these are the curves where ξ and η are constant.

4.8.2 D'Alembert's formula

We know what any solution must look like, but we need to solve for the given side conditions. We will just give the formula and see that it works. Let $F(x)$ denote the odd periodic extension of $f(x)$, and let $G(x)$ denote the odd periodic extension of $g(x)$. Define

$$A(x) = \frac{1}{2}F(x) - \frac{1}{2a} \int_0^x G(s) ds, \quad B(x) = \frac{1}{2}F(x) + \frac{1}{2a} \int_0^x G(s) ds.$$

We claim these $A(x)$ and $B(x)$ give the solution. Explicitly, the solution is $y(x, t) = A(x - at) + B(x + at)$ or in other words:

$$\begin{aligned} y(x, t) &= \frac{1}{2}F(x - at) - \frac{1}{2a} \int_0^{x-at} G(s) ds + \frac{1}{2}F(x + at) + \frac{1}{2a} \int_0^{x+at} G(s) ds \\ &= \frac{F(x - at) + F(x + at)}{2} + \frac{1}{2a} \int_{x-at}^{x+at} G(s) ds. \end{aligned} \quad (4.17)$$

Let us check that the d'Alembert formula really works. First,

$$y(x, 0) = \frac{F(x) + F(x)}{2} + \frac{1}{2a} \int_x^x G(s) ds = F(x).$$

So far so good. Assume for simplicity F is differentiable. And we use the first form of (4.17) as it is easier to differentiate. By the fundamental theorem of calculus we have

$$y_t(x, t) = \frac{-a}{2}F'(x - at) + \frac{1}{2}G(x - at) + \frac{a}{2}F'(x + at) + \frac{1}{2}G(x + at).$$

So

$$y_t(x, 0) = \frac{-a}{2}F'(x) + \frac{1}{2}G(x) + \frac{a}{2}F'(x) + \frac{1}{2}G(x) = G(x).$$

*There is nothing special about η , you can integrate with ξ first, if you wish.

Yay! We're smoking now. OK, now the boundary conditions. Note that $F(x)$ and $G(x)$ are odd. So

$$y(0, t) = \frac{F(-at) + F(at)}{2} + \frac{1}{2a} \int_{-at}^{at} G(s) ds = \frac{-F(at) + F(at)}{2} + \frac{1}{2a} \int_{-at}^{at} G(s) ds = 0 + 0 = 0.$$

Now $F(x)$ is odd and $2L$ -periodic, so

$$F(L - at) + F(L + at) = F(-L - at) + F(L + at) = -F(L + at) + F(L + at) = 0.$$

Next, $G(s)$ is odd and $2L$ -periodic, so we change variables $v = s - L$. We then notice that $G(v + L) = G(v - L) = -G(-v + L)$, so $G(v + L)$ is odd as a function of v :

$$\int_{L-at}^{L+at} G(s) ds = \int_{-at}^{at} G(v + L) dv = 0.$$

Hence

$$y(L, t) = \frac{F(L - at) + F(L + at)}{2} + \frac{1}{2a} \int_{L-at}^{L+at} G(s) ds = 0 + 0 = 0.$$

And voilà, it works.

Example 4.8.1: D'Alembert says that the solution is a superposition of two functions (waves) moving in opposite directions at "speed" a . To get an idea of how it works, we work out the following example. Consider the simpler setup

$$\begin{aligned} y_{tt} &= y_{xx} && \text{for } 0 < x < 1 \text{ and } t > 0, \\ y(0, t) &= y(1, t) = 0 && \text{for } t > 0, \\ y(x, 0) &= f(x) && \text{for } 0 < x < 1, \\ y_t(x, 0) &= 0 && \text{for } 0 < x < 1. \end{aligned}$$

Let $f(x)$ be the following triangular impulse of height 1 centered at $x = 0.5$:

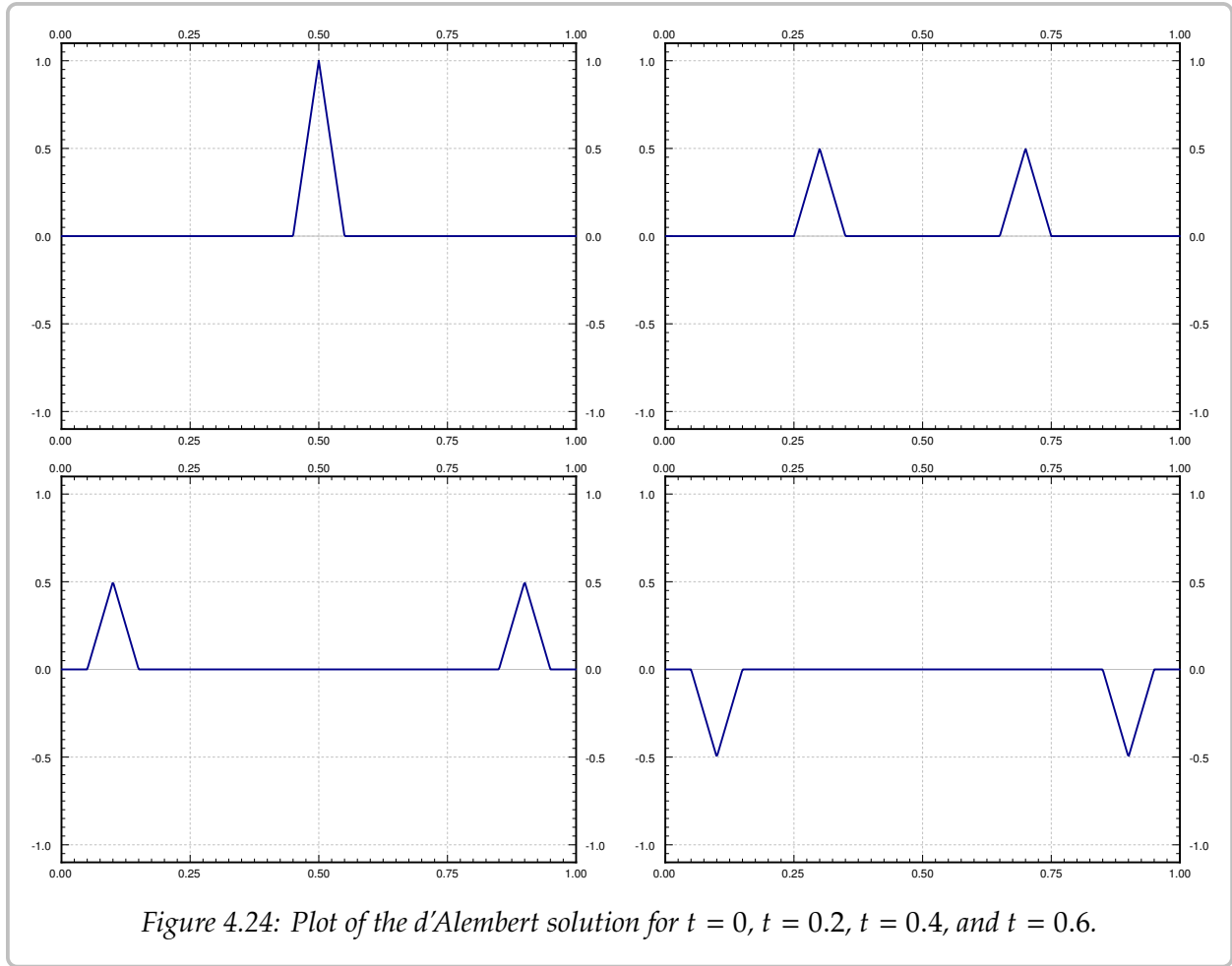
$$f(x) = \begin{cases} 0 & \text{if } 0 \leq x < 0.45, \\ 20(x - 0.45) & \text{if } 0.45 \leq x < 0.5, \\ 20(0.55 - x) & \text{if } 0.5 \leq x < 0.55, \\ 0 & \text{if } 0.55 \leq x \leq 1. \end{cases}$$

The graph of this impulse is the top left plot in [Figure 4.24](#) on the facing page.

Let $F(x)$ be the odd periodic extension of $f(x)$. Then (4.17) says that the solution is

$$y(x, t) = \frac{F(x - t) + F(x + t)}{2}.$$

It is not hard to compute specific values of $y(x, t)$. For example, to compute $y(0.1, 0.6)$, we notice $x - t = -0.5$ and $x + t = 0.7$. Now $F(-0.5) = -f(0.5) = -20(0.55 - 0.5) = -1$ and $F(0.7) = f(0.7) = 0$. Hence $y(0.1, 0.6) = \frac{-1+0}{2} = -0.5$. As you can see, the d'Alembert solution is much easier to actually compute and to plot than the Fourier series solution. See [Figure 4.24](#) on the next page for plots of the solution y for several different t .



4.8.3 Another way to solve for the side conditions

It is perhaps easier and more useful to remember the procedure rather than memorizing the formula itself. The important thing is that a solution to the wave equation is a superposition of two waves traveling in opposite directions. That is,

$$y(x, t) = A(x - at) + B(x + at).$$

If you think about it, the exact formulas for A and B are not hard to guess once you realize what kind of side conditions $y(x, t)$ is supposed to satisfy. Let us walk through the formula again, but slightly differently. The best approach is to do it in stages. When $g(x) = 0$ (and hence $G(x) = 0$), the solution is

$$\frac{F(x - at) + F(x + at)}{2}.$$

On the other hand, when $f(x) = 0$ (and hence $F(x) = 0$), we let

$$H(x) = \int_0^x G(s) ds.$$

The solution in this case is

$$\frac{1}{2a} \int_{x-at}^{x+at} G(s) ds = \frac{-H(x-at) + H(x+at)}{2a}.$$

By superposition, we get a solution for the general side conditions (4.16) (when neither $f(x)$ nor $g(x)$ are identically zero).

$$y(x, t) = \frac{F(x-at) + F(x+at)}{2} + \frac{-H(x-at) + H(x+at)}{2a}. \quad (4.18)$$

Do note the minus sign before the H and the a in the second denominator.

Exercise 4.8.1: Check that the new formula (4.18) satisfies the side conditions (4.16).

Warning: Make sure you use the odd periodic extensions $F(x)$ and $G(x)$, when you have formulas for $f(x)$ and $g(x)$. The thing is, those formulas in general hold only for $0 < x < L$ and are not usually equal to $F(x)$ and $G(x)$ for other x .

4.8.4 Some remarks

We remark that the formula $y(x, t) = A(x-at) + B(x+at)$ is the reason why the solution of the wave equation does not get “nicer” as time goes on, that is, why in the examples where the initial conditions had corners, the solution also has corners at every time t .

The corners bring us to another interesting remark. Nobody ever notices at first that our example solutions are not even differentiable (they have corners): In Example 4.8.1 above, the solution is not differentiable whenever $x = t + 0.5$ or $x = -t + 0.5$ for example. Really, to be able to compute y_{xx} or y_{tt} , you need not one, but two derivatives. Fear not, we could think of a shape that is very nearly $F(x)$ but does have two derivatives by rounding the corners a little bit, and then the solution would be very nearly $\frac{F(x-t)+F(x+t)}{2}$ and nobody would notice the switch.

One final remark is what the d’Alembert solution tells us about what bits of the initial conditions influence the solution at a certain point. We can figure this out by “traveling backwards along the characteristics.” Suppose that the string is very long (perhaps infinite) for simplicity. Since the solution at time t is

$$y(x, t) = \frac{F(x-at) + F(x+at)}{2} + \frac{1}{2a} \int_{x-at}^{x+at} G(s) ds,$$

we notice that we have only used the initial conditions in the interval $[x-at, x+at]$. The endpoints of this interval are called the *wavefronts*, as that is where the wave front is, given an initial ($t = 0$) disturbance at x . If $a = 1$, an observer sitting at $x = 0$ at time $t = 1$ has only seen the initial conditions for x in the interval $[-1, 1]$ and is blissfully unaware of anything else. This is why, for example, we do not know that a supernova has occurred in the universe until we see its light, millions of years from the time when it happened.

4.8.5 Exercises

Exercise 4.8.2: Using the d'Alembert solution solve $y_{tt} = 4y_{xx}$, $0 < x < \pi$, $t > 0$, $y(0, t) = y(\pi, t) = 0$, $y(x, 0) = \sin x$, and $y_t(x, 0) = \sin x$. Hint: Note that $\sin x$ is the odd periodic extension of $y(x, 0)$ and $y_t(x, 0)$.

Exercise 4.8.3: Using the d'Alembert solution solve $y_{tt} = 2y_{xx}$, $0 < x < 1$, $t > 0$, $y(0, t) = y(1, t) = 0$, $y(x, 0) = \sin^5(\pi x)$, and $y_t(x, 0) = \sin^3(\pi x)$.

Exercise 4.8.4: Take $y_{tt} = 4y_{xx}$, $0 < x < \pi$, $t > 0$, $y(0, t) = y(\pi, t) = 0$, $y(x, 0) = x(\pi - x)$, and $y_t(x, 0) = 0$.

- Solve using the d'Alembert formula. Hint: You can use the sine series for $y(x, 0)$.
- Find the solution as a function of x for a fixed $t = 0.5$, $t = 1$, and $t = 2$. Do not use the sine series here.

Exercise 4.8.5: Derive the d'Alembert solution for $y_{tt} = a^2y_{xx}$, $0 < x < \pi$, $t > 0$, $y(0, t) = y(\pi, t) = 0$, $y(x, 0) = f(x)$, and $y_t(x, 0) = 0$, using the Fourier series solution of the wave equation, by applying an appropriate trigonometric identity. Hint: Do it first for a single term of the Fourier series solution, in particular do it when y is $\sin(nx) \cos(nat)$.

Exercise 4.8.6: The d'Alembert solution still works if there are no boundary conditions and the initial condition is defined on the whole real line. Suppose that $y_{tt} = y_{xx}$ (for all x on the real line and $t \geq 0$), $y(x, 0) = f(x)$, and $y_t(x, 0) = 0$, where

$$f(x) = \begin{cases} 0 & \text{if } x < -1, \\ x + 1 & \text{if } -1 \leq x < 0, \\ -x + 1 & \text{if } 0 \leq x < 1, \\ 0 & \text{if } 1 < x. \end{cases}$$

Solve using the d'Alembert solution. That is, write down a piecewise definition for the solution. Then sketch the solution for $t = 0$, $t = 1/2$, $t = 1$, and $t = 2$.

Exercise 4.8.101: Using the d'Alembert solution solve $y_{tt} = 9y_{xx}$, $0 < x < 1$, $t > 0$, $y(0, t) = y(1, t) = 0$, $y(x, 0) = \sin(2\pi x)$, and $y_t(x, 0) = \sin(3\pi x)$.

Exercise 4.8.102: Take $y_{tt} = 4y_{xx}$, $0 < x < 1$, $t > 0$, $y(0, t) = y(1, t) = 0$, $y(x, 0) = x - x^2$, and $y_t(x, 0) = 0$. Using the d'Alembert solution find the solution at

- $t = 0.1$,
- $t = 1/2$,
- $t = 1$.

You may have to split your answer up by cases.

Exercise 4.8.103: Take $y_{tt} = 100y_{xx}$, $0 < x < 4$, $t > 0$, $y(0, t) = y(4, t) = 0$, $y(x, 0) = F(x)$, and $y_t(x, 0) = 0$. Suppose that $F(0) = 0$, $F(1) = 2$, $F(2) = 3$, $F(3) = 1$. Using the d'Alembert solution find

- $y(1, 1)$,
- $y(4, 3)$,
- $y(3, 9)$.

4.9 Steady state temperature and the Laplacian

Note: 1 lecture, §9.7 in [EP], §10.8 in [BD]

Consider an insulated wire, a plate, or a 3-dimensional object. We apply certain fixed temperatures on the ends of the wire, the edges of the plate, or on all sides of the 3-dimensional object. We wish to find the *steady-state temperature* distribution. That is, we wish to know the temperature after a long enough period of time.

We are really seeking a solution to the heat equation that is not dependent on time. We first solve the problem in one space variable. We are looking for a function u that satisfies

$$u_t = ku_{xx},$$

but such that $u_t = 0$ for all x and t . Hence, we are looking for a function of x alone that satisfies $u_{xx} = 0$. It is easy to solve this equation by integration, and we see that $u = Ax + B$ for some constants A and B .

Consider an insulated wire where we apply constant temperature T_1 at one end (say where $x = 0$) and T_2 on the other end (at $x = L$ where L is the length of the wire). Our steady-state solution is

$$u(x) = \frac{T_2 - T_1}{L}x + T_1.$$

It is simply a straight line from one end to the other. This solution agrees with the common-sense intuition about how heat should be distributed in the wire. So in one dimension, the steady-state solutions are just straight lines.

Things are more complicated in two or more space dimensions. We restrict ourselves to two space dimensions for simplicity. The heat equation in two space variables is

$$u_t = k(u_{xx} + u_{yy}), \quad (4.19)$$

or more commonly written as $u_t = k\Delta u$ or $u_t = k\nabla^2 u$. The Δ and ∇^2 symbols both mean $\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}$. We will use Δ here. The reason for using such a notation is that you can define Δ to be the right thing for any number of space dimensions and then the heat equation is always $u_t = k\Delta u$. The operator Δ is called the *Laplacian*.

OK, now that we have notation out of the way, let us see what an equation for the steady-state solution looks like. We are looking for a solution to (4.19) that does not depend on t , that is, $u_t = 0$. Hence, we are looking for a function $u(x, y)$ such that

$$\Delta u = u_{xx} + u_{yy} = 0.$$

This equation is called the *Laplace equation*^{*} and is an example of an elliptic equation. Solutions to the Laplace equation are called *harmonic functions* and have many nice properties and applications far beyond the steady-state heat problem.

^{*}Named after the French mathematician [Pierre-Simon, marquis de Laplace](#) (1749–1827).

Harmonic functions in two variables are no longer just linear (graphs are not just planes). For example, you can check that the functions $x^2 - y^2$ and xy are harmonic. However, note that if u_{xx} is positive, u is concave up in the x direction, then u_{yy} must be negative and u must be concave down in the y direction. A harmonic function can never have a “hilltop” or “valley” on the graph. This observation is consistent with our intuitive idea of steady-state heat distribution; the hottest or coldest spot should not be inside.

Commonly the Laplace equation is part of a so-called *Dirichlet problem**. That is, we have a region in the xy -plane and we specify certain values along the boundaries of the region. We then try to find a solution u to the Laplace equation defined on this region such that u agrees with the values we specified on the boundary.

In this section, we consider a rectangular region. For simplicity, we specify boundary values to be zero at three of the four edges and only specify an arbitrary function at one edge. With the principle of superposition, we can use this simpler solution to derive the general solution for arbitrary boundary values by solving four different problems, one for each edge, and adding those solutions together. This setup is left as an exercise.

We wish to solve the following problem. Let h and w be the height and width of our rectangle, with one corner at the origin and lying in the first quadrant, so that the region is given by $0 < x < w$ and $0 < y < h$. Consider the problem

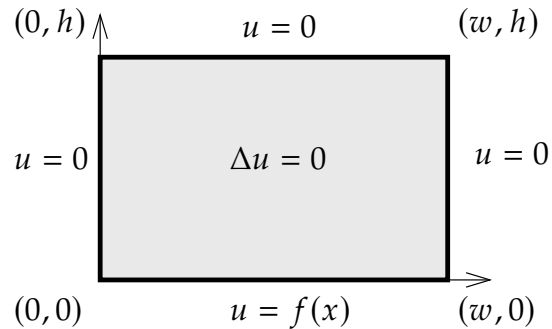
$$\Delta u = 0, \quad (4.20)$$

$$u(0, y) = 0 \quad \text{for } 0 < y < h, \quad (4.21)$$

$$u(x, h) = 0 \quad \text{for } 0 < x < w, \quad (4.22)$$

$$u(w, y) = 0 \quad \text{for } 0 < y < h, \quad (4.23)$$

$$u(x, 0) = f(x) \quad \text{for } 0 < x < w. \quad (4.24)$$



The method we apply is separation of variables. Again, we will come up with enough building-block solutions satisfying all the homogeneous boundary conditions (all conditions except (4.24)). We notice that superposition still works for the equation and all the homogeneous conditions. Therefore, we can use the Fourier series for $f(x)$ to find a solution u that also solves (4.24).

We try $u(x, y) = X(x)Y(y)$. We plug u into the equation to get

$$X''Y + XY'' = 0.$$

We put the X s on one side and the Y s on the other to get

$$-\frac{X''}{X} = \frac{Y''}{Y}.$$

*Named after the German mathematician [Johann Peter Gustav Lejeune Dirichlet](#) (1805–1859).

The left-hand side only depends on x and the right-hand side only depends on y . Therefore, there is some constant λ such that $\lambda = \frac{-X''}{X} = \frac{Y''}{Y}$. And we get two equations

$$X'' + \lambda X = 0,$$

$$Y'' - \lambda Y = 0.$$

Furthermore, the homogeneous boundary conditions imply that $X(0) = X(w) = 0$ and $Y(h) = 0$. Using the equation for X , we have already seen that there is a nontrivial solution if and only if $\lambda = \lambda_n = \frac{n^2\pi^2}{w^2}$ and the solution is a multiple of

$$X_n(x) = \sin\left(\frac{n\pi}{w}x\right).$$

For these given λ_n , the general solution for Y (one for each n) is

$$Y_n(y) = A_n \cosh\left(\frac{n\pi}{w}y\right) + B_n \sinh\left(\frac{n\pi}{w}y\right). \quad (4.25)$$

There is only one condition on Y_n and hence we can pick one of A_n or B_n to be something convenient. It will be useful to have $Y_n(0) = 1$, so let $A_n = 1$. Setting $Y_n(h) = 0$ and solving for B_n , we get

$$B_n = \frac{-\cosh\left(\frac{n\pi h}{w}\right)}{\sinh\left(\frac{n\pi h}{w}\right)}.$$

After we plug the A_n and B_n into (4.25) and simplify by using the identity $\sinh(\alpha - \beta) = \sinh(\alpha)\cosh(\beta) - \cosh(\alpha)\sinh(\beta)$, we find

$$Y_n(y) = \frac{\sinh\left(\frac{n\pi(h-y)}{w}\right)}{\sinh\left(\frac{n\pi h}{w}\right)}.$$

We define $u_n(x, y) = X_n(x)Y_n(y)$. Note that u_n satisfies (4.20)–(4.23). Observe

$$u_n(x, 0) = X_n(x)Y_n(0) = \sin\left(\frac{n\pi}{w}x\right).$$

Suppose

$$f(x) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi x}{w}\right).$$

Then we get a solution of (4.20)–(4.24) of the following form.

$$u(x, y) = \sum_{n=1}^{\infty} b_n u_n(x, y) = \sum_{n=1}^{\infty} b_n \sin\left(\frac{n\pi}{w}x\right) \left(\frac{\sinh\left(\frac{n\pi(h-y)}{w}\right)}{\sinh\left(\frac{n\pi h}{w}\right)} \right).$$

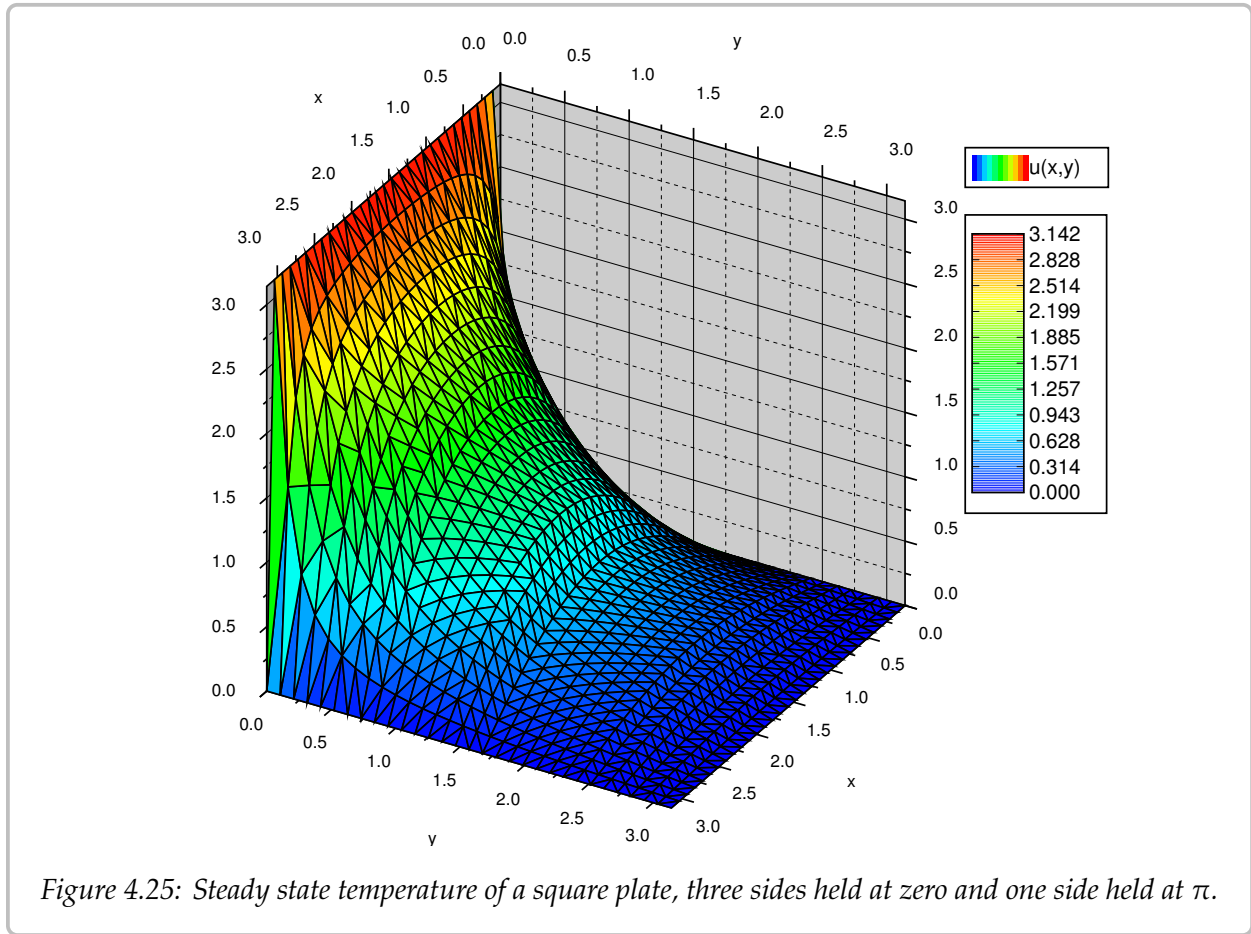
As u_n satisfies (4.20)–(4.23), a linear combination (finite or infinite) of u_n , and hence u , satisfies (4.20)–(4.23). We plug in $y = 0$ to see that u satisfies (4.24) as well.

Example 4.9.1: Take $w = h = \pi$ and let $f(x) = \pi$. Let us compute the sine series for the function π (same as the series for the square wave). For $0 < x < \pi$, we have

$$f(x) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{4}{n} \sin(nx).$$

The solution $u(x, y)$, see Figure 4.25, to the corresponding Dirichlet problem is given as

$$u(x, y) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{4}{n} \sin(nx) \left(\frac{\sinh(n(\pi - y))}{\sinh(n\pi)} \right).$$



This scenario corresponds to the steady-state temperature on a square plate of width π with three sides held at 0 degrees and one side held at π degrees. If we have arbitrary initial data on all sides, then we solve four problems, each using one piece of nonhomogeneous data. Then we use the principle of superposition to add up all four solutions to have a solution to the original problem.

A different way to visualize solutions of the Laplace equation is to take a wire and bend it so that it corresponds to the graph of the temperature above the boundary of your region. Cut a rubber sheet in the shape of your region—a square in our case—and stretch it, fixing the edges of the sheet to the wire. The rubber sheet is a good approximation of the graph of the solution to the Laplace equation with the given boundary data.

4.9.1 Exercises

Exercise 4.9.1: In the region described by $0 < x < \pi$ and $0 < y < \pi$, solve the problem

$$\Delta u = 0, \quad u(x, 0) = \sin x, \quad u(x, \pi) = 0, \quad u(0, y) = 0, \quad u(\pi, y) = 0.$$

Exercise 4.9.2: In the region described by $0 < x < 1$ and $0 < y < 1$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= \sin(\pi x) - \sin(2\pi x), & u(x, 1) &= 0, \\ u(0, y) &= 0, & u(1, y) &= 0. \end{aligned}$$

Exercise 4.9.3: In the region described by $0 < x < 1$ and $0 < y < 1$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= u(x, 1) = u(0, y) = u(1, y) = C. \end{aligned}$$

for some constant C . Hint: Guess, then check your intuition.

Exercise 4.9.4: In the region described by $0 < x < \pi$ and $0 < y < \pi$, solve

$$\Delta u = 0, \quad u(x, 0) = 0, \quad u(x, \pi) = \pi, \quad u(0, y) = y, \quad u(\pi, y) = y.$$

Hint: Try a solution of the form $u(x, y) = X(x) + Y(y)$ (different separation of variables).

Exercise 4.9.5: Use the solution of [Exercise 4.9.4](#) to solve

$$\Delta u = 0, \quad u(x, 0) = \sin x, \quad u(x, \pi) = \pi, \quad u(0, y) = y, \quad u(\pi, y) = y.$$

Hint: Use superposition.

Exercise 4.9.6: In the region described by $0 < x < w$ and $0 < y < h$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= 0, & u(x, h) &= f(x), \\ u(0, y) &= 0, & u(w, y) &= 0. \end{aligned}$$

The solution should be in series form using the Fourier-series coefficients of $f(x)$.

Exercise 4.9.7: In the region described by $0 < x < w$ and $0 < y < h$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= 0, & u(x, h) &= 0, \\ u(0, y) &= f(y), & u(w, y) &= 0. \end{aligned}$$

The solution should be in series form using the Fourier-series coefficients of $f(y)$.

Exercise 4.9.8: In the region described by $0 < x < w$ and $0 < y < h$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= 0, & u(x, h) &= 0, \\ u(0, y) &= 0, & u(w, y) &= f(y). \end{aligned}$$

The solution should be in series form using the Fourier-series coefficients of $f(y)$.

Exercise 4.9.9: In the region described by $0 < x < 1$ and $0 < y < 1$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= \sin(9\pi x), & u(x, 1) &= \sin(2\pi x), \\ u(0, y) &= 0, & u(1, y) &= 0. \end{aligned}$$

Hint: Use superposition.

Exercise 4.9.10: In the region described by $0 < x < 1$ and $0 < y < 1$, solve the problem

$$\begin{aligned} u_{xx} + u_{yy} &= 0, \\ u(x, 0) &= \sin(\pi x), & u(x, 1) &= \sin(\pi x), \\ u(0, y) &= \sin(\pi y), & u(1, y) &= \sin(\pi y). \end{aligned}$$

Hint: Use superposition.

Exercise 4.9.11 (challenging): Using only your intuition find $u(1/2, 1/2)$, for the problem $\Delta u = 0$, where $u(0, y) = u(1, y) = 100$ for $0 < y < 1$, and $u(x, 0) = u(x, 1) = 0$ for $0 < x < 1$. Explain.

Exercise 4.9.101: In the region described by $0 < x < 1$ and $0 < y < 1$, solve the problem

$$\Delta u = 0, \quad u(x, 0) = \sum_{n=1}^{\infty} \frac{1}{n^2} \sin(n\pi x), \quad u(x, 1) = 0, \quad u(0, y) = 0, \quad u(1, y) = 0.$$

Exercise 4.9.102: In the region described by $0 < x < 1$ and $0 < y < 2$, solve the problem

$$\Delta u = 0, \quad u(x, 0) = 0.1 \sin(\pi x), \quad u(x, 2) = 0, \quad u(0, y) = 0, \quad u(1, y) = 0.$$

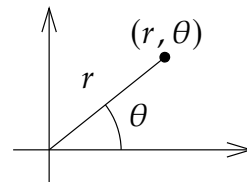
4.10 Dirichlet problem in the circle and the Poisson kernel

Note: 2 lectures, §9.7 in [EP], §10.8 in [BD]

4.10.1 Laplace in polar coordinates

A more natural setting for the Laplace equation $\Delta u = 0$ is a circle rather than a rectangle. On the other hand, what makes the problem somewhat more difficult is that we need polar coordinates. Recall that the polar coordinates for the (x, y) -plane are (r, θ) :

$$x = r \cos \theta, \quad y = r \sin \theta,$$

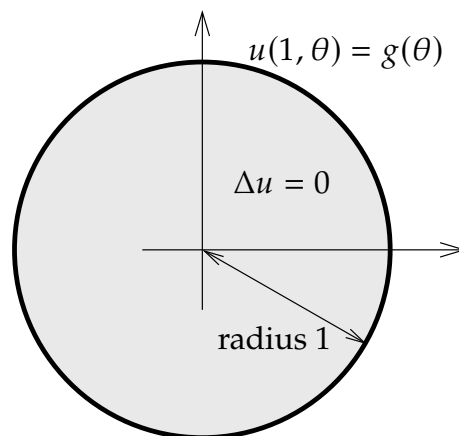


where $r \geq 0$ and $-\pi < \theta \leq \pi$. So the point (x, y) is distance r from the origin at an angle θ from the positive x -axis.

Now that we know our coordinates, let us give the problem we wish to solve. We have a circular region of radius 1, and we are interested in the Dirichlet problem for the Laplace equation for this region. Let $u(r, \theta)$ denote the temperature at the point (r, θ) in polar coordinates. We will prescribe the temperature on the circle as $g(\theta)$.

We have the problem:

$$\begin{aligned} \Delta u &= 0 & \text{for } r < 1, \\ u(1, \theta) &= g(\theta) & \text{for } -\pi < \theta \leq \pi. \end{aligned} \quad (4.26)$$



The first issue we face is that we do not know the Laplacian in polar coordinates. Normally, we would find u_{xx} and u_{yy} in terms of the derivatives in r and θ . We would need to solve for r and θ in terms of x and y . In this case, it is more convenient to work in reverse. We compute derivatives in r and θ in terms of derivatives in x and y and then we solve. The computations are easier this way. First,

$$\begin{aligned} x_r &= \cos \theta, & x_\theta &= -r \sin \theta, \\ y_r &= \sin \theta, & y_\theta &= r \cos \theta. \end{aligned}$$

By chain rule, we obtain

$$\begin{aligned} u_r &= u_x x_r + u_y y_r = \cos(\theta) u_x + \sin(\theta) u_y, \\ u_{rr} &= \cos(\theta)(u_{xx} x_r + u_{xy} y_r) + \sin(\theta)(u_{yx} x_r + u_{yy} y_r) \\ &= \cos^2(\theta) u_{xx} + 2 \cos(\theta) \sin(\theta) u_{xy} + \sin^2(\theta) u_{yy}. \end{aligned}$$

Similarly for the θ derivative. Note that we have to use the product rule for the second derivative.

$$\begin{aligned} u_\theta &= u_x x_\theta + u_y y_\theta = -r \sin(\theta) u_x + r \cos(\theta) u_y, \\ u_{\theta\theta} &= -r \cos(\theta) u_x - r \sin(\theta) (u_{xx} x_\theta + u_{xy} y_\theta) - r \sin(\theta) u_y + r \cos(\theta) (u_{yx} x_\theta + u_{yy} y_\theta) \\ &= -r \cos(\theta) u_x - r \sin(\theta) u_y + r^2 \sin^2(\theta) u_{xx} - r^2 2 \sin(\theta) \cos(\theta) u_{xy} + r^2 \cos^2(\theta) u_{yy}. \end{aligned}$$

Let us now try to find $u_{xx} + u_{yy}$. We start with $\frac{1}{r^2} u_{\theta\theta}$ to get rid of those pesky r^2 s. If we add u_{rr} and use the fact that $\cos^2(\theta) + \sin^2(\theta) = 1$, we get

$$\frac{1}{r^2} u_{\theta\theta} + u_{rr} = u_{xx} + u_{yy} - \frac{1}{r} \cos(\theta) u_x - \frac{1}{r} \sin(\theta) u_y.$$

We are not quite there yet, but all we are lacking is $\frac{1}{r} u_r$. We add it to obtain the *Laplacian in polar coordinates*:

$$\Delta u = u_{xx} + u_{yy} = \frac{1}{r^2} u_{\theta\theta} + \frac{1}{r} u_r + u_{rr}.$$

Notice that the Laplacian in polar coordinates no longer has constant coefficients.

4.10.2 Series solution

Let us separate variables as usual. That is, we try $u(r, \theta) = R(r)\Theta(\theta)$. Then

$$0 = \Delta u = \frac{1}{r^2} R \Theta'' + \frac{1}{r} R' \Theta + R'' \Theta.$$

We put R on one side and Θ on the other and conclude that both sides must be constant.

$$\begin{aligned} \frac{1}{r^2} R \Theta'' &= - \left(\frac{1}{r} R' + R'' \right) \Theta \\ \frac{\Theta''}{\Theta} &= - \frac{r R' + r^2 R''}{R} = -\lambda \end{aligned}$$

We get two equations:

$$\begin{aligned} \Theta'' + \lambda \Theta &= 0, \\ r^2 R'' + r R' - \lambda R &= 0. \end{aligned}$$

We first focus on Θ . We know that $u(r, \theta)$ ought to be 2π -periodic in θ , that is, $u(r, \theta) = u(r, \theta + 2\pi)$. Therefore, the solution to $\Theta'' + \lambda \Theta = 0$ must be 2π -periodic. We have seen such a problem in [Example 4.1.5](#). We conclude that $\lambda = n^2$ for a nonnegative integer $n = 0, 1, 2, 3, \dots$. The equation becomes $\Theta'' + n^2 \Theta = 0$. When $n = 0$ the equation is just $\Theta'' = 0$, so we have the general solution $A\theta + B$. As Θ is periodic, $A = 0$. For convenience, we write this solution as

$$\Theta_0 = \frac{a_0}{2}$$

for some constant a_0 . For positive n , the solution to $\Theta'' + n^2\Theta = 0$ is

$$\Theta_n = a_n \cos(n\theta) + b_n \sin(n\theta),$$

for some constants a_n and b_n .

Next, we consider the equation for R ,

$$r^2 R'' + rR' - n^2 R = 0.$$

This equation appeared in exercises before—we solved it in [Exercise 2.1.6](#) and [Exercise 2.1.7](#) on page 83. The idea is to try a solution r^s and if that does not give us two solutions, also try a solution of the form $r^s \ln r$. We name the solution R_n as usual. When $n = 0$ we obtain

$$R_0 = Ar^0 + Br^0 \ln r = A + B \ln r,$$

and if $n > 0$, we get

$$R_n = Ar^n + Br^{-n}.$$

The function $u(r, \theta)$ must be finite (it cannot blow up) at the origin, that is, when $r = 0$. So $B = 0$ in both cases as otherwise r^{-n} or $\ln r$ does blow up as $r \rightarrow 0$. Set $A = 1$ in both cases; the constants in Θ_n will pick up the slack so nothing is lost. That is,

$$R_0 = 1 \quad \text{and} \quad R_n = r^n.$$

Our building block solutions are

$$u_0(r, \theta) = \frac{a_0}{2} \quad \text{and} \quad u_n(r, \theta) = a_n r^n \cos(n\theta) + b_n r^n \sin(n\theta).$$

Putting everything together, our solution is

$$u(r, \theta) = \frac{a_0}{2} + \sum_{n=1}^{\infty} [a_n r^n \cos(n\theta) + b_n r^n \sin(n\theta)].$$

We look at the boundary condition in (4.26),

$$g(\theta) = u(1, \theta) = \frac{a_0}{2} + \sum_{n=1}^{\infty} [a_n \cos(n\theta) + b_n \sin(n\theta)].$$

Therefore, to solve (4.26) we expand $g(\theta)$, which is a 2π -periodic function, as a Fourier series, and then multiply the n^{th} term by r^n . To find the a_n and the b_n , we compute

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} g(\theta) \cos(n\theta) d\theta \quad \text{and} \quad b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} g(\theta) \sin(n\theta) d\theta.$$

Example 4.10.1: Suppose we wish to solve

$$\begin{aligned} \Delta u &= 0, & 0 \leq r < 1, & \quad -\pi < \theta \leq \pi, \\ u(1, \theta) &= \cos(10\theta), & -\pi < \theta \leq \pi. \end{aligned}$$

The $g(\theta)$ is already expanded, so the solution is

$$u(r, \theta) = r^{10} \cos(10 \theta).$$

See the plot in [Figure 4.26](#). The thing to notice in this example is that the effect of a high frequency is mostly felt at the boundary. In the middle of the disc, the solution is very close to zero. That is because r^{10} is rather small when r is close to 0.

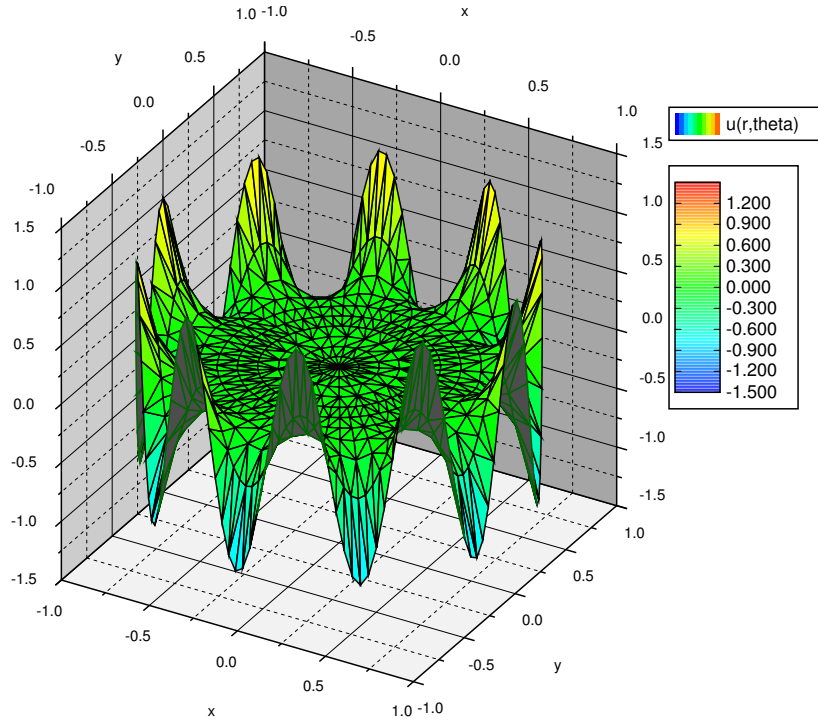


Figure 4.26: The solution of the Dirichlet problem in the disc with $\cos(10 \theta)$ as boundary data.

Example 4.10.2: Let us solve a more difficult problem. Consider a long rod with a circular cross section of radius 1. Suppose we wish to solve the steady-state heat problem in the rod. If the rod is long enough, we simply need to solve the Laplace equation in two dimensions. We put the center of the rod at the origin, and we have exactly the region we are currently studying—a circle of radius 1. For the boundary conditions, suppose in Cartesian coordinates x and y , the temperature on the boundary is 0 when $y < 0$, and it is $2y$ when $y > 0$.

Let us set the problem up. As $y = r \sin(\theta)$, then on the circle of radius 1, that is, where $r = 1$, we have $2y = 2 \sin(\theta)$. So

$$\begin{aligned} \Delta u &= 0, & 0 \leq r < 1, & \quad -\pi < \theta \leq \pi, \\ u(1, \theta) &= \begin{cases} 2 \sin(\theta) & \text{if } 0 \leq \theta \leq \pi, \\ 0 & \text{if } -\pi < \theta < 0. \end{cases} \end{aligned}$$

We must now compute the Fourier series for the boundary condition. By now the reader has plentiful experience in computing Fourier series, and so we simply state that

$$u(1, \theta) = \frac{2}{\pi} + \sin(\theta) + \sum_{n=1}^{\infty} \frac{-4}{\pi(4n^2 - 1)} \cos(2n\theta).$$

Exercise 4.10.1: Compute the series for $u(1, \theta)$ and verify that it really is what we have just claimed. Hint: Be careful, make sure not to divide by zero.

To obtain the solution (see Figure 4.27), we write its series by multiplying terms in the series for $g(\theta)$ by r^n in the right places:

$$u(r, \theta) = \frac{2}{\pi} + r \sin(\theta) + \sum_{n=1}^{\infty} \frac{-4r^{2n}}{\pi(4n^2 - 1)} \cos(2n\theta).$$

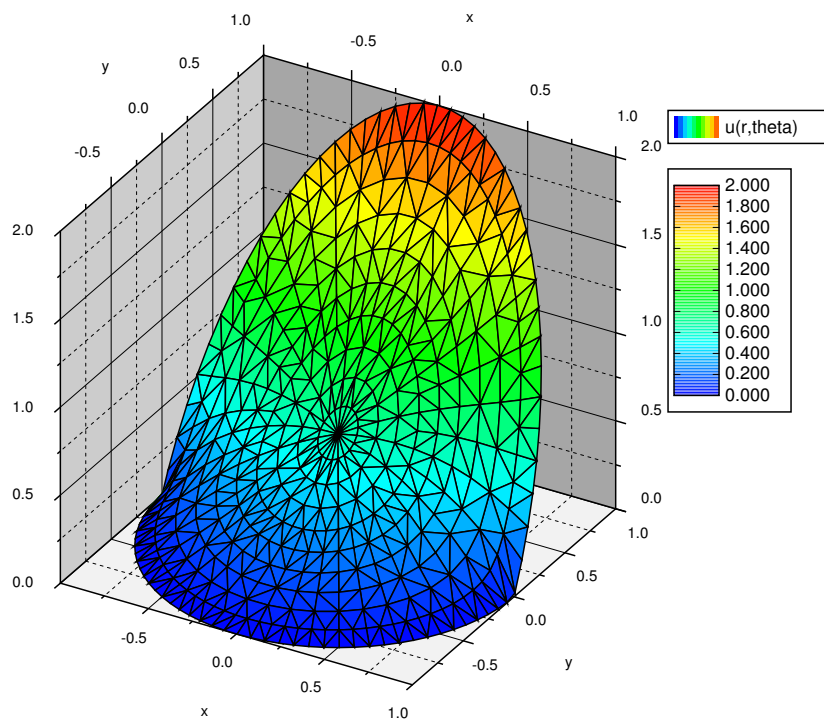


Figure 4.27: The solution of the Dirichlet problem with boundary data 0 for $y < 0$ and $2y$ for $y > 0$.

4.10.3 Poisson kernel

There is another way to solve the Dirichlet problem—with the help of an integral kernel. That is, we will find a function $P(r, \theta, \alpha)$ called the *Poisson kernel*^{*} such that

$$u(r, \theta) = \frac{1}{2\pi} \int_{-\pi}^{\pi} P(r, \theta, \alpha) g(\alpha) d\alpha.$$

While the integral will generally not be solvable analytically, it can be evaluated numerically. In fact, unless the boundary data is given as a Fourier series already, it may be much easier to numerically evaluate this formula as there is only one integral to evaluate.

The formula also has theoretical applications. For instance, as $P(r, \theta, \alpha)$ will have infinitely many derivatives, then via differentiating under the integral, we find that the solution $u(r, \theta)$ has infinitely many derivatives, at least when inside the circle, $r < 1$. By “having infinitely many derivatives,” what you should think of is that $u(r, \theta)$ has “no corners” and all of its partial derivatives of all orders exist and also have “no corners.”

We will compute the formula for $P(r, \theta, \alpha)$ from the series solution, and this idea can be applied anytime you have a convenient series solution where the coefficients are obtained via integration. Hence you can apply this reasoning to obtain such integral kernels for other equations, such as the heat equation. The computation is long and tedious, but not overly difficult. Since the ideas are often applied in similar contexts, it is good to understand how this computation works.

What we do is start with the series solution and replace the coefficients with the integrals that compute them. Then we try to write everything as a single integral. We must use a different dummy variable for the integration and hence we use α instead of θ .

$$\begin{aligned} u(r, \theta) &= \frac{a_0}{2} + \sum_{n=1}^{\infty} \left[a_n r^n \cos(n\theta) + b_n r^n \sin(n\theta) \right] \\ &= \underbrace{\left(\frac{1}{2\pi} \int_{-\pi}^{\pi} g(\alpha) d\alpha \right)}_{\frac{a_0}{2}} + \sum_{n=1}^{\infty} \left[\underbrace{\left(\frac{1}{\pi} \int_{-\pi}^{\pi} g(\alpha) \cos(n\alpha) d\alpha \right)}_{a_n} r^n \cos(n\theta) + \right. \\ &\quad \left. + \underbrace{\left(\frac{1}{\pi} \int_{-\pi}^{\pi} g(\alpha) \sin(n\alpha) d\alpha \right)}_{b_n} r^n \sin(n\theta) \right] \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \left(g(\alpha) + 2 \sum_{n=1}^{\infty} g(\alpha) \cos(n\alpha) r^n \cos(n\theta) + g(\alpha) \sin(n\alpha) r^n \sin(n\theta) \right) d\alpha \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \underbrace{\left(1 + 2 \sum_{n=1}^{\infty} r^n (\cos(n\alpha) \cos(n\theta) + \sin(n\alpha) \sin(n\theta)) \right)}_{P(r, \theta, \alpha)} g(\alpha) d\alpha. \end{aligned}$$

^{*}Named for the French mathematician [Siméon Denis Poisson](#) (1781–1840).

OK, so we have what we wanted: The expression in the parentheses is the Poisson kernel, $P(r, \theta, \alpha)$. However, we can do a lot better. It is still given as a series, and we would really like to have a nice simple expression for it. We must work a little harder. The trick is to rewrite everything in terms of complex exponentials. Let us work just on the kernel.

$$\begin{aligned}
 P(r, \theta, \alpha) &= 1 + 2 \sum_{n=1}^{\infty} r^n (\cos(n\alpha) \cos(n\theta) + \sin(n\alpha) \sin(n\theta)) \\
 &= 1 + 2 \sum_{n=1}^{\infty} r^n \cos(n(\theta - \alpha)) \\
 &= 1 + \sum_{n=1}^{\infty} r^n (e^{in(\theta-\alpha)} + e^{-in(\theta-\alpha)}) \\
 &= 1 + \sum_{n=1}^{\infty} (re^{i(\theta-\alpha)})^n + \sum_{n=1}^{\infty} (re^{-i(\theta-\alpha)})^n.
 \end{aligned}$$

In the expression above, we recognize the *geometric series*. Recall from calculus that if z is a complex number where $|z| < 1$, then

$$\sum_{n=1}^{\infty} z^n = \frac{z}{1-z}.$$

Note that n starts at 1, and that is why we have the z in the numerator. It is the standard geometric series multiplied by z . We can use $z = re^{i(\theta-\alpha)}$, as lo and behold $|re^{i(\theta-\alpha)}| = r < 1$. We continue with the computation.

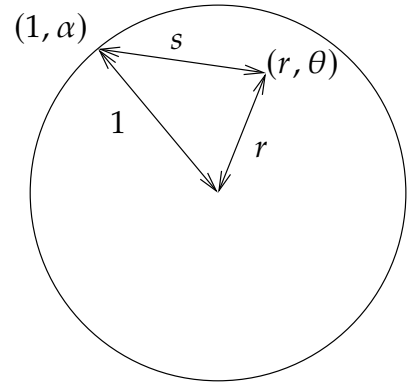
$$\begin{aligned}
 P(r, \theta, \alpha) &= 1 + \sum_{n=1}^{\infty} (re^{i(\theta-\alpha)})^n + \sum_{n=1}^{\infty} (re^{-i(\theta-\alpha)})^n \\
 &= 1 + \frac{re^{i(\theta-\alpha)}}{1 - re^{i(\theta-\alpha)}} + \frac{re^{-i(\theta-\alpha)}}{1 - re^{-i(\theta-\alpha)}} \\
 &= \frac{(1 - re^{i(\theta-\alpha)})(1 - re^{-i(\theta-\alpha)}) + (1 - re^{-i(\theta-\alpha)})re^{i(\theta-\alpha)} + (1 - re^{i(\theta-\alpha)})re^{-i(\theta-\alpha)}}{(1 - re^{i(\theta-\alpha)})(1 - re^{-i(\theta-\alpha)})} \\
 &= \frac{1 - r^2}{1 - re^{i(\theta-\alpha)} - re^{-i(\theta-\alpha)} + r^2} \\
 &= \frac{1 - r^2}{1 - 2r \cos(\theta - \alpha) + r^2}.
 \end{aligned}$$

That is a formula we can live with. The solution to the Dirichlet problem using the Poisson kernel is

$$u(r, \theta) = \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1 - r^2}{1 - 2r \cos(\theta - \alpha) + r^2} g(\alpha) d\alpha.$$

Sometimes the formula for the Poisson kernel is given together with the constant $\frac{1}{2\pi}$, in which case we should, of course, not leave it in front of the integral. Sometimes the limits of the integral are given as 0 to 2π ; everything inside is 2π -periodic in α , so this does not change the integral.

Let us not leave the Poisson kernel without explaining its geometric meaning. The kernel corresponding to the point (r, θ) in the disc and the point $(1, \alpha)$ on the boundary depends only on the distance of (r, θ) from the center and the distance between (r, θ) and $(1, \alpha)$. Let s be the distance from (r, θ) to $(1, \alpha)$. This distance s in polar coordinates is given precisely by the square root of $1 - 2r \cos(\theta - \alpha) + r^2$. That is, the Poisson kernel is really the formula



$$\frac{1 - r^2}{s^2}.$$

One final note we make about the formula is that it is really a weighted average of the boundary values. First, we look at what happens at the origin, that is, when $r = 0$:

$$\begin{aligned} u(0, 0) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1 - 0^2}{1 - 2(0) \cos(0 - \alpha) + 0^2} g(\alpha) d\alpha \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} g(\alpha) d\alpha. \end{aligned}$$

So $u(0, 0)$ is precisely the average value of $g(\theta)$ and therefore the average value of u on the boundary. This is a general feature of harmonic functions, the value at some point p is equal to the average of the values on a circle centered at p .

What the formula says at other points inside the circle is that the value of the solution is a weighted average of the boundary data $g(\theta)$. The kernel is bigger when $(1, \alpha)$ is closer to (r, θ) . Therefore, when computing $u(r, \theta)$, we give more weight to the values $g(\alpha)$ when $(1, \alpha)$ is closer to (r, θ) and less weight to the values $g(\alpha)$ when $(1, \alpha)$ is far from (r, θ) .

4.10.4 Exercises

Exercise 4.10.2: Using series, solve $\Delta u = 0$, $u(1, \theta) = |\theta|$, for $-\pi < \theta \leq \pi$.

Exercise 4.10.3: Using series solve $\Delta u = 0$, $u(1, \theta) = g(\theta)$ for the following data. Hint: trig identities.

a) $g(\theta) = 1/2 + 3 \sin(\theta) + \cos(3\theta)$

b) $g(\theta) = 3 \cos(3\theta) + 3 \sin(3\theta) + \sin(9\theta)$

c) $g(\theta) = 2 \cos(\theta + 1)$

d) $g(\theta) = \sin^2(\theta)$

Exercise 4.10.4: Using the Poisson kernel, give the solution to $\Delta u = 0$, where $u(1, \theta)$ is zero for θ outside the interval $[-\pi/4, \pi/4]$ and $u(1, \theta)$ is 1 for θ on the interval $[-\pi/4, \pi/4]$.

Exercise 4.10.5:

- a) Draw a graph for the Poisson kernel as a function of α when $r = 1/2$ and $\theta = 0$.
- b) Describe what happens to the graph when you make r bigger (as it approaches 1).
- c) Knowing that the solution $u(r, \theta)$ is the weighted average of $g(\theta)$ with the Poisson kernel as the weight, explain what your answer to part b) means.

Exercise 4.10.6: Let $g(\theta)$ be the function $xy = \cos \theta \sin \theta$ on the boundary. Use the series solution to find a solution to the Dirichlet problem $\Delta u = 0$, $u(1, \theta) = g(\theta)$. Now convert the solution to Cartesian coordinates x and y . Is this solution surprising? Hint: Use your trig identities.

Exercise 4.10.7: Carry out the computation we needed in the separation of variables and solve $r^2 R'' + rR' - n^2 R = 0$, for $n = 0, 1, 2, 3, \dots$

Exercise 4.10.8 (challenging): Derive the series solution to the Dirichlet problem if the region is a circle of radius ρ rather than 1. That is, solve $\Delta u = 0$, $u(\rho, \theta) = g(\theta)$.

Exercise 4.10.9 (challenging):

- a) Find the solution for $\Delta u = 0$, $u(1, \theta) = x^2 y^3 + 5x^2$. Write the answer in Cartesian coordinates.
- b) Now solve $\Delta u = 0$, $u(1, \theta) = x^k y^\ell$. Write the solution in Cartesian coordinates.
- c) Given a polynomial $P(x, y) = \sum_{j=0}^m \sum_{k=0}^n c_{j,k} x^j y^k$, solve $\Delta u = 0$, $u(1, \theta) = P(x, y)$ (that is, write down the formula for the answer). Write the answer in Cartesian coordinates.

Notice the answer is again a polynomial in x and y . See also [Exercise 4.10.6](#).

Exercise 4.10.101: Using series solve $\Delta u = 0$, $u(1, \theta) = 1 + \sum_{n=1}^{\infty} \frac{1}{n^2} \sin(n\theta)$.

Exercise 4.10.102: Using the series solution find the solution to $\Delta u = 0$, $u(1, \theta) = 1 - \cos(\theta)$. Express the solution in Cartesian coordinates (that is, using x and y).

Exercise 4.10.103:

- a) Try and guess a solution to $\Delta u = -1$, $u(1, \theta) = 0$. Hint: Try a solution that only depends on r . Also, don't worry about the boundary condition at first.
- b) Now solve $\Delta u = -1$, $u(1, \theta) = \sin(2\theta)$ using superposition.

Exercise 4.10.104 (challenging): Derive the Poisson kernel solution if the region is a circle of radius ρ rather than 1. That is, solve $\Delta u = 0$, $u(\rho, \theta) = g(\theta)$.

Chapter 5

More on eigenvalue problems

5.1 Sturm–Liouville problems

Note: 2 lectures, §10.1 in [EP], §11.2 in [BD]

5.1.1 Boundary value problems

In [chapter 4](#), we encountered several different eigenvalue problems such as:

$$X''(x) + \lambda X(x) = 0,$$

with different boundary conditions

$$\begin{array}{lll} X(0) = 0 & X(L) = 0 & \text{(Dirichlet), or} \\ X'(0) = 0 & X'(L) = 0 & \text{(Neumann), or} \\ X'(0) = 0 & X(L) = 0 & \text{(Mixed), or} \\ X(0) = 0 & X'(L) = 0 & \text{(Mixed), } \dots \end{array}$$

These boundary problems came up in the study of the heat equation $u_t = ku_{xx}$ when we were trying to solve the equation by the method of separation of variables in [§ 4.6](#). Dirichlet conditions correspond to applying a zero temperature at the ends, Neumann means insulating the ends, etc. Other types of endpoint conditions also arise naturally, such as the *Robin boundary conditions*

$$hX(0) - X'(0) = 0, \quad hX(L) + X'(L) = 0,$$

for some constant h . These conditions come up when the ends are immersed in some medium.

In the separation of variables computation we encountered an eigenvalue problem and found the eigenfunctions $X_n(x)$. We then found the *eigenfunction decomposition* of the initial temperature $f(x) = u(x, 0)$,

$$f(x) = \sum_{n=1}^{\infty} c_n X_n(x).$$

Once we had this decomposition and found suitable $T_n(t)$ such that $T_n(0) = 1$ and such that $T_n(t)X_n(x)$ were solutions to the heat equation, we wrote the solution to the original problem, including the initial condition, as

$$u(x, t) = \sum_{n=1}^{\infty} c_n T_n(t) X_n(x).$$

To study more general problems with this method, we must study more general eigenvalue problems. First, we study second-order linear equations of the form

$$\frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - q(x)y + \lambda r(x)y = 0. \quad (5.1)$$

Essentially any second-order linear equation of the form $a(x)y'' + b(x)y' + c(x)y + \lambda d(x)y = 0$ can be written as (5.1) after multiplying by a proper factor.

Example 5.1.1 (Bessel): Put the following equation into the form (5.1):

$$x^2 y'' + x y' + (\lambda x^2 - n^2) y = 0.$$

Multiply both sides by $\frac{1}{x}$ to obtain

$$\begin{aligned} \frac{1}{x} (x^2 y'' + x y' + (\lambda x^2 - n^2) y) &= x y'' + y' + \left(\lambda x - \frac{n^2}{x} \right) y \\ &= \frac{d}{dx} \left(x \frac{dy}{dx} \right) - \frac{n^2}{x} y + \lambda x y = 0. \end{aligned}$$

The Bessel equation turns up for example in the solution of the two-dimensional wave equation. If you want to see how one solves the equation, you can look at [subsection 7.3.3](#).

The so-called *Sturm–Liouville problem*^{*} is to seek nontrivial solutions to

$$\begin{aligned} \frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - q(x)y + \lambda r(x)y &= 0, & a < x < b, \\ \alpha_1 y(a) - \alpha_2 y'(a) &= 0, \\ \beta_1 y(b) + \beta_2 y'(b) &= 0. \end{aligned} \quad (5.2)$$

By nontrivial, we again mean a y that is not simply the constant zero (which is always a solution). In particular, we ask which λ s allow such nontrivial solutions. The λ s that admit nontrivial solutions are called the *eigenvalues* and the corresponding nontrivial solutions are called *eigenfunctions*. The constants α_1 and α_2 should not both be zero; the same applies to β_1 and β_2 .

^{*}Named after the French mathematicians [Jacques Charles François Sturm](#) (1803–1855) and [Joseph Liouville](#) (1809–1882).

Theorem 5.1.1. Suppose $p(x)$, $p'(x)$, $q(x)$ and $r(x)$ are continuous on $[a, b]$ and suppose $p(x) > 0$ and $r(x) > 0$ for all x in $[a, b]$. Then the Sturm–Liouville problem (5.2) has an increasing sequence of eigenvalues

$$\lambda_1 < \lambda_2 < \lambda_3 < \cdots$$

such that

$$\lim_{n \rightarrow \infty} \lambda_n = +\infty$$

and such that to each λ_n there is (up to a constant multiple) a single eigenfunction $y_n(x)$.

Moreover, if $q(x) \geq 0$ and $\alpha_1, \alpha_2, \beta_1, \beta_2 \geq 0$, then $\lambda_n \geq 0$ for all n .

Problems satisfying the hypothesis of the theorem (including the “Moreover”) are called *regular Sturm–Liouville problems*, and we will only consider such problems here. That is, a regular problem is one where $p(x)$, $p'(x)$, $q(x)$ and $r(x)$ are continuous, $p(x) > 0$, $r(x) > 0$, $q(x) \geq 0$, and $\alpha_1, \alpha_2, \beta_1, \beta_2 \geq 0$, where neither $\alpha_1 = \alpha_2 = 0$ nor $\beta_1 = \beta_2 = 0$. Note: Be careful about the signs. Also be careful about the inequalities for r and p . They must be strict for all x in the interval $[a, b]$, including the endpoints!

When zero is an eigenvalue, we usually start labeling the eigenvalues from 0 rather than from 1 for convenience. That is, we label the eigenvalues $\lambda_0 < \lambda_1 < \lambda_2 < \cdots$.

Example 5.1.2: The problem $y'' + \lambda y = 0$, $0 < x < L$, $y(0) = 0$, and $y(L) = 0$ is a regular Sturm–Liouville problem: $p(x) = 1$, $q(x) = 0$, $r(x) = 1$, and we have $p(x) = 1 > 0$ and $r(x) = 1 > 0$. We also have $a = 0$, $b = L$, $\alpha_1 = \beta_1 = 1$, $\alpha_2 = \beta_2 = 0$. The eigenvalues are $\lambda_n = \frac{n^2\pi^2}{L^2}$ and eigenfunctions are $y_n(x) = \sin(\frac{n\pi}{L}x)$. All eigenvalues are nonnegative as predicted by the theorem.

Exercise 5.1.1: Find eigenvalues and eigenfunctions for

$$y'' + \lambda y = 0, \quad y'(0) = 0, \quad y'(1) = 0.$$

Identify the $p, q, r, \alpha_j, \beta_j$. Can you use the theorem above to make the search for eigenvalues easier? Hint: Consider the condition $-y'(0) = 0$.

Example 5.1.3: Find eigenvalues and eigenfunctions of the problem

$$\begin{aligned} y'' + \lambda y &= 0, & 0 < x < 1, \\ hy(0) - y'(0) &= 0, & y'(1) = 0, & h > 0. \end{aligned}$$

These equations give a regular Sturm–Liouville problem.

Exercise 5.1.2: Identify $p, q, r, \alpha_j, \beta_j$ in the example above.

By Theorem 5.1.1, $\lambda \geq 0$. So the general solution (without boundary conditions) is

$$\begin{aligned} y(x) &= A \cos(\sqrt{\lambda} x) + B \sin(\sqrt{\lambda} x) & \text{if } \lambda > 0, \\ y(x) &= Ax + B & \text{if } \lambda = 0. \end{aligned}$$

Let us see if $\lambda = 0$ is an eigenvalue: We must satisfy $0 = hB - A$ and $A = 0$, hence $B = 0$ (as $h > 0$). Therefore, 0 is not an eigenvalue (no nonzero solution, so no eigenfunction).

Now we try $\lambda > 0$. We plug in the boundary conditions:

$$\begin{aligned} 0 &= hA - \sqrt{\lambda} B, \\ 0 &= -A\sqrt{\lambda} \sin(\sqrt{\lambda}) + B\sqrt{\lambda} \cos(\sqrt{\lambda}). \end{aligned}$$

If $A = 0$, then $B = 0$ and vice versa; therefore, both are nonzero. So $B = \frac{hA}{\sqrt{\lambda}}$, and $0 = -A\sqrt{\lambda} \sin(\sqrt{\lambda}) + \frac{hA}{\sqrt{\lambda}} \sqrt{\lambda} \cos(\sqrt{\lambda})$. As $A \neq 0$, we get

$$0 = -\sqrt{\lambda} \sin(\sqrt{\lambda}) + h \cos(\sqrt{\lambda}),$$

or

$$\frac{h}{\sqrt{\lambda}} = \tan \sqrt{\lambda}.$$

We use a computer to find λ_n . There are tables available, though using a computer or a graphing calculator is far more convenient nowadays. The easiest method is to plot the functions h/x and $\tan x$ and see for which x they intersect. There is an infinite number of intersections. Denote the first intersection by $\sqrt{\lambda_1}$, the second intersection by $\sqrt{\lambda_2}$, etc. For example, when $h = 1$, we get $\sqrt{\lambda_1} \approx 0.86$, $\sqrt{\lambda_2} \approx 3.43$, That is $\lambda_1 \approx 0.74$, $\lambda_2 \approx 11.73$, A plot for $h = 1$ is given in [Figure 5.1](#) on the facing page. The appropriate eigenfunction (let $A = 1$ for convenience, then $B = h/\sqrt{\lambda}$) is

$$y_n(x) = \cos(\sqrt{\lambda_n} x) + \frac{h}{\sqrt{\lambda_n}} \sin(\sqrt{\lambda_n} x).$$

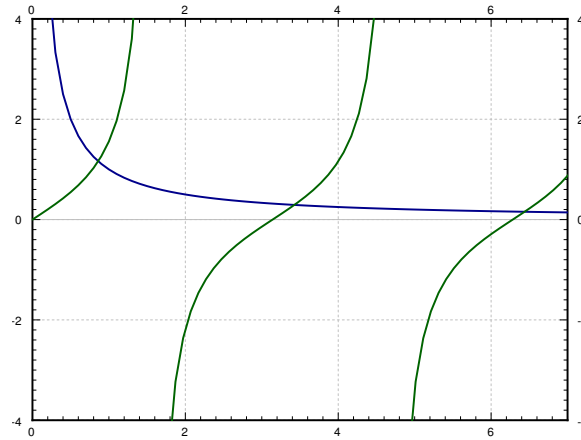
When $h = 1$ we get (approximately)

$$y_1(x) \approx \cos(0.86 x) + \frac{1}{0.86} \sin(0.86 x), \quad y_2(x) \approx \cos(3.43 x) + \frac{1}{3.43} \sin(3.43 x), \quad \dots$$

5.1.2 Orthogonality

We have seen the notion of orthogonality before. For example, we have shown that $\sin(nx)$ are orthogonal for distinct n on $[0, \pi]$. For general Sturm–Liouville problems we need a more general setup. Let $r(x)$ be a *weight function* (any function, though generally we assume it is positive) on $[a, b]$. Two functions $f(x)$ and $g(x)$ are said to be *orthogonal* with respect to the weight function $r(x)$ when

$$\int_a^b f(x) g(x) r(x) dx = 0.$$

Figure 5.1: Plot of $\frac{1}{x}$ and $\tan x$.

In this setting, we define the *inner product* as

$$\langle f, g \rangle \stackrel{\text{def}}{=} \int_a^b f(x) g(x) r(x) dx,$$

and then say f and g are orthogonal whenever $\langle f, g \rangle = 0$. The results and concepts are again analogous to finite-dimensional linear algebra.

The idea of the given inner product is that those x where $r(x)$ is greater have more weight. Nontrivial (nonconstant) $r(x)$ arise naturally, for example from a change of variables. Hence, you could think of a change of variables such that $d\xi = r(x) dx$.

Eigenfunctions of a regular Sturm–Liouville problem satisfy an orthogonality property, just like the eigenfunctions in § 4.1. Its proof is very similar to that of the analogous Theorem 4.1.1 on page 193.

Theorem 5.1.2. *Suppose we have a regular Sturm–Liouville problem*

$$\begin{aligned} \frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - q(x)y + \lambda r(x)y &= 0, \\ \alpha_1 y(a) - \alpha_2 y'(a) &= 0, \\ \beta_1 y(b) + \beta_2 y'(b) &= 0. \end{aligned}$$

Let y_j and y_k be two distinct eigenfunctions for two distinct eigenvalues λ_j and λ_k . Then

$$\int_a^b y_j(x) y_k(x) r(x) dx = 0,$$

that is, y_j and y_k are orthogonal with respect to the weight function r .

5.1.3 Fredholm alternative

The *Fredholm alternative* theorem we talked about before ([Theorem 4.1.2](#) on page 194) holds for all regular Sturm–Liouville problems. We state it here for completeness.

Theorem 5.1.3 (Fredholm alternative). *Suppose that we have a regular Sturm–Liouville problem. Then either*

$$\begin{aligned} \frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - q(x)y + \lambda r(x)y &= 0, \\ \alpha_1 y(a) - \alpha_2 y'(a) &= 0, \\ \beta_1 y(b) + \beta_2 y'(b) &= 0, \end{aligned}$$

has a nonzero solution (λ is an eigenvalue), or

$$\begin{aligned} \frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - q(x)y + \lambda r(x)y &= f(x), \\ \alpha_1 y(a) - \alpha_2 y'(a) &= 0, \\ \beta_1 y(b) + \beta_2 y'(b) &= 0, \end{aligned}$$

has a unique solution for any $f(x)$ continuous on $[a, b]$.

This theorem is used in much the same way as we did before in [§ 4.4](#). It is used when solving more general nonhomogeneous boundary value problems. The theorem does not help us solve the problem. It tells us when a unique solution exists, so that we know when to spend time looking for it. To solve the problem, we decompose $f(x)$ and $y(x)$ in terms of eigenfunctions of the homogeneous problem, and then solve for the coefficients of the series for $y(x)$.

5.1.4 Eigenfunction series

What we want to do with the eigenfunctions once we have them is to compute the *eigenfunction decomposition* of an arbitrary function $f(x)$. That is, we wish to write

$$f(x) = \sum_{n=1}^{\infty} c_n y_n(x), \tag{5.3}$$

where $y_n(x)$ are eigenfunctions. We wish to find out if we can represent any function $f(x)$ in this way, and if so, we wish to calculate the c_n (and of course we would want to know if the sum converges). OK, so imagine we could write $f(x)$ as (5.3). We will assume convergence and the ability to integrate the series term by term. Because of orthogonality

we have

$$\begin{aligned}
 \langle f, y_m \rangle &= \int_a^b f(x) y_m(x) r(x) dx = \int_a^b \left(\sum_{n=1}^{\infty} c_n y_n(x) \right) y_m(x) r(x) dx \\
 &= \sum_{n=1}^{\infty} c_n \int_a^b y_n(x) y_m(x) r(x) dx \\
 &= c_m \int_a^b y_m(x) y_m(x) r(x) dx = c_m \langle y_m, y_m \rangle.
 \end{aligned}$$

Hence,

$$c_m = \frac{\langle f, y_m \rangle}{\langle y_m, y_m \rangle} = \frac{\int_a^b f(x) y_m(x) r(x) dx}{\int_a^b (y_m(x))^2 r(x) dx}. \quad (5.4)$$

Note that y_m are known up to a constant multiple, so we could have picked a scalar multiple of an eigenfunction such that $\langle y_m, y_m \rangle = 1$ (if we had an arbitrary eigenfunction $\tilde{y}_m$, divide it by $\sqrt{\langle \tilde{y}_m, \tilde{y}_m \rangle}$). When $\langle y_m, y_m \rangle = 1$, we have the simpler form $c_m = \langle f, y_m \rangle$. The following theorem holds more generally, but the statement given is enough for our purposes.

Theorem 5.1.4. Suppose f is a piecewise smooth continuous function on $[a, b]$. If $y_1, y_2, \dots$ are eigenfunctions of a regular Sturm–Liouville problem, one for each eigenvalue, then there exist real constants $c_1, c_2, \dots$ given by (5.4) such that (5.3) converges and holds for $a < x < b$.

Example 5.1.4: Consider

$$\begin{aligned}
 y'' + \lambda y &= 0, \quad 0 < x < \pi/2, \\
 y(0) &= 0, \quad y'(\pi/2) = 0.
 \end{aligned}$$

The above is a regular Sturm–Liouville problem, and Theorem 5.1.1 on page 275 says that if λ is an eigenvalue then $\lambda \geq 0$.

Suppose $\lambda = 0$. The general solution is $y(x) = Ax + B$. We plug in the initial conditions to get $0 = y(0) = B$, and $0 = y'(\pi/2) = A$. Hence $\lambda = 0$ is not an eigenvalue.

So let us consider $\lambda > 0$, where the general solution is

$$y(x) = A \cos(\sqrt{\lambda} x) + B \sin(\sqrt{\lambda} x).$$

Plugging in the boundary conditions we get $0 = y(0) = A$ and $0 = y'(\pi/2) = \sqrt{\lambda} B \cos(\sqrt{\lambda} \frac{\pi}{2})$. Since A is zero, B cannot be zero. Hence $\cos(\sqrt{\lambda} \frac{\pi}{2}) = 0$. This means that $\sqrt{\lambda} \frac{\pi}{2}$ is an odd integral multiple of $\pi/2$, i.e. $(2n - 1)\frac{\pi}{2} = \sqrt{\lambda_n} \frac{\pi}{2}$. Solving for λ_n we get

$$\lambda_n = (2n - 1)^2.$$

We can take $B = 1$. Our eigenfunctions are

$$y_n(x) = \sin((2n-1)x).$$

A little bit of calculus shows

$$\int_0^{\pi/2} \left(\sin((2n-1)x) \right)^2 dx = \frac{\pi}{4}.$$

So any piecewise smooth function $f(x)$ on $[0, \pi/2]$ can be written as

$$f(x) = \sum_{n=1}^{\infty} c_n \sin((2n-1)x),$$

where

$$c_n = \frac{\langle f, y_n \rangle}{\langle y_n, y_n \rangle} = \frac{\int_0^{\pi/2} f(x) \sin((2n-1)x) dx}{\int_0^{\pi/2} \left(\sin((2n-1)x) \right)^2 dx} = \frac{4}{\pi} \int_0^{\pi/2} f(x) \sin((2n-1)x) dx.$$

Note that the series converges to an odd 2π -periodic extension of $f(x)$. With the regular sine series, we would expect a function with period $2 \frac{\pi}{2} = \pi$.

Exercise 5.1.3 (challenging): *In the example above, the function is defined on $0 < x < \pi/2$, yet the series with respect to the eigenfunctions $\sin((2n-1)x)$ converges to an odd 2π -periodic extension of $f(x)$. Find out how the extension is defined for $\pi/2 < x < \pi$.*

Let us compute an example. Consider $f(x) = x$ for $0 < x < \pi/2$. Some calculus later we find

$$c_n = \frac{4}{\pi} \int_0^{\pi/2} f(x) \sin((2n-1)x) dx = \frac{4(-1)^{n+1}}{\pi(2n-1)^2},$$

and so for x in $[0, \pi/2]$,

$$f(x) = \sum_{n=1}^{\infty} \frac{4(-1)^{n+1}}{\pi(2n-1)^2} \sin((2n-1)x).$$

This is different from the π -periodic regular sine series which can be computed to be

$$f(x) = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} \sin(2nx).$$

Both series converge to $f(x)$ for $0 < x < \pi/2$, but the eigenfunctions involved come from different eigenvalue problems.

5.1.5 Exercises

Exercise 5.1.4: Find eigenvalues and eigenfunctions of

$$y'' + \lambda y = 0, \quad y(0) - y'(0) = 0, \quad y(1) = 0.$$

Exercise 5.1.5: Expand the function $f(x) = x$ on $0 \leq x \leq 1$ using eigenfunctions of the system

$$y'' + \lambda y = 0, \quad y'(0) = 0, \quad y(1) = 0.$$

Exercise 5.1.6: Suppose that you had a Sturm–Liouville problem on the interval $[0, 1]$ and came up with $y_n(x) = \sin(\gamma n x)$, where $\gamma > 0$ is some constant. Decompose $f(x) = x$, $0 < x < 1$ in terms of these eigenfunctions.

Exercise 5.1.7: Find eigenvalues and eigenfunctions of

$$y^{(4)} + \lambda y = 0, \quad y(0) = 0, \quad y'(0) = 0, \quad y(1) = 0, \quad y'(1) = 0.$$

This problem is not a Sturm–Liouville problem, but the idea is the same.

Exercise 5.1.8 (more challenging): Find eigenvalues and eigenfunctions for

$$\frac{d}{dx}(e^x y') + \lambda e^x y = 0, \quad y(0) = 0, \quad y(1) = 0.$$

Hint: First write the system as a constant-coefficient system to find general solutions. Do note that [Theorem 5.1.1](#) on page 275 guarantees $\lambda \geq 0$.

Exercise 5.1.101: Find eigenvalues and eigenfunctions of

$$y'' + \lambda y = 0, \quad y(-1) = 0, \quad y(1) = 0.$$

Exercise 5.1.102: Put the following problems into the standard form for Sturm–Liouville problems, that is, find $p(x)$, $q(x)$, $r(x)$, α_1 , α_2 , β_1 , and β_2 , and decide if the problems are regular or not.

a) $xy'' + \lambda y = 0$ for $0 < x < 1$, $y(0) = 0$, $y(1) = 0$.

b) $(1 + x^2)y'' + 2xy' + (\lambda - x^2)y = 0$ for $-1 < x < 1$, $y(-1) = 0$, $y(1) + y'(1) = 0$.*

*In an earlier version of this book, a typo rendered the equation as $(1 + x^2)y'' - 2xy' + (\lambda - x^2)y = 0$ ending up with something harder than intended. Try this equation for a further challenge.

5.2 Higher-order eigenvalue problems

Note: 1 lecture, §10.2 in [EP], exercises in §11.2 in [BD]

The eigenfunction series can arise even from higher-order equations. Consider an elastic beam (say made of steel). We will study the *transversal vibrations* (vibrations where the beam “bends”) of the beam. Suppose the beam lies along the x -axis and let $y(x, t)$ measure the displacement of the point x on the beam at time t . See Figure 5.2.

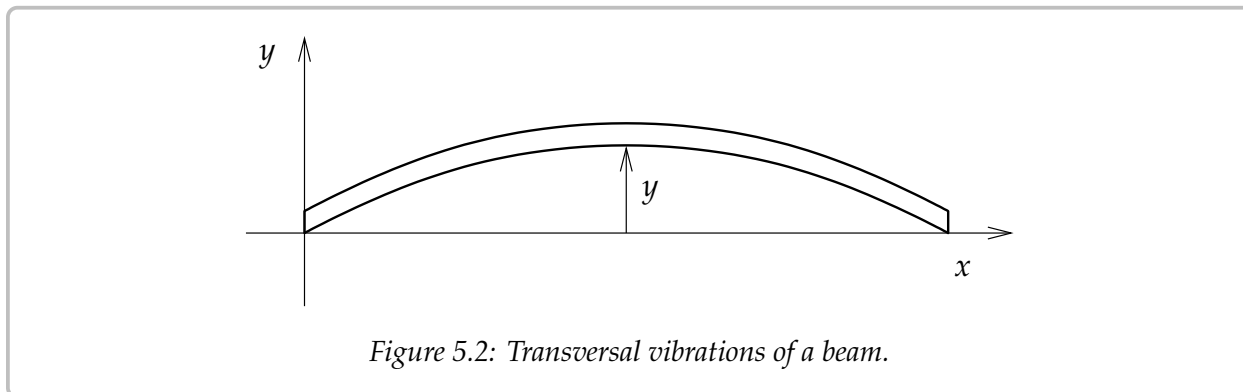


Figure 5.2: Transversal vibrations of a beam.

The equation that governs this setup is

$$a^4 \frac{\partial^4 y}{\partial x^4} + \frac{\partial^2 y}{\partial t^2} = 0,$$

for some constant $a > 0$. Let us not worry about the physics*.

Suppose the beam is of length 1 and simply supported (hinged) at the ends. The beam is displaced by some function $f(x)$ at time $t = 0$ and then let go (initial velocity is 0). Then y satisfies:

$$\begin{aligned} a^4 y_{xxxx} + y_{tt} &= 0 & (0 < x < 1, \ t > 0), \\ y(0, t) &= y_{xx}(0, t) = 0, \\ y(1, t) &= y_{xx}(1, t) = 0, \\ y(x, 0) &= f(x), \quad y_t(x, 0) = 0. \end{aligned} \tag{5.5}$$

Again we try $y(x, t) = X(x)T(t)$ and plug in to get $a^4 X^{(4)}T + XT'' = 0$ or

$$\frac{X^{(4)}}{X} = \frac{-T''}{a^4 T} = \lambda.$$

The equations are

$$T'' + \lambda a^4 T = 0 \quad \text{and} \quad X^{(4)} - \lambda X = 0.$$

*If you are interested, $a^4 = \frac{EI}{\rho}$, where E is the elastic modulus, I is the second moment of area of the cross section, and ρ is linear density.

The boundary conditions $y(0, t) = y_{xx}(0, t) = 0$ and $y(1, t) = y_{xx}(1, t) = 0$ imply

$$X(0) = X''(0) = 0 \quad \text{and} \quad X(1) = X''(1) = 0.$$

The initial homogeneous condition $y_t(x, 0) = 0$ implies

$$T'(0) = 0.$$

As usual, we leave the nonhomogeneous $y(x, 0) = f(x)$ for later.

Considering the equation for T , that is, $T'' + \lambda a^4 T = 0$, physical intuition leads us to the fact that if λ is an eigenvalue then $\lambda > 0$: We expect vibration and not exponential growth or decay in the t direction (there is no friction in our model for instance). So there are no negative eigenvalues. Similarly, $\lambda = 0$ is not an eigenvalue.

Exercise 5.2.1: Justify $\lambda > 0$ just from the equation for X and the boundary conditions.

Let $\omega = \sqrt[4]{\lambda}$, that is, $\omega^4 = \lambda$, to avoid writing the fourth root all the time. Notice $\omega > 0$. The general solution to $X^{(4)} - \omega^4 X = 0$ is

$$X(x) = Ae^{\omega x} + Be^{-\omega x} + C \sin(\omega x) + D \cos(\omega x).$$

Now $0 = X(0) = A + B + D$, $0 = X''(0) = \omega^2(A + B - D)$. Solving, $D = 0$ and $B = -A$. So

$$X(x) = Ae^{\omega x} - Ae^{-\omega x} + C \sin(\omega x).$$

Also $0 = X(1) = A(e^\omega - e^{-\omega}) + C \sin \omega$, and $0 = X''(1) = A\omega^2(e^\omega - e^{-\omega}) - C\omega^2 \sin \omega$. This means that $C \sin \omega = 0$ and $A(e^\omega - e^{-\omega}) = 2A \sinh \omega = 0$. If $\omega > 0$, then $\sinh \omega \neq 0$ and so $A = 0$. Thus $C \neq 0$, otherwise λ is not an eigenvalue. Also, ω must be an integer multiple of π . Hence $\omega = n\pi$ and $n \geq 1$ (as $\omega > 0$). We can take $C = 1$. So the eigenvalues are $\lambda_n = n^4\pi^4$, and corresponding eigenfunctions are $\sin(n\pi x)$.

Now $T'' + n^4\pi^4 a^4 T = 0$. The general solution is $T(t) = A \sin(n^2\pi^2 a^2 t) + B \cos(n^2\pi^2 a^2 t)$. But $T'(0) = 0$ and hence $A = 0$. We take $B = 1$ to make $T(0) = 1$ for convenience. So our solutions are $T_n(t) = \cos(n^2\pi^2 a^2 t)$.

The eigenfunctions are just the sines, so we decompose the function $f(x)$ using the sine series. That is, we find numbers b_n such that for $0 < x < 1$,

$$f(x) = \sum_{n=1}^{\infty} b_n \sin(n\pi x).$$

Then the solution to (5.5) is

$$y(x, t) = \sum_{n=1}^{\infty} b_n X_n(x) T_n(t) = \sum_{n=1}^{\infty} b_n \sin(n\pi x) \cos(n^2\pi^2 a^2 t).$$

The point is that $X_n T_n$ is a solution that satisfies all the homogeneous conditions (all conditions except the initial position). And since $T_n(0) = 1$,

$$y(x, 0) = \sum_{n=1}^{\infty} b_n X_n(x) T_n(0) = \sum_{n=1}^{\infty} b_n X_n(x) = \sum_{n=1}^{\infty} b_n \sin(n\pi x) = f(x).$$

So $y(x, t)$ solves (5.5).

The natural (angular) frequencies of the system are $n^2\pi^2 a^2$. These frequencies are all integer multiples of the fundamental frequency $\pi^2 a^2$, so we get a nice musical note. The exact frequencies and their amplitudes are what musicians call the *timbre* of the note (outside of music it is called the spectrum).

The timbre of a simply supported beam is different from that of a vibrating string, where we get “more” of the lower frequencies since we get all integer multiples, $1, 2, 3, 4, 5, \dots$. For a simply supported beam, we get only the square multiples $1, 4, 9, 16, 25, \dots$, and so we get something much closer to a pure sound. The sound of a xylophone or vibraphone is, therefore, very different from that of a guitar or piano.

Example 5.2.1: Consider $f(x) = \frac{x(x-1)}{10}$. On $0 < x < 1$ (you know how to do this by now)

$$f(x) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{5\pi^3 n^3} \sin(n\pi x).$$

Hence, the solution to (5.5) with the given initial position $f(x)$ is

$$y(x, t) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{5\pi^3 n^3} \sin(n\pi x) \cos(n^2\pi^2 a^2 t).$$

There are other boundary conditions than just hinged ends. There are three basic possibilities: hinged, free, or fixed. Consider the end at $x = 0$. For the other end, it is the same idea. If the end is *hinged*, then

$$y(0, t) = y_{xx}(0, t) = 0.$$

If the end is *free*, that is, it is just floating in air, then

$$y_{xx}(0, t) = y_{xxx}(0, t) = 0.$$

And finally, if the end is *clamped* or *fixed*, for example it is welded to a wall, then

$$y(0, t) = y_x(0, t) = 0.$$

5.2.1 Exercises

Exercise 5.2.2: Suppose you have a beam of length 5 with free ends. Let y be the transverse deviation of the beam at position x on the beam ($0 < x < 5$). You know that the constants are such that this satisfies the equation $y_{tt} + 4y_{xxxx} = 0$. Suppose you know that the initial shape of the beam is the graph of $x(5 - x)$, and the initial velocity is uniformly equal to 2 (same for each x) in the positive y direction. Set up the equation together with the boundary and initial conditions. Just set up, do not solve.

Exercise 5.2.3: Suppose you have a beam of length 5 with one end free and one end fixed (the fixed end is at $x = 5$). Let u be the longitudinal deviation of the beam at position x on the beam ($0 < x < 5$). You know that the constants are such that this satisfies the equation $u_{tt} = 4u_{xx}$. Suppose you know that the initial displacement of the beam is $\frac{x-5}{50}$, and the initial velocity is $\frac{-(x-5)}{100}$ in the positive u direction. Set up the equation together with the boundary and initial conditions. Just set up, do not solve.

Exercise 5.2.4: Suppose the beam is L units long, everything else kept the same as in (5.5). What is the equation and the series solution?

Exercise 5.2.5: Suppose you have

$$\begin{aligned} a^4 y_{xxxx} + y_{tt} &= 0 \quad (0 < x < 1, t > 0), \\ y(0, t) &= y_{xx}(0, t) = 0, \\ y(1, t) &= y_{xx}(1, t) = 0, \\ y(x, 0) &= f(x), \quad y_t(x, 0) = g(x). \end{aligned}$$

That is, you also have an initial velocity. Find a series solution. Hint: Use the same idea as we did for the wave equation.

Exercise 5.2.101: Suppose you have a beam of length 1 with hinged ends. Let y be the transverse deviation of the beam at position x on the beam ($0 < x < 1$). You know that the constants are such that this satisfies the equation $y_{tt} + 4y_{xxxx} = 0$. Suppose you know that the initial shape of the beam is the graph of $\sin(\pi x)$, and the initial velocity is 0. Solve for y .

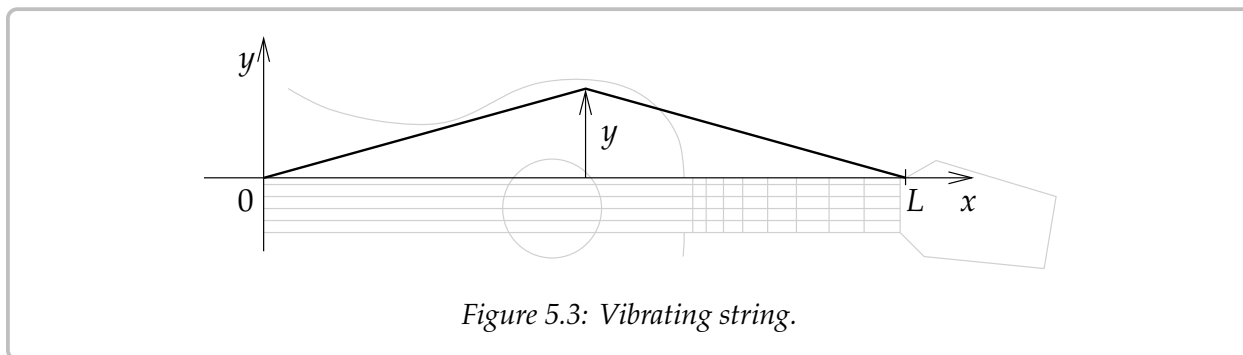
Exercise 5.2.102: Suppose you have a beam of length 10 with two fixed ends. Let y be the transverse deviation of the beam at position x on the beam ($0 < x < 10$). You know that the constants are such that this satisfies the equation $y_{tt} + 9y_{xxxx} = 0$. Suppose you know that the initial shape of the beam is the graph of $\sin^2(\pi x)$, and the initial velocity is equal to $x(10 - x)$. Set up the equation together with the boundary and initial conditions. Just set up, do not solve.

5.3 Steady periodic solutions

Note: 1–2 lectures, §10.3 in [EP], not in [BD]

5.3.1 Forced vibrating string

Consider a guitar string of length L . We studied this setup in § 4.7. Let x be the position on the string, t the time, and y the displacement of the string. See Figure 5.3.



The problem is governed by the wave equation

$$\begin{aligned} y_{tt} &= a^2 y_{xx}, \\ y(0, t) &= 0, \quad y(L, t) = 0, \\ y(x, 0) &= f(x), \quad y_t(x, 0) = g(x). \end{aligned} \tag{5.6}$$

We found that the solution is of the form

$$y = \sum_{n=1}^{\infty} \left(A_n \cos\left(\frac{n\pi a}{L}t\right) + B_n \sin\left(\frac{n\pi a}{L}t\right) \right) \sin\left(\frac{n\pi}{L}x\right),$$

where A_n and B_n are determined by the initial conditions. The natural frequencies of the system are the (angular) frequencies $\frac{n\pi a}{L}$ for integers $n \geq 1$.

But these are free vibrations. What if there is an external force acting on the string? Let us assume, say, air vibrations (noise), for example from a second string. Or perhaps from a jet engine. For simplicity, assume a nice pure sound and assume the force is uniform at every position on the string: The external force as a function of time t is given as $F_0 \cos(\omega t)$ as force per unit mass. We will call this the forcing function. Then our wave equation becomes (remember force is mass times acceleration)

$$y_{tt} = a^2 y_{xx} + F_0 \cos(\omega t), \tag{5.7}$$

with the same boundary conditions of course.

Suppose we want to find the solution here that satisfies the equation (5.7) and

$$y(0, t) = 0, \quad y(L, t) = 0, \quad y(x, 0) = 0, \quad y_t(x, 0) = 0. \quad (5.8)$$

That is, the string is initially at rest. First, we find a particular solution y_p of (5.7) that satisfies only the conditions $y(0, t) = y(L, t) = 0$. We define the functions f and g as

$$f(x) = -y_p(x, 0), \quad g(x) = -\frac{\partial y_p}{\partial t}(x, 0).$$

We then find a solution y_c of (5.6). If we add the two solutions, we find that $y = y_c + y_p$ solves (5.7) with the initial conditions.

Exercise 5.3.1: Check that $y = y_c + y_p$ solves (5.7) and the side conditions (5.8).

So the big issue here is to find the particular solution y_p . We look at the equation and we make an educated guess

$$y_p(x, t) = X(x) \cos(\omega t).$$

We plug in to get

$$-\omega^2 X \cos(\omega t) = a^2 X'' \cos(\omega t) + F_0 \cos(\omega t),$$

or $-\omega^2 X = a^2 X'' + F_0$ after canceling the cosine. We know how to find a general solution to this equation (it is a nonhomogeneous constant-coefficient equation). The general solution is

$$X(x) = A \cos\left(\frac{\omega}{a}x\right) + B \sin\left(\frac{\omega}{a}x\right) - \frac{F_0}{\omega^2}.$$

The endpoint conditions imply $X(0) = X(L) = 0$. So

$$0 = X(0) = A - \frac{F_0}{\omega^2},$$

or $A = \frac{F_0}{\omega^2}$, and also

$$0 = X(L) = \frac{F_0}{\omega^2} \cos\left(\frac{\omega L}{a}\right) + B \sin\left(\frac{\omega L}{a}\right) - \frac{F_0}{\omega^2}.$$

Assuming that $\sin(\frac{\omega L}{a})$ is not zero, we can solve for B to get

$$B = \frac{-F_0 \left(\cos\left(\frac{\omega L}{a}\right) - 1 \right)}{\omega^2 \sin\left(\frac{\omega L}{a}\right)}. \quad (5.9)$$

Therefore,

$$X(x) = \frac{F_0}{\omega^2} \left(\cos\left(\frac{\omega}{a}x\right) - \frac{\cos\left(\frac{\omega L}{a}\right) - 1}{\sin\left(\frac{\omega L}{a}\right)} \sin\left(\frac{\omega}{a}x\right) - 1 \right).$$

The particular solution y_p we are looking for is

$$y_p(x, t) = \frac{F_0}{\omega^2} \left(\cos\left(\frac{\omega}{a}x\right) - \frac{\cos\left(\frac{\omega L}{a}\right) - 1}{\sin\left(\frac{\omega L}{a}\right)} \sin\left(\frac{\omega}{a}x\right) - 1 \right) \cos(\omega t).$$

Exercise 5.3.2: Check that y_p works.

We get to the point that we skipped. Suppose $\sin(\frac{\omega L}{a}) = 0$. This means that ω is equal to one of the natural frequencies of the system, i.e. a multiple of the base frequency $\frac{\pi a}{L}$. If ω is not equal to a multiple of $\frac{\pi a}{L}$, but is very close, then the coefficient B in (5.9) seems to become very large as the denominator goes to zero. But let us not jump to conclusions just yet because the numerator may also go to zero. When $\omega = \frac{n\pi a}{L}$ for n even, then $\cos(\frac{\omega L}{a}) = 1$. We could simply take $B = 0$ to obtain a nice bounded solution of the same form, $y_p = \frac{F_0}{\omega^2} (\cos(\frac{\omega}{a}x) - 1) \cos(\omega t)$, and so no resonance is happening. Resonance occurs only when both $\cos(\frac{\omega L}{a}) = -1$ and $\sin(\frac{\omega L}{a}) = 0$, that is, when $\omega = \frac{n\pi a}{L}$ for odd n . When n is odd, if we take ω approaching $\frac{n\pi a}{L}$, then the numerator in the coefficient B is not going to zero, but the denominator is, so B really does “blow up.” We could again try to explicitly solve for the resonance solution if we wanted to, but it is, in the right sense, the limit of solutions as ω gets close to a resonance frequency. In real life, pure resonance never occurs anyway.

The calculation of resonance frequencies above explains why a string begins to vibrate if an identical string is plucked close by. In the absence of friction this vibration would get louder and louder as time goes on. On the other hand, you are unlikely to get large vibration if the forcing frequency is not close to a resonance frequency even if you have a jet engine running close to the string. That is, the amplitude does not keep increasing unless you tune to just the right frequency.

Similar resonance phenomena occur when you break a wine glass using human voice (yes, this is possible, but not easy*) if you happen to hit just the right frequency. However, a glass has a much purer sound, i.e. it is more like a vibraphone, so there are far fewer resonance frequencies to hit.

When the forcing function is more complicated, you decompose it in terms of its Fourier series and apply the result above. You may also need to solve the problem above if the forcing function is a sine rather than a cosine, but if you think about it, the solution is almost the same.

Example 5.3.1: Suppose $F_0 = 1$, $\omega = 1$, $L = 1$, and $a = 1$. Let us find a solution y to (5.7) satisfying the zero boundary conditions (5.8). First we find

$$y_p(x, t) = \left(\cos(x) - \frac{\cos(1) - 1}{\sin(1)} \sin(x) - 1 \right) \cos(t).$$

Write $B = \frac{\cos(1)-1}{\sin(1)}$ for simplicity.

Plug in $t = 0$ into y_p to get

$$f(x) = -y_p(x, 0) = -\cos x + B \sin x + 1,$$

and after differentiating y_p in t , we see that $g(x) = -\frac{\partial y_p}{\partial t}(x, 0) = 0$.

**Mythbusters*, episode 31, Discovery Channel, originally aired May 18th 2005.

Hence to find y_c , we need to solve the problem

$$\begin{aligned} w_{tt} &= w_{xx}, \\ w(0, t) &= 0, \quad w(1, t) = 0, \\ w(x, 0) &= f(x) = -\cos x + B \sin x + 1, \\ w_t(x, 0) &= g(x) = 0, \end{aligned}$$

and set $y_c = w$. If we use d'Alembert, we note that the formula that we use to define $w(x, 0)$ is not odd, hence it is not a simple matter of plugging in the expression for $w(x, 0)$ to the d'Alembert formula directly! You must define F to be the odd, 2-periodic extension of $w(x, 0)$. Then our solution $y = y_p + y_c$ to (5.7) subject to (5.8) is

$$y(x, t) = \frac{F(x+t) + F(x-t)}{2} + \left(\cos(x) - \frac{\cos(1)-1}{\sin(1)} \sin(x) - 1 \right) \cos(t). \quad (5.10)$$

It is not hard to compute specific values for an odd periodic extension of a function and hence (5.10) is a wonderful solution to the problem. For example, it is very easy to have a computer do it, unlike a series solution. A plot is given in Figure 5.4.

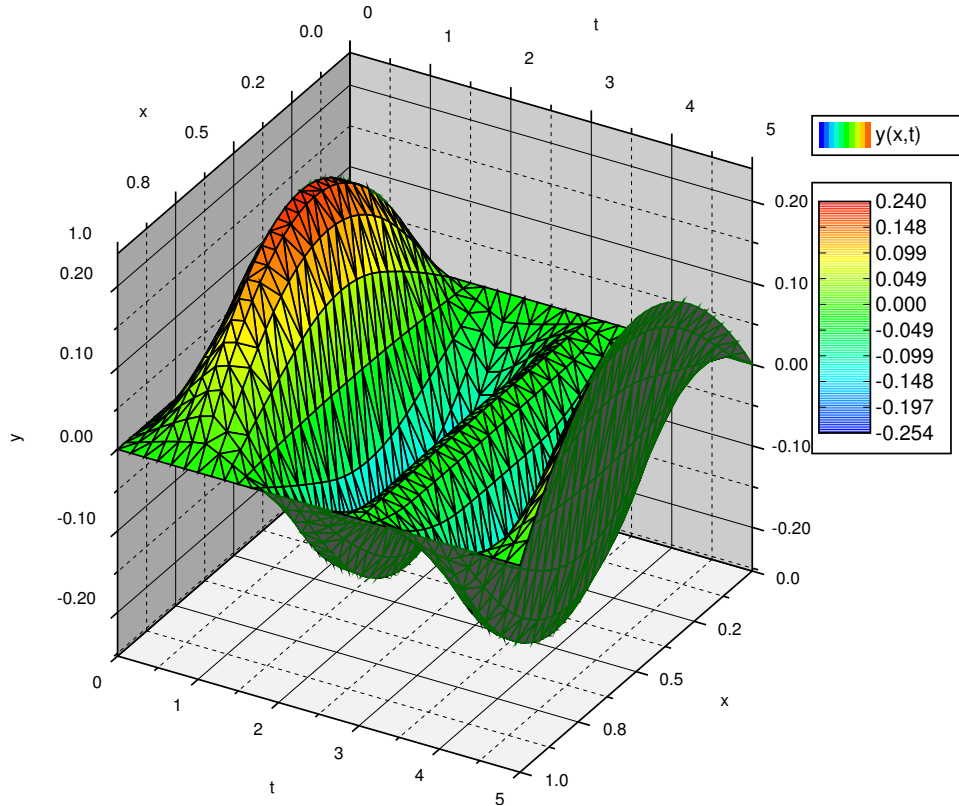


Figure 5.4: Plot of $y(x, t) = \frac{F(x+t)+F(x-t)}{2} + \left(\cos(x) - \frac{\cos(1)-1}{\sin(1)} \sin(x) - 1 \right) \cos(t)$.

5.3.2 Underground temperature oscillations

Let $u(x, t)$ be the temperature at a certain location at depth x underground at time t . See Figure 5.5.

The temperature u satisfies the heat equation $u_t = ku_{xx}$, where k is the diffusivity of the soil. We know the temperature at the surface $u(0, t)$ from weather records. Let us assume for simplicity that

$$u(0, t) = T_0 + A_0 \cos(\omega t),$$

where T_0 is the yearly mean temperature, and $t = 0$ is midsummer (put a negative sign above to make it midwinter if you wish). The constant A_0 gives the typical variation for the year: The hottest temperature is $T_0 + A_0$ and the coldest is $T_0 - A_0$. For simplicity, we assume that $T_0 = 0$. The frequency ω is picked depending on the units of t , such that when $t = 1$ year, then $\omega t = 2\pi$. For example, if t is in years, then $\omega = 2\pi$.

It seems reasonable that the temperature at depth x also oscillates with the same frequency. This, in fact, is the steady periodic solution, a solution independent of the initial conditions. So we are looking for a solution of the form

$$u(x, t) = V(x) \cos(\omega t) + W(x) \sin(\omega t)$$

for the problem

$$u_t = ku_{xx}, \quad u(0, t) = A_0 \cos(\omega t). \quad (5.11)$$

We employ the complex exponential here to make calculations simpler. Suppose we have a complex-valued function

$$h(x, t) = X(x) e^{i\omega t}.$$

We look for an h such that $\text{Re } h = u$. To find an h whose real part satisfies (5.11), we look for an h such that

$$h_t = kh_{xx}, \quad h(0, t) = A_0 e^{i\omega t}. \quad (5.12)$$

Exercise 5.3.3: Suppose h satisfies (5.12). Use *Euler's formula* for the complex exponential to check that $u = \text{Re } h$ satisfies (5.11).

Substitute h into (5.12).

$$i\omega X e^{i\omega t} = kX'' e^{i\omega t}.$$

Hence,

$$kX'' - i\omega X = 0,$$

or

$$X'' - \alpha^2 X = 0,$$

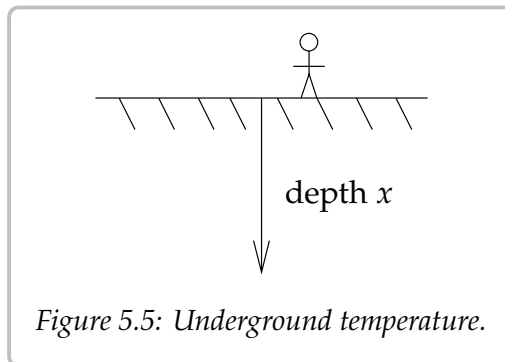


Figure 5.5: Underground temperature.

where $\alpha = \pm\sqrt{\frac{i\omega}{k}}$. Note that $\pm\sqrt{i} = \pm\frac{1+i}{\sqrt{2}}$ so you could simplify to $\alpha = \pm(1+i)\sqrt{\frac{\omega}{2k}}$. Hence the general solution is

$$X(x) = Ae^{-(1+i)\sqrt{\frac{\omega}{2k}}x} + Be^{(1+i)\sqrt{\frac{\omega}{2k}}x}.$$

We assume that an $X(x)$ that solves the problem must be bounded as $x \rightarrow \infty$ since $u(x, t)$ should be bounded (we are not worrying about Earth's core!). If you use [Euler's formula](#) to expand the complex exponentials, note that the second term is unbounded (if $B \neq 0$), while the first term is always bounded. Hence $B = 0$.

Exercise 5.3.4: Use [Euler's formula](#) to show that $e^{(1+i)\sqrt{\frac{\omega}{2k}}x}$ is unbounded as $x \rightarrow \infty$, while $e^{-(1+i)\sqrt{\frac{\omega}{2k}}x}$ is bounded as $x \rightarrow \infty$.

Furthermore, $X(0) = A_0$ since $h(0, t) = A_0e^{i\omega t}$. Thus $A = A_0$. This means that

$$h(x, t) = A_0e^{-(1+i)\sqrt{\frac{\omega}{2k}}x}e^{i\omega t} = A_0e^{-(1+i)\sqrt{\frac{\omega}{2k}}x+i\omega t} = A_0e^{-\sqrt{\frac{\omega}{2k}}x}e^{i(\omega t - \sqrt{\frac{\omega}{2k}}x)}.$$

We need to get the real part of h , so we apply [Euler's formula](#) to get

$$h(x, t) = A_0e^{-\sqrt{\frac{\omega}{2k}}x} \left(\cos \left(\omega t - \sqrt{\frac{\omega}{2k}}x \right) + i \sin \left(\omega t - \sqrt{\frac{\omega}{2k}}x \right) \right).$$

Then finally

$$u(x, t) = \operatorname{Re} h(x, t) = A_0e^{-\sqrt{\frac{\omega}{2k}}x} \cos \left(\omega t - \sqrt{\frac{\omega}{2k}}x \right).$$

Yay!

Notice the phase is different at different depths. At depth x the phase is delayed by $x\sqrt{\frac{\omega}{2k}}$. For example, in cgs units (centimeters-grams-seconds) we have $k = 0.005$ (typical value for soil), $\omega = \frac{2\pi}{\text{seconds in a year}} = \frac{2\pi}{31,557,341} \approx 1.99 \times 10^{-7}$. Then, if we compute where the phase shift $x\sqrt{\frac{\omega}{2k}} = \pi$, we find the depth in centimeters where the seasons are reversed. That is, we get the depth at which summer is the coldest and winter is the warmest. We get approximately 700 centimeters, which is approximately 23 feet below ground.

Be careful not to jump to conclusions. The temperature swings decay rapidly as you dig deeper. The amplitude of the temperature swings is $A_0e^{-\sqrt{\frac{\omega}{2k}}x}$. This function decays *very* quickly as x (the depth) grows. Let us again take typical parameters as above. We also assume that our surface temperature swing is $\pm 15^\circ$ Celsius, that is, $A_0 = 15$. Then the maximum temperature variation at 700 centimeters is only $\pm 0.66^\circ$ Celsius.

You need not dig very deep to get an effective “refrigerator,” with nearly constant temperature. That is why wines are kept in a cellar; you need consistent temperature. The temperature differential could also be used for energy. A home could be heated or cooled by taking advantage of the fact above. Even without Earth's core, you could heat a home in the winter and cool it in the summer. Earth's core makes the temperature higher the deeper you dig, although you need to dig somewhat deep to feel a difference. We did not take that into account above.

5.3.3 Exercises

Exercise 5.3.5: Suppose that the forcing function for the vibrating string is $F_0 \sin(\omega t)$. Derive the particular solution y_p .

Exercise 5.3.6: Take the forced vibrating string. Suppose that $L = 1$, $a = 1$. Suppose that the forcing function is the square wave that is 1 on the interval $0 < t < \pi$ and -1 on the interval $-\pi < t < 0$. Find the particular solution. Hint: You may want to use result of [Exercise 5.3.5](#).

Exercise 5.3.7: The units are cgs (centimeters-grams-seconds). For $k = 0.005$, $\omega = 1.991 \times 10^{-7}$, $A_0 = 20$. Find the depth at which the temperature variation is half (± 10 degrees) of what it is on the surface.

Exercise 5.3.8: Derive the solution for underground temperature oscillation without assuming that $T_0 = 0$.

Exercise 5.3.101: Take the forced vibrating string. Suppose that $L = 1$, $a = 1$. Suppose that the forcing function is a triangle wave, $|t| - \frac{\pi}{2}$ on $-\pi < t < \pi$ extended periodically. Find the particular solution.

Exercise 5.3.102: The units are cgs (centimeters-grams-seconds). For $k = 0.01$, $\omega = 1.991 \times 10^{-7}$, $A_0 = 25$. Find the depth at which the summer is again the hottest point.

Chapter 6

The Laplace transform

6.1 The Laplace transform

Note: 1.5–2 lectures, §7.1 in [EP], §6.1 and parts of §6.2 in [BD]

6.1.1 The transform

In this chapter, we will discuss the Laplace transform[†]. The Laplace transform is a very efficient method to solve certain ODE or PDE problems. The transform takes a differential equation and turns it into an algebraic equation. If the algebraic equation can be solved, applying the inverse transform gives us our desired solution. The Laplace transform also has applications in the analysis of electrical circuits, NMR spectroscopy, signal processing, and elsewhere. Finally, understanding the Laplace transform will also help with understanding the related Fourier transform, which, however, requires more understanding of complex numbers. We will not cover the Fourier transform.

The Laplace transform also gives a lot of insight into the nature of the equations we are dealing with. It can be seen as converting between the time and the frequency domain. For example, take the standard equation

$$mx''(t) + cx'(t) + kx(t) = f(t).$$

We can think of t as time and $f(t)$ as incoming signal. The Laplace transform will convert the equation from a differential equation in time to an algebraic (no derivatives) equation, where the new independent variable s is thought of as the frequency.[‡]

We can think of the *Laplace transform* as a black box. It eats functions and spits out functions in a new variable. We write $\mathcal{L}\{f(t)\} = F(s)$ for the Laplace transform of $f(t)$. It is common to write lower case letters for functions in the time domain and upper case letters for functions in the frequency domain. We use the same letter to denote that one

[†]Just like the Laplace equation and the Laplacian, the Laplace transform is also named after [Pierre-Simon, marquis de Laplace](#) (1749–1827).

[‡]Really, it is the “frequency” in terms of complex numbers, but we digress.

function is the Laplace transform of the other. For example, $F(s)$ is the Laplace transform of $f(t)$. Let us define the transform.

$$\mathcal{L}\{f(t)\} = F(s) \stackrel{\text{def}}{=} \int_0^{\infty} e^{-st} f(t) dt.$$

We note that we are only considering $t \geq 0$ in the transform. Of course, if we think of t as time, there is no problem; we are generally interested in finding out what will happen in the future (the Laplace transform is one place where it is safe to ignore the past). Let us compute some simple transforms.

Example 6.1.1: Suppose $f(t) = 1$, then

$$\mathcal{L}\{1\} = \int_0^{\infty} e^{-st} dt = \left[\frac{e^{-st}}{-s} \right]_{t=0}^{\infty} = \lim_{h \rightarrow \infty} \left[\frac{e^{-sh}}{-s} \right]_{t=0}^h = \lim_{h \rightarrow \infty} \left(\frac{e^{-sh}}{-s} - \frac{1}{-s} \right) = \frac{1}{s}.$$

The limit (the improper integral) only exists if $s > 0$. So $\mathcal{L}\{1\}$ is only defined for $s > 0$.

Example 6.1.2: Suppose $f(t) = e^{-at}$, then

$$\mathcal{L}\{e^{-at}\} = \int_0^{\infty} e^{-st} e^{-at} dt = \int_0^{\infty} e^{-(s+a)t} dt = \left[\frac{e^{-(s+a)t}}{-(s+a)} \right]_{t=0}^{\infty} = \frac{1}{s+a}.$$

The limit only exists if $s + a > 0$. So $\mathcal{L}\{e^{-at}\}$ is only defined for $s + a > 0$.

Example 6.1.3: Suppose $f(t) = t$, then using integration by parts

$$\begin{aligned} \mathcal{L}\{t\} &= \int_0^{\infty} e^{-st} t dt \\ &= \left[\frac{-te^{-st}}{s} \right]_{t=0}^{\infty} + \frac{1}{s} \int_0^{\infty} e^{-st} dt \\ &= 0 + \frac{1}{s} \left[\frac{e^{-st}}{-s} \right]_{t=0}^{\infty} \\ &= \frac{1}{s^2}. \end{aligned}$$

Again, the limit only exists if $s > 0$.

Example 6.1.4: A common function is the *unit step function*, which is sometimes called the *Heaviside function**. This function is generally given as

$$u(t) = \begin{cases} 0 & \text{if } t < 0, \\ 1 & \text{if } t \geq 0. \end{cases}$$

*The function is named after the English mathematician, engineer, and physicist [Oliver Heaviside](#) (1850–1925). Only by coincidence is the function “heavy” on “one side.”

Let us find the Laplace transform of $u(t - a)$, where $a \geq 0$ is some constant. That is, it is the function that is 0 for $t < a$ and 1 for $t \geq a$.

$$\mathcal{L}\{u(t - a)\} = \int_0^{\infty} e^{-st} u(t - a) dt = \int_a^{\infty} e^{-st} dt = \left[\frac{e^{-st}}{-s} \right]_{t=a}^{\infty} = \frac{e^{-as}}{s},$$

where of course $s > 0$ (and $a \geq 0$ as we said before).

By applying similar procedures, we can compute the transforms of many elementary functions. Many basic transforms are listed in [Table 6.1](#) (see also [appendix B](#)).

$f(t)$	$\mathcal{L}\{f(t)\}$	$f(t)$	$\mathcal{L}\{f(t)\}$
C	$\frac{C}{s}$	$\sin(\omega t)$	$\frac{\omega}{s^2 + \omega^2}$
t	$\frac{1}{s^2}$	$\cos(\omega t)$	$\frac{s}{s^2 + \omega^2}$
t^2	$\frac{2}{s^3}$	$\sinh(\omega t)$	$\frac{\omega}{s^2 - \omega^2}$
t^3	$\frac{6}{s^4}$	$\cosh(\omega t)$	$\frac{s}{s^2 - \omega^2}$
t^n	$\frac{n!}{s^{n+1}}$	$u(t - a) \quad (a \geq 0)$	$\frac{e^{-as}}{s}$
e^{-at}	$\frac{1}{s + a}$		

Table 6.1: Some Laplace transforms (C , ω , and a are constants).

Exercise 6.1.1: Verify [Table 6.1](#).

Since the transform is defined by an integral, we can use the linearity properties of the integral. For example, if C is a constant, then

$$\mathcal{L}\{C f(t)\} = \int_0^{\infty} e^{-st} C f(t) dt = C \int_0^{\infty} e^{-st} f(t) dt = C \mathcal{L}\{f(t)\}.$$

So we can “pull out” a constant from the transform. Similarly, with addition. As linearity is important, we state it as a theorem.

Theorem 6.1.1 (Linearity of the Laplace transform). *If A , B , and C are constants, then*

$$\mathcal{L}\{A f(t) + B g(t)\} = A \mathcal{L}\{f(t)\} + B \mathcal{L}\{g(t)\},$$

and in particular,

$$\mathcal{L}\{C f(t)\} = C \mathcal{L}\{f(t)\}.$$

Exercise 6.1.2: Verify the theorem. That is, show that $\mathcal{L}\{A f(t) + B g(t)\} = A \mathcal{L}\{f(t)\} + B \mathcal{L}\{g(t)\}$.

These rules together with Table 6.1 on the previous page make it easy to find the Laplace transform of a whole lot of functions already. For example:

$$\mathcal{L}\{2 + 5t + 9e^{-2t}\} = \mathcal{L}\{2\} + 5\mathcal{L}\{t\} + 9\mathcal{L}\{e^{-2t}\} = \frac{2}{s} + \frac{5}{s^2} + \frac{9}{s+2}.$$

Be careful! The Laplace transform of a product is *not* the product of the transforms. In general

$$\mathcal{L}\{f(t)g(t)\} \neq \mathcal{L}\{f(t)\}\mathcal{L}\{g(t)\}.$$

Moreover, not all functions have a Laplace transform. For example, the function $\frac{1}{t}$ does not have a Laplace transform as the integral diverges for all s . Similarly, $\tan t$ or e^{t^2} do not have Laplace transforms.

6.1.2 Existence and uniqueness

When does the Laplace transform exist? A function $f(t)$ is of *exponential order* as t goes to infinity if

$$|f(t)| \leq Me^{ct},$$

for some constants M and c , for sufficiently large t (say for all $t > t_0$ for some t_0). The simplest way to check this condition is to try and compute

$$\lim_{t \rightarrow \infty} \frac{f(t)}{e^{ct}}.$$

If the limit exists and is finite (usually zero), then $f(t)$ is of exponential order.

Exercise 6.1.3: Use L'Hopital's rule from calculus to show that a polynomial is of exponential order. Hint: Note that a sum of two exponential-order functions is also of exponential order. Then show that t^n is of exponential order for any n .

For an exponential-order function, we have existence and uniqueness of the Laplace transform.

Theorem 6.1.2 (Existence). Let $f(t)$ be continuous on the interval $[0, \infty)$ and of exponential order for a certain constant c . Then $F(s) = \mathcal{L}\{f(t)\}$ is defined for all $s > c$.

The existence is not difficult to see. Let $f(t)$ be of exponential order, that is, $|f(t)| \leq Me^{ct}$ for all $t > 0$ (for simplicity $t_0 = 0$). Let $s > c$, or in other words $(s - c) > 0$. By the comparison theorem from calculus, the improper integral defining $\mathcal{L}\{f(t)\}$ exists because the following integral exists

$$\int_0^\infty e^{-st}(Me^{ct}) dt = M \int_0^\infty e^{-(s-c)t} dt = M \left[\frac{e^{-(s-c)t}}{-(s-c)} \right]_{t=0}^\infty = \frac{M}{s-c}.$$

The transform also exists for some other functions that are not of exponential order, but that will not be relevant to us. Before dealing with uniqueness, we note that for functions of exponential order, their Laplace transform decays at infinity:

$$\lim_{s \rightarrow \infty} F(s) = 0.$$

Theorem 6.1.3 (Uniqueness). *Let $f(t)$ and $g(t)$ be continuous and of exponential order. Suppose that there exists a constant C , such that $F(s) = G(s)$ for all $s > C$. Then $f(t) = g(t)$ for all $t \geq 0$.*

Both theorems hold for piecewise continuous functions as well. Recall that piecewise continuous means that the function is continuous except perhaps at a discrete set of points, where it has jump discontinuities like the Heaviside function. The Laplace transform, however, does not “see” values at the discontinuities. So we can only conclude that $f(t) = g(t)$ outside of discontinuities. For example, the unit step function is sometimes defined using $u(0) = 1/2$. This new step function, however, has the exact same Laplace transform as the one we defined earlier, where $u(0) = 1$.

6.1.3 The inverse transform

As we said, the Laplace transform will allow us to convert a differential equation into an algebraic equation. Once we solve the algebraic equation in the frequency domain, we will want to get back to the time domain, as that is what we are interested in. Given a function $F(s)$, we wish to find a function $f(t)$ such that $\mathcal{L}\{f(t)\} = F(s)$. [Theorem 6.1.3](#) says that the solution $f(t)$ is unique. So we can without fear make the following definition.

Suppose $F(s) = \mathcal{L}\{f(t)\}$ for some function $f(t)$. Define the *inverse Laplace transform* as

$$\mathcal{L}^{-1}\{F(s)\} \stackrel{\text{def}}{=} f(t).$$

There is an integral formula for the inverse, but it is not as simple as the transform itself—it requires complex numbers and path integrals. For us it will suffice to compute the inverse using [Table 6.1](#) on page 295.

Example 6.1.5: Take $F(s) = \frac{1}{s+1}$. Find the inverse Laplace transform.

We look at the table to find

$$\mathcal{L}^{-1}\left\{\frac{1}{s+1}\right\} = e^{-t}.$$

As the Laplace transform is linear, the inverse Laplace transform is also linear. That is,

$$\mathcal{L}^{-1}\{AF(s) + BG(s)\} = A\mathcal{L}^{-1}\{F(s)\} + B\mathcal{L}^{-1}\{G(s)\}.$$

Of course, we also have $\mathcal{L}^{-1}\{AF(s)\} = A\mathcal{L}^{-1}\{F(s)\}$.

Example 6.1.6: Take $F(s) = \frac{s^2+s+1}{s^3+s}$. Find the inverse Laplace transform.

First we use the *method of partial fractions* to write F in a form where we can use [Table 6.1](#) on page 295. We factor the denominator as $s(s^2 + 1)$ and write

$$\frac{s^2 + s + 1}{s^3 + s} = \frac{A}{s} + \frac{Bs + C}{s^2 + 1}.$$

Putting the right-hand side over a common denominator and equating the numerators, we get $A(s^2 + 1) + s(Bs + C) = s^2 + s + 1$. Expanding and equating coefficients, we obtain

$A + B = 1$, $C = 1$, $A = 1$, and thus $B = 0$. In other words,

$$F(s) = \frac{s^2 + s + 1}{s^3 + s} = \frac{1}{s} + \frac{1}{s^2 + 1}.$$

By linearity of the inverse Laplace transform, we get

$$\mathcal{L}^{-1} \left\{ \frac{s^2 + s + 1}{s^3 + s} \right\} = \mathcal{L}^{-1} \left\{ \frac{1}{s} \right\} + \mathcal{L}^{-1} \left\{ \frac{1}{s^2 + 1} \right\} = 1 + \sin t.$$

Another useful property is the so-called *shifting property* or the *first shifting property*

$$\mathcal{L}\{e^{-at} f(t)\} = F(s + a),$$

where $F(s)$ is the Laplace transform of $f(t)$.

Exercise 6.1.4: Derive the first shifting property from the definition of the Laplace transform.

The shifting property can be used, for example, when the denominator is a more complicated quadratic that may come up in the method of partial fractions. We complete the square and write such quadratics as $(s + a)^2 + b$ and then use the shifting property.

Example 6.1.7: Find $\mathcal{L}^{-1} \left\{ \frac{1}{s^2 + 4s + 8} \right\}$.

First, we complete the square to make the denominator $(s + 2)^2 + 4$. We find

$$\mathcal{L}^{-1} \left\{ \frac{1}{s^2 + 4} \right\} = \frac{1}{2} \mathcal{L}^{-1} \left\{ \frac{2}{s^2 + 2^2} \right\} = \frac{1}{2} \sin(2t).$$

Finally, we put it all together with the shifting property to find

$$\mathcal{L}^{-1} \left\{ \frac{1}{s^2 + 4s + 8} \right\} = \mathcal{L}^{-1} \left\{ \frac{1}{(s + 2)^2 + 4} \right\} = \frac{1}{2} e^{-2t} \sin(2t).$$

Often, we want to be able to apply the inverse Laplace transform to rational functions, that is, functions of the form

$$\frac{F(s)}{G(s)}$$

where $F(s)$ and $G(s)$ are polynomials. If $\frac{F(s)}{G(s)}$ is the Laplace transform of an exponential-order function, it goes to zero as $s \rightarrow \infty$, and so the degree of $F(s)$ must be smaller than that of $G(s)$. Such rational functions are called *proper rational functions* and we can always apply the method of partial fractions without polynomial division. Of course, we still need to be able to factor the denominator into linear and quadratic terms, which involves finding the roots of the denominator.

6.1.4 Exercises

Exercise 6.1.5: Find the Laplace transform of $3 + t^5 + \sin(\pi t)$.

Exercise 6.1.6: Find the Laplace transform of $a + bt + ct^2$ for some constants a , b , and c .

Exercise 6.1.7: Find the Laplace transform of $A \cos(\omega t) + B \sin(\omega t)$.

Exercise 6.1.8: Find the Laplace transform of $\cos^2(\omega t)$.

Exercise 6.1.9: Find the inverse Laplace transform of $\frac{4}{s^2-9}$.

Exercise 6.1.10: Find the inverse Laplace transform of $\frac{2s}{s^2-1}$.

Exercise 6.1.11: Find the inverse Laplace transform of $\frac{1}{(s-1)^2(s+1)}$.

Exercise 6.1.12: Find the Laplace transform of $f(t) = \begin{cases} t & \text{if } t \geq 1, \\ 0 & \text{if } t < 1. \end{cases}$

Exercise 6.1.13: Find the inverse Laplace transform of $\frac{s}{(s^2+s+2)(s+4)}$.

Exercise 6.1.14: Find the Laplace transform of $\sin(\omega(t-a))$.

Exercise 6.1.15: Find the Laplace transform of $t \sin(\omega t)$. Hint: Several integrations by parts.

Exercise 6.1.101: Find the Laplace transform of $4(t+1)^2$.

Exercise 6.1.102: Find the inverse Laplace transform of $\frac{8}{s^3(s+2)}$.

Exercise 6.1.103: Find the Laplace transform of te^{-t} . Hint: Shifting property.

Exercise 6.1.104: Find the Laplace transform of $\sin(t)e^{-t}$. Hint: Integrate by parts.

6.2 Transforms of derivatives and ODEs

Note: 2 lectures, §7.2–7.3 in [EP], §6.2 and §6.3 in [BD]

6.2.1 Transforms of derivatives

Let us see how the Laplace transform is used for differential equations. First we find the Laplace transform of the derivative of a function. Suppose $g(t)$ is a differentiable function of exponential order, that is, $|g(t)| \leq Me^{ct}$ for some M and c . So $\mathcal{L}\{g(t)\}$ exists, and what is more, $\lim_{t \rightarrow \infty} e^{-st}g(t) = 0$ when $s > c$. Then

$$\mathcal{L}\{g'(t)\} = \int_0^\infty e^{-st}g'(t)dt = \left[e^{-st}g(t)\right]_{t=0}^\infty - \int_0^\infty (-s)e^{-st}g(t)dt = -g(0) + s\mathcal{L}\{g(t)\}.$$

We repeat this procedure for higher derivatives. The results are listed in Table 6.2. The procedure also works for continuous piecewise smooth functions, that is, functions that are continuous with a piecewise continuous derivative.

$f(t)$	$\mathcal{L}\{f(t)\} = F(s)$
$g'(t)$	$sG(s) - g(0)$
$g''(t)$	$s^2G(s) - sg(0) - g'(0)$
$g'''(t)$	$s^3G(s) - s^2g(0) - sg'(0) - g''(0)$

Table 6.2: Laplace transforms of derivatives ($G(s) = \mathcal{L}\{g(t)\}$ as usual).

Exercise 6.2.1: Verify Table 6.2.

6.2.2 Solving ODEs with the Laplace transform

Notice that the Laplace transform turns differentiation into multiplication by s . It is this property that makes it useful to apply the transform to differential equations.

Example 6.2.1: Consider the problem

$$x''(t) + x(t) = \cos(2t), \quad x(0) = 0, \quad x'(0) = 1.$$

We will take the Laplace transform of both sides of the equation. By $X(s)$, we will, as usual, denote the Laplace transform of $x(t)$.

$$\begin{aligned} \mathcal{L}\{x''(t) + x(t)\} &= \mathcal{L}\{\cos(2t)\}, \\ s^2X(s) - sx(0) - x'(0) + X(s) &= \frac{s}{s^2 + 4}. \end{aligned}$$

We plug in the initial conditions now—making computations more streamlined—to find

$$s^2X(s) - 1 + X(s) = \frac{s}{s^2 + 4}.$$

We solve for $X(s)$,

$$X(s) = \frac{s}{(s^2 + 1)(s^2 + 4)} + \frac{1}{s^2 + 1}.$$

We use partial fractions (exercise) to write

$$X(s) = \frac{1}{3} \frac{s}{s^2 + 1} - \frac{1}{3} \frac{s}{s^2 + 4} + \frac{1}{s^2 + 1}.$$

Now take the inverse Laplace transform to obtain

$$x(t) = \frac{1}{3} \cos(t) - \frac{1}{3} \cos(2t) + \sin(t).$$

The procedure for linear constant-coefficient equations is as follows: Take an ordinary differential equation in the time variable t . Apply the Laplace transform to transform the equation into an algebraic (non differential) equation in the frequency domain. All the $x(t)$, $x'(t)$, $x''(t)$, and so on, will be converted to $X(s)$, $sX(s) - x(0)$, $s^2X(s) - sx(0) - x'(0)$, and so on. Solve the equation for $X(s)$. Then taking the inverse transform, if possible, find $x(t)$.

It should be noted that since not every function has a Laplace transform, not every equation can be solved in this manner. Also if the equation is not a linear constant-coefficient ODE, then applying the Laplace transform may not give us an algebraic equation.

6.2.3 Using the Heaviside function

Before we move on to more general equations than those we could solve before, we want to consider the Heaviside function. See [Figure 6.1](#) on the following page for the graph.

$$u(t) = \begin{cases} 0 & \text{if } t < 0, \\ 1 & \text{if } t \geq 0. \end{cases}$$

The Heaviside function is useful for putting other functions together or cutting functions off. Usually we use $u(t - a)$ for some constant a ; we just shift the graph to the right by a . That is, $u(t - a) = 0$ when $t < a$ and $u(t - a) = 1$ when $t \geq a$. For example, suppose $f(t)$ is a “signal”, and we started receiving $\sin t$ at time $t = \pi$. The function $f(t)$ is then defined as

$$f(t) = \begin{cases} 0 & \text{if } t < \pi, \\ \sin t & \text{if } t \geq \pi. \end{cases}$$

Using the Heaviside function, $f(t)$ can be written as

$$f(t) = u(t - \pi) \sin t.$$

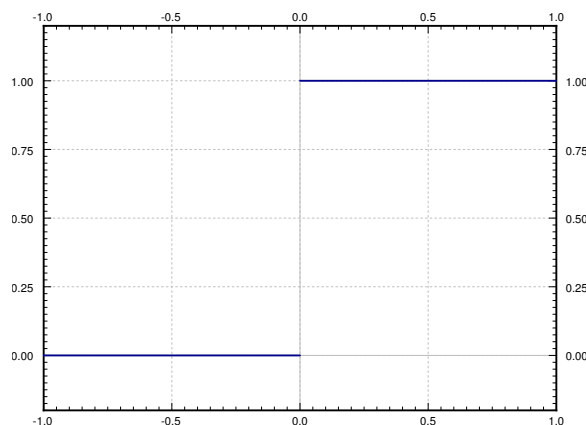


Figure 6.1: Plot of the Heaviside (unit step) function $u(t)$.

Similarly, the step function that is 1 on the interval $[1, 2)$ and 0 elsewhere is written as

$$u(t - 1) - u(t - 2).$$

With the Heaviside function we can express functions defined piecewise. If $f(t) = t$ when t is in $[0, 1]$, $f(t) = -t + 2$ when t is in $[1, 2]$, and $f(t) = 0$ otherwise, then you write

$$f(t) = t(u(t) - u(t - 1)) + (-t + 2)(u(t - 1) - u(t - 2)).$$

How does the Heaviside function interact with the Laplace transform? We saw that (for $a \geq 0$)

$$\mathcal{L}\{u(t - a)\} = \frac{e^{-as}}{s}.$$

This computation can be generalized into a *shifting property* or *second shifting property*:

$$\mathcal{L}\{f(t - a)u(t - a)\} = e^{-as} \mathcal{L}\{f(t)\} \quad (a \geq 0). \quad (6.1)$$

For example,

$$\mathcal{L}\{t u(t - 1)\} = \mathcal{L}\{((t - 1) + 1) u(t - 1)\} = e^{-s} \mathcal{L}\{t + 1\} = e^{-s} \left(\frac{1}{s^2} + \frac{1}{s} \right).$$

Example 6.2.2: Consider the mass-spring system

$$x''(t) + x(t) = f(t), \quad x(0) = 0, \quad x'(0) = 0,$$

where $f(t) = 1$ if $1 \leq t < 5$ and zero otherwise. The function $f(t)$ is not periodic and is defined piecewise; no problem for Laplace. Imagine a rocket attached to the mass is fired for 4 seconds starting at $t = 1$. Or perhaps imagine an RLC circuit, where the voltage is raised at a constant rate for 4 seconds starting at $t = 1$ and then held steady again starting at $t = 5$ (recall that $f(t)$ represents the derivative of the voltage in the RLC circuit).

We write $f(t) = u(t-1) - u(t-5)$. We transform the equation and we plug in the initial conditions as before

$$s^2 X(s) + X(s) = \frac{e^{-s}}{s} - \frac{e^{-5s}}{s}.$$

We solve for $X(s)$,

$$X(s) = \frac{e^{-s}}{s(s^2 + 1)} - \frac{e^{-5s}}{s(s^2 + 1)}.$$

We leave it as an exercise to the reader to show that

$$\mathcal{L}^{-1} \left\{ \frac{1}{s(s^2 + 1)} \right\} = 1 - \cos t.$$

In other words, $\mathcal{L}\{1 - \cos t\} = \frac{1}{s(s^2 + 1)}$. Using the second shifting property (6.1), we find

$$\mathcal{L}^{-1} \left\{ \frac{e^{-s}}{s(s^2 + 1)} \right\} = \mathcal{L}^{-1} \{ e^{-s} \mathcal{L}\{1 - \cos t\} \} = (1 - \cos(t-1)) u(t-1).$$

Similarly,

$$\mathcal{L}^{-1} \left\{ \frac{e^{-5s}}{s(s^2 + 1)} \right\} = \mathcal{L}^{-1} \{ e^{-5s} \mathcal{L}\{1 - \cos t\} \} = (1 - \cos(t-5)) u(t-5).$$

Hence, the solution is

$$x(t) = (1 - \cos(t-1)) u(t-1) - (1 - \cos(t-5)) u(t-5).$$

The plot of this solution is given in [Figure 6.2](#).

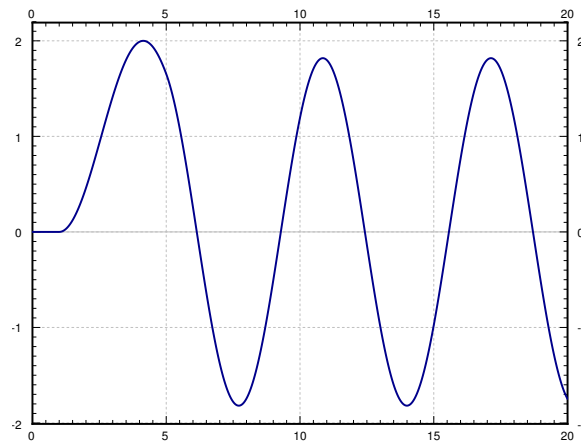


Figure 6.2: Plot of $x(t)$.

6.2.4 Transfer functions

The Laplace transform leads to the following useful concept for studying the steady-state behavior of a linear system. Consider an equation of the form

$$Lx = f(t),$$

where L is a linear constant-coefficient differential operator. Then $f(t)$ is usually thought of as the input of the system and $x(t)$ is thought of as the output of the system. For example, for a mass-spring system the input is the forcing function and the output is the behavior of the mass. We would like to have a convenient way to study the behavior of the system for different inputs.

Let us suppose that all the initial conditions are zero. We take the Laplace transform of the equation to obtain, for some $A(s)$,

$$A(s)X(s) = F(s).$$

Solving for the ratio $X(s)/F(s)$ gives the so-called *transfer function* $H(s) = 1/A(s)$, that is,

$$H(s) = \frac{X(s)}{F(s)}.$$

In other words, $X(s) = H(s)F(s)$. We get an algebraic dependence of the output of the system based on the input. We can now easily study the steady-state behavior of the system given different inputs by simply multiplying by the transfer function. Moreover, it is possible to compute the $H(s)$ without knowing exactly what the equation is by observing the output $X(s)$ for a given input $F(s)$. Once $H(s)$ is known, you can find the output for any input.

Example 6.2.3: Given $x'' + \omega_0^2 x = f(t)$ (assume the initial conditions are zero), let us find the transfer function.

First, we take the Laplace transform of the equation,

$$s^2 X(s) + \omega_0^2 X(s) = F(s).$$

Now we solve for the transfer function $X(s)/F(s)$,

$$H(s) = \frac{X(s)}{F(s)} = \frac{1}{s^2 + \omega_0^2}.$$

Let us see how to use the transfer function. Suppose we have the constant input $f(t) = 1$. Hence $F(s) = 1/s$, and

$$X(s) = H(s)F(s) = \frac{1}{s^2 + \omega_0^2} \frac{1}{s}.$$

Taking the inverse Laplace transform of $X(s)$, we obtain

$$x(t) = \frac{1 - \cos(\omega_0 t)}{\omega_0^2}.$$

Similarly, for any other input $F(s)$, the output is $X(s) = H(s)F(s) = \frac{1}{s^2 + \omega_0^2} F(s)$.

6.2.5 Transforms of integrals

A feature of the Laplace transform is that it is able to deal with integral equations. These are equations in which integrals rather than derivatives of functions appear. The basic property, which can be proved by applying the definition and doing integration by parts, is

$$\mathcal{L} \left\{ \int_0^t f(\tau) d\tau \right\} = \frac{1}{s} F(s).$$

It is sometimes useful (e.g. for computing the inverse transform) to write this as

$$\int_0^t f(\tau) d\tau = \mathcal{L}^{-1} \left\{ \frac{1}{s} F(s) \right\}.$$

Example 6.2.4: To compute $\mathcal{L}^{-1} \left\{ \frac{1}{s(s^2+1)} \right\}$, we could apply this integration rule:

$$\mathcal{L}^{-1} \left\{ \frac{1}{s} \frac{1}{s^2+1} \right\} = \int_0^t \mathcal{L}^{-1} \left\{ \frac{1}{s^2+1} \right\} d\tau = \int_0^t \sin \tau d\tau = 1 - \cos t.$$

Example 6.2.5: An equation containing an integral of the unknown function is called an *integral equation*. Consider

$$x(t) - t = \int_0^t x(\tau) d\tau,$$

where we wish to solve for $x(t)$. We apply the Laplace transform to get

$$X(s) - \frac{1}{s^2} = \frac{1}{s} X(s),$$

where $X(s) = \mathcal{L}\{x(t)\}$. Thus

$$X(s) = \frac{1}{s(s-1)} = \frac{1}{s-1} - \frac{1}{s}.$$

The inverse Laplace transform gives

$$x(t) = e^t - 1.$$

6.2.6 Periodic functions

The reader might ask: What about periodic functions as our input $f(t)$? That is, a function $f(t)$ where $f(t) = f(t+P)$ for some constant P (the period). Well, let us compute $F(s)$:

$$\begin{aligned} F(s) &= \int_0^\infty e^{-st} f(t) dt = \int_0^P e^{-st} f(t) dt + \int_P^\infty e^{-st} f(t) dt \\ &= \int_0^P e^{-st} f(t) dt + \int_0^\infty e^{-s(t+P)} f(t+P) dt = \int_0^P e^{-st} f(t) dt + e^{-Ps} \int_0^\infty e^{-st} f(t) dt \\ &= \int_0^P e^{-st} f(t) dt + e^{-Ps} F(s). \end{aligned}$$

Solving for $F(s)$, we get

$$F(s) = \frac{1}{1 - e^{-Ps}} \int_0^P e^{-st} f(t) dt.$$

As before, computing the inverse would be more complex and possibly involve consulting a table. Let us not worry about computing the inverse here.

Example 6.2.6: Suppose $f(t)$ is a version of the *sawtooth*, that is, let $f(t) = t$ for $0 \leq t < 1$ and use $f(t) = f(t + 1)$ to extend it periodically. So $f(t) = t - 1$ for $1 \leq t < 2$, $f(t) = t - 2$ for $2 \leq t < 3$, etc. Then $P = 1$ and a short computation with integration by parts gets

$$F(s) = \frac{1}{1 - e^{-s}} \int_0^1 e^{-st} t dt = \frac{1}{1 - e^{-s}} \left(\frac{-e^{-s}}{s} - \frac{e^{-s}}{s^2} + \frac{1}{s^2} \right) = \frac{-e^{-s}}{(1 - e^{-s})s} + \frac{1}{s^2}.$$

6.2.7 Exercises

Exercise 6.2.2: Using the Heaviside function write down the piecewise function that is 0 for $t < 0$, t^2 for t in $[0, 1]$ and t for $t > 1$.

Exercise 6.2.3: Using the Laplace transform solve

$$mx'' + cx' + kx = 0, \quad x(0) = a, \quad x'(0) = b,$$

where $m > 0$, $c > 0$, $k > 0$, and $c^2 - 4km > 0$ (system is overdamped).

Exercise 6.2.4: Using the Laplace transform solve

$$mx'' + cx' + kx = 0, \quad x(0) = a, \quad x'(0) = b,$$

where $m > 0$, $c > 0$, $k > 0$, and $c^2 - 4km < 0$ (system is underdamped).

Exercise 6.2.5: Using the Laplace transform solve

$$mx'' + cx' + kx = 0, \quad x(0) = a, \quad x'(0) = b,$$

where $m > 0$, $c > 0$, $k > 0$, and $c^2 = 4km$ (system is critically damped).

Exercise 6.2.6: Solve $x'' + x = u(t - 1)$ for initial conditions $x(0) = 0$ and $x'(0) = 0$.

Exercise 6.2.7: Show the “differentiation of the transform” property. Suppose $\mathcal{L}\{f(t)\} = F(s)$, then show

$$\mathcal{L}\{-t f(t)\} = F'(s).$$

Hint: Differentiate under the integral sign.

Exercise 6.2.8: Solve $x''' + x = t^3 u(t - 1)$ for initial conditions $x(0) = 1$ and $x'(0) = 0$, $x''(0) = 0$.

Exercise 6.2.9: Using the definition of the Laplace transform, justify the second shifting property: $\mathcal{L}\{f(t - a) u(t - a)\} = e^{-as} \mathcal{L}\{f(t)\}$.

Exercise 6.2.10: Consider the mass-spring system with a rocket from [Example 6.2.2](#). We noticed that the solution kept oscillating after the rocket stopped running. The amplitude of the oscillation depends on the time that the rocket was fired (for 4 seconds in the example).

- Find a formula for the amplitude of the resulting oscillation in terms of the amount of time the rocket is fired.
- Is there a nonzero time (if so, what is it?) for which the rocket fires and the resulting oscillation has amplitude 0 (the mass is not moving)?

Exercise 6.2.11: Define

$$f(t) = \begin{cases} (t-1)^2 & \text{if } 1 \leq t < 2, \\ 3-t & \text{if } 2 \leq t < 3, \\ 0 & \text{otherwise.} \end{cases}$$

- Sketch the graph of $f(t)$.
- Write down $f(t)$ using the Heaviside function.
- Solve $x'' + x = f(t)$, $x(0) = 0$, $x'(0) = 0$ using the Laplace transform.

Exercise 6.2.12: Find the transfer function for $mx'' + cx' + kx = f(t)$, $x(0) = 0$, $x'(0) = 0$.

Exercise 6.2.13: Suppose $Lx = f(t)$, $x(0) = 0$, $x'(0) = 0$ for the input function $f(t) = t$ gives the output $x(t) = 2e^{-t} + te^{-t} + (t-2)$.

- Find $F(s)$, $X(s)$, and the transfer function $H(s)$.
- If the input is instead $f(t) = \sin(t)$, find the new output $x(t)$.

Exercise 6.2.14: Suppose $f(t) = 1$ if $0 \leq t < 1$ and $f(t) = 0$ if $1 \leq t < 2$, and then extend periodically for all $t \geq 0$ so that $f(t) = f(t+2)$. Compute the Laplace transform $F(s)$.

Exercise 6.2.101: Using the Heaviside function $u(t)$, write down the function

$$f(t) = \begin{cases} 0 & \text{if } t < 1, \\ t-1 & \text{if } 1 \leq t < 2, \\ 1 & \text{if } 2 \leq t. \end{cases}$$

Exercise 6.2.102: Solve $x'' - x = (t^2 - 1)u(t-1)$ for initial conditions $x(0) = 1$, $x'(0) = 2$ using the Laplace transform.

Exercise 6.2.103: Find the transfer function for $x' + x = f(t)$, $x(0) = 0$.

Exercise 6.2.104: Suppose $Lx = f(t)$, $x(0) = 0$, $x'(0) = 0$ for the input function $f(t) = 4$ gives the output $x(t) = 1 - \cos(2t)$. Find $F(s)$, $X(s)$, and the transfer function $H(s)$.

6.3 Convolution

Note: 1 or 1.5 lectures, §7.2 in [EP], §6.6 in [BD]

6.3.1 The convolution

The Laplace transformation of a product is not the product of the transforms. All hope is not lost, however. We simply have to use a different type of “product.” Take two functions $f(t)$ and $g(t)$ defined for $t \geq 0$, and define the *convolution** of $f(t)$ and $g(t)$ as

$$(f * g)(t) \stackrel{\text{def}}{=} \int_0^t f(\tau)g(t - \tau) d\tau. \quad (6.2)$$

As you can see, the convolution of two functions of t is another function of t .

Example 6.3.1: Take $f(t) = e^t$ and $g(t) = t$ for $t \geq 0$. Then

$$(f * g)(t) = \int_0^t e^\tau(t - \tau) d\tau = e^t - t - 1.$$

To solve the integral we did one integration by parts.

Example 6.3.2: Take $f(t) = \sin(\omega t)$ and $g(t) = \cos(\omega t)$ for $t \geq 0$. Then

$$(f * g)(t) = \int_0^t \sin(\omega \tau) \cos(\omega(t - \tau)) d\tau.$$

Apply the identity

$$\cos(\theta) \sin(\psi) = \frac{1}{2} (\sin(\theta + \psi) - \sin(\theta - \psi)),$$

to get

$$\begin{aligned} (f * g)(t) &= \int_0^t \frac{1}{2} (\sin(\omega t) - \sin(\omega t - 2\omega \tau)) d\tau \\ &= \left[\frac{1}{2} \tau \sin(\omega t) - \frac{1}{4\omega} \cos(\omega t - 2\omega \tau) \right]_{\tau=0}^t \\ &= \frac{1}{2} t \sin(\omega t). \end{aligned}$$

The formula holds only for $t \geq 0$. The functions f , g , and $f * g$ are undefined for $t < 0$.

*For those who have seen convolution before, you may have seen it defined as $(f * g)(t) = \int_{-\infty}^{\infty} f(\tau)g(t - \tau) d\tau$. This definition agrees with (6.2) if you define $f(t)$ and $g(t)$ to be zero for $t < 0$. When discussing the Laplace transform, the definition we gave is sufficient. Convolution does occur in many other applications, however, where you may have to use the more general definition with infinities.

Convolution has many properties that make it behave like a product. Let c be a constant and f , g , and h be functions. It is a calculus exercise to verify that

$$\begin{aligned} f * g &= g * f, \\ (cf) * g &= f * (cg) = c(f * g), \\ (f + g) * h &= f * h + g * h, \\ (f * g) * h &= f * (g * h). \end{aligned}$$

The first property means that $(f * g)(t) = \int_0^t f(\tau)g(t - \tau) d\tau = \int_0^t f(t - \tau)g(\tau) d\tau$. The most interesting property for us is the following theorem.

Theorem 6.3.1. *If $f(t)$ and $g(t)$ are of exponential order, then so is $(f * g)(t)$ and*

$$\mathcal{L}\{(f * g)(t)\} = \mathcal{L}\left\{\int_0^t f(\tau)g(t - \tau) d\tau\right\} = \mathcal{L}\{f(t)\}\mathcal{L}\{g(t)\}.$$

In other words, the Laplace transform of a convolution is the product of the Laplace transforms: $\mathcal{L}\{(f * g)(t)\} = F(s)G(s)$, or in reverse, $\mathcal{L}^{-1}\{F(s)G(s)\} = (f * g)(t)$.

Example 6.3.3: Suppose we wish to find the inverse Laplace transform of

$$\frac{1}{s^2(s + 1)} = \frac{1}{s^2} \frac{1}{s + 1}.$$

We recognize the two entries of [Table 6.2](#):

$$\mathcal{L}^{-1}\left\{\frac{1}{s^2}\right\} = t \quad \text{and} \quad \mathcal{L}^{-1}\left\{\frac{1}{s + 1}\right\} = e^{-t}.$$

Therefore, we convolve t and e^{-t} ,

$$\mathcal{L}^{-1}\left\{\frac{1}{s^2} \frac{1}{s + 1}\right\} = \int_0^t \tau e^{-(t-\tau)} d\tau = e^{-t} + t - 1.$$

The calculation of the integral involved an integration by parts.

6.3.2 Solving ODEs

The next example demonstrates the full power of the convolution and the Laplace transform. We can give the solution to the forced oscillation problem for any forcing function as a definite integral.

Example 6.3.4: Find the solution to

$$x'' + \omega_0^2 x = f(t), \quad x(0) = 0, \quad x'(0) = 0,$$

for an arbitrary function $f(t)$.

We first apply the Laplace transform to the equation. Denote the transform of $x(t)$ by $X(s)$ and the transform of $f(t)$ by $F(s)$ as usual. We get

$$s^2 X(s) + \omega_0^2 X(s) = F(s),$$

or in other words,

$$X(s) = \frac{1}{s^2 + \omega_0^2} F(s).$$

Recall that $H(s) = \frac{1}{s^2 + \omega_0^2}$ is the transfer function. We know

$$\mathcal{L}^{-1} \left\{ \frac{1}{s^2 + \omega_0^2} \right\} = \frac{\sin(\omega_0 t)}{\omega_0}.$$

Therefore,

$$x(t) = \int_0^t \frac{\sin(\omega_0 \tau)}{\omega_0} f(t - \tau) d\tau,$$

or, if we reverse the order,

$$x(t) = \int_0^t f(\tau) \frac{\sin(\omega_0(t - \tau))}{\omega_0} d\tau.$$

Notice one more feature of the example above. We can now see how the Laplace transform handles resonance. Suppose that $f(t) = \cos(\omega_0 t)$. Then

$$x(t) = \int_0^t \frac{\sin(\omega_0 \tau)}{\omega_0} \cos(\omega_0(t - \tau)) d\tau = \frac{1}{\omega_0} \int_0^t \sin(\omega_0 \tau) \cos(\omega_0(t - \tau)) d\tau.$$

We have computed the convolution of sine and cosine in [Example 6.3.2](#). Hence

$$x(t) = \left(\frac{1}{\omega_0} \right) \left(\frac{1}{2} t \sin(\omega_0 t) \right) = \frac{1}{2\omega_0} t \sin(\omega_0 t).$$

Note the t in front of the sine. The solution, therefore, grows without bound as t gets large, meaning we get resonance.

The general idea here is that if $H(s)$ is the transfer function, then $X(s) = H(s)F(s)$. If we find the $h(t) = \mathcal{L}^{-1}\{H(s)\}$, then

$$x(t) = \mathcal{L}^{-1}\{X(s)\} = \mathcal{L}^{-1}\{F(s)H(s)\} = (f * h)(t) = \int_0^t f(\tau)h(t - \tau) d\tau.$$

Hence, we can solve any constant-coefficient equation with an arbitrary forcing function $f(t)$ as a definite integral using convolution. A definite integral, rather than a closed-form solution, is usually enough for most practical purposes. It is not hard to numerically evaluate a definite integral.

6.3.3 Volterra integral equation

An integral equation involving convolution that comes up, for example, when studying systems with memory effects is the *Volterra integral equation**

$$x(t) = f(t) + \int_0^t g(t - \tau)x(\tau) d\tau.$$

Here $f(t)$ and $g(t)$ are known functions and $x(t)$ is an unknown we wish to solve for. We have $x(t) = f(t) + (g * x)(t)$, so the question is: If convolution is like a product, can we also “divide” to solve this equation? We apply the Laplace transform to the equation to obtain

$$X(s) = F(s) + G(s)X(s),$$

where $X(s)$, $F(s)$, and $G(s)$ are the Laplace transforms of $x(t)$, $f(t)$, and $g(t)$, respectively. We find

$$X(s) = \frac{F(s)}{1 - G(s)}.$$

To find $x(t)$, we now need to find the inverse Laplace transform of $X(s)$.

Example 6.3.5: Solve

$$x(t) = e^{-t} + \int_0^t \sinh(t - \tau)x(\tau) d\tau.$$

We apply the Laplace transform to obtain

$$X(s) = \frac{1}{s+1} + \frac{1}{s^2-1}X(s),$$

or

$$X(s) = \frac{\frac{1}{s+1}}{1 - \frac{1}{s^2-1}} = \frac{s-1}{s^2-2} = \frac{s}{s^2-2} - \frac{1}{s^2-2}.$$

It is not hard to apply [Table 6.1](#) on page 295 to find

$$x(t) = \cosh(\sqrt{2}t) - \frac{1}{\sqrt{2}} \sinh(\sqrt{2}t).$$

6.3.4 Exercises

Exercise 6.3.1: Let $f(t) = t^2$ for $t \geq 0$, and $g(t) = u(t-1)$. Compute $f * g$.

Exercise 6.3.2: Let $f(t) = t$ for $t \geq 0$, and $g(t) = \sin t$ for $t \geq 0$. Compute $f * g$.

Exercise 6.3.3: Find the solution to

$$mx'' + cx' + kx = f(t), \quad x(0) = 0, \quad x'(0) = 0,$$

for an arbitrary function $f(t)$, where $m > 0$, $c > 0$, $k > 0$, and $c^2 - 4km > 0$ (the system is overdamped). Write the solution as a definite integral.

*Named for the Italian mathematician [Vito Volterra](#) (1860–1940).

Exercise 6.3.4: Find the solution to

$$mx'' + cx' + kx = f(t), \quad x(0) = 0, \quad x'(0) = 0,$$

for an arbitrary function $f(t)$, where $m > 0$, $c > 0$, $k > 0$, and $c^2 - 4km < 0$ (the system is underdamped). Write the solution as a definite integral.

Exercise 6.3.5: Find the solution to

$$mx'' + cx' + kx = f(t), \quad x(0) = 0, \quad x'(0) = 0,$$

for an arbitrary function $f(t)$, where $m > 0$, $c > 0$, $k > 0$, and $c^2 = 4km$ (the system is critically damped). Write the solution as a definite integral.

Exercise 6.3.6: Solve

$$x(t) = e^{-t} + \int_0^t \cos(t - \tau)x(\tau) d\tau.$$

Exercise 6.3.7: Solve

$$x(t) = \cos t + \int_0^t \cos(t - \tau)x(\tau) d\tau.$$

Exercise 6.3.8: Compute $\mathcal{L}^{-1} \left\{ \frac{s}{(s^2+4)^2} \right\}$ using convolution.

Exercise 6.3.9: Write down the solution to $x'' - 2x = e^{-t^2}$, $x(0) = 0$, $x'(0) = 0$ as a definite integral. Hint: Do not try to compute the Laplace transform of e^{-t^2} .

Exercise 6.3.101: Let $f(t) = \cos t$ for $t \geq 0$, and $g(t) = e^{-t}$. Compute $f * g$.

Exercise 6.3.102: Compute $\mathcal{L}^{-1} \left\{ \frac{5}{s^4+s^2} \right\}$ using convolution.

Exercise 6.3.103: Solve $x'' + x = \sin t$, $x(0) = 0$, $x'(0) = 0$ using convolution.

Exercise 6.3.104: Solve $x''' + x' = f(t)$, $x(0) = 0$, $x'(0) = 0$, $x''(0) = 0$ using convolution. Write the result as a definite integral.

6.4 Dirac delta and impulse response

Note: 1 or 1.5 lecture, §7.6 in [EP], §6.5 in [BD]

6.4.1 Rectangular pulse

Often we study a physical system by putting in a short pulse and then seeing what the system does. The resulting behavior is often called *impulse response*, and understanding it tells us how the system responds to any input. Let us see what we mean by a pulse. The simplest kind of a pulse is a *rectangular pulse* defined by

$$\varphi(t) = \begin{cases} 0 & \text{if } t < a, \\ M & \text{if } a \leq t < b, \\ 0 & \text{if } b \leq t. \end{cases}$$

See Figure 6.3 for a graph.

Notice that

$$\varphi(t) = M(u(t - a) - u(t - b)),$$

where $u(t)$ is the unit step function. The Laplace transform of a rectangular pulse is

$$\begin{aligned} \mathcal{L}\{\varphi(t)\} &= \mathcal{L}\{M(u(t - a) - u(t - b))\} \\ &= M \frac{e^{-as} - e^{-bs}}{s}. \end{aligned}$$

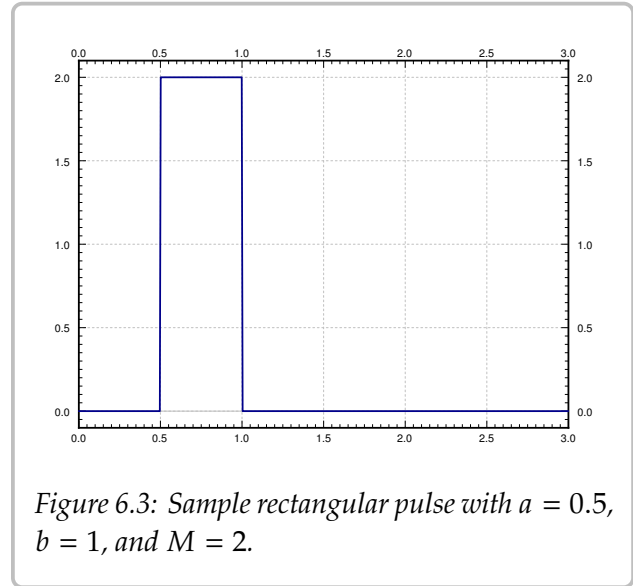
For simplicity, let $a = 0$. We also set $M = 1/b$ so that

$$\int_0^\infty \varphi(t) dt = 1.$$

That is, we have a pulse of “unit mass.” This way for any b the pulse has the same “oomph.” For such a pulse,

$$\mathcal{L}\{\varphi(t)\} = \mathcal{L}\left\{\frac{u(t) - u(t - b)}{b}\right\} = \frac{1 - e^{-bs}}{bs}.$$

We want b to be very small; we wish to have the pulse be very short and very tall. By letting b go to zero, we arrive at the concept of the Dirac delta function. The limit of $\frac{1 - e^{-bs}}{bs}$ as $b \rightarrow 0$ is 1, so we are looking for a “function” whose Laplace transform is 1.



6.4.2 The delta function

The *Dirac delta function*^{*} is not exactly a function; it is sometimes called a *generalized function*. We avoid unnecessary details and simply say that it is an object that does not really make sense unless we integrate it. The motivation is that we would like a “function” $\delta(t)$ such that for any continuous function $f(t)$,

$$\int_{-\infty}^{\infty} \delta(t) f(t) dt = f(0).$$

The formula should hold if we integrate over any interval that contains 0, not just $(-\infty, \infty)$. So $\delta(t)$ is a “function” with all its “mass” at the single point $t = 0$. For any interval[†] $[c, d]$,

$$\int_c^d \delta(t) dt = \begin{cases} 1 & \text{if the interval } [c, d] \text{ contains } 0, \text{ i.e. } c \leq 0 \leq d, \\ 0 & \text{otherwise.} \end{cases}$$

Unfortunately, there is no such function in the classical sense. You could informally think that $\delta(t)$ is zero for $t \neq 0$ and somehow infinite at $t = 0$.

A good way to think about $\delta(t)$ is as a limit of pulses of decreasing length, each having an integral of 1. For example, consider a rectangular pulse $\varphi(t)$ as above with $a = 0$ and $M = 1/b$, that is, $\varphi(t) = \frac{u(t) - u(t-b)}{b}$. Compute

$$\int_{-\infty}^{\infty} \varphi(t) f(t) dt = \int_{-\infty}^{\infty} \frac{u(t) - u(t-b)}{b} f(t) dt = \frac{1}{b} \int_0^b f(t) dt.$$

If $f(t)$ is continuous at $t = 0$, then for very small b , the function $f(t)$ is approximately equal to $f(0)$ on the interval $[0, b]$. We approximate the integral

$$\frac{1}{b} \int_0^b f(t) dt \approx \frac{1}{b} \int_0^b f(0) dt = f(0).$$

Hence,

$$\lim_{b \rightarrow 0} \int_{-\infty}^{\infty} \varphi(t) f(t) dt = \lim_{b \rightarrow 0} \frac{1}{b} \int_0^b f(t) dt = f(0).$$

Let us therefore accept $\delta(t)$ as an object that is possible to integrate. We often want to shift δ to another point, for example $\delta(t - a)$. In that case,

$$\int_{-\infty}^{\infty} \delta(t - a) f(t) dt = f(a).$$

Note that $\delta(a - t)$ is the same object as $\delta(t - a)$. With that, the convolution of $\delta(t)$ with $f(t)$ is again $f(t)$,

$$(f * \delta)(t) = \int_0^t f(\tau) \delta(t - \tau) d\tau = f(t).$$

^{*}Named after the English physicist and mathematician [Paul Adrien Maurice Dirac](#) (1902–1984).

[†]It is important that we consider c and d as part of the interval. One could write this integral as $\int_{c-}^{d+}$.

As we can integrate $\delta(t)$, we compute its Laplace transform:

$$\mathcal{L}\{\delta(t-a)\} = \int_0^{\infty} e^{-st} \delta(t-a) dt = e^{-as}.$$

In particular,

$$\mathcal{L}\{\delta(t)\} = 1.$$

Remark 6.4.1: The Laplace transform of $\delta(t-a)$ would be the Laplace transform of the derivative of the Heaviside function $u(t-a)$, if the Heaviside function had a derivative. First,

$$\mathcal{L}\{u(t-a)\} = \frac{e^{-as}}{s}.$$

To obtain what the Laplace transform of the derivative would be, we multiply by s to obtain e^{-as} , which is the Laplace transform of $\delta(t-a)$. We see the same thing using integration,

$$\int_0^t \delta(s-a) ds = u(t-a).$$

So in a certain sense

$$\frac{d}{dt} [u(t-a)] = \delta(t-a).$$

This line of reasoning allows us to talk about derivatives of functions with jump discontinuities. We can think of the derivative of the Heaviside function $u(t-a)$ as being somehow infinite at a , which is precisely our intuitive understanding of the delta function.

Example 6.4.1: Let us compute $\mathcal{L}^{-1}\left\{\frac{s+1}{s}\right\}$. So far, we only computed the inverse transform of proper rational functions in the s variable. That is, the numerator was of lower degree than the denominator. Not so with $\frac{s+1}{s}$. We can use the delta function to compute,

$$\mathcal{L}^{-1}\left\{\frac{s+1}{s}\right\} = \mathcal{L}^{-1}\left\{1 + \frac{1}{s}\right\} = \mathcal{L}^{-1}\{1\} + \mathcal{L}^{-1}\left\{\frac{1}{s}\right\} = \delta(t) + 1.$$

The resulting object is a generalized function and only makes sense when put underneath an integral.

6.4.3 Impulse response

As we said before, in the differential equation $Lx = f(t)$, we think of $f(t)$ as the input, and $x(t)$ as the output. We think of the delta function as an impulse, and so to find the response to an impulse, we use the delta function in place of $f(t)$. The solution to

$$Lx = \delta(t)$$

is called the *impulse response*.

Example 6.4.2: Solve (find the impulse response)

$$x'' + \omega_0^2 x = \delta(t), \quad x(0) = 0, \quad x'(0) = 0. \quad (6.3)$$

Apply the Laplace transform to the equation, and denote the transform of $x(t)$ by $X(s)$:

$$s^2 X(s) + \omega_0^2 X(s) = 1, \quad \text{and so} \quad X(s) = \frac{1}{s^2 + \omega_0^2}.$$

The inverse Laplace transform produces (for $t > 0$)

$$x(t) = \frac{\sin(\omega_0 t)}{\omega_0}.$$

Remark 6.4.2: Perhaps an astute reader will notice that it does not seem like $x'(0) = 0$. However, we really want to think that $x(t) = 0$ for $t \leq 0$, so $x(t)$ has a “corner” at $t = 0$, that is, x' has a jump discontinuity there, which is what produces the $\delta(t)$ when we take x'' . See [Remark 6.4.1](#). The initial condition really is $x'(0-) = \lim_{t \uparrow 0} x'(t) = 0$.

Let us notice something about the example above. In [Example 6.3.4](#), we found that when the input is $f(t)$, the solution to $Lx = f(t)$ is given by

$$x(t) = \int_0^t f(\tau) \frac{\sin(\omega_0(t - \tau))}{\omega_0} d\tau.$$

That is, *the solution for an arbitrary input is given as convolution with the impulse response*. Let us see why. For functions $f(t)$ and $h(t)$ that are zero for negative t ,

$$(f * h)''(t) = \frac{d^2}{dt^2} \left[\int_0^t f(\tau) h(t - \tau) d\tau \right] = \int_0^t f(\tau) h''(t - \tau) d\tau = (f * h'')(t).$$

We differentiate twice under the integral*. Suppose that $h(t)$ is the impulse response (solution to $h'' + \omega_0^2 h = \delta(t)$). If we convolve the entire equation (6.3), the left-hand side becomes

$$f * (h'' + \omega_0^2 h) = (f * h'') + \omega_0^2 (f * h) = (f * h)'' + \omega_0^2 (f * h).$$

The right-hand side becomes

$$(f * \delta)(t) = f(t).$$

Therefore, if h is the impulse response, then $x(t) = (f * h)(t)$ is the solution to

$$x'' + \omega_0^2 x = f(t).$$

The procedure works for any constant-coefficient linear equation $Lx = f(t)$. If you find the impulse response h (solution to $Lh = \delta(t)$), you also know how to obtain the output $x(t)$ for any input $f(t)$. We simply convolve, $x(t) = (f * h)(t)$. As you may have noticed in the example, the impulse response is in fact just the inverse Laplace transform of the transfer function, that is, $h(t) = \mathcal{L}^{-1}\{H(s)\}$.

*You may be worried about the fact that our h is not really twice differentiable, in fact, $h'(0)$ does not exist as a normal derivative, and h'' involves the δ function. You can pretend that h is twice differentiable for this computation and either apply the standard Leibniz rule (that $h(t) = 0$ for $t < 0$ is important) or note that for such functions, the integral is actually over $(-\infty, \infty)$ rather than over $[0, t]$.

6.4.4 Three-point beam bending

A quite different example application where the delta function appears is in representing point loads on a steel beam. Consider a beam of length L , resting on two simple supports at the ends. Let x denote the position on the beam, and let $y(x)$ denote the deflection of the beam in the vertical direction. The deflection $y(x)$ satisfies the *Euler–Bernoulli equation*^{*},

$$EI \frac{d^4 y}{dx^4} = F(x),$$

where E and I are constants[†] and $F(x)$ is the force applied per unit length at position x . The situation we are interested in is when the force is applied at a single point as in [Figure 6.4](#).

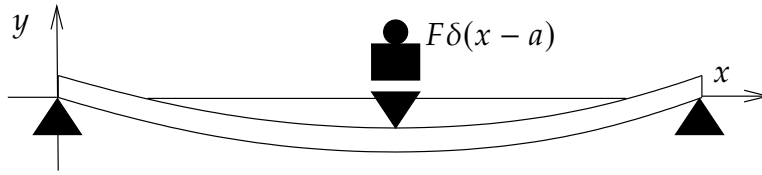


Figure 6.4: Three-point bending.

The equation becomes

$$EI \frac{d^4 y}{dx^4} = -F\delta(x - a),$$

where $x = a$ is the point where the mass is applied. The constant F is the force applied and the minus sign indicates that the force is downward, that is, in the negative y direction. The end points of the beam satisfy the conditions,

$$\begin{aligned} y(0) &= 0, & y''(0) &= 0, \\ y(L) &= 0, & y''(L) &= 0. \end{aligned}$$

See [§ 5.2](#) for further information about endpoint conditions applied to beams.

Example 6.4.3: Suppose that the length of the beam is 2, and $EI = 1$ for simplicity. Further suppose that a force $F = 1$ is applied at $x = 1$. That is, we have the equation

$$\frac{d^4 y}{dx^4} = -\delta(x - 1),$$

and the endpoint conditions are

$$y(0) = 0, \quad y''(0) = 0, \quad y(2) = 0, \quad y''(2) = 0.$$

^{*}Named for the Swiss mathematicians [Jacob Bernoulli](#) (1654–1705), [Daniel Bernoulli](#) (1700–1782), the nephew of Jacob, and [Leonhard Paul Euler](#) (1707–1783).

[†] E is the elastic modulus and I is the second moment of area. Let us not worry about the details and simply think of these as some given constants.

We could integrate, but using the Laplace transform is even easier. We apply the transform in the x variable rather than the t variable. We again denote the transform of $y(x)$ by $Y(s)$.

$$s^4 Y(s) - s^3 y(0) - s^2 y'(0) - s y''(0) - y'''(0) = -e^{-s}.$$

We notice that $y(0) = 0$ and $y''(0) = 0$. Let us call $C_1 = y'(0)$ and $C_2 = y'''(0)$. We solve for $Y(s)$,

$$Y(s) = \frac{-e^{-s}}{s^4} + \frac{C_1}{s^2} + \frac{C_2}{s^4}.$$

We take the inverse Laplace transform utilizing the second shifting property (6.1) for the first term:

$$y(x) = \frac{-(x-1)^3}{6} u(x-1) + C_1 x + \frac{C_2}{6} x^3.$$

We still need to apply two of the endpoint conditions. As the conditions are at $x = 2$ we can simply replace $u(x-1) = 1$ when taking the derivatives. Therefore,

$$0 = y(2) = \frac{-(2-1)^3}{6} + C_1(2) + \frac{C_2}{6} 2^3 = \frac{-1}{6} + 2C_1 + \frac{4}{3}C_2,$$

and

$$0 = y''(2) = \frac{-3 \cdot 2 \cdot (2-1)}{6} + \frac{C_2}{6} 3 \cdot 2 \cdot 2 = -1 + 2C_2.$$

Hence, $C_2 = \frac{1}{2}$. Solving for C_1 using the first equation, we obtain $C_1 = \frac{-1}{4}$. Our solution for the beam deflection is

$$y(x) = \frac{-(x-1)^3}{6} u(x-1) - \frac{x}{4} + \frac{x^3}{12}.$$

6.4.5 Exercises

Exercise 6.4.1: Solve (find the impulse response) $x'' + x' + x = \delta(t)$, $x(0) = 0$, $x'(0) = 0$.

Exercise 6.4.2: Solve (find the impulse response) $x'' + 2x' + x = \delta(t)$, $x(0) = 0$, $x'(0) = 0$.

Exercise 6.4.3: A pulse can come later and can be bigger. Solve $x'' + 4x = 4\delta(t-1)$, $x(0) = 0$, $x'(0) = 0$.

Exercise 6.4.4: Suppose that $f(t)$ and $g(t)$ are differentiable functions (and the derivatives are continuous) and suppose that $f(t) = g(t) = 0$ for all $t \leq 0$. Show that

$$(f * g)'(t) = (f' * g)(t) = (f * g')(t).$$

Exercise 6.4.5: Suppose that $Lx = \delta(t)$, $x(0) = 0$, $x'(0) = 0$, has the solution $x(t) = te^{-t}$ for $t > 0$. Find the solution to $Lx = t^2$, $x(0) = 0$, $x'(0) = 0$ for $t > 0$.

Exercise 6.4.6: Compute $\mathcal{L}^{-1} \left\{ \frac{s^2 + s + 1}{s^2} \right\}$.

Exercise 6.4.7 (challenging): Solve [Example 6.4.3](#) via integrating 4 times in the x variable.

Exercise 6.4.8: Suppose we have a beam of length 1 simply supported at the ends and suppose that a force $F = 1$ is applied at $x = 3/4$ in the downward direction. Suppose that $EI = 1$ for simplicity. Find the beam deflection $y(x)$.

Exercise 6.4.101: Solve (find the impulse response) $x'' = \delta(t)$, $x(0) = 0$, $x'(0) = 0$.

Exercise 6.4.102: Solve (find the impulse response) $x' + ax = \delta(t)$, $x(0) = 0$.

Exercise 6.4.103: Suppose that $Lx = \delta(t)$, $x(0) = 0$, $x'(0) = 0$, has the solution $x(t) = e^t \sin(t)$ for $t > 0$. Find (in closed form) the solution to $Lx = e^t$, $x(0) = 0$, $x'(0) = 0$ for $t > 0$.

Exercise 6.4.104: Compute $\mathcal{L}^{-1} \left\{ \frac{s^2}{s^2+1} \right\}$.

Exercise 6.4.105: Compute $\mathcal{L}^{-1} \left\{ \frac{3s^2 e^{-s} + 2}{s^2} \right\}$.

6.5 Solving PDEs with the Laplace transform

Note: 1–1.5 lecture, can be skipped

The Laplace transform comes from the same family of transforms as does the Fourier series*, which we used in [chapter 4](#) to solve partial differential equations (PDEs). It is therefore not surprising that we can also solve PDEs with the Laplace transform.

Given a PDE in two independent variables x and t , we use the Laplace transform on one of the variables (taking the transform of everything in sight), and derivatives in that variable become multiplications by the transformed variable s . The PDE becomes an ODE, which we solve. Afterwards, we invert the transform to find a solution to the original problem. It is best to see the procedure in an example.

Example 6.5.1: Consider the first-order PDE

$$y_t = -\alpha y_x, \quad \text{for } x > 0, \ t > 0,$$

with side conditions

$$y(0, t) = C, \quad y(x, 0) = 0.$$

We will assume $\alpha > 0$ is a constant. This equation is called the *convection equation* or sometimes the *transport equation*, and it already made an appearance in [§ 1.9](#), with different conditions. See [Figure 6.5](#) for a diagram of the setup.

A physical setup of this equation is a river of viscous goo, as we do not want anything to diffuse. The function y is the concentration of some toxic substance†. The variable x denotes position, where $x = 0$ is the location of a factory spewing the toxic substance into the river. The toxic substance flows into the river so that at $x = 0$ the concentration is always C . We wish to see what happens past the factory, that is, at $x > 0$. Let t be the time, and assume the factory started operations at $t = 0$, so that at $t = 0$ the river is just pure goo.

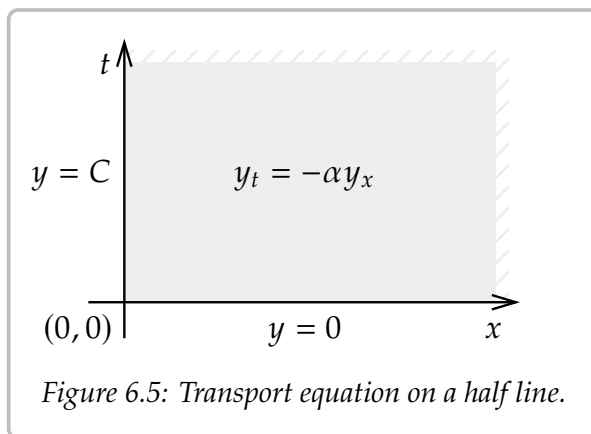


Figure 6.5: Transport equation on a half line.

Consider a function of two variables $y(x, t)$. Let us fix x and transform the t variable. For convenience, we treat the transformed s variable as a parameter, since there are no derivatives in s . That is, we write $Y(x)$ for the transformed function, and treat it as a function of x , leaving s as a parameter.

$$Y(x) = \mathcal{L}\{y(x, t)\} = \int_0^\infty y(x, t)e^{-st} dt.$$

*There is also a Fourier transform on the real line that looks sort of like the Laplace transform.

†It's a river of goo already. We're not hurting the environment much more.

The transform of a derivative with respect to x is just the derivative of the transformed function:

$$\mathcal{L}\{y_x(x, t)\} = \int_0^\infty y_x(x, t)e^{-st} dt = \frac{d}{dx} \left[\int_0^\infty y(x, t)e^{-st} dt \right] = Y'(x).$$

To transform the derivative in t (the variable being transformed), we use the rules from § 6.2:

$$\mathcal{L}\{y_t(x, t)\} = sY(x) - y(x, 0).$$

In our specific case, $y(x, 0) = 0$, and so $\mathcal{L}\{y_t(x, t)\} = sY(x)$. We transform the equation to find

$$sY(x) = -\alpha Y'(x).$$

This ODE needs an initial condition. The initial condition is the other side condition of the PDE, the one that depends on x . Everything is transformed, so we must also transform this condition

$$Y(0) = \mathcal{L}\{y(0, t)\} = \mathcal{L}\{C\} = \frac{C}{s}.$$

We solve the ODE problem $sY(x) = -\alpha Y'(x)$, $Y(0) = \frac{C}{s}$, to find

$$Y(x) = \frac{C}{s} e^{-\frac{s}{\alpha}x}.$$

We are not done. We have $Y(x)$, but we really want $y(x, t)$. We transform the s variable back to t . Let

$$u(t) = \begin{cases} 0 & \text{if } t < 0, \\ 1 & \text{otherwise} \end{cases}$$

be the Heaviside function. As

$$\mathcal{L}\{u(t-a)\} = \int_0^\infty u(t-a)e^{-st} dt = \int_a^\infty e^{-st} dt = \frac{e^{-as}}{s},$$

then

$$y(x, t) = \mathcal{L}^{-1} \left\{ \frac{C}{s} e^{-\frac{s}{\alpha}x} \right\} = Cu(t - x/\alpha).$$

In other words,

$$y(x, t) = \begin{cases} 0 & \text{if } t < x/\alpha, \\ C & \text{otherwise.} \end{cases}$$

See Figure 6.6 on the next page for a diagram of this solution. The line of slope $1/\alpha$ indicates the wavefront of the toxic substance in the picture as it is leaving the factory. What the equation does is to try to move the initial condition to the right at speed α , and so it moves the boundary condition (the wavefront) too.

Shhh. . . y is not differentiable. It is not even continuous (nobody ever seems to notice). How could we plug something that's not differentiable into the equation? Well, just think

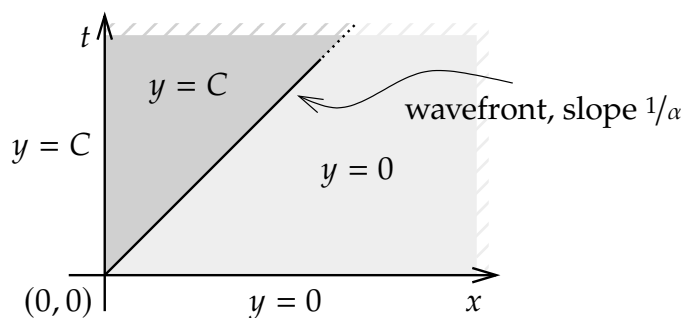


Figure 6.6: Wavefront of toxic substance is a line of slope $1/\alpha$.

of a differentiable function very, very close to y . Or, if you recognize the derivative of the Heaviside function as the delta function, then all is well too:

$$y_t(x, t) = \frac{\partial}{\partial t} [Cu(t - x/\alpha)] = Cu'(t - x/\alpha) = C\delta(t - x/\alpha)$$

and

$$y_x(x, t) = \frac{\partial}{\partial x} [Cu(t - x/\alpha)] = -\frac{C}{\alpha}u'(t - x/\alpha) = -\frac{C}{\alpha}\delta(t - x/\alpha).$$

So $y_t = -\alpha y_x$.

The Laplace transform is very good with constant-coefficient equations. One advantage of Laplace is that it easily handles nonhomogeneous side conditions. Let us try a more complicated example.

Example 6.5.2: Consider

$$\begin{aligned} y_t + y_x + y &= 0, & \text{for } x > 0, \ t > 0, \\ y(0, t) &= \sin(t), & y(x, 0) &= 0. \end{aligned}$$

Again, we transform t , and we write $Y(x)$ for the transformed function. As $y(x, 0) = 0$, we find

$$sY(x) + Y'(x) + Y(x) = 0, \quad Y(0) = \frac{1}{s^2 + 1}.$$

The solution of the transformed equation is

$$Y(x) = \frac{1}{s^2 + 1}e^{-(s+1)x} = \frac{1}{s^2 + 1}e^{-xs}e^{-x}.$$

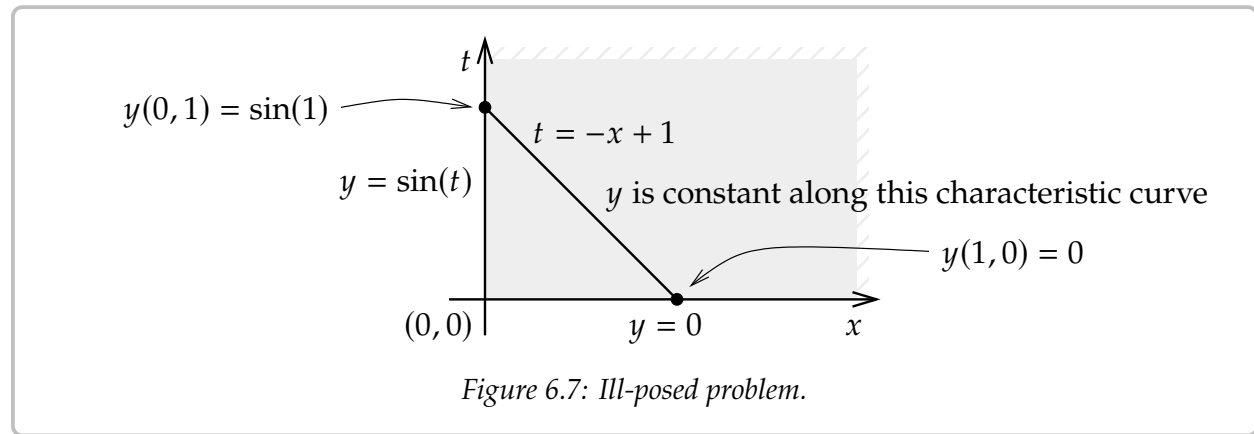
Using the second shifting property (6.1) and linearity of the transform, we obtain the solution

$$y(x, t) = e^{-x} \sin(t - x)u(t - x).$$

We can also detect when the problem is *ill-posed* in the sense that it has no solution. Let us change the equation to

$$\begin{aligned} -y_t + y_x &= 0, & \text{for } x > 0, \ t > 0, \\ y(0, t) &= \sin(t), & y(x, 0) = 0. \end{aligned}$$

Then the problem has no solution. First, let us see this in the language of § 1.9. The characteristic curves are $t = -x + C$. If τ is the characteristic coordinate, then we find the equation $y_\tau = 0$ along the curve, meaning a solution is constant along characteristic curves. But these curves intersect both the x -axis and the t -axis. For example, the curve $t = -x + 1$ intersects at $(1, 0)$ and $(0, 1)$. The solution is constant along the curve so $y(1, 0)$ should equal $y(0, 1)$. But $y(1, 0) = 0$ and $y(0, 1) = \sin(1) \neq 0$. See Figure 6.7.



Now consider the transform. The transformed problem is

$$-sY(x) + Y'(x) = 0, \quad Y(0) = \frac{1}{s^2 + 1},$$

and the solution ought to be

$$Y(x) = \frac{1}{s^2 + 1} e^{sx}.$$

Importantly, this Laplace transform does not decay to zero at infinity! That is, since $x > 0$ in the region of interest,

$$\lim_{s \rightarrow \infty} \frac{1}{s^2 + 1} e^{sx} = \infty \neq 0.$$

It almost looks as if we could use the shifting property, but notice that the shift is in the wrong direction.

Of course, we need not restrict ourselves to first-order equations, although the computations become more involved for higher-order equations.

Example 6.5.3: Let us use Laplace for the following problem:

$$\begin{aligned} y_t &= y_{xx}, & 0 < x < \infty, \quad t > 0, \\ y_x(0, t) &= f(t), \\ y(x, 0) &= 0. \end{aligned}$$

This problem corresponds to a half-infinite insulated rod with a given heat flux at one end. The setup would come up in real life in the case of a sufficiently long rod, where we would not consider a long enough time for our solution to notice that we do not have an infinite rod but just a very long one. In any case, the real-life situation imposes other conditions on our solution y . For example, we will assume that the solution is bounded. This boundedness condition will stand in for a boundary condition at the infinite end. Boundedness is also sufficient for the Laplace transform in the t variable to exist.

Transform the equation in the t variable to find

$$sY(x) = Y''(x).$$

The general solution to this ODE is

$$Y(x) = Ae^{\sqrt{s}x} + Be^{-\sqrt{s}x}.$$

Note that A and B depend on s . As y is bounded, then $Y(x)$ is bounded for any fixed $s > 0$, so for $Y(x)$ to stay bounded as $x \rightarrow \infty$, we must have $A = 0$.

Now consider the boundary condition at $x = 0$. Transform $Y'(0) = F(s)$ and so $-\sqrt{s}B = F(s)$. In other words,

$$Y(x) = F(s) \frac{-1}{\sqrt{s}} e^{-\sqrt{s}x}.$$

If we look up the inverse transform in a table such as the one in [Appendix B](#) (or we spend the afternoon doing calculus), we find

$$\mathcal{L}^{-1} \left[\frac{1}{\sqrt{s}} e^{-\sqrt{s}x} \right] = \frac{1}{\sqrt{\pi t}} e^{-\frac{x^2}{4t}}.$$

So

$$y(x, t) = \mathcal{L}^{-1} \left[F(s) \frac{-1}{\sqrt{s}} e^{-\sqrt{s}x} \right] = \int_0^t f(\tau) \frac{-1}{\sqrt{\pi(t-\tau)}} e^{-\frac{x^2}{4(t-\tau)}} d\tau.$$

Laplace can solve problems where separation of variables fails. Laplace does not mind nonhomogeneity, but it is essentially only useful for constant-coefficient equations.

6.5.1 Exercises

Exercise 6.5.1: Solve

$$\begin{aligned} y_t + y_x &= 1, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= 1, & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.2: Solve

$$\begin{aligned} y_t + \alpha y_x &= 0, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= t, & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.3: Solve

$$\begin{aligned} y_t + 2y_x &= x + t, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= 0, & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.4: For an $\alpha > 0$, solve

$$\begin{aligned} y_t + \alpha y_x + y &= 0, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= \sin(t), & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.5: Find the corresponding ODE problem for $Y(x)$, after transforming the t variable

$$\begin{aligned} y_{tt} + 3y_{xx} + y_{xt} + 3y_x + y &= \sin(x) + t, & 0 < x < 1, \quad t > 0, \\ y(0, t) &= 1, & y(1, t) &= t, & y(x, 0) &= 1 - x, & y_t(x, 0) &= 1. \end{aligned}$$

Do not solve the problem.

Exercise 6.5.6: Write down a solution to

$$\begin{aligned} y_t &= y_{xx}, & 0 < x < \infty, \quad t > 0, \\ y_x(0, t) &= e^{-t}, & y(x, 0) &= 0, \end{aligned}$$

as a definite integral (convolution). Assume that y is bounded.

Exercise 6.5.7: Use the Laplace transform in t to solve

$$\begin{aligned} y_{tt} &= y_{xx}, & -\infty < x < \infty, \quad t > 0, \\ y(x, 0) &= 0, & y_t(x, 0) &= \sin(x). \end{aligned}$$

Assume that y is bounded. Hint: Note that for any fixed $s > 0$, e^{sx} blows up as $x \rightarrow +\infty$ and e^{-sx} blows up as $x \rightarrow -\infty$.

Exercise 6.5.101: Solve

$$\begin{aligned} y_t + y_x &= 1, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= 0, & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.102: For a $c > 0$, solve

$$\begin{aligned} y_t + y_x + cy &= 0, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= \sin(t), & y(x, 0) &= 0. \end{aligned}$$

Exercise 6.5.103: Find the corresponding ODE problem for $Y(x)$, after transforming the t variable

$$\begin{aligned} y_{tt} + 3y_{xx} + y &= x + t, & -1 < x < 1, \quad t > 0, \\ y(-1, t) &= 0, \quad y(1, t) = 0, & y(x, 0) = (1 - x^2), \quad y_t(x, 0) = 0. \end{aligned}$$

Do not solve the problem.

Exercise 6.5.104: Use the Laplace transform in t to solve

$$\begin{aligned} y_{tt} &= y_{xx}, & -\infty < x < \infty, \quad t > 0, \\ y(x, 0) &= \sin(x), \quad y_t(x, 0) = 0. \end{aligned}$$

Assume that y is bounded. Hint: Note that for any fixed $s > 0$, e^{sx} blows up as $x \rightarrow +\infty$ and e^{-sx} blows up as $x \rightarrow -\infty$.

Exercise 6.5.105: Use the Laplace transform in t to solve

$$\begin{aligned} y_t &= y_{xx}, & 0 < x < \infty, \quad t > 0, \\ y(0, t) &= f(t), \quad y(x, 0) = 0, \end{aligned}$$

where $f(t)$ is some function. Assume that y is bounded. Give the answer as a convolution.

Chapter 7

Power-series methods

7.1 Power series

Note: 1.5 or 2 lectures, §8.1 in [EP], §5.1 in [BD]

Many functions can be written in terms of a power series

$$\sum_{k=0}^{\infty} a_k (x - x_0)^k.$$

Assuming a solution of a differential equation is a power series, we can perhaps use a method reminiscent of undetermined coefficients—we try to solve for the numbers a_k . Before we carry out this process, we review some results and concepts about power series.

7.1.1 Definition

As we said, a *power series* is an expression such as

$$\sum_{k=0}^{\infty} a_k (x - x_0)^k = a_0 + a_1(x - x_0) + a_2(x - x_0)^2 + a_3(x - x_0)^3 + \cdots, \quad (7.1)$$

where $a_0, a_1, a_2, \dots, a_k, \dots$ and x_0 are constants. Let

$$S_n(x) = \sum_{k=0}^n a_k (x - x_0)^k = a_0 + a_1(x - x_0) + a_2(x - x_0)^2 + a_3(x - x_0)^3 + \cdots + a_n(x - x_0)^n,$$

denote the so-called *partial sum*. If for some x , the limit

$$\lim_{n \rightarrow \infty} S_n(x) = \lim_{n \rightarrow \infty} \sum_{k=0}^n a_k (x - x_0)^k$$

exists, we say the series (7.1) *converges* at x . At $x = x_0$, the series always converges to a_0 . When (7.1) converges at any other $x \neq x_0$, we say (7.1) is a *convergent power series*, and we write

$$\sum_{k=0}^{\infty} a_k(x - x_0)^k = \lim_{n \rightarrow \infty} \sum_{k=0}^n a_k(x - x_0)^k.$$

If the series does not converge for any point $x \neq x_0$, we say that the series is *divergent*.

Example 7.1.1: The series

$$\sum_{k=0}^{\infty} \frac{1}{k!} x^k = 1 + x + \frac{x^2}{2} + \frac{x^3}{6} + \cdots$$

is convergent for any x . Recall that $k! = 1 \cdot 2 \cdot 3 \cdots k$ is the factorial. By convention, we define $0! = 1$. You may recall that this series converges to e^x .

We say that (7.1) *converges absolutely* at x whenever the limit

$$\lim_{n \rightarrow \infty} \sum_{k=0}^n |a_k| |x - x_0|^k$$

exists. That is, the series $\sum_{k=0}^{\infty} |a_k| |x - x_0|^k$ is convergent. If (7.1) converges absolutely at x , then it converges at x . However, the opposite implication is not true.

Example 7.1.2: The series

$$\sum_{k=1}^{\infty} \frac{1}{k} x^k$$

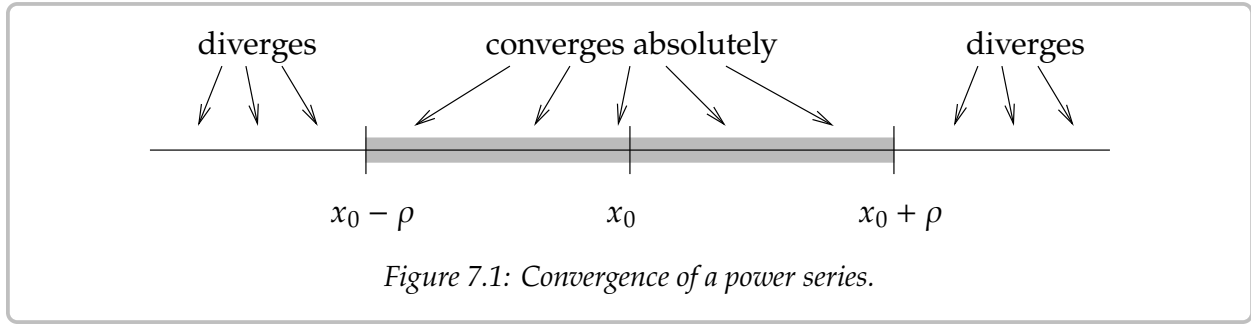
converges absolutely for all x in the interval $(-1, 1)$. It converges at $x = -1$, as $\sum_{k=1}^{\infty} \frac{(-1)^k}{k}$ converges (conditionally) by the alternating series test. The power series does not converge absolutely at $x = -1$, because $\sum_{k=1}^{\infty} \frac{1}{k}$ does not converge. The series diverges at $x = 1$.

7.1.2 Radius of convergence

If a power series converges absolutely at some x_1 , then for all x such that $|x - x_0| \leq |x_1 - x_0|$ (that is, x is no further than x_1 to x_0), we have $|a_k(x - x_0)^k| \leq |a_k(x_1 - x_0)^k|$ for all k . As the numbers $|a_k(x_1 - x_0)^k|$ sum to some finite limit, summing smaller positive numbers $|a_k(x - x_0)^k|$ must also have a finite limit. Hence, the series must converge absolutely at x .

Theorem 7.1.1. For a power series (7.1), there exists a number ρ (we allow $\rho = \infty$) called the radius of convergence such that the series converges absolutely on the interval $(x_0 - \rho, x_0 + \rho)$ and diverges for $x < x_0 - \rho$ and $x > x_0 + \rho$. We write $\rho = \infty$ if the series converges for all x .

See Figure 7.1 on the next page. In Example 7.1.1, the radius of convergence is $\rho = \infty$ as the series converges everywhere. In Example 7.1.2, the radius of convergence is $\rho = 1$. We note that $\rho = 0$ is another way of saying that the series is divergent.



A useful test for convergence of a series is the *ratio test*. Suppose that

$$\sum_{k=0}^{\infty} c_k$$

is a series and the limit

$$L = \lim_{k \rightarrow \infty} \left| \frac{c_{k+1}}{c_k} \right|$$

exists. Then the series converges absolutely if $L < 1$ and diverges if $L > 1$.

We apply this test to the power series (7.1). Let $c_k = a_k(x - x_0)^k$ in the test. Compute

$$L = \lim_{k \rightarrow \infty} \left| \frac{c_{k+1}}{c_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}(x - x_0)^{k+1}}{a_k(x - x_0)^k} \right| = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| |x - x_0|.$$

Define A by

$$A = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right|.$$

Then the series (7.1) converges absolutely if $1 > L = A|x - x_0|$. If $A > 0$, then the series converges absolutely if $|x - x_0| < 1/A$, and diverges if $|x - x_0| > 1/A$. That is, the radius of convergence is $1/A$. If $A = 0$, then the series always converges.

The *root test* is similar. Suppose

$$L = \lim_{k \rightarrow \infty} \sqrt[k]{|c_k|}$$

exists. Then $\sum_{k=0}^{\infty} c_k$ converges absolutely if $L < 1$ and diverges if $L > 1$. We can use the same calculation as above to find A . Let us summarize.

Theorem 7.1.2 (Ratio and root tests for power series). *Consider a power series*

$$\sum_{k=0}^{\infty} a_k(x - x_0)^k$$

such that

$$A = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| \quad \text{or} \quad A = \lim_{k \rightarrow \infty} \sqrt[k]{|a_k|}$$

exists. If $A = 0$, then the radius of convergence of the series is ∞ . Otherwise, the radius of convergence is $1/A$. Moreover, if either limit for A diverges to ∞ , then the series is divergent.

Example 7.1.3: Consider

$$\sum_{k=0}^{\infty} 2^{-k}(x-1)^k.$$

We compute the limit in the ratio test,

$$A = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{2^{-k-1}}{2^{-k}} \right| = \lim_{k \rightarrow \infty} 2^{-1} = \frac{1}{2}.$$

Therefore, the radius of convergence is 2, and the series converges absolutely on the interval $(-1, 3)$. We could just as well have used the root test:

$$A = \lim_{k \rightarrow \infty} \sqrt[k]{|a_k|} = \lim_{k \rightarrow \infty} \sqrt[k]{2^{-k}} = \lim_{k \rightarrow \infty} 2^{-1} = \frac{1}{2}.$$

Example 7.1.4: Consider

$$\sum_{k=1}^{\infty} \frac{1}{k^k} x^k.$$

Compute the limit for the root test (ratio test would also work),

$$A = \lim_{k \rightarrow \infty} \sqrt[k]{|a_k|} = \lim_{k \rightarrow \infty} \sqrt[k]{\left| \frac{1}{k^k} \right|} = \lim_{k \rightarrow \infty} \sqrt[k]{\frac{1}{k}} = \lim_{k \rightarrow \infty} \frac{1}{k} = 0.$$

So the radius of convergence is ∞ : The series converges everywhere.

The root or the ratio test as given does not always apply. That is, the limit of $\left| \frac{a_{k+1}}{a_k} \right|$ or $\sqrt[k]{|a_k|}$ might not exist. Sometimes the test must be applied to the series itself—to find the L rather than the A . For example, for the series $\sum_{k=1}^{\infty} 2^k x^{2k}$, we cannot apply the root with the A , we must work with the limit of $\sqrt[k]{2^k x^{2k}}$, which equals $L = 2|x|^2$. This L is less than 1 when $|x| < \frac{1}{\sqrt{2}}$ and so $\frac{1}{\sqrt{2}}$ is the radius of convergence. Other ways exist to find the radius, but ratio and root tests cover many series arising in practice.

At the endpoints, $x = x_0 - \rho$ and $x = x_0 + \rho$, the series may or may not converge, and the tests above say nothing about convergence there. Sometimes convergence at the endpoints is important, but for our purposes, we will not worry about it much.

7.1.3 Analytic functions

Functions represented by power series are called *analytic functions*. Not all functions are analytic, but the functions you have seen in calculus likely all are. An analytic function $f(x)$ equals its *Taylor series** (a power series computed from f) for x near a given point x_0 :

$$f(x) = \sum_{k=0}^{\infty} \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k = f(x_0) + f'(x_0)(x - x_0) + \frac{f''(x_0)}{2}(x - x_0)^2 + \cdots, \quad (7.2)$$

where $f^{(k)}(x_0)$ denotes the k^{th} derivative of $f(x)$ at the point x_0 .

*Named after the English mathematician [Sir Brook Taylor](#) (1685–1731).

For example, sine is an analytic function and its Taylor series around $x_0 = 0$ is given by

$$\sin(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}.$$

In [Figure 7.2](#), we plot $\sin(x)$ and the truncations of the series up to degree 5 and 9. You can see that the approximation is very good for x near 0, but gets worse for x further away from 0. This is what happens in general. To get a good approximation far away from x_0 you need to take more and more terms of the Taylor series.

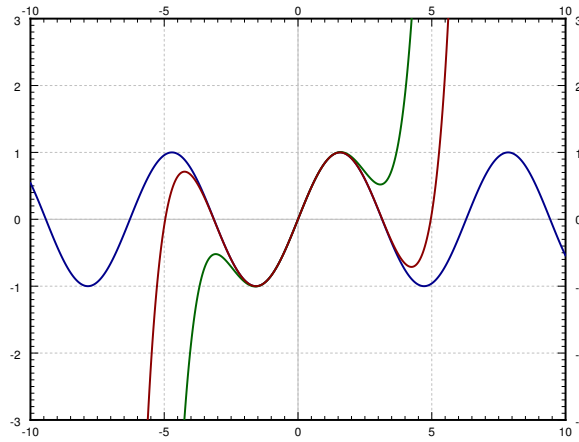


Figure 7.2: The sine function and its Taylor approximations around $x_0 = 0$ of 5th and 9th degree.

7.1.4 Manipulating power series

One of the main properties of power series that we will use is that we can differentiate them term by term. That is, suppose that $\sum a_k(x - x_0)^k$ is a convergent power series. Then for x in the radius of convergence, we have

$$\frac{d}{dx} \left[\sum_{k=0}^{\infty} a_k(x - x_0)^k \right] = \sum_{k=1}^{\infty} k a_k(x - x_0)^{k-1} = a_1 + 2a_2(x - x_0) + 3a_3(x - x_0)^2 + \cdots.$$

Notice that the term corresponding to $k = 0$ disappeared as it was constant. The radius of convergence of the differentiated series is the same as that of the original.

Example 7.1.5: Let us show that the exponential $y = e^x$ solves $y' = y$. Suppose we did not know that. Write

$$y = e^x = \sum_{k=0}^{\infty} \frac{1}{k!} x^k.$$

Differentiate y ,

$$y' = \sum_{k=1}^{\infty} k \frac{1}{k!} x^{k-1} = \sum_{k=1}^{\infty} \frac{1}{(k-1)!} x^{k-1}.$$

We *reindex* the series by simply replacing k with $k + 1$. The series does not change; what changes is simply how we write it. After reindexing, the series starts at $k = 0$ again.

$$\sum_{k=1}^{\infty} \frac{1}{(k-1)!} x^{k-1} = \sum_{k+1=1}^{\infty} \frac{1}{((k+1)-1)!} x^{(k+1)-1} = \sum_{k=0}^{\infty} \frac{1}{k!} x^k.$$

That was precisely the power series for e^x we started with, so we showed that $\frac{d}{dx}[e^x] = e^x$.

Convergent power series can be added and multiplied together, and multiplied by constants using the following rules. First, we can add series by adding term by term,

$$\left(\sum_{k=0}^{\infty} a_k (x - x_0)^k \right) + \left(\sum_{k=0}^{\infty} b_k (x - x_0)^k \right) = \sum_{k=0}^{\infty} (a_k + b_k) (x - x_0)^k.$$

We can multiply by constants,

$$\alpha \left(\sum_{k=0}^{\infty} a_k (x - x_0)^k \right) = \sum_{k=0}^{\infty} \alpha a_k (x - x_0)^k.$$

We can also multiply series together,

$$\left(\sum_{k=0}^{\infty} a_k (x - x_0)^k \right) \left(\sum_{k=0}^{\infty} b_k (x - x_0)^k \right) = \sum_{k=0}^{\infty} c_k (x - x_0)^k,$$

where $c_k = a_0 b_k + a_1 b_{k-1} + \cdots + a_k b_0$. The radius of convergence of the sum or the product is at least the minimum of the radii of convergence of the two series involved.

7.1.5 Power series for rational functions

Polynomials are simply finite power series. That is, a polynomial is a power series where the a_k are zero for all k large enough. We can always expand a polynomial as a power series about any point x_0 by writing the polynomial as a polynomial in $(x - x_0)$. For example, let us write $2x^2 - 3x + 4$ as a power series around $x_0 = 1$:

$$2x^2 - 3x + 4 = 3 + (x - 1) + 2(x - 1)^2.$$

In other words, $a_0 = 3$, $a_1 = 1$, $a_2 = 2$, and all other $a_k = 0$. To do this, we know that $a_k = 0$ for all $k \geq 3$ as the polynomial is of degree 2. We write $a_0 + a_1(x - 1) + a_2(x - 1)^2$, we expand, and we solve for a_0 , a_1 , and a_2 . We could have also differentiated at $x = 1$ and used the Taylor series formula (7.2).

Let us look at rational functions, that is, ratios of polynomials. An important fact is that a series for a function only defines the function on an interval even if the function is defined elsewhere. For example, for $-1 < x < 1$,

$$\frac{1}{1-x} = \sum_{k=0}^{\infty} x^k = 1 + x + x^2 + \cdots$$

This series is called the *geometric series*. The ratio test tells us that the radius of convergence is 1. The series diverges for $x \leq -1$ and $x \geq 1$, even though $\frac{1}{1-x}$ is defined for all $x \neq 1$.

Instead of applying the Taylor series expansion (7.2), the geometric series together with rules for addition and multiplication of power series can be used to expand any rational function around a point x_0 , as long as the denominator is not zero at x_0 .

Example 7.1.6: Expand $\frac{x}{1+2x+x^2}$ as a power series around the origin ($x_0 = 0$) and find the radius of convergence.

First, write $1 + 2x + x^2 = (1 + x)^2 = (1 - (-x))^2$. Compute

$$\begin{aligned} \frac{x}{1+2x+x^2} &= x \left(\frac{1}{1-(-x)} \right)^2 \\ &= x \left(\sum_{k=0}^{\infty} (-1)^k x^k \right)^2 \\ &= x \left(\sum_{k=0}^{\infty} c_k x^k \right) \\ &= \sum_{k=0}^{\infty} c_k x^{k+1}, \end{aligned}$$

where to get c_k , we use the formula for the product of series: $c_0 = 1$, $c_1 = -1 - 1 = -2$, $c_2 = 1 + 1 + 1 = 3$, etc. Therefore

$$\frac{x}{1+2x+x^2} = \sum_{k=1}^{\infty} (-1)^{k+1} k x^k = x - 2x^2 + 3x^3 - 4x^4 + \cdots$$

The radius of convergence is at least 1. We use the ratio test

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{(-1)^{k+2}(k+1)}{(-1)^{k+1}k} \right| = \lim_{k \rightarrow \infty} \frac{k+1}{k} = 1.$$

So the radius of convergence is actually equal to 1.

When the rational function is more complicated, it is also possible to use the method of partial fractions. For example, to find the Taylor series for $\frac{x^3+x}{x^2-1}$, we write

$$\frac{x^3+x}{x^2-1} = x + \frac{1}{1+x} - \frac{1}{1-x} = x + \sum_{k=0}^{\infty} (-1)^k x^k - \sum_{k=0}^{\infty} x^k = -x + \sum_{\substack{k=3 \\ k \text{ odd}}}^{\infty} (-2)x^k.$$

7.1.6 Exercises

Exercise 7.1.1: Is $\sum_{k=0}^{\infty} e^k x^k$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.2: Is $\sum_{k=0}^{\infty} kx^k$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.3: Is $\sum_{k=0}^{\infty} k!x^k$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.4: Is $\sum_{k=0}^{\infty} \frac{1}{(2k)!} (x-10)^k$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.5: Determine the Taylor series for $\sin x$ around the point $x_0 = \pi$.

Exercise 7.1.6: Find the Taylor series for $\ln x$ around the point $x_0 = 1$, and the radius of convergence.

Exercise 7.1.7: Determine the Taylor series and its radius of convergence of $\frac{1}{1+x}$ around $x_0 = 0$.

Exercise 7.1.8: Determine the Taylor series and its radius of convergence of $\frac{x}{4-x^2}$ around $x_0 = 0$.

Hint: You will not be able to use the ratio test.

Exercise 7.1.9: Expand $x^5 + 5x + 1$ as a power series around $x_0 = 5$.

Exercise 7.1.10: Suppose that the ratio test applies to a series $\sum_{k=0}^{\infty} a_k x^k$. Show, using the ratio test, that the radius of convergence of the differentiated series is the same as that of the original series.

Exercise 7.1.11: Suppose that f is an analytic function such that $f^{(n)}(0) = n$. Find $f(1)$.

Exercise 7.1.101: Is $\sum_{n=1}^{\infty} (0.1)^n x^n$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.102 (challenging): Is $\sum_{n=1}^{\infty} \frac{n!}{n^n} x^n$ a convergent power series? If so, find the radius of convergence.

Exercise 7.1.103: Using the geometric series, expand $\frac{1}{1-x}$ around $x_0 = 2$. For what x does the series converge?

Exercise 7.1.104 (challenging): Find the Taylor series for $x^7 e^x$ around $x_0 = 0$.

Exercise 7.1.105 (challenging): Imagine f and g are analytic functions such that $f^{(k)}(0) = g^{(k)}(0)$ for all large enough k . What can you say about $f(x) - g(x)$?

Exercise 7.1.106: Reindex the following series to have the powers be x^k .

a) $\sum_{k=3}^{\infty} k(k-1)x^{k+3}$

b) $\sum_{k=1}^{\infty} (k+1)x^{k-1}$

c) $\sum_{k=5}^{\infty} 2kx^{k-2}$

7.2 Series solutions of linear second-order ODEs

Note: 1.5 or 2 lectures, §8.2 in [EP], §5.2 and §5.3 in [BD]

Consider a linear second-order homogeneous ODE of the form

$$P(x)y'' + Q(x)y' + R(x)y = 0.$$

Suppose that $P(x)$, $Q(x)$, and $R(x)$ are polynomials. We will try a solution of the form

$$y = \sum_{k=0}^{\infty} a_k (x - x_0)^k,$$

and solve for the a_k to obtain a solution defined in some interval around x_0 .

The point x_0 is called an *ordinary point* if $P(x_0) \neq 0$. That is, the functions

$$\frac{Q(x)}{P(x)} \quad \text{and} \quad \frac{R(x)}{P(x)}$$

are defined for x near x_0 . If $P(x_0) = 0$, then x_0 is a *singular point*. Handling singular points is harder than handling ordinary points and so we now focus only on ordinary points.

Example 7.2.1: We start with a very simple example

$$y'' - y = 0.$$

Let us try a power series solution near $x_0 = 0$, which is an ordinary point. Every point is an ordinary point in fact, as the equation is a constant-coefficient one. We already know we should obtain exponentials or the hyperbolic sine and cosine, but let us pretend we do not know this fact.

We try

$$y = \sum_{k=0}^{\infty} a_k x^k.$$

If we differentiate, the $k = 0$ term is a constant and hence disappears. We get

$$y' = \sum_{k=1}^{\infty} k a_k x^{k-1}.$$

We differentiate yet again to obtain (now the $k = 1$ term disappears)

$$y'' = \sum_{k=2}^{\infty} k(k-1) a_k x^{k-2}.$$

We need to subtract y from y'' , so we reindex the series for y'' (replace k with $k+2$):

$$y'' = \sum_{k=0}^{\infty} (k+2)(k+1) a_{k+2} x^k.$$

Now we plug y and y'' into the differential equation

$$\begin{aligned} 0 = y'' - y &= \left(\sum_{k=0}^{\infty} (k+2)(k+1) a_{k+2} x^k \right) - \left(\sum_{k=0}^{\infty} a_k x^k \right) \\ &= \sum_{k=0}^{\infty} \left((k+2)(k+1) a_{k+2} - a_k \right) x^k \\ &= \sum_{k=0}^{\infty} \left((k+2)(k+1) a_{k+2} - a_k \right) x^k. \end{aligned}$$

As $y'' - y$ is supposed to be equal to 0, we know that the coefficients of the resulting series must be equal to 0. Therefore, for every $k = 0, 1, 2, 3, \dots$, we have

$$(k+2)(k+1) a_{k+2} - a_k = 0, \quad \text{or} \quad a_{k+2} = \frac{a_k}{(k+2)(k+1)}.$$

The equation above is called a *recurrence relation* for the coefficients of the power series. It does not matter what a_0 or a_1 are. They can be arbitrary. But once we pick a_0 and a_1 , all other coefficients are determined by the recurrence relation.

Let us see what the coefficients must be. First, a_0 and a_1 are arbitrary. Then,

$$a_2 = \frac{a_0}{2}, \quad a_3 = \frac{a_1}{3 \cdot 2}, \quad a_4 = \frac{a_2}{4 \cdot 3} = \frac{a_0}{4 \cdot 3 \cdot 2}, \quad a_5 = \frac{a_3}{5 \cdot 4} = \frac{a_1}{5 \cdot 4 \cdot 3 \cdot 2}, \quad \dots$$

For even k , that is, $k = 2n$, we have

$$a_k = a_{2n} = \frac{a_0}{(2n)!}.$$

For odd k , that is, $k = 2n + 1$, we have

$$a_k = a_{2n+1} = \frac{a_1}{(2n+1)!}.$$

We write down the series for the solution

$$y = \sum_{k=0}^{\infty} a_k x^k = \sum_{n=0}^{\infty} \left(\frac{a_0}{(2n)!} x^{2n} + \frac{a_1}{(2n+1)!} x^{2n+1} \right) = a_0 \sum_{n=0}^{\infty} \frac{1}{(2n)!} x^{2n} + a_1 \sum_{n=0}^{\infty} \frac{1}{(2n+1)!} x^{2n+1}.$$

We recognize the two series as the hyperbolic sine and cosine. Therefore,

$$y = a_0 \cosh x + a_1 \sinh x.$$

Of course, in general we will not be able to recognize the series that appears, since usually there will not be any elementary function that matches it. In that case, we will be content with the series.

Example 7.2.2: Let us do a more complex example. Consider *Airy's equation**:

$$y'' - xy = 0,$$

near the point $x_0 = 0$. Note that $x_0 = 0$ is an ordinary point.

*Named after the English mathematician [Sir George Biddell Airy](#) (1801–1892).

We try

$$y = \sum_{k=0}^{\infty} a_k x^k.$$

We differentiate twice (as above) to obtain

$$y'' = \sum_{k=2}^{\infty} k(k-1) a_k x^{k-2}.$$

We plug y into the equation:

$$\begin{aligned} 0 = y'' - xy &= \left(\sum_{k=2}^{\infty} k(k-1) a_k x^{k-2} \right) - x \left(\sum_{k=0}^{\infty} a_k x^k \right) \\ &= \left(\sum_{k=2}^{\infty} k(k-1) a_k x^{k-2} \right) - \left(\sum_{k=0}^{\infty} a_k x^{k+1} \right). \end{aligned}$$

We reindex and synchronize the start indexes to make the series possible to sum:

$$\begin{aligned} 0 = y'' - xy &= \left(2a_2 + \sum_{k=1}^{\infty} (k+2)(k+1) a_{k+2} x^k \right) - \left(\sum_{k=1}^{\infty} a_{k-1} x^k \right) \\ &= 2a_2 + \sum_{k=1}^{\infty} \left((k+2)(k+1) a_{k+2} - a_{k-1} \right) x^k. \end{aligned}$$

As $y'' - xy$ is supposed to be 0, we have $a_2 = 0$, and

$$(k+2)(k+1) a_{k+2} - a_{k-1} = 0, \quad \text{or} \quad a_{k+2} = \frac{a_{k-1}}{(k+2)(k+1)}.$$

We have arbitrary constants a_0 and a_1 , and $a_2 = 0$. We compute a few higher coefficients:

$$\begin{aligned} a_3 &= \frac{a_0}{3 \cdot 2}, \quad a_4 = \frac{a_1}{4 \cdot 3}, \quad a_5 = \frac{a_2}{5 \cdot 4} = 0, \quad a_6 = \frac{a_3}{6 \cdot 5} = \frac{a_0}{6 \cdot 5 \cdot 3 \cdot 2}, \\ a_7 &= \frac{a_4}{7 \cdot 6} = \frac{a_1}{7 \cdot 6 \cdot 4 \cdot 3}, \quad a_8 = \frac{a_5}{8 \cdot 7} = 0, \quad a_9 = \frac{a_6}{9 \cdot 8} = \frac{a_0}{9 \cdot 8 \cdot 6 \cdot 5 \cdot 3 \cdot 2}, \quad \dots \end{aligned}$$

We jump in steps of three. For a_k where k is a multiple of 3, that is, $k = 3n$, we notice that

$$a_{3n} = \frac{a_0}{2 \cdot 3 \cdot 5 \cdot 6 \cdots (3n-1) \cdot (3n)}.$$

For a_k where $k = 3n + 1$, we notice

$$a_{3n+1} = \frac{a_1}{3 \cdot 4 \cdot 6 \cdot 7 \cdots (3n) \cdot (3n+1)}.$$

Finally, since $a_2 = 0$, we have $a_5 = 0$, $a_8 = 0$, $a_{11} = 0$, etc. In general, $a_{3n+2} = 0$.

In other words, if we write down the series for y , it has two parts

$$\begin{aligned}
 y &= \left(a_0 + \frac{a_0}{6}x^3 + \frac{a_0}{180}x^6 + \cdots + \frac{a_0}{2 \cdot 3 \cdot 5 \cdot 6 \cdots (3n-1) \cdot (3n)}x^{3n} + \cdots \right) \\
 &\quad + \left(a_1x + \frac{a_1}{12}x^4 + \frac{a_1}{504}x^7 + \cdots + \frac{a_1}{3 \cdot 4 \cdot 6 \cdot 7 \cdots (3n) \cdot (3n+1)}x^{3n+1} + \cdots \right) \\
 &= a_0 \left(1 + \frac{1}{6}x^3 + \frac{1}{180}x^6 + \cdots + \frac{1}{2 \cdot 3 \cdot 5 \cdot 6 \cdots (3n-1) \cdot (3n)}x^{3n} + \cdots \right) \\
 &\quad + a_1 \left(x + \frac{1}{12}x^4 + \frac{1}{504}x^7 + \cdots + \frac{1}{3 \cdot 4 \cdot 6 \cdot 7 \cdots (3n) \cdot (3n+1)}x^{3n+1} + \cdots \right).
 \end{aligned}$$

We define

$$\begin{aligned}
 y_1(x) &= 1 + \frac{1}{6}x^3 + \frac{1}{180}x^6 + \cdots + \frac{1}{2 \cdot 3 \cdot 5 \cdot 6 \cdots (3n-1) \cdot (3n)}x^{3n} + \cdots, \\
 y_2(x) &= x + \frac{1}{12}x^4 + \frac{1}{504}x^7 + \cdots + \frac{1}{3 \cdot 4 \cdot 6 \cdot 7 \cdots (3n) \cdot (3n+1)}x^{3n+1} + \cdots,
 \end{aligned}$$

and write the general solution to the equation as $y(x) = a_0y_1(x) + a_1y_2(x)$. If we plug $x = 0$ into the power series for y_1 and y_2 , we find $y_1(0) = 1$ and $y_2(0) = 0$. Similarly, $y_1'(0) = 0$ and $y_2'(0) = 1$. Therefore $y = a_0y_1 + a_1y_2$ is a solution that satisfies the initial conditions $y(0) = a_0$ and $y'(0) = a_1$.

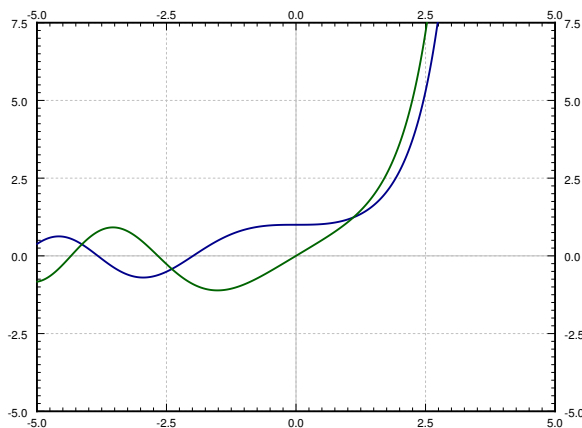


Figure 7.3: The two solutions y_1 and y_2 to Airy's equation.

The functions y_1 and y_2 cannot be written in terms of the elementary functions that you know. See Figure 7.3 for the plot of the solutions y_1 and y_2 . These functions have many interesting properties. For example, they are oscillatory for negative x (like solutions to $y'' + y = 0$) and for positive x they grow without bound (like solutions to $y'' - y = 0$).

Sometimes a series solution may turn out to be a polynomial. Let us see an example.

Example 7.2.3: Let us find a solution to the so-called *Hermite's equation of order n^** :

$$y'' - 2xy' + 2ny = 0.$$

Let us find a solution around the point $x_0 = 0$. We try

$$y = \sum_{k=0}^{\infty} a_k x^k.$$

We differentiate (as above) to obtain

$$y' = \sum_{k=1}^{\infty} k a_k x^{k-1} \quad \text{and} \quad y'' = \sum_{k=2}^{\infty} k(k-1) a_k x^{k-2}.$$

Now we plug into the equation

$$\begin{aligned} 0 &= y'' - 2xy' + 2ny \\ &= \left(\sum_{k=2}^{\infty} k(k-1) a_k x^{k-2} \right) - 2x \left(\sum_{k=1}^{\infty} k a_k x^{k-1} \right) + 2n \left(\sum_{k=0}^{\infty} a_k x^k \right) \\ &= \left(\sum_{k=2}^{\infty} k(k-1) a_k x^{k-2} \right) - \left(\sum_{k=1}^{\infty} 2k a_k x^k \right) + \left(\sum_{k=0}^{\infty} 2n a_k x^k \right) \\ &= \left(2a_2 + \sum_{k=1}^{\infty} (k+2)(k+1) a_{k+2} x^k \right) - \left(\sum_{k=1}^{\infty} 2k a_k x^k \right) + \left(2n a_0 + \sum_{k=1}^{\infty} 2n a_k x^k \right) \\ &= 2a_2 + 2n a_0 + \sum_{k=1}^{\infty} ((k+2)(k+1) a_{k+2} - 2k a_k + 2n a_k) x^k. \end{aligned}$$

As $y'' - 2xy' + 2ny = 0$, we have

$$(k+2)(k+1) a_{k+2} + (-2k + 2n) a_k = 0, \quad \text{or} \quad a_{k+2} = \frac{(2k - 2n)}{(k+2)(k+1)} a_k.$$

This recurrence relation actually includes $a_2 = -n a_0$ (which comes about from the constant term $2a_2 + 2n a_0 = 0$). Again, a_0 and a_1 are arbitrary.

$$\begin{aligned} a_2 &= \frac{-2n}{2 \cdot 1} a_0, \\ a_3 &= \frac{2(1-n)}{3 \cdot 2} a_1, \\ a_4 &= \frac{2(2-n)}{4 \cdot 3} a_2 = \frac{2^2(2-n)(-n)}{4 \cdot 3 \cdot 2 \cdot 1} a_0, \end{aligned}$$

*Named after the French mathematician [Charles Hermite](#) (1822–1901).

$$a_5 = \frac{2(3-n)}{5 \cdot 4} a_3 = \frac{2^2(3-n)(1-n)}{5 \cdot 4 \cdot 3 \cdot 2} a_1, \quad \dots$$

We separate the even and odd coefficients to find that

$$a_{2m} = \frac{2^m(-n)(2-n) \cdots (2m-2-n)}{(2m)!} a_0,$$

$$a_{2m+1} = \frac{2^m(1-n)(3-n) \cdots (2m-1-n)}{(2m+1)!} a_1.$$

We write down the two series, one with the even powers and one with the odd:

$$y_1(x) = 1 + \frac{2(-n)}{2!}x^2 + \frac{2^2(-n)(2-n)}{4!}x^4 + \frac{2^3(-n)(2-n)(4-n)}{6!}x^6 + \dots,$$

$$y_2(x) = x + \frac{2(1-n)}{3!}x^3 + \frac{2^2(1-n)(3-n)}{5!}x^5 + \frac{2^3(1-n)(3-n)(5-n)}{7!}x^7 + \dots.$$

Then

$$y(x) = a_0 y_1(x) + a_1 y_2(x).$$

We remark that if n is a positive even integer, then $y_1(x)$ is a polynomial as all the coefficients in the series beyond degree n are zero. If n is a positive odd integer, then $y_2(x)$ is a polynomial. For example, if $n = 4$, then

$$y_1(x) = 1 + \frac{2(-4)}{2!}x^2 + \frac{2^2(-4)(2-4)}{4!}x^4 = 1 - 4x^2 + \frac{4}{3}x^4.$$

7.2.1 Exercises

In the following exercises, when asked to solve an equation using power series methods, you should find the first few terms of the series, and if possible find a general formula for the k^{th} coefficient.

Exercise 7.2.1: Use power series methods to solve $y'' + y = 0$ at the point $x_0 = 1$.

Exercise 7.2.2: Use power series methods to solve $y'' + 4xy = 0$ at the point $x_0 = 0$.

Exercise 7.2.3: Use power series methods to solve $y'' - xy = 0$ at the point $x_0 = 1$.

Exercise 7.2.4: Use power series methods to solve $y'' + x^2y = 0$ at the point $x_0 = 0$.

Exercise 7.2.5: The methods work for other orders than second order. Try the methods of this section to solve the first-order system $y' - xy = 0$ at the point $x_0 = 0$.

Exercise 7.2.6 (Chebyshev's equation of order p):

- Solve $(1 - x^2)y'' - xy' + p^2y = 0$ using power series methods at $x_0 = 0$.
- For what p is there a polynomial solution?

Exercise 7.2.7: Find a polynomial solution to $(x^2 + 1)y'' - 2xy' + 2y = 0$ using power series methods.

Exercise 7.2.8:

- a) Use power series methods to solve $(1 - x)y'' + y = 0$ at the point $x_0 = 0$.
- b) Use the solution to part a) to find a solution for $xy'' + y = 0$ around the point $x_0 = 1$.

Exercise 7.2.101: Use power series methods to solve $y'' + 2x^3y = 0$ at the point $x_0 = 0$.

Exercise 7.2.102 (challenging): Power-series methods also work for nonhomogeneous equations.

- a) Use power series methods to solve $y'' - xy = \frac{1}{1-x}$ at the point $x_0 = 0$. Hint: Recall the geometric series.
- b) Now solve for the initial condition $y(0) = 0, y'(0) = 0$.

Exercise 7.2.103: Attempt to solve $x^2y'' - y = 0$ at $x_0 = 0$ using the power series method of this section (x_0 is a singular point). Can you find at least one solution? Can you find more than one solution?

7.3 Singular points and the method of Frobenius

Note: 1.5 or 2 lectures, §8.3 and §8.4 in [EP], §5.4–§5.7 in [BD]

The behavior of ODEs at singular points can be complicated. For certain singular points, we can find a solution on at least one side of the singular point using a modification of the power series. Let us look at some examples before giving a general method.

7.3.1 Examples

Example 7.3.1: Consider the simple first-order equation

$$2xy' - y = 0.$$

Note that $x = 0$ is a singular point. Setting $x = 0$ in the equation, we find that any solution defined near zero satisfies $y(0) = 0$, but it is even worse. If we try to plug in

$$y = \sum_{k=0}^{\infty} a_k x^k,$$

we obtain

$$\begin{aligned} 0 = 2xy' - y &= 2x \left(\sum_{k=1}^{\infty} k a_k x^{k-1} \right) - \left(\sum_{k=0}^{\infty} a_k x^k \right) \\ &= -a_0 + \sum_{k=1}^{\infty} (2k a_k - a_k) x^k. \end{aligned}$$

First, $a_0 = 0$. Next, the only way to solve $0 = 2k a_k - a_k = (2k - 1) a_k$ for $k = 1, 2, 3, \dots$ is for $a_k = 0$ for all k . Therefore, in this manner we only get the trivial solution $y = 0$. We need a nonzero solution to get the general solution to the equation.

Let us try $y = x^r$ for some real number r . Consequently our solution—if we can find one—may only make sense for positive x . Then $y' = r x^{r-1}$. So

$$0 = 2xy' - y = 2xr x^{r-1} - x^r = (2r - 1)x^r.$$

Thus $r = 1/2$ and so $y = x^{1/2}$. As the equation is linear, the general solution for positive x is

$$y = Cx^{1/2}.$$

If $C \neq 0$, then the derivative of the solution “blows up” at $x = 0$ (the singular point). There is only one solution that is differentiable at $x = 0$ and that’s the trivial solution $y = 0$.

Not every problem with a singular point has a solution of the form $y = x^r$, of course. But perhaps we can combine the methods. What we will do is to try a solution of the form

$$y = x^r f(x)$$

for positive x , where $f(x)$ is an analytic function (a power series).

Example 7.3.2: Consider the equation

$$4x^2y'' - 4x^2y' + (1 - 2x)y = 0,$$

and again note that $x = 0$ is a singular point.

Let us try

$$y = x^r \sum_{k=0}^{\infty} a_k x^k = \sum_{k=0}^{\infty} a_k x^{k+r},$$

where r is a real number, not necessarily an integer. Again if such a solution exists, it may only exist for positive x . First we find the derivatives

$$y' = \sum_{k=0}^{\infty} (k+r) a_k x^{k+r-1},$$

$$y'' = \sum_{k=0}^{\infty} (k+r)(k+r-1) a_k x^{k+r-2}.$$

We plug those into our equation:

$$\begin{aligned} 0 &= 4x^2y'' - 4x^2y' + (1 - 2x)y \\ &= 4x^2 \left(\sum_{k=0}^{\infty} (k+r)(k+r-1) a_k x^{k+r-2} \right) - 4x^2 \left(\sum_{k=0}^{\infty} (k+r) a_k x^{k+r-1} \right) + (1 - 2x) \left(\sum_{k=0}^{\infty} a_k x^{k+r} \right) \\ &= \left(\sum_{k=0}^{\infty} 4(k+r)(k+r-1) a_k x^{k+r} \right) \\ &\quad - \left(\sum_{k=0}^{\infty} 4(k+r) a_k x^{k+r+1} \right) + \left(\sum_{k=0}^{\infty} a_k x^{k+r} \right) - \left(\sum_{k=0}^{\infty} 2a_k x^{k+r+1} \right) \\ &= \left(\sum_{k=0}^{\infty} 4(k+r)(k+r-1) a_k x^{k+r} \right) \\ &\quad - \left(\sum_{k=1}^{\infty} 4(k+r-1) a_{k-1} x^{k+r} \right) + \left(\sum_{k=0}^{\infty} a_k x^{k+r} \right) - \left(\sum_{k=1}^{\infty} 2a_{k-1} x^{k+r} \right) \\ &= 4r(r-1) a_0 x^r + a_0 x^r + \sum_{k=1}^{\infty} \left(4(k+r)(k+r-1) a_k - 4(k+r-1) a_{k-1} + a_k - 2a_{k-1} \right) x^{k+r} \\ &= (4r(r-1) + 1) a_0 x^r + \sum_{k=1}^{\infty} \left((4(k+r)(k+r-1) + 1) a_k - (4(k+r-1) + 2) a_{k-1} \right) x^{k+r}. \end{aligned}$$

First, to have a solution we must have $(4r(r-1) + 1) a_0 = 0$. Supposing $a_0 \neq 0$,

$$4r(r-1) + 1 = 0.$$

This equation is called the *indicial equation*. This particular indicial equation has a double root at $r = 1/2$.

OK, so we know what r has to be. That knowledge we obtained simply by looking at the coefficient of x^r . All other coefficients of x^{k+r} also have to be zero so

$$(4(k+r)(k+r-1)+1)a_k - (4(k+r-1)+2)a_{k-1} = 0.$$

If we plug in $r = 1/2$ and solve for a_k , we get

$$a_k = \frac{4(k+1/2-1)+2}{4(k+1/2)(k+1/2-1)+1} a_{k-1} = \frac{1}{k} a_{k-1}.$$

Let us set $a_0 = 1$. Then

$$a_1 = \frac{1}{1}a_0 = 1, \quad a_2 = \frac{1}{2}a_1 = \frac{1}{2}, \quad a_3 = \frac{1}{3}a_2 = \frac{1}{3 \cdot 2}, \quad a_4 = \frac{1}{4}a_3 = \frac{1}{4 \cdot 3 \cdot 2}, \quad \dots$$

Extrapolating, we notice that

$$a_k = \frac{1}{k(k-1)(k-2)\cdots 3 \cdot 2} = \frac{1}{k!}.$$

In other words,

$$y = \sum_{k=0}^{\infty} a_k x^{k+r} = \sum_{k=0}^{\infty} \frac{1}{k!} x^{k+1/2} = x^{1/2} \sum_{k=0}^{\infty} \frac{1}{k!} x^k = x^{1/2} e^x.$$

That was lucky! In general, we will not be able to write the series in terms of elementary functions.

We have one solution, let us call it $y_1 = x^{1/2}e^x$. But what about a second solution? If we want a general solution, we need two linearly independent solutions. Picking a_0 to be a different constant only gets us a constant multiple of y_1 , and we do not have any other r to try; we only have one solution to the indicial equation. Well, there are powers of x floating around and we are taking derivatives, perhaps the logarithm (the antiderivative of x^{-1}) is around as well. It turns out we want to try for another solution of the form

$$y_2 = \sum_{k=0}^{\infty} b_k x^{k+r} + (\ln x)y_1,$$

which in our case is

$$y_2 = \sum_{k=0}^{\infty} b_k x^{k+1/2} + (\ln x)x^{1/2}e^x.$$

We now differentiate this equation, substitute into the differential equation and solve for b_k . A long computation ensues and we obtain some recursion relation for b_k . The reader can (and should) try this to obtain for example the first three terms

$$b_1 = b_0 - 1, \quad b_2 = \frac{2b_1 - 1}{4}, \quad b_3 = \frac{6b_2 - 1}{18}, \quad \dots$$

We then fix b_0 and obtain a solution y_2 . Then we write the general solution as $y = Ay_1 + By_2$.

7.3.2 The method of Frobenius

Before giving the general method, let us clarify when the method applies. Let

$$P(x)y'' + Q(x)y' + R(x)y = 0$$

be an ODE for polynomials $P(x)$, $Q(x)$, and $R(x)$. As before, if $P(x_0) = 0$, then x_0 is a singular point. If we divide by $P(x)$ to put the equation in standard form $y'' + \frac{Q(x)}{P(x)}y' + \frac{R(x)}{P(x)}y = 0$, perhaps the singularities introduced are not too bad. More specifically, if the limits

$$\lim_{x \rightarrow x_0} (x - x_0) \frac{Q(x)}{P(x)} \quad \text{and} \quad \lim_{x \rightarrow x_0} (x - x_0)^2 \frac{R(x)}{P(x)}$$

both exist and are finite, then we say that x_0 is a *regular singular point*.

Example 7.3.3: Often, and for the rest of this section, $x_0 = 0$. Consider

$$x^2y'' + x(1+x)y' + (\pi + x^2)y = 0.$$

Write

$$\begin{aligned} \lim_{x \rightarrow 0} x \frac{Q(x)}{P(x)} &= \lim_{x \rightarrow 0} x \frac{x(1+x)}{x^2} = \lim_{x \rightarrow 0} (1+x) = 1, \\ \lim_{x \rightarrow 0} x^2 \frac{R(x)}{P(x)} &= \lim_{x \rightarrow 0} x^2 \frac{(\pi + x^2)}{x^2} = \lim_{x \rightarrow 0} (\pi + x^2) = \pi. \end{aligned}$$

So $x = 0$ is a regular singular point.

On the other hand, if we make the slight change

$$x^2y'' + (1+x)y' + (\pi + x^2)y = 0,$$

then

$$\lim_{x \rightarrow 0} x \frac{Q(x)}{P(x)} = \lim_{x \rightarrow 0} x \frac{(1+x)}{x^2} = \lim_{x \rightarrow 0} \frac{1+x}{x} = \text{DNE}.$$

Here DNE stands for *does not exist*. The point 0 is singular, but not a regular singular point.

We now discuss the general *Method of Frobenius**. We only consider the method at the point $x = 0$ for simplicity. If $x_0 \neq 0$, then in the solution, we would replace every x with $(x - x_0)$. The main idea is the following theorem.

Theorem 7.3.1 (Method of Frobenius). *Suppose that*

$$P(x)y'' + Q(x)y' + R(x)y = 0 \tag{7.3}$$

has a regular singular point at $x = 0$, then there exists at least one solution of the form

$$y = x^r \sum_{k=0}^{\infty} a_k x^k.$$

A solution of this form is called a Frobenius-type solution.

*Named after the German mathematician [Ferdinand Georg Frobenius](#) (1849–1917).

The method usually breaks down like this:

- (i) We seek a Frobenius-type solution of the form

$$y = \sum_{k=0}^{\infty} a_k x^{k+r}.$$

We plug this y into the left-hand side of equation (7.3), collect terms, and write everything as a single series.

- (ii) The obtained series must be zero. Setting the first coefficient (the one with the smallest power) in the series to zero we obtain the *indicial equation*, which is a quadratic polynomial in r .
- (iii) If the indicial equation has two real roots r_1 and r_2 such that $r_1 - r_2$ is not an integer, then we have two linearly independent Frobenius-type solutions. Using the first root, we plug in

$$y_1 = x^{r_1} \sum_{k=0}^{\infty} a_k x^k,$$

and we solve for all a_k to obtain the first solution. Then using the second root, we plug in

$$y_2 = x^{r_2} \sum_{k=0}^{\infty} b_k x^k,$$

and solve for all b_k to obtain the second solution.

- (iv) If the indicial equation has a doubled root r , then there we find one solution

$$y_1 = x^r \sum_{k=0}^{\infty} a_k x^k,$$

and then we obtain a new solution by plugging

$$y_2 = x^r \sum_{k=0}^{\infty} b_k x^k + (\ln x)y_1,$$

into equation (7.3) and solving for the constants b_k .

- (v) If the indicial equation has two real roots such that $r_1 - r_2$ is an integer, then, assuming r_1 is the larger root, one solution is

$$y_1 = x^{r_1} \sum_{k=0}^{\infty} a_k x^k,$$

and the second linearly independent solution is of the form

$$y_2 = x^{r_2} \sum_{k=0}^{\infty} b_k x^k + C(\ln x)y_1,$$

where we plug y_2 into (7.3) and solve for the constants b_k and C .

(vi) Finally, if the indicial equation has complex roots, then solving for a_k in the solution

$$y = x^{r_1} \sum_{k=0}^{\infty} a_k x^k$$

results in a complex-valued function—the a_k are complex numbers. We obtain our two linearly independent solutions* by taking the real and imaginary parts of y .

The main idea is to find at least one Frobenius-type solution. If we are lucky and find two, we are done. If we only get one, we either use the ideas above or even a different method such as reduction of order (see § 2.1) to obtain a second solution.

7.3.3 Bessel functions

An important class of functions that arise commonly in physics are the *Bessel functions*[†]. For example, these functions appear when solving the wave equation in two and three dimensions. First consider *Bessel's equation* of order p :

$$x^2 y'' + x y' + (x^2 - p^2) y = 0.$$

We allow p to be any number, not just an integer, although integers and multiples of $1/2$ are most important in applications.

When we plug

$$y = \sum_{k=0}^{\infty} a_k x^{k+r}$$

into Bessel's equation of order p , we obtain the indicial equation

$$r(r-1) + r - p^2 = (r-p)(r+p) = 0.$$

Unless $p = 0$, we obtain two roots, $r_1 = p$ and $r_2 = -p$. If p is not an integer, then following the method of Frobenius and setting $a_0 = 1$, we find linearly independent solutions of the form

$$y_1 = x^p \sum_{k=0}^{\infty} \frac{(-1)^k x^{2k}}{2^{2k} k! (k+p)(k-1+p) \cdots (2+p)(1+p)},$$

$$y_2 = x^{-p} \sum_{k=0}^{\infty} \frac{(-1)^k x^{2k}}{2^{2k} k! (k-p)(k-1-p) \cdots (2-p)(1-p)}.$$

*See Joseph L. Neuringer, *The Frobenius method for complex roots of the indicial equation*, International Journal of Mathematical Education in Science and Technology, Volume 9, Issue 1, 1978, 71–77.

[†]Named after the German astronomer and mathematician [Friedrich Wilhelm Bessel](#) (1784–1846).

Exercise 7.3.1:

- a) Verify that the indicial equation of Bessel's equation of order p is $(r - p)(r + p) = 0$.
 b) Suppose p is not an integer. Carry out the computation to find the solutions y_1 and y_2 above.

Bessel functions are convenient constant multiples of y_1 and y_2 . First, we define the *gamma function*

$$\Gamma(x) = \int_0^{\infty} t^{x-1} e^{-t} dt.$$

Notice that $\Gamma(1) = 1$. The gamma function also has a wonderful property

$$\Gamma(x + 1) = x\Gamma(x).$$

From this property, it follows that $\Gamma(n) = (n - 1)!$ when n is a positive integer. So the gamma function is a continuous version of the factorial. We compute:

$$\begin{aligned}\Gamma(k + p + 1) &= (k + p)(k - 1 + p) \cdots (2 + p)(1 + p)\Gamma(1 + p), \\ \Gamma(k - p + 1) &= (k - p)(k - 1 - p) \cdots (2 - p)(1 - p)\Gamma(1 - p).\end{aligned}$$

Exercise 7.3.2: Verify the identities above using $\Gamma(x + 1) = x\Gamma(x)$.

Using $\Gamma(x + 1) = x\Gamma(x)$, we define the gamma function for all x except nonpositive integers. The reciprocal $\frac{1}{\Gamma(x)}$ can be defined for all x by letting $\frac{1}{\Gamma(x)}$ be zero for integers $x \leq 0$.

We define the *Bessel functions of the first kind* of order p and $-p$ as

$$\begin{aligned}J_p(x) &= \frac{1}{2^p \Gamma(1 + p)} y_1 = \sum_{k=0}^{\infty} \frac{(-1)^k}{k! \Gamma(k + p + 1)} \left(\frac{x}{2}\right)^{2k+p}, \\ J_{-p}(x) &= \frac{1}{2^{-p} \Gamma(1 - p)} y_2 = \sum_{k=0}^{\infty} \frac{(-1)^k}{k! \Gamma(k - p + 1)} \left(\frac{x}{2}\right)^{2k-p}.\end{aligned}$$

As these are constant multiples of the solutions we found above, these are both solutions to Bessel's equation of order p . The constants are picked for convenience. As the reciprocal of Γ is defined for all x , the formulas on the right hold even for integer p . When n is a nonnegative integer,

$$J_n(x) = \sum_{k=0}^{\infty} \frac{(-1)^k}{k! (k + n)!} \left(\frac{x}{2}\right)^{2k+n}.$$

When p is not an integer, J_p and J_{-p} are linearly independent. For integer orders, one can check that

$$J_n(x) = (-1)^n J_{-n}(x).$$

So J_{-n} is not a second linearly independent solution. The other solution is the so-called *Bessel function of second kind*. These make sense only for integer orders n and are defined as limits of linear combinations of $J_p(x)$ and $J_{-p}(x)$, as p approaches n in the following way:

$$Y_n(x) = \lim_{p \rightarrow n} \frac{\cos(p\pi)J_p(x) - J_{-p}(x)}{\sin(p\pi)}.$$

Each linear combination of $J_p(x)$ and $J_{-p}(x)$ is a solution to Bessel's equation of order p . Then as we take the limit as p goes to n , we see that $Y_n(x)$ is a solution to Bessel's equation of order n . It also turns out that $Y_n(x)$ and $J_n(x)$ are linearly independent. Therefore, when n is an integer, we have the general solution to Bessel's equation of order n :

$$y = AJ_n(x) + BY_n(x),$$

for arbitrary constants A and B . Note that $Y_n(x)$ goes to negative infinity at $x = 0$. Many mathematical software packages have these functions $J_n(x)$ and $Y_n(x)$ defined, so they can be used just like say $\sin(x)$ and $\cos(x)$. In fact, Bessel functions have some similar properties. For example, $-J_1(x)$ is a derivative of $J_0(x)$, and in general the derivative of $J_n(x)$ can be written as a linear combination of $J_{n-1}(x)$ and $J_{n+1}(x)$. Furthermore, these functions oscillate, although they are not periodic. See Figure 7.4 for graphs of Bessel functions.

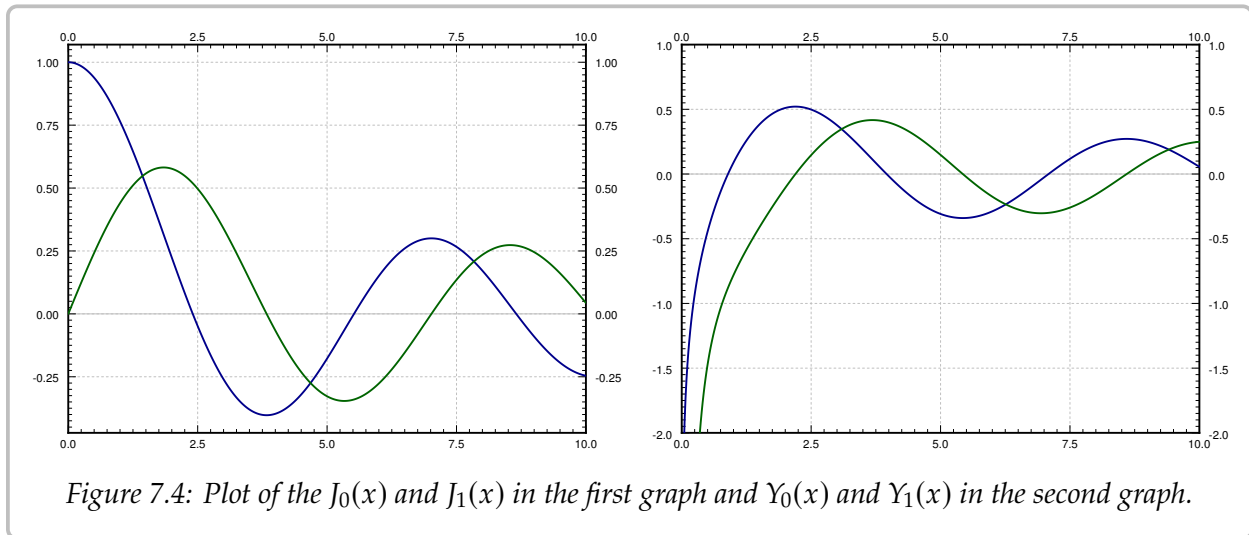


Figure 7.4: Plot of the $J_0(x)$ and $J_1(x)$ in the first graph and $Y_0(x)$ and $Y_1(x)$ in the second graph.

Example 7.3.4: Other equations can sometimes be solved in terms of the Bessel functions. For example, given a positive constant λ ,

$$xy'' + y' + \lambda^2 xy = 0,$$

can be changed to $x^2 y'' + xy' + \lambda^2 x^2 y = 0$. Changing variables $t = \lambda x$, we obtain, via the chain rule, the equation in y and t :

$$t^2 y'' + ty' + t^2 y = 0,$$

which we recognize as Bessel's equation of order 0. Therefore the general solution is $y(t) = AJ_0(t) + BY_0(t)$, or in terms of x :

$$y = AJ_0(\lambda x) + BY_0(\lambda x).$$

This equation comes up, for example, when finding the fundamental modes of vibration of a circular drum, but we digress.

7.3.4 Exercises

Exercise 7.3.3: Find a particular (Frobenius-type) solution of $x^2y'' + xy' + (1+x)y = 0$.

Exercise 7.3.4: Find a particular (Frobenius-type) solution of $xy'' - y = 0$.

Exercise 7.3.5: Find a particular (Frobenius-type) solution of $y'' + \frac{1}{x}y' - xy = 0$.

Exercise 7.3.6: Find the general solution of $2xy'' + y' - x^2y = 0$.

Exercise 7.3.7: Find the general solution of $x^2y'' - xy' - y = 0$.

Exercise 7.3.8: In the following equations classify the point $x = 0$ as ordinary, regular singular, or singular but not regular singular.

a) $x^2(1+x^2)y'' + xy = 0$

b) $x^2y'' + y' + y = 0$

c) $xy'' + x^3y' + y = 0$

d) $xy'' + xy' - e^x y = 0$

e) $x^2y'' + x^2y' + x^2y = 0$

Exercise 7.3.101: In the following equations classify the point $x = 0$ as ordinary, regular singular, or singular but not regular singular.

a) $y'' + y = 0$

b) $x^3y'' + (1+x)y = 0$

c) $xy'' + x^5y' + y = 0$

d) $\sin(x)y'' - y = 0$

e) $\cos(x)y'' - \sin(x)y = 0$

Exercise 7.3.102: Find the general solution of $x^2y'' - y = 0$.

Exercise 7.3.103: Find a particular solution of $x^2y'' + (x - 3/4)y = 0$.

Exercise 7.3.104 (tricky): Find the general solution of $x^2y'' - xy' + y = 0$.

Chapter 8

Nonlinear systems

8.1 Linearization, critical points, and equilibria

Note: 1 lecture, §6.1–§6.2 in [EP], §9.2–§9.3 in [BD]

Except for a few brief detours in [chapter 1](#), we considered mostly linear equations. Linear equations suffice in many applications, but in reality, most phenomena require nonlinear equations. Nonlinear equations, however, are notoriously more difficult to understand than linear ones, and many strange new phenomena appear when we allow our equations to be nonlinear. Not to worry, we did not waste all this time studying linear equations. Nonlinear equations can often be approximated by linear ones if we only need a solution “locally”—only for a short period of time or only for certain parameters. Understanding linear equations can also give us qualitative understanding about a more general nonlinear problem. The idea is similar to what you did in calculus in trying to approximate a function by a line with the right slope.

In § 2.4 we looked at the pendulum of length L . The goal was to solve for the angle $\theta(t)$ as a function of the time t . The equation for the setup is the nonlinear equation

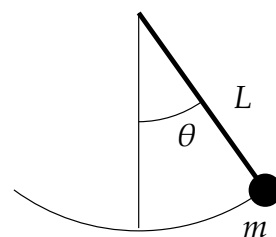
$$\theta'' + \frac{g}{L} \sin \theta = 0.$$

Instead of solving this equation, we solved the rather easier linear equation

$$\theta'' + \frac{g}{L} \theta = 0.$$

While the solution to the linear equation is not exactly what we were looking for, it is rather close to the original, as long as the angle θ is small and the time period involved is short.

You might ask: Why don't we just solve the nonlinear problem? Well, it might be very difficult, impractical, or impossible to solve analytically, depending on the equation in question. We may not even be interested in the actual solution. We might only be interested in some qualitative idea of what the solution is doing. For example, what happens as time goes to infinity?



8.1.1 Autonomous systems and phase plane analysis

We restrict our attention to a two-dimensional autonomous system

$$x' = f(x, y), \quad y' = g(x, y),$$

where $f(x, y)$ and $g(x, y)$ are functions of two variables, and the derivatives are taken with respect to time t . Solutions are functions $x(t)$ and $y(t)$ such that

$$x'(t) = f(x(t), y(t)), \quad y'(t) = g(x(t), y(t)).$$

The way we will analyze the system is very similar to § 1.6, where we studied a single autonomous equation. The ideas in two dimensions are the same, but the behavior can be far more complicated.

It may be best to think of the system of equations as the single vector equation

$$\begin{bmatrix} x \\ y \end{bmatrix}' = \begin{bmatrix} f(x, y) \\ g(x, y) \end{bmatrix}. \quad (8.1)$$

As in § 3.1, we draw the *phase portrait* (or *phase diagram*), where each point (x, y) corresponds to a specific state of the system. We draw the *vector field* given at each point (x, y) by the vector $\begin{bmatrix} f(x, y) \\ g(x, y) \end{bmatrix}$. And as before, if we find solutions, we draw the trajectories by plotting all points $(x(t), y(t))$ for a certain range of t .

Example 8.1.1: Consider the second-order equation $x'' = -x + x^2$. Write this equation as a first-order nonlinear system

$$x' = y, \quad y' = -x + x^2.$$

The phase portrait with some trajectories is drawn in Figure 8.1 on the next page.

From the phase portrait it should be clear that even this simple system has fairly complicated behavior. Some trajectories keep oscillating around the origin, and some go off towards infinity. We will return to this example often, and analyze it completely in this (and the next) section.

If we zoom into the diagram near a point where $\begin{bmatrix} f(x, y) \\ g(x, y) \end{bmatrix}$ is not zero, then nearby the arrows point generally in essentially that same direction and have essentially the same magnitude. In other words, the behavior is not that interesting near such a point. We are of course assuming that $f(x, y)$ and $g(x, y)$ are continuous.

Let us concentrate on those points in the phase diagram above where the trajectories seem to start, end, or go around. We see two such points: $(0, 0)$ and $(1, 0)$. The trajectories seem to go around the point $(0, 0)$, and they seem to either go in or out of the point $(1, 0)$. These points are precisely those points where the derivatives of both x and y are zero. Let us define the *critical points* as the points (x, y) such that

$$\begin{bmatrix} f(x, y) \\ g(x, y) \end{bmatrix} = \vec{0}.$$

In other words, these are the points where both $f(x, y) = 0$ and $g(x, y) = 0$.

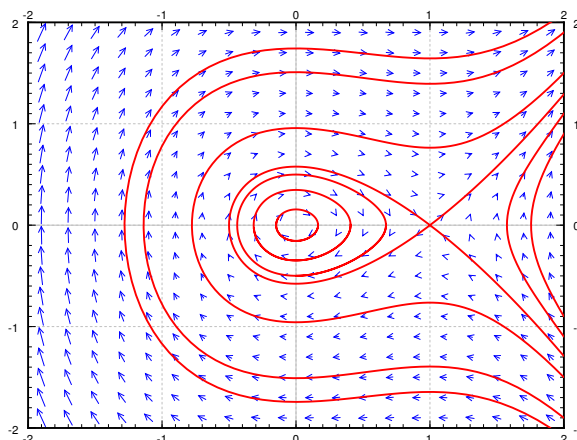


Figure 8.1: Phase portrait with some trajectories of $x' = y$, $y' = -x + x^2$.

The critical points are where the behavior of the system is in some sense the most complicated. If $\begin{bmatrix} f(x,y) \\ g(x,y) \end{bmatrix}$ is zero, then nearby, the vector can point in any direction whatsoever. Also, the trajectories are either going towards, away from, or around these points, so if we are looking for long-term qualitative behavior of the system, we should look at what is happening near the critical points.

Critical points are also sometimes called *equilibria*, since we have so-called *equilibrium solutions* at critical points. If (x_0, y_0) is a critical point, then we have the solutions

$$x(t) = x_0, \quad y(t) = y_0.$$

In [Example 8.1.1](#) on the facing page, there are two equilibrium solutions:

$$x(t) = 0, \quad y(t) = 0, \quad \text{and} \quad x(t) = 1, \quad y(t) = 0.$$

Compare this discussion on equilibria to the discussion in [§ 1.6](#). The underlying concept is exactly the same.

8.1.2 Linearization

In [§ 3.5](#) we studied the behavior of a homogeneous linear system of two equations near a critical point. For a linear system of two variables given by an invertible matrix, the only critical point is the origin $(0, 0)$. Let us put the understanding we gained in that section to good use understanding what happens near critical points of nonlinear systems.

In calculus, we learned to estimate a function by taking its derivative and linearizing. We work similarly with nonlinear systems of ODEs. Suppose (x_0, y_0) is a critical point. First change variables to (u, v) , so that $(u, v) = (0, 0)$ corresponds to (x_0, y_0) . That is,

$$u = x - x_0, \quad v = y - y_0.$$

Next we need to find the derivative. In multivariable calculus you may have seen that the several-variables version of the derivative is the *Jacobian matrix**. The Jacobian matrix of the vector-valued function $\begin{bmatrix} f(x,y) \\ g(x,y) \end{bmatrix}$ at (x_0, y_0) is

$$\begin{bmatrix} \frac{\partial f}{\partial x}(x_0, y_0) & \frac{\partial f}{\partial y}(x_0, y_0) \\ \frac{\partial g}{\partial x}(x_0, y_0) & \frac{\partial g}{\partial y}(x_0, y_0) \end{bmatrix}.$$

This matrix gives the best linear approximation as u and v (and therefore x and y) vary. We define the *linearization* of the equation (8.1) as the linear system

$$\begin{bmatrix} u \\ v \end{bmatrix}' = \begin{bmatrix} \frac{\partial f}{\partial x}(x_0, y_0) & \frac{\partial f}{\partial y}(x_0, y_0) \\ \frac{\partial g}{\partial x}(x_0, y_0) & \frac{\partial g}{\partial y}(x_0, y_0) \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix}.$$

Example 8.1.2: Let us keep with the same equations as [Example 8.1.1](#): $x' = y$, $y' = -x + x^2$. There are two critical points, $(0, 0)$ and $(1, 0)$. The Jacobian matrix at any point is

$$\begin{bmatrix} \frac{\partial f}{\partial x}(x, y) & \frac{\partial f}{\partial y}(x, y) \\ \frac{\partial g}{\partial x}(x, y) & \frac{\partial g}{\partial y}(x, y) \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -1 + 2x & 0 \end{bmatrix}.$$

At $(0, 0)$, that is, when $x = 0$ and $y = 0$, the Jacobian matrix is

$$\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix},$$

and our new variables are $u = x$ and $v = y$. The linearization at $(0, 0)$ is

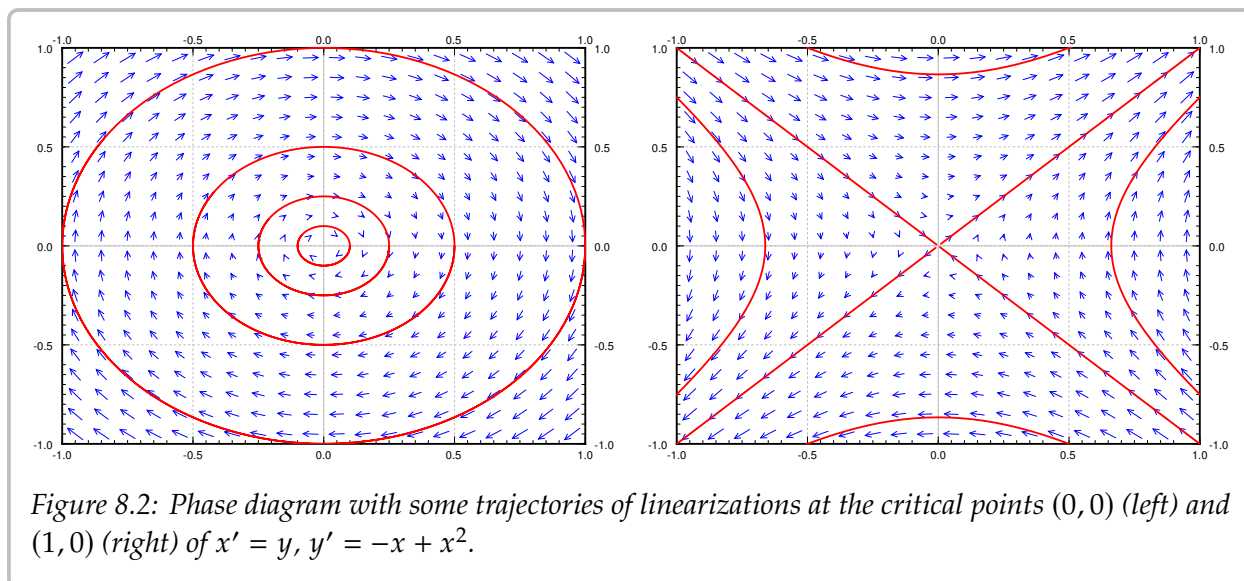
$$\begin{bmatrix} u \\ v \end{bmatrix}' = \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix}.$$

At the point $(1, 0)$, the Jacobian matrix is $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$, and our new variables are $u = x - 1$ and $v = y$. The linearization at $(1, 0)$ is

$$\begin{bmatrix} u \\ v \end{bmatrix}' = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix}.$$

The phase diagrams of the two linearizations at the critical points $(0, 0)$ and $(1, 0)$ are given in [Figure 8.2](#) on the next page. Note that the variables are now u and v , which are different at each critical point. Compare [Figure 8.2](#) with [Figure 8.1](#) on the preceding page, and look especially at the behavior near the critical points.

*Named for the German mathematician [Carl Gustav Jacob Jacobi](#) (1804–1851).



8.1.3 Exercises

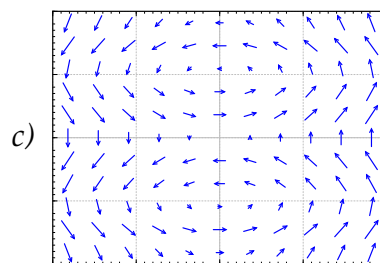
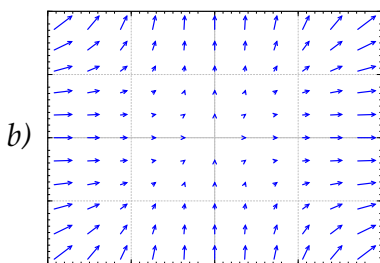
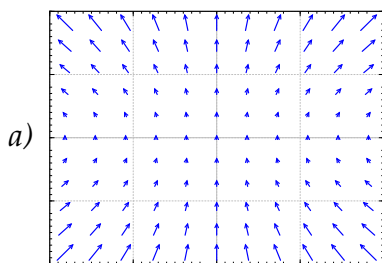
Exercise 8.1.1: Sketch the phase plane vector field for:

a) $x' = x^2$, $y' = y^2$, b) $x' = (x - y)^2$, $y' = -x$, c) $x' = e^y$, $y' = e^x$.

Exercise 8.1.2: Match systems

1) $x' = x^2$, $y' = y^2$, 2) $x' = xy$, $y' = 1 + y^2$, 3) $x' = \sin(\pi y)$, $y' = x$,

to the vector fields below. Justify.



Exercise 8.1.3: Find the critical points and linearizations of the following systems.

a) $x' = x^2 - y^2$, $y' = x^2 + y^2 - 1$, b) $x' = -y$, $y' = 3x + yx^2$,
c) $x' = x^2 + y$, $y' = y^2 + x$.

Exercise 8.1.4: For the following systems, verify they have critical point at $(0,0)$, and find the linearization at $(0,0)$.

a) $x' = x + 2y + x^2 - y^2$, $y' = 2y - x^2$ b) $x' = -y$, $y' = x - y^3$
c) $x' = ax + by + f(x, y)$, $y' = cx + dy + g(x, y)$, where $f(0,0) = 0$, $g(0,0) = 0$, and all first partial derivatives of f and g are also zero at $(0,0)$, that is, $\frac{\partial f}{\partial x}(0,0) = \frac{\partial f}{\partial y}(0,0) = \frac{\partial g}{\partial x}(0,0) = \frac{\partial g}{\partial y}(0,0) = 0$.

Exercise 8.1.5: Take $x' = (x - y)^2$, $y' = (x + y)^2$.

- Find the set of critical points.
- Sketch a phase diagram and describe the behavior near the critical point(s).
- Find the linearization. Is it helpful in understanding the system?

Exercise 8.1.6: Take $x' = x^2$, $y' = x^3$.

- Find the set of critical points.
- Sketch a phase diagram and describe the behavior near the critical point(s).
- Find the linearization. Is it helpful in understanding the system?

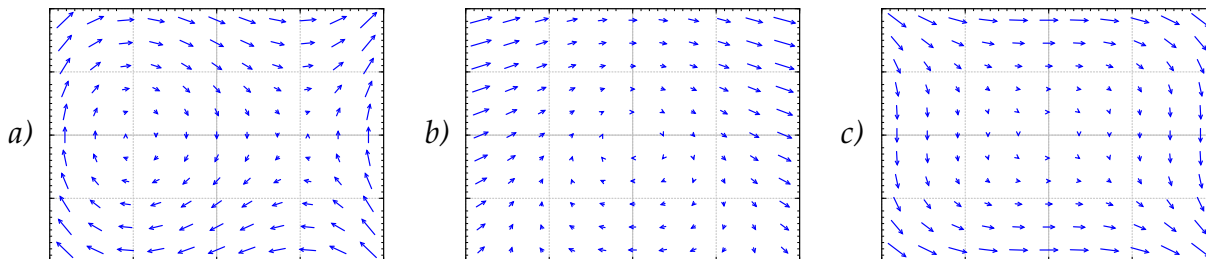
Exercise 8.1.101: Find the critical points and linearizations of the following systems.

- $x' = \sin(\pi y) + (x - 1)^2$, $y' = y^2 - y$,
- $x' = x + y + y^2$, $y' = x$,
- $x' = (x - 1)^2 + y$, $y' = x^2 + y$.

Exercise 8.1.102: Match systems

- $x' = y^2$, $y' = -x^2$,
- $x' = y$, $y' = (x - 1)(x + 1)$,
- $x' = y + x^2$, $y' = -x$,

to the vector fields below. Justify.



Exercise 8.1.103: The idea of critical points and linearization works in higher dimensions as well. You simply make the Jacobian matrix bigger by adding more functions and more variables. For the following system of 3 equations find the critical points and their linearizations:

$$x' = x + z^2, \quad y' = z^2 - y, \quad z' = z + x^2.$$

Exercise 8.1.104: Any two-dimensional non-autonomous system $x' = f(x, y, t)$, $y' = g(x, y, t)$ can be written as a three-dimensional autonomous system (three equations). Write down this autonomous system using the variables u, v, w .

8.2 Stability and classification of isolated critical points

Note: 1.5–2 lectures, §6.1–§6.2 in [EP], §9.2–§9.3 in [BD]

8.2.1 Isolated critical points and almost linear systems

A critical point is *isolated* if it is the only critical point in some small “neighborhood” of the point. That is, if we zoom in enough, it is the only critical point we see. In [Example 8.1.1](#) above, both critical points are isolated. If, on the other hand, there were a whole curve of critical points, then it would not be isolated.

A system is called *almost linear* at a critical point (x_0, y_0) if the critical point is isolated and the Jacobian matrix at the point is invertible, or equivalently if the linearized system has an isolated critical point. In such a case, the nonlinear terms are very small and the system behaves like its linearization, at least if we are close to the critical point.

For example, the system in [Examples 8.1.1](#) and [8.1.2](#) has two isolated critical points $(0, 0)$ and $(1, 0)$, and is almost linear at both critical points as the Jacobian matrices at both points, $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ and $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$, are invertible.

On the other hand, the system $x' = x^2$, $y' = y^2$ has an isolated critical point at $(0, 0)$, however the Jacobian matrix

$$\begin{bmatrix} 2x & 0 \\ 0 & 2y \end{bmatrix}$$

is zero when $(x, y) = (0, 0)$. So the system is not almost linear. An even worse example is the system $x' = x$, $y' = x^2$, which does not have isolated critical points; x' and y' are both zero whenever $x = 0$, that is, the entire y -axis.

Fortunately, most often critical points are isolated, and the system is almost linear at the critical points. So if we learn what happens there, we will have figured out the majority of situations that arise in applications.

8.2.2 Stability and classification of isolated critical points

Once we have an isolated critical point, the system is almost linear at that critical point, and we have computed the associated linearized system, we can classify what happens to the solutions. We more or less use the classification for linear two-variable systems from [§ 3.5](#), with one minor caveat. Let us list the behaviors depending on the eigenvalues of the Jacobian matrix at the critical point in [Table 8.1](#) on the next page. This table is very similar to [Table 3.1](#) on page 150, with the exception of missing “center” points. We will discuss centers later, as they are more complicated.

In the third column, we mark points as *asymptotically stable* or *unstable*. Formally, a *stable critical point* (x_0, y_0) is one where given any small distance ϵ to (x_0, y_0) , and any initial condition within a perhaps smaller radius around (x_0, y_0) , the trajectory of the system never goes further away from (x_0, y_0) than ϵ . An *unstable critical point* is one that is not

Eigenvalues of the Jacobian matrix	Behavior	Stability
real and both positive	source / unstable node	unstable
real and both negative	sink / stable node	asymptotically stable
real and opposite signs	saddle	unstable
complex with positive real part	spiral source	unstable
complex with negative real part	spiral sink	asymptotically stable

Table 8.1: Behavior of an almost linear system near an isolated critical point.

stable. Informally, a point is stable if we start close to a critical point and follow a trajectory, we either go towards, or at least not away from, this critical point.

A stable critical point (x_0, y_0) is called *asymptotically stable* if given any initial condition sufficiently close to (x_0, y_0) and any solution $(x(t), y(t))$ satisfying that condition, then

$$\lim_{t \rightarrow \infty} (x(t), y(t)) = (x_0, y_0).$$

That is, the critical point is asymptotically stable if any trajectory for a sufficiently close initial condition goes towards the critical point (x_0, y_0) .

Example 8.2.1: Consider $x' = -y - x^2$, $y' = -x + y^2$. See Figure 8.3 on the next page for the phase diagram. Let us find the critical points. These are the points where $-y - x^2 = 0$ and $-x + y^2 = 0$. The first equation means $y = -x^2$, and so $y^2 = x^4$. Plugging into the second equation, we obtain $-x + x^4 = 0$. Factoring, we obtain $x(x^3 - 1) = 0$. Since we are looking only for real solutions, we get either $x = 0$ or $x = 1$. Solving for the corresponding y using $y = -x^2$, we get two critical points, one being $(0, 0)$ and the other being $(1, -1)$. Clearly the critical points are isolated.

The Jacobian matrix is

$$\begin{bmatrix} -2x & -1 \\ -1 & 2y \end{bmatrix}.$$

At the point $(0, 0)$, we get the matrix $\begin{bmatrix} 0 & -1 \\ -1 & 0 \end{bmatrix}$ whose eigenvalues are 1 and -1 . As the matrix is invertible, the system is almost linear at $(0, 0)$. As the eigenvalues are real and of opposite signs, we get a saddle point, which is an unstable equilibrium point.

At the point $(1, -1)$, we get the matrix $\begin{bmatrix} -2 & -1 \\ -1 & -2 \end{bmatrix}$ whose eigenvalues are -1 and -3 . The matrix is invertible, and so the system is almost linear at $(1, -1)$. As both eigenvalues are real and negative, the critical point is a sink, and therefore an asymptotically stable equilibrium point. That is, if we start with any point (x_i, y_i) close to $(1, -1)$ as an initial condition and plot a trajectory, it approaches $(1, -1)$. In other words,

$$\lim_{t \rightarrow \infty} (x(t), y(t)) = (1, -1).$$

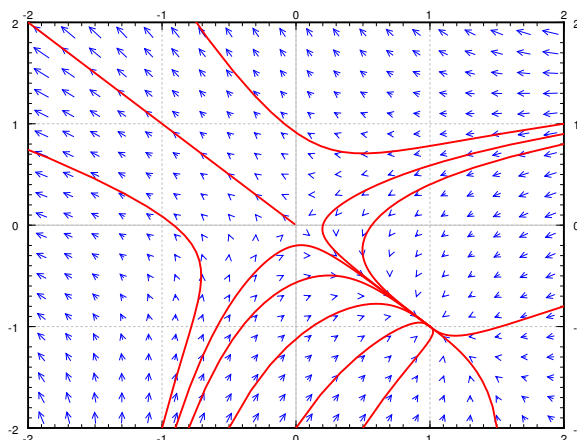


Figure 8.3: The phase portrait with few sample trajectories of $x' = -y - x^2$, $y' = -x + y^2$.

As you can see from the diagram, this behavior is true even for some initial points quite far from $(1, -1)$, but it is definitely not true for all initial points.

Example 8.2.2: Consider $x' = y + y^2e^x$, $y' = x$. First, we find the critical points. These are the points where $y + y^2e^x = 0$ and $x = 0$. Simplifying we get $0 = y + y^2 = y(y + 1)$. So the critical points are $(0, 0)$ and $(0, -1)$, and hence are isolated. Let us compute the Jacobian matrix:

$$\begin{bmatrix} y^2e^x & 1 + 2ye^x \\ 1 & 0 \end{bmatrix}.$$

At the point $(0, 0)$ we get the matrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$ and so the two eigenvalues are 1 and -1 . As the matrix is invertible, the system is almost linear at $(0, 0)$. And, as the eigenvalues are real and of opposite signs, we get a saddle point, which is an unstable equilibrium point.

At the point $(0, -1)$ we get the matrix $\begin{bmatrix} 1 & -1 \\ 1 & 0 \end{bmatrix}$ whose eigenvalues are $\frac{1}{2} \pm i\frac{\sqrt{3}}{2}$. The matrix is invertible, and so the system is almost linear at $(0, -1)$. As we have complex eigenvalues with positive real part, the critical point is a spiral source, and therefore an unstable equilibrium point.

See [Figure 8.4](#) on the following page for the phase diagram. Notice the two critical points, and the behavior of the arrows in the vector field around these points.

8.2.3 The trouble with centers

Recall, a linear system with a center means that trajectories travel in closed elliptical orbits in some direction around the critical point. Such a critical point we call a *center* or a *stable center*. It is not an asymptotically stable critical point, as the trajectories never approach the critical point, but at least if you start sufficiently close to the critical point, you stay close to the critical point. The simplest example of such behavior is the linear system with a center. Another example is the critical point $(0, 0)$ in [Example 8.1.1](#) on page 352.

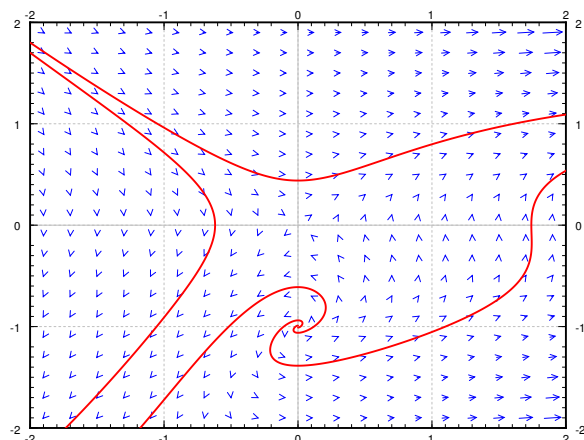


Figure 8.4: The phase portrait with few sample trajectories of $x' = y + y^2 e^x$, $y' = x$.

The trouble with a center in a nonlinear system is that whether the trajectory goes towards or away from the critical point is governed by the sign of the real part of the eigenvalues of the Jacobian matrix, and the Jacobian matrix in a nonlinear system changes from point to point. Since this real part is zero at the critical point itself, it can have either sign nearby, meaning the trajectory could be pulled towards or away from the critical point.

Example 8.2.3: An example of such a problematic behavior is the system $x' = y$, $y' = -x + y^3$. The only critical point is the origin $(0, 0)$. The Jacobian matrix is

$$\begin{bmatrix} 0 & 1 \\ -1 & 3y^2 \end{bmatrix}.$$

At $(0, 0)$ the Jacobian matrix is $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$, which has eigenvalues $\pm i$. So the linearization has a center.

Via the quadratic equation, the eigenvalues of the Jacobian matrix at any point (x, y) are

$$\lambda = \frac{3}{2}y^2 \pm i \frac{\sqrt{4 - 9y^4}}{2}.$$

At any point where $y \neq 0$ (so at most points near the origin), the eigenvalues have a positive real part (y^2 can never be negative). This positive real part pulls the trajectory away from the origin. A sample trajectory for an initial condition near the origin is given in [Figure 8.5](#) on the facing page.

The moral of the example is that further analysis is needed when the linearization has a center. The analysis will in general be more complicated than in the example above and is more likely to involve case-by-case consideration. Such a complication should not be surprising to you. By now in your mathematical career, you have seen many places where a simple test is inconclusive, and more careful, perhaps ad hoc, analysis is required. Recall for example when the second derivative test for maxima or minima is inconclusive.

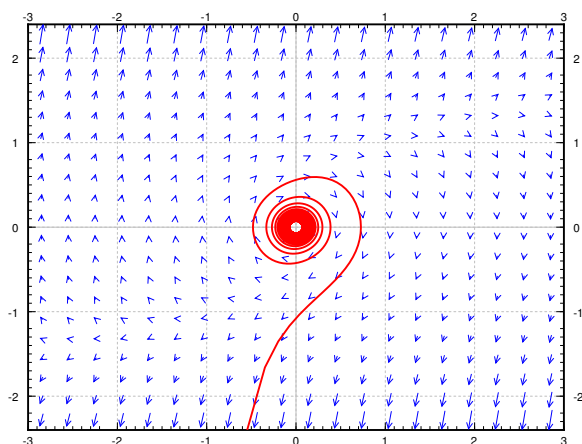


Figure 8.5: An unstable critical point (spiral source) at the origin for $x' = y, y' = -x + y^3$, even if the linearization has a center.

8.2.4 Conservative equations

An equation of the form

$$x'' + f(x) = 0,$$

where $f(x)$ is an arbitrary function, is called a *conservative equation*. For example, the pendulum equation is a conservative equation. The equations are conservative as there is no friction in the system so the energy in the system is “conserved.” Let us write this equation as a system of nonlinear ODEs:

$$x' = y, \quad y' = -f(x).$$

These types of equations have the advantage that we can solve for their trajectories easily.

The trick is to first think of y as a function of x for a moment. Then use the chain rule

$$x'' = y' = \frac{dy}{dx}x' = y \frac{dy}{dx},$$

where the prime indicates a derivative with respect to t . We obtain $y \frac{dy}{dx} + f(x) = 0$. We integrate with respect to x to get $\int y \frac{dy}{dx} dx + \int f(x) dx = C$. In other words,

$$\frac{1}{2}y^2 + \int f(x) dx = C.$$

We obtained an implicit equation for the trajectories, with different C giving different trajectories. The value of C is conserved on any trajectory. This expression is sometimes called the *Hamiltonian* or the energy of the system. If you look back to § 1.8, you will notice that $y \frac{dy}{dx} + f(x) = 0$ is an exact equation, and we just found a potential function.

Example 8.2.4: Let us find the trajectories for the equation $x'' + x - x^2 = 0$, which is the equation from [Example 8.1.1](#) on page 352. The corresponding first-order system is

$$x' = y, \quad y' = -x + x^2.$$

Trajectories satisfy

$$\frac{1}{2}y^2 + \frac{1}{2}x^2 - \frac{1}{3}x^3 = C.$$

We solve for y

$$y = \pm \sqrt{-x^2 + \frac{2}{3}x^3 + 2C}.$$

Plotting these graphs we get exactly the trajectories in [Figure 8.1](#) on page 353. In particular we notice that near the origin the trajectories are *closed curves*: they keep going around the origin, never spiraling in or out. Therefore we discovered a way to verify that the critical point at $(0, 0)$ is a stable center. The critical point at $(1, 0)$ is a saddle as we already noticed. This example is typical for conservative equations.

Consider an arbitrary conservative equation $x'' + f(x) = 0$. All critical points occur when $y = 0$ (the x -axis), that is, when $x' = 0$. The critical points are those points on the x -axis where $f(x) = 0$. The trajectories are given by

$$y = \pm \sqrt{-2 \int f(x) dx + 2C}.$$

So all trajectories are mirrored across the x -axis. In particular, there can be no spiral sources or sinks. The Jacobian matrix is

$$\begin{bmatrix} 0 & 1 \\ -f'(x) & 0 \end{bmatrix}.$$

The critical point is almost linear if $f'(x) \neq 0$ at the critical point. Let J denote the Jacobian matrix. The eigenvalues of J are solutions to

$$0 = \det(J - \lambda I) = \lambda^2 + f'(x).$$

Therefore $\lambda = \pm \sqrt{-f'(x)}$. In other words, either we get real eigenvalues of opposite signs (if $f'(x) < 0$), or we get purely imaginary eigenvalues (if $f'(x) > 0$). There are only two possibilities for critical points, either an *unstable saddle point*, or a *stable center*. There are never any sinks or sources.

8.2.5 Exercises

Exercise 8.2.1: For the systems below, find and classify the critical points. Also indicate if the equilibria are stable, asymptotically stable, or unstable.

a) $x' = -x + 3x^2, y' = -y$

b) $x' = x^2 + y^2 - 1, y' = x$

c) $x' = ye^x, y' = y - x + y^2$

Exercise 8.2.2: Find the implicit equations of the trajectories of the following conservative systems. Next find their critical points (if any) and classify them.

a) $x'' + x + x^3 = 0$

b) $\theta'' + \sin \theta = 0$

c) $z'' + (z - 1)(z + 1) = 0$

d) $x'' + x^2 + 1 = 0$

Exercise 8.2.3: Find and classify the critical point(s) of $x' = -x^2$, $y' = -y^2$.

Exercise 8.2.4: Suppose $x' = -xy$, $y' = x^2 - 1 - y$.

a) Show there are two spiral sinks at $(-1, 0)$ and $(1, 0)$.

b) For any initial point of the form $(0, y_0)$, find the trajectory.

c) Can a trajectory starting at (x_0, y_0) where $x_0 > 0$ spiral into the critical point at $(-1, 0)$? Why or why not?

Exercise 8.2.5: In the example $x' = y$, $y' = y^3 - x$, show that for any trajectory, the distance from the origin is an increasing function. Conclude that at the origin the system behaves like a spiral source. Hint: Consider $f(t) = (x(t))^2 + (y(t))^2$ and show it has positive derivative.

Exercise 8.2.6: Suppose f is always positive. Find the trajectories of $x'' + f(x') = 0$. Are there any critical points?

Exercise 8.2.7: Suppose that $x' = f(x, y)$, $y' = g(x, y)$. Suppose that $g(x, y) > 1$ for all x and y . Are there any critical points? What can we say about the trajectories as t goes to infinity?

Exercise 8.2.101: For the systems below, find and classify the critical points.

a) $x' = -x + x^2$, $y' = y$

b) $x' = y - y^2 - x$, $y' = -x$

c) $x' = xy$, $y' = x + y - 1$

Exercise 8.2.102: Find the implicit equations of the trajectories of the following conservative systems. Next find their critical points (if any) and classify them.

a) $x'' + x^2 = 4$

b) $x'' + e^x = 0$

c) $x'' + (x + 1)e^x = 0$

Exercise 8.2.103: The conservative system $x'' + x^3 = 0$ is not almost linear. Classify its critical point(s) nonetheless.

Exercise 8.2.104: Derive an analogous classification of critical points for equations in one dimension, such as $x' = f(x)$ based on the derivative. A point x_0 is critical when $f(x_0) = 0$ and almost linear if in addition $f'(x_0) \neq 0$. Figure out if the critical point is stable or unstable depending on the sign of $f'(x_0)$. Explain. Hint: see § 1.6.

8.3 Applications of nonlinear systems

Note: 2 lectures, §6.3–§6.4 in [EP], §9.3, §9.5 in [BD]

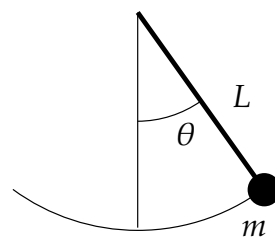
In this section we study two very standard examples of nonlinear systems. First, we look at the nonlinear pendulum equation. We saw the pendulum equation's linearization before, but we noted it was only valid for small angles and short times. Now we find out what happens for large angles. Next, we look at the predator-prey equation, which finds various applications in modeling problems in biology, chemistry, economics, and elsewhere.

8.3.1 Pendulum

The first example we study is the pendulum equation $\theta'' + \frac{g}{L} \sin \theta = 0$. Here, θ is the angular displacement, g is the gravitational acceleration, and L is the length of the pendulum. In this equation we disregard friction, so we are talking about an idealized pendulum.

This equation is a conservative equation, so we can use our analysis of conservative equations from the previous section. We change the equation to a two-dimensional system in variables (θ, ω) by introducing the new variable $\omega = \theta'$:

$$\begin{bmatrix} \theta \\ \omega \end{bmatrix}' = \begin{bmatrix} \omega \\ -\frac{g}{L} \sin \theta \end{bmatrix}.$$



The critical points of this system are when $\omega = 0$ and $-\frac{g}{L} \sin \theta = 0$, or in other words if $\sin \theta = 0$. So the critical points are when $\omega = 0$ and θ is a multiple of π . That is, the points are $\dots (-2\pi, 0), (-\pi, 0), (0, 0), (\pi, 0), (2\pi, 0) \dots$. While there are infinitely many critical points, they are all isolated. For conservative equations, there are two types of critical points: either stable centers or saddle points.

We compute the Jacobian matrix:

$$\begin{bmatrix} \frac{\partial}{\partial \theta}(\omega) & \frac{\partial}{\partial \omega}(\omega) \\ \frac{\partial}{\partial \theta}(-\frac{g}{L} \sin \theta) & \frac{\partial}{\partial \omega}(-\frac{g}{L} \sin \theta) \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -\frac{g}{L} \cos \theta & 0 \end{bmatrix}.$$

The eigenvalues of this matrix are $\lambda = \pm \sqrt{-\frac{g}{L} \cos \theta}$. The eigenvalues are going to be real when $\cos \theta < 0$. This happens at the odd multiples of π . The eigenvalues are going to be purely imaginary when $\cos \theta > 0$. This happens at the even multiples of π . Therefore the system has a stable center at the points $\dots (-2\pi, 0), (0, 0), (2\pi, 0) \dots$, and it has an unstable saddle at the points $\dots (-3\pi, 0), (-\pi, 0), (\pi, 0), (3\pi, 0) \dots$. Look at the phase diagram in [Figure 8.6](#) on the next page, where for simplicity we let $\frac{g}{L} = 1$.

In the linearized equation, we have only a single critical point, the center at $(0, 0)$. Now we see more clearly what we meant when we said the linearization is good for small

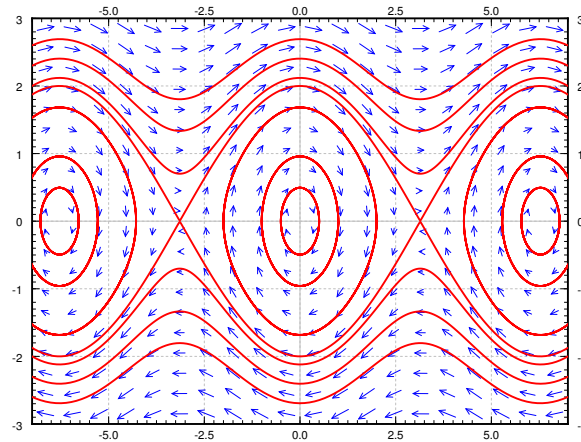


Figure 8.6: Phase plane diagram and some trajectories of the nonlinear pendulum equation.

angles. The horizontal axis is the deflection angle. The vertical axis is the angular velocity of the pendulum. Suppose we start at $\theta = 0$ (no deflection), and we start with a small angular velocity ω . Then the trajectory keeps going around the critical point $(0, 0)$ in an approximate circle. This corresponds to short swings of the pendulum back and forth. When θ stays small, the trajectories really look like circles and hence are very close to our linearization.

When we give the pendulum a big enough push, it goes across the top and keeps spinning about its axis. This behavior corresponds to the wavy curves that do not cross the horizontal axis in the phase diagram. Let us suppose we look at the top curves, when the angular velocity ω is large and positive. Then the pendulum is going around and around its axis. The velocity is going to be large when the pendulum is near the bottom, and the velocity is the smallest when the pendulum is close to the top of its loop.

At each critical point, there is an equilibrium solution. Consider the solution $\theta = 0$; the pendulum is not moving and is hanging straight down. This is a stable place for the pendulum to be, hence this is a *stable* equilibrium.

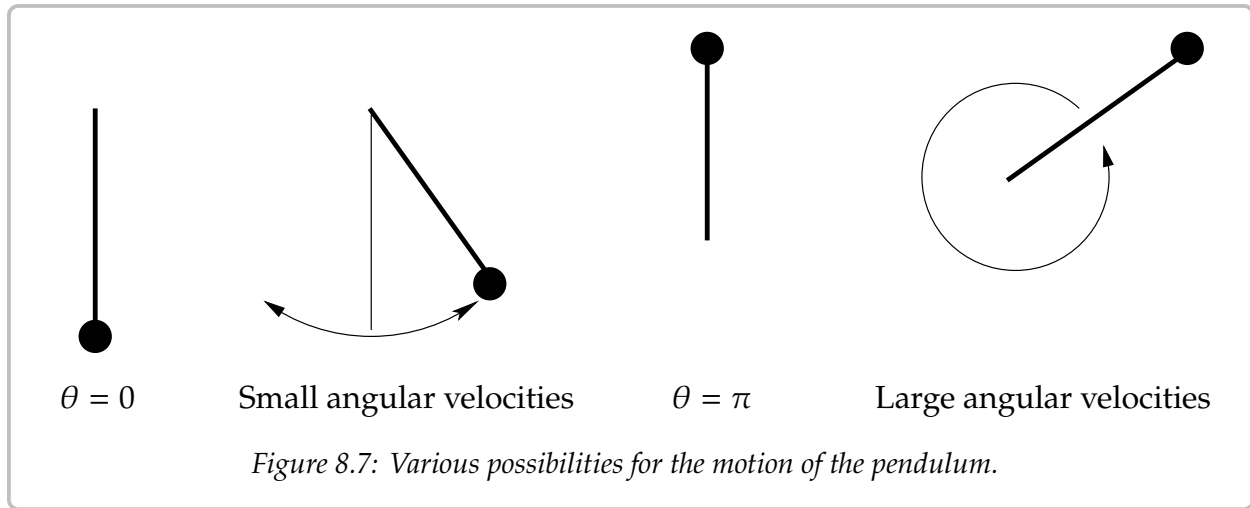
The other type of equilibrium solution is at the unstable point, for example $\theta = \pi$. Here the pendulum is upside down. Sure, you can balance the pendulum this way and it will stay, but this is an *unstable* equilibrium. Even the tiniest push will make the pendulum start swinging wildly.

See Figure 8.7 on the following page for a diagram. The first picture is the stable equilibrium $\theta = 0$. The second corresponds to those “almost circles” in the phase diagram around $\theta = 0$ when the angular velocity is small. The third is the unstable equilibrium $\theta = \pi$. The last picture corresponds to the wavy lines for large angular velocities.

The quantity

$$\frac{1}{2}\omega^2 - \frac{g}{L}\cos\theta$$

is conserved by any solution. This is the energy or the Hamiltonian of the system.



We have a conservative equation and so (exercise) the trajectories are given by

$$\omega = \pm \sqrt{\frac{2g}{L} \cos \theta + C},$$

for various values of C . Let us look at the initial condition of $(\theta_0, 0)$, that is, we take the pendulum to angle θ_0 , and just let it go (initial angular velocity 0). We plug the initial conditions into the above and solve for C to obtain

$$C = -\frac{2g}{L} \cos \theta_0.$$

Thus the expression for the trajectory is

$$\omega = \pm \sqrt{\frac{2g}{L}} \sqrt{\cos \theta - \cos \theta_0}.$$

Let us figure out the period. That is, the time it takes for the pendulum to swing back and forth. We notice that the trajectory about the origin in the phase plane is symmetric about both the θ and the ω -axis. That is, in terms of θ , the time it takes from θ_0 to $-\theta_0$ is the same as it takes from $-\theta_0$ back to θ_0 . Furthermore, the time it takes from $-\theta_0$ to 0 is the same as to go from 0 to θ_0 . Therefore, let us find how long it takes for the pendulum to go from angle 0 to angle θ_0 , which is a quarter of the full oscillation and then multiply by 4.

We figure out this time by finding $\frac{dt}{d\theta}$ and integrating from 0 to θ_0 . The period is four times this integral. Let us stay in the region where ω is positive. Since $\omega = \frac{d\theta}{dt}$, inverting we get

$$\frac{dt}{d\theta} = \sqrt{\frac{L}{2g}} \frac{1}{\sqrt{\cos \theta - \cos \theta_0}}.$$

Therefore the period T is given by

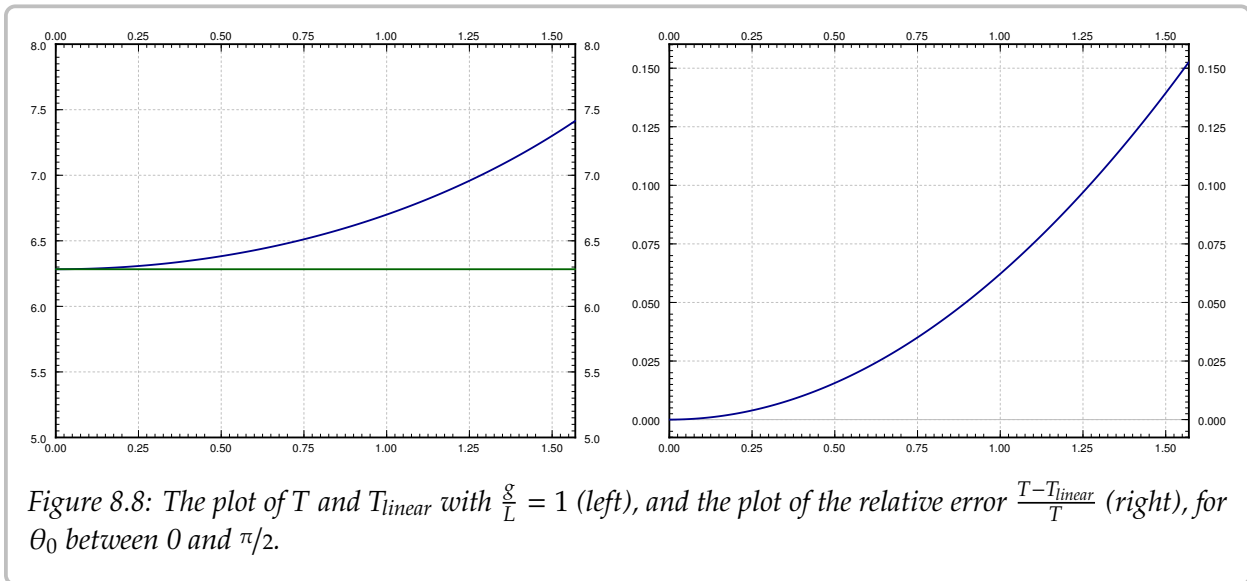
$$T = 4 \sqrt{\frac{L}{2g}} \int_0^{\theta_0} \frac{1}{\sqrt{\cos \theta - \cos \theta_0}} d\theta.$$

The integral is an improper integral, and we cannot in general evaluate it symbolically. We must resort to numerical approximation if we want to compute a particular T .

Recall from § 2.4, the linearized equation $\theta'' + \frac{g}{L}\theta = 0$ has period

$$T_{\text{linear}} = 2\pi\sqrt{\frac{L}{g}}.$$

We plot T , T_{linear} , and the relative error $\frac{T - T_{\text{linear}}}{T}$ in Figure 8.8. The relative error says how far our approximation is from the real period percentage-wise. Note that T_{linear} is simply a constant. It does not change with the initial angle θ_0 . The actual period T gets larger and larger as θ_0 gets larger. Notice how the relative error is small when θ_0 is small. It is still only 15% when $\theta_0 = \pi/2$, that is, a 90 degree angle. The error is 3.8% when starting at $\pi/4$, a 45 degree angle. At a 5 degree initial angle, the error is only 0.048%.



While it is not immediately obvious from the formula, it is true that

$$\lim_{\theta_0 \uparrow \pi} T = \infty.$$

That is, the period goes to infinity as the initial angle approaches the unstable equilibrium point. So if we put the pendulum almost upside down it may take a very long time before it gets down. This is consistent with the limiting behavior, where the exactly upside down pendulum never makes an oscillation, so we could think of that as infinite period.

8.3.2 Predator-prey or Lotka–Volterra systems

One of the most common simple applications of nonlinear systems is the so-called *predator-prey* or *Lotka–Volterra*^{*} systems. For example, such a system arises when two species interact, one as the prey and one as the predator. It is then no surprise that the equations also see applications in economics. The system also arises in chemical reactions. In biology, this system of equations explains the natural periodic variations of populations of different species. Before the application of differential equations, these periodic variations in the population baffled biologists.

We keep with the classical example of hares and foxes in a forest. It is the easiest to understand.

$$\begin{aligned}x &= \# \text{ of hares (the prey),} \\ y &= \# \text{ of foxes (the predator).}\end{aligned}$$

When there are a lot of hares, there is plenty of food for the foxes, so the fox population grows. However, when the fox population grows, the foxes eat more hares, so when there are lots of foxes, the hare population should go down, and vice versa. The Lotka–Volterra model proposes that this behavior is described by the system of equations

$$\begin{aligned}x' &= (a - by)x, \\ y' &= (cx - d)y,\end{aligned}$$

where a, b, c, d are some parameters that describe the interaction of the foxes and hares[†]. In this model, these are all positive numbers.

Let us analyze the idea behind this model. The model is a slightly more complicated idea based on the exponential population model. First expand,

$$x' = (a - by)x = ax - byx.$$

The hares are expected to simply grow exponentially in the absence of foxes. That is where the ax term comes in, the growth in population is proportional to the population itself. We are assuming the hares always find enough food and have enough space to reproduce. However, there is another component $-byx$, that is, the population also is decreasing proportionally to the number of foxes. Together we can write the equation as $(a - by)x$, so it is like exponential growth or decay but the constant depends on the number of foxes.

The equation for foxes is very similar, expand again

$$y' = (cx - d)y = cxy - dy.$$

The foxes need food (hares) to reproduce: the more food, the bigger the rate of growth, hence the cxy term. On the other hand, there are natural deaths in the fox population, and hence the $-dy$ term.

^{*}Named for the American mathematician, chemist, and statistician [Alfred James Lotka](#) (1880–1949) and the Italian mathematician and physicist [Vito Volterra](#) (1860–1940).

[†]This interaction does not end well for the hare.

Without further delay, let us start with an explicit example. Suppose the equations are

$$x' = (0.4 - 0.01y)x, \quad y' = (0.003x - 0.3)y.$$

See Figure 8.9 for the phase portrait. In this example it makes sense to also plot x and y as graphs with respect to time. Therefore the second graph in Figure 8.9 is the graph of x and y on the vertical axis (the prey x is the thinner line with taller peaks), against time on the horizontal axis. The particular solution graphed was with initial conditions of 20 foxes and 50 hares.

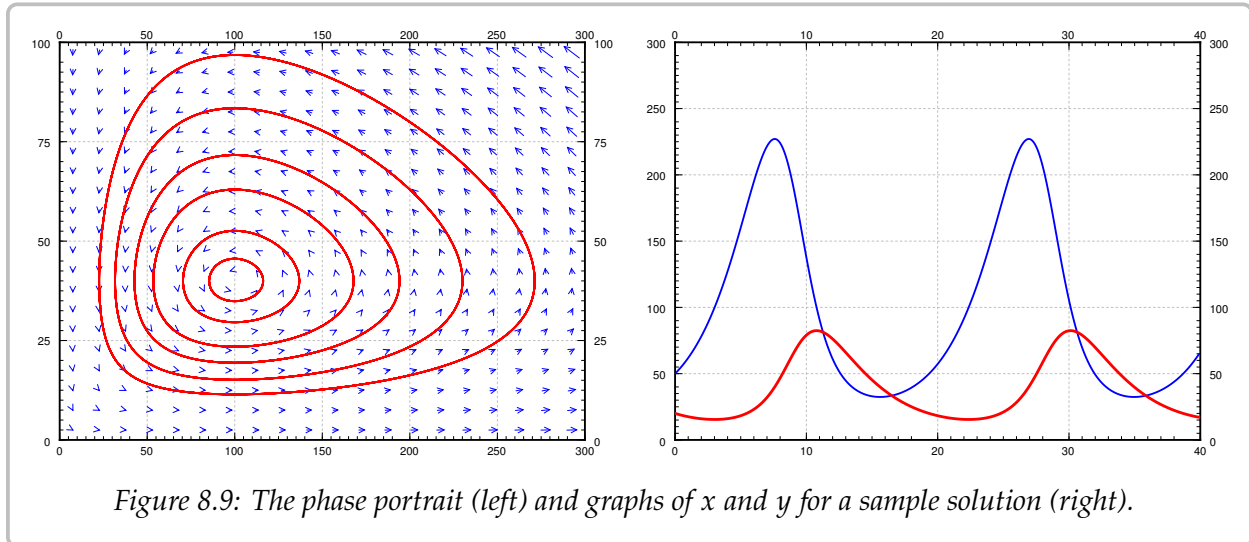


Figure 8.9: The phase portrait (left) and graphs of x and y for a sample solution (right).

Let us analyze what we see on the graphs. We work in the general setting rather than putting in specific numbers. We start with finding the critical points. Set $(a - by)x = 0$, and $(cx - d)y = 0$. The first equation is satisfied if either $x = 0$ or $y = a/b$. If $x = 0$, the second equation implies $y = 0$. If $y = a/b$, the second equation implies $x = d/c$. There are two equilibria: at $(0, 0)$ when there are no animals at all, and at $(d/c, a/b)$. In our specific example $x = d/c = 100$, and $y = a/b = 40$. This is the point where there are 100 hares and 40 foxes.

We compute the Jacobian matrix:

$$\begin{bmatrix} a - by & -bx \\ cy & cx - d \end{bmatrix}.$$

At the origin $(0, 0)$ we get the matrix $\begin{bmatrix} a & 0 \\ 0 & -d \end{bmatrix}$, so the eigenvalues are a and $-d$, hence real and of opposite signs. So the critical point at the origin is a saddle. This makes sense. If you started with some foxes but no hares, then the foxes would go extinct, that is, you would approach the origin. If you started with no foxes and a few hares, then the hares would keep multiplying without check, and so you would go away from the origin.

OK, how about the other critical point at $(d/c, a/b)$. Here the Jacobian matrix becomes

$$\begin{bmatrix} 0 & -\frac{bd}{c} \\ \frac{ac}{b} & 0 \end{bmatrix}.$$

The eigenvalues satisfy $\lambda^2 + ad = 0$. In other words, $\lambda = \pm i\sqrt{ad}$. The eigenvalues being purely imaginary, we are in the case where we cannot quite decide using only linearization. We could have a stable center, spiral sink, or a spiral source. That is, the equilibrium could be asymptotically stable, stable, or unstable. Of course I gave you a picture above that seems to imply it is a stable center. But never trust a picture only. Perhaps the oscillations are getting larger and larger, but only *very* slowly. Of course this would be bad as it would imply something will go wrong with our population sooner or later. And I only graphed a very specific example with very specific trajectories.

How can we be sure we are in the stable situation? As we said before, in the case of purely imaginary eigenvalues, we have to do a bit more work. Previously we found that for conservative systems, there was a certain quantity that was conserved on the trajectories, and hence the trajectories had to go in closed loops. We can use a similar technique here. We just have to figure out what is the conserved quantity. After some trial and error we find the constant

$$C = \frac{y^a x^d}{e^{cx+by}} = y^a x^d e^{-cx-by}$$

is conserved. Such a quantity is called the *constant of motion*. Let us check C really is a constant of motion. How do we check, you say? Well, a constant is something that does not change with time, so let us compute the derivative with respect to time:

$$C' = ay^{a-1}y'x^de^{-cx-by} + y^adx^{d-1}x'e^{-cx-by} + y^ax^de^{-cx-by}(-cx' - by').$$

Our equations give us what x' and y' are so let us plug those in:

$$\begin{aligned} C' &= ay^{a-1}(cx - d)yx^de^{-cx-by} + y^adx^{d-1}(a - by)xe^{-cx-by} \\ &\quad + y^ax^de^{-cx-by}(-c(a - by)x - b(cx - d)y) \\ &= y^ax^de^{-cx-by}\left(a(cx - d) + d(a - by) + (-c(a - by)x - b(cx - d)y)\right) \\ &= 0. \end{aligned}$$

So C is constant along the trajectories. In fact, the expression $C = \frac{y^a x^d}{e^{cx+by}}$ gives us an implicit equation for the trajectories. In any case, once we have found this constant of motion, it must be true that the trajectories are simple curves—the level curves of $\frac{y^a x^d}{e^{cx+by}}$. It turns out, the critical point at $(d/c, a/b)$ is a maximum for C (left as an exercise). So $(d/c, a/b)$ is a stable equilibrium point, and we do not have to worry about the foxes and hares going extinct or their populations exploding.

One blemish on this wonderful model is that the number of foxes and hares are discrete quantities and we are modeling with continuous variables. Our model has no problem with there being 0.1 fox in the forest for example, while in reality that makes no sense. The approximation is a reasonable one as long as the number of foxes and hares are large, but it does not make much sense for small numbers. One must be careful in interpreting any results from such a model.

An interesting consequence (perhaps counterintuitive) of this model is that adding animals to the forest might lead to extinction, because the variations will get too big, and one of the populations will get close to zero. For example, suppose there are 20 foxes and 50 hares as before, but now we bring in more foxes, bringing their number to 200. If we run the computation, we find the number of hares will plummet to just slightly more than 1 hare in the whole forest. In reality that most likely means the hares die out, and then the foxes will die out as well as they will have nothing to eat.

Showing that a system of equations has a stable solution can be a very difficult problem. When Isaac Newton put forth his laws of planetary motions, he proved that a single planet orbiting a single sun is a stable system. But any solar system with more than 1 planet proved very difficult indeed. In fact, such a system behaves chaotically (see § 8.5), meaning small changes in initial conditions lead to very different long-term outcomes. From numerical experimentation and measurements, we know the Earth will not fly out into empty space or crash into the Sun, for at least some millions of years or so. But we do not know what happens beyond that.

8.3.3 Exercises

Exercise 8.3.1: Take the damped nonlinear pendulum equation $\theta'' + \mu\theta' + (g/L)\sin\theta = 0$ for some $\mu > 0$ (that is, there is some friction).

- Suppose $\mu = 1$ and $g/L = 1$ for simplicity. Find and classify the critical points.
- Do the same for any $\mu > 0$ and any g and L , but such that the damping is small, in particular, $\mu^2 < 4(g/L)$.
- Explain what your findings mean, and if it agrees with what you expect in reality.

Exercise 8.3.2: Suppose the hares do not grow exponentially, but logistically. In particular consider

$$x' = (0.4 - 0.01y)x - \gamma x^2, \quad y' = (0.003x - 0.3)y.$$

For the following two values of γ , find and classify all the critical points in the positive quadrant, that is, for $x \geq 0$ and $y \geq 0$. Then sketch the phase diagram. Discuss the implication for the long-term behavior of the population.

- $\gamma = 0.001$,
- $\gamma = 0.01$.

Exercise 8.3.3:

- Suppose x and y are positive variables. Show $\frac{yx}{e^{x+y}}$ attains a maximum at $(1, 1)$.
- Suppose a, b, c, d are positive constants, and also suppose x and y are positive variables. Show $\frac{y^a x^d}{e^{cx+by}}$ attains a maximum at $(d/c, a/b)$.

Exercise 8.3.4: Suppose that for the pendulum equation we take a trajectory giving the spinning-around motion, for example $\omega = \sqrt{\frac{2g}{L} \cos\theta + \frac{2g}{L} + \omega_0^2}$. This is the trajectory where the lowest angular velocity is ω_0 . Find an integral expression for how long it takes the pendulum to go all the way around.

Exercise 8.3.5 (challenging): Take the pendulum, suppose the initial position is $\theta = 0$.

- a) Find the expression for ω giving the trajectory with initial condition $(0, \omega_0)$. Hint: Figure out what C should be in terms of ω_0 .
- b) Find the crucial angular velocity ω_1 , such that for any higher initial angular velocity, the pendulum will keep going around its axis, and for any lower initial angular velocity, the pendulum will simply swing back and forth. Hint: When the pendulum doesn't go over the top, the expression for ω will be undefined for some θ s.
- c) What do you think happens if the initial condition is $(0, \omega_1)$, that is, the initial angle is 0, and the initial angular velocity is exactly ω_1 ?

Exercise 8.3.101: Take the damped nonlinear pendulum equation $\theta'' + \mu\theta' + (g/L)\sin\theta = 0$ for some $\mu > 0$ (that is, there is friction). Suppose the friction is large, in particular $\mu^2 > 4(g/L)$.

- a) Find and classify the critical points.
- b) Explain what your findings mean, and if it agrees with what you expect in reality.

Exercise 8.3.102: Suppose we have the predator-prey system where the foxes are also killed at a constant rate h (h foxes killed per unit time): $x' = (a - by)x$, $y' = (cx - d)y - h$.

- a) Find the critical points and the Jacobian matrices of the system.
- b) Put in the constants $a = 0.4$, $b = 0.01$, $c = 0.003$, $d = 0.3$, $h = 10$. Analyze the critical points. What do you think it says about the forest?

Exercise 8.3.103 (challenging): Suppose the foxes never die. That is, we have the system $x' = (a - by)x$, $y' = cxy$. Find the critical points and notice they are not isolated. What will happen to the population in the forest if it starts at some positive numbers. Hint: Think of the constant of motion.

8.4 Limit cycles

Note: less than 1 lecture, discussed in §6.1 and §6.4 in [EP], §9.7 in [BD]

For nonlinear systems, trajectories do not simply need to approach or leave a single point. They may in fact approach a larger set, such as a circle or another closed curve.

Example 8.4.1: The *Van der Pol oscillator** is the following equation

$$x'' - \mu(1 - x^2)x' + x = 0,$$

where μ is some positive constant. The Van der Pol oscillator originated with electrical circuits, but finds applications in diverse fields such as biology, seismology, and other physical sciences.

For simplicity, we use $\mu = 1$. A phase diagram is given in the left-hand plot in Figure 8.10, where $y = x'$ as usual. Notice how the trajectories seem to very quickly settle on a closed curve. On the right-hand side is the plot of a single solution for $t = 0$ to $t = 30$ with initial conditions $x(0) = 0.1$ and $x'(0) = 0.1$. The solution quickly tends to a periodic solution.

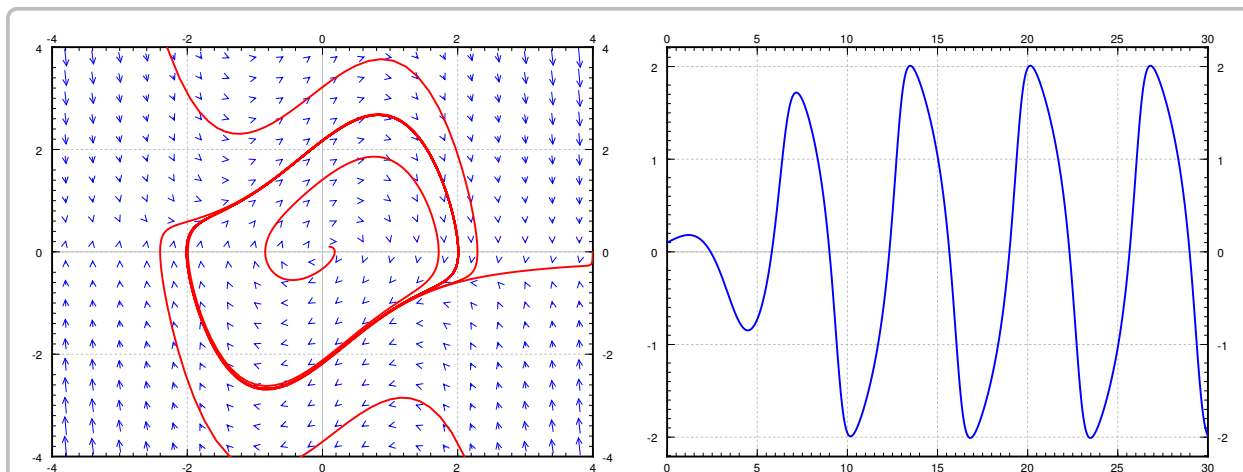


Figure 8.10: The phase portrait (left) and a graph of a sample solution of the Van der Pol oscillator.

The Van der Pol oscillator is an example of *relaxation oscillation*. The word relaxation refers to the sudden jump (the very steep part of the solution). For larger μ , the steep part is more pronounced. For smaller μ , the limit cycle looks more like a circle. In fact, if $\mu = 0$, then $x'' + x = 0$, which is a linear system with a center and all trajectories are circles.

A trajectory in the phase portrait that is a closed curve (a curve that is a loop) is called a *closed trajectory*. A *limit cycle* is a closed trajectory such that at least one other trajectory spirals into it (or spirals out of it). For example, the closed curve in the phase portrait for the Van der Pol equation is a limit cycle. If all trajectories that start near the limit cycle

*Named for the Dutch physicist [Balthasar van der Pol](#) (1889–1959).

spiral into it, the limit cycle is called *asymptotically stable*. The limit cycle in the Van der Pol oscillator is asymptotically stable.

Given a closed trajectory in an autonomous system, any solution that starts on it is periodic. Such a curve is called a *periodic orbit*. More precisely, if $(x(t), y(t))$ is a solution such that for some t_0 the point $(x(t_0), y(t_0))$ lies on a periodic orbit, then both $x(t)$ and $y(t)$ are periodic functions (with the same period). That is, there is some number P such that $x(t) = x(t + P)$ and $y(t) = y(t + P)$.

Consider the system

$$x' = f(x, y), \quad y' = g(x, y), \quad (8.2)$$

where the functions f and g have continuous derivatives in some region R in the plane.

Theorem 8.4.1 (Poincaré–Bendixson*). *Suppose R is a closed bounded region (a region in the plane that includes its boundary and does not have points arbitrarily far from the origin) and (8.2) has no critical points in R . Suppose $(x(t), y(t))$ is a solution of (8.2) in R that exists for all $t \geq t_0$. Then either the solution is a periodic function, or the solution tends towards a periodic solution in R .*

The point of the theorem is that if you find one solution that exists for all t large enough (that is, as t goes to infinity) and stays within a bounded region without critical points, then you have found either a periodic orbit or a solution that spirals towards a limit cycle—in the long term, the solution is very close to a periodic function. If R contains critical points, the behavior may be more complicated: The solution could tend towards a loop composed of critical points connected by trajectories, or the solution could simply tend towards a critical point. The theorem is more a qualitative statement rather than something to help us in computations. In practice, it is hard to either find analytic solutions or to show that they exist for all time. But if we think the solution exists, we numerically solve for a large time to approximate the limit cycle. Another caveat is that the theorem only works in two dimensions. In three dimensions and higher, there is simply too much room.

The theorem applies to all solutions in the Van der Pol oscillator. Solutions that start at any point except the origin $(0, 0)$ tend to the periodic solution around the limit cycle, and the initial condition of $(0, 0)$ gives the constant solution $x = 0, y = 0$.

Example 8.4.2: Consider

$$x' = y + (x^2 + y^2 - 1)^2 x, \quad y' = -x + (x^2 + y^2 - 1)^2 y.$$

A vector field and solutions with initial conditions $(1.02, 0)$, $(0.9, 0)$, and $(0.1, 0)$ are drawn in [Figure 8.11](#) on the facing page.

Notice that points on the unit circle (distance one from the origin) satisfy $x^2 + y^2 - 1 = 0$. And $x(t) = \sin(t)$, $y = \cos(t)$ is a solution of the system. Therefore we have a closed trajectory. For points off the unit circle, the second term in x' pushes the solution further away from the y -axis than the system $x' = y$, $y' = -x$, and y' pushes the solution further away from the x -axis than the linear system $x' = y$, $y' = -x$. In other words, for all other initial conditions the trajectory will spiral out.

*Ivar Otto Bendixson (1861–1935) was a Swedish mathematician.

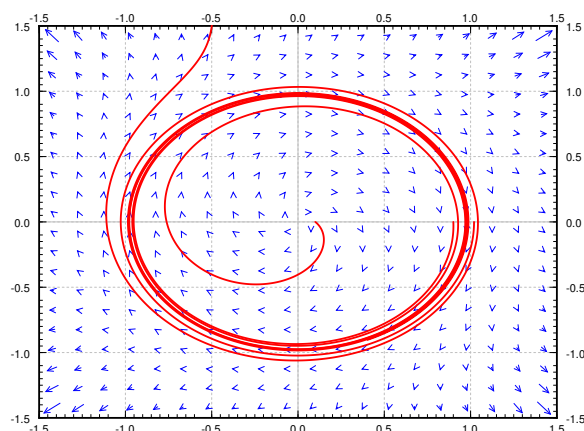


Figure 8.11: Unstable limit cycle example.

This means that for initial conditions inside the unit circle, the solution spirals out towards the periodic solution on the unit circle, and for initial conditions outside the unit circle the solutions spiral off towards infinity. Therefore the unit circle is a limit cycle, but not an asymptotically stable one. The Poincaré–Bendixson Theorem applies to the initial points inside the unit circle, as those solutions stay bounded, but not to those outside, as those solutions go off to infinity.

A similar analysis applies to the system

$$x' = y + (x^2 + y^2 - 1)x, \quad y' = -x + (x^2 + y^2 - 1)y.$$

The unit circle is again a closed trajectory, points outside the unit circle spiral out to infinity, but now points inside the unit circle spiral towards the critical point at the origin. The unit circle is an asymptotically unstable limit cycle where all trajectories spiral away. If we consider the linear system $x' = y$, $y' = -x$, then all circles centered at the origin (not just the unit circle) are closed trajectories, but no trajectory spirals onto any other, so none of the circles are limit cycles. All these solutions are periodic.

By Picard's theorem ([Theorem 3.1.1](#) on page 124), at any point in the plane, we can always find a solution to (8.2) a little bit forward or backwards in time, as long as f and g have continuous derivatives. So if we find a closed trajectory in an autonomous system, then for every initial point inside the closed trajectory, the solution will exist for all time and it will stay bounded (it will stay inside the closed trajectory). The moment we found the solution above going around the unit circle, we knew that for every initial point inside the circle, the solution exists for all time and the Poincaré–Bendixson theorem applies.

We next look for conditions when limit cycles (or periodic orbits) do not exist. We assume the equation (8.2) is defined on a *simply connected region*, that is, a region with no holes we can go around. For example, the entire plane is a simply connected region. So is the inside of the unit disc. However, the entire plane minus a point is not a simply connected region as it has a “hole” at the origin.

Theorem 8.4.2 (Bendixson–Dulac^{*}). Suppose R is a simply connected region, and the expression[†]

$$\frac{\partial f}{\partial x} + \frac{\partial g}{\partial y}$$

is either always positive or always negative on R (except perhaps a small set such as on isolated points or curves) then the system (8.2) has no closed trajectory inside R .

The theorem gives a way of ruling out the existence of a closed trajectory, and hence a way of ruling out limit cycles. The expression $\frac{\partial f}{\partial x} + \frac{\partial g}{\partial y}$ may seem random, but it is called the *divergence* of the vector field. Divergence measures how much the vector field is “expanding” at any point and comes up in many other contexts. The theorem says that if the vector field is either expanding everywhere on R or contracting everywhere on R , then there is no closed trajectory and so no limit cycle. Perhaps this is intuitive—if particles travel along the vector fields and are getting further and further apart, then we do not expect any particles to travel in loops. The exception about points or curves means that divergence can be zero at a few points, or on a curve, but not on any larger set.

Example 8.4.3: Consider $x' = y + y^2e^x$, $y' = x$ in the entire plane (see [Example 8.2.2](#) on page 359). The plane is simply connected and the theorem applies. We compute $\frac{\partial f}{\partial x} + \frac{\partial g}{\partial y} = y^2e^x + 0$. The function y^2e^x is always positive except on the line $y = 0$. Therefore, via the theorem, the system has no closed trajectories.

A common informal, but not quite correct, way to state the theorem is to conclude there are no periodic solutions that stay in R . The example above has two critical points and hence it has constant solutions. Constant functions are periodic. The conclusion of the theorem is that there exist no trajectories that form closed curves, or in other words, that there exist no nonconstant periodic solutions that stay in R .

Example 8.4.4: Consider a somewhat more complicated example. Take the system $x' = -y - x^2$, $y' = -x + y^2$ (see [Example 8.2.1](#) on page 358). We compute $\frac{\partial f}{\partial x} + \frac{\partial g}{\partial y} = -2x + 2y = 2(-x + y)$. This expression takes on both signs, so if we are talking about the whole plane we cannot simply apply the theorem. However, we could apply it on the set where $-x + y \geq 0$. Via the theorem, there is no closed trajectory in that set. Similarly, there is no closed trajectory in the set $-x + y \leq 0$. We cannot conclude (yet) that there is no closed trajectory in the entire plane. For all we know, perhaps half of it is in the set where $-x + y \geq 0$ and the other half is in the set where $-x + y \leq 0$.

The key is to look at the line where $-x + y = 0$, or $x = y$. On this line, $x' = -y - x^2 = -x - x^2$ and $y' = -x + y^2 = -x + x^2$. In particular, if $x = y$, then $x' \leq y'$. So the arrows, the vectors (x', y') , always point into the set where $-x + y \geq 0$. There is no way we can start in the set where $-x + y \geq 0$ and go into the set where $-x + y \leq 0$. Once we are in the set where $-x + y \geq 0$, we stay there. So no closed trajectory can have points in both sets.

^{*}Henri Dulac (1870–1955) was a French mathematician.

[†]Usually the expression in the Bendixson–Dulac Theorem is $\frac{\partial(\varphi f)}{\partial x} + \frac{\partial(\varphi g)}{\partial y}$ for some continuously differentiable function φ . For simplicity, let us just consider the case $\varphi = 1$.

Example 8.4.5: Consider $x' = y + (x^2 + y^2 - 1)x$, $y' = -x + (x^2 + y^2 - 1)y$ in the region R given by $x^2 + y^2 > \frac{1}{2}$. That is, R is the region outside a circle of radius $\frac{1}{\sqrt{2}}$ centered at the origin. Then there is a closed trajectory in R , namely $x = \cos(t)$, $y = \sin(t)$. Furthermore,

$$\frac{\partial f}{\partial x} + \frac{\partial g}{\partial y} = 4x^2 + 4y^2 - 2,$$

which is always positive on R . So what is going on? The Bendixson–Dulac theorem does not apply since the region R is not simply connected—it has a hole, the circle we cut out!

8.4.1 Exercises

Exercise 8.4.1: Show that the following systems have no closed trajectories.

- a) $x' = x^3 + y$, $y' = y^3 + x^2$, b) $x' = e^{x-y}$, $y' = e^{x+y}$,
 c) $x' = x + 3y^2 - y^3$, $y' = y^3 + x^2$.

Exercise 8.4.2: Formulate a condition for a 2-by-2 linear system $\vec{x}' = A\vec{x}$ to not be a center using the Bendixson–Dulac theorem. That is, the theorem says something about certain elements of A .

Exercise 8.4.3: Explain why the Bendixson–Dulac Theorem does not apply for any conservative system $x'' + h(x) = 0$.

Exercise 8.4.4: A system such as $x' = x$, $y' = y$ has solutions that exist for all time t , yet there are no closed trajectories. Explain why the Poincaré–Bendixson Theorem does not apply.

Exercise 8.4.5: Differential equations can also be given in different coordinate systems. Suppose we have the system $r' = 1 - r^2$, $\theta' = 1$ given in polar coordinates. Find all the closed trajectories and check if they are limit cycles and if so, if they are asymptotically stable or not.

Exercise 8.4.101: Show that the following systems have no closed trajectories.

- a) $x' = x + y^2$, $y' = y + x^2$, b) $x' = -x \sin^2(y)$, $y' = e^x$,
 c) $x' = xy^2$, $y' = x + x^2$.

Exercise 8.4.102: Suppose an autonomous system in the plane has a solution $x = \cos(t)(1 + e^{-t})$, $y = \sin(t)(1 + e^{-t})$. What can you say about the system (in particular about limit cycles and periodic solutions)?

Exercise 8.4.103: Show that the limit cycle of the Van der Pol oscillator (for $\mu > 0$) must not lie completely in the set where $-1 < x < 1$. Compare with [Figure 8.10](#) on page 373.

Exercise 8.4.104: Consider the system $r' = \sin(r)$, $\theta' = 1$ given in polar coordinates. Find all the closed trajectories.

8.5 Chaos

Note: 1 lecture, §6.5 in [EP], §9.8 in [BD]

You have surely heard the story about the flap of a butterfly wing in the Amazon causing hurricanes in the North Atlantic. In a prior section, we mentioned that a small change in initial conditions of the planets can lead to very different configurations of the planets in the long term. These are examples of *chaotic systems*. Mathematical chaos is not really chaos—there is precise order behind the scenes. Everything is still deterministic. However, a chaotic system is extremely sensitive to initial conditions. This also means even small errors induced via numerical approximation create large errors very quickly, so it is almost impossible to numerically approximate for long times. This is a large part of the trouble, as chaotic systems cannot be in general solved analytically.

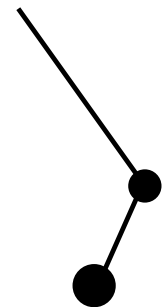
Take the weather, perhaps the most well-known chaotic system. A small change in the initial conditions (the temperature at every point of the atmosphere for example) produces drastically different predictions in a relatively short time, and so we cannot accurately predict weather. And we do not actually know the exact initial conditions. We measure temperatures at a few points with some error, and then we somehow estimate what is in between. There is no way we can accurately measure the effects of every butterfly wing. Then we solve the equations numerically, introducing new errors. You should not trust weather prediction more than a few days out. But we will see we can still get some information about a chaotic system on a longer time scale. In the context of weather, we can study the climate.

Chaotic behavior was first noticed by Edward Lorenz* in the 1960s when trying to model thermally induced air convection (movement). Lorenz was looking at the relatively simple system:

$$x' = -10x + 10y, \quad y' = 28x - y - xz, \quad z' = -\frac{8}{3}z + xy.$$

A small change in the initial conditions yields a very different solution after a reasonably short time.

A simple example the reader can experiment with, and which displays chaotic behavior, is a double pendulum. The equations for this setup are somewhat complicated, and their derivation is quite tedious, so we will not bother to write them down. The idea is to put a pendulum on the end of another pendulum. The movement of the bottom mass will appear chaotic. This type of chaotic system is the basis for a number of office novelty desk toys. It is simple to build a version. Take a piece of string. Tie two heavy nuts at different points of the string; one at the end, and one a bit above. Now give the bottom nut a little push. As long as the swings are not too big and the string stays tight, you have a double pendulum system.



*Edward Norton Lorenz (1917–2008) was an American mathematician and meteorologist.

8.5.1 Duffing equation and strange attractors

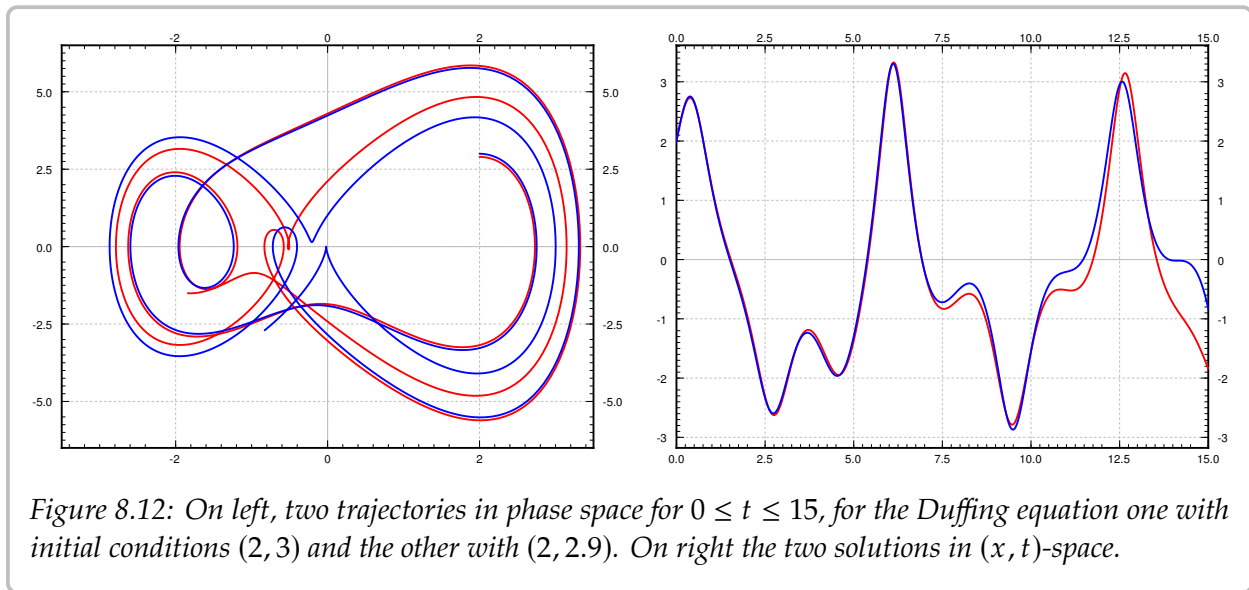
Let us study the so-called *Duffing equation*:

$$x'' + ax' + bx + cx^3 = C \cos(\omega t).$$

Here a, b, c, C , and ω are constants. Except for the cx^3 term, this equation looks like a forced mass-spring system. The cx^3 means the spring does not exactly obey Hooke's law (which no real-world spring actually obeys exactly). When c is not zero, the equation does not have a closed-form solution, so we must resort to numerical solutions, as is usual for nonlinear systems. Not all choices of constants and initial conditions exhibit chaotic behavior. Let us study

$$x'' + 0.05x' + x^3 = 8 \cos(t).$$

The equation is not autonomous, so we cannot draw the vector field in the phase plane. We can still draw trajectories, where $y = x'$ as usual. In Figure 8.12, we plot trajectories for t going from 0 to 15 for two very close initial conditions $(2, 3)$ and $(2, 2.9)$, and also the solutions in the (x, t) space. The two trajectories are close at first, but soon diverge significantly. This sensitivity to initial conditions is precisely what we mean by the system behaving chaotically.



Let us consider the long-term behavior. In Figure 8.13 on the following page, we plot the solution x with initial conditions $(2, 3)$ for a longer period of time. It is hard to see any pattern in the shape of the solution except that it seems to oscillate, but each oscillation appears quite unique. Oscillation is expected due to the forcing term. We mention that to produce the picture accurately, a ridiculously large number of steps* were used in the numerical algorithm, as even small errors quickly propagate in a chaotic system.

*In fact for reference, 30,000 steps were used with the Runge–Kutta algorithm, see exercises in § 1.7.

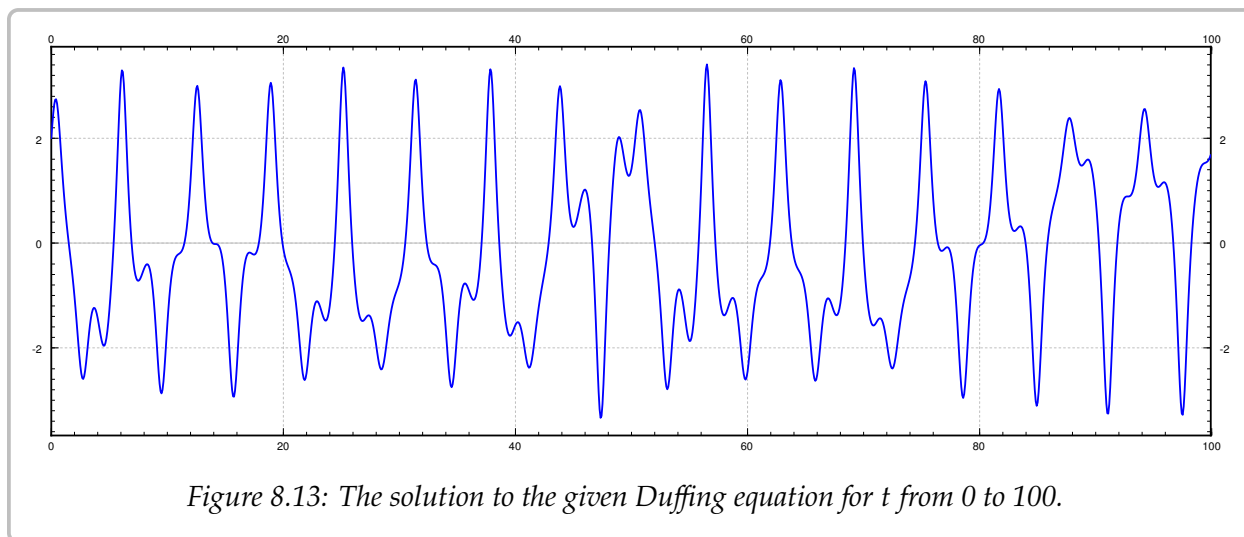


Figure 8.13: The solution to the given Duffing equation for t from 0 to 100.

It is very difficult to analyze chaotic systems, or to find the order behind the madness, but let us try to do something that we did for the standard mass-spring system. One way we analyzed the system is that we figured out what the long-term behavior was (not dependent on initial conditions). From the figure above, it is clear that we will not get a nice exact description of the long-term behavior for this chaotic system, but perhaps we can find some order to what happens on each “oscillation” and what these oscillations have in common.

The concept we explore is that of a *Poincaré section**. Instead of looking at t in a certain interval, we look at where the system is at a certain sequence of points in time. Imagine flashing a strobe at a fixed frequency and drawing the points where the solution is during the flashes. The right strobing frequency depends on the system in question. The correct frequency for the forced Duffing equation (and other similar systems) is the frequency of the forcing term. For the Duffing equation above, find a solution $(x(t), y(t))$, and look at the points

$$(x(0), y(0)), \quad (x(2\pi), y(2\pi)), \quad (x(4\pi), y(4\pi)), \quad (x(6\pi), y(6\pi)), \quad \dots$$

As we are really not interested in the transient part of the solution, that is, the part of the solution that depends on the initial condition, we skip some number of steps in the beginning. For example, we might skip the first 100 such steps and start plotting points at $t = 100(2\pi)$, that is,

$$(x(200\pi), y(200\pi)), \quad (x(202\pi), y(202\pi)), \quad (x(204\pi), y(204\pi)), \quad \dots$$

The plot of these points is the Poincaré section. After plotting enough points, a curious pattern emerges in Figure 8.14 on the next page (the left-hand picture), a so-called *strange attractor*.

*Named for the French polymath [Jules Henri Poincaré](#) (1854–1912).

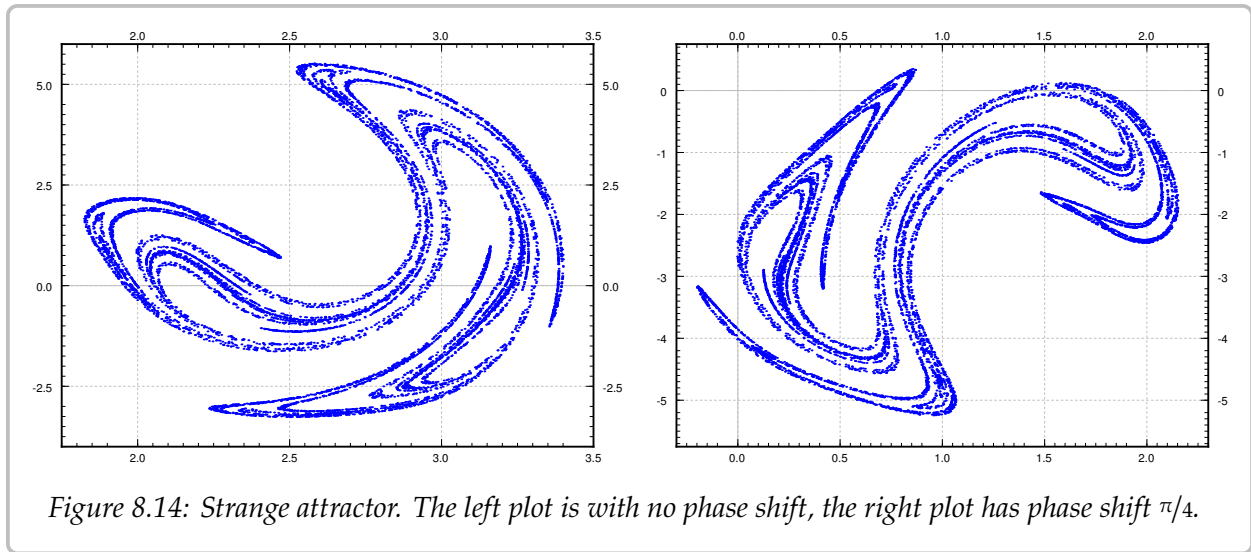


Figure 8.14: Strange attractor. The left plot is with no phase shift, the right plot has phase shift $\pi/4$.

Given a sequence of points, an *attractor* is a set towards which the points in the sequence eventually get closer and closer, that is, they are attracted. The Poincaré section is not really the attractor itself, but as the points are very close to it, we see its shape. The strange attractor is a very complicated set. It has fractal structure, that is, if you zoom in as far as you want, you keep seeing the same complicated structure.

The initial condition makes no difference. If we start with a different initial condition, the points eventually gravitate towards the attractor, and so as long as we throw away the first few points, we get the same picture. Similarly, small errors in the numerical approximations do not matter here.

An amazing thing is that a chaotic system such as the Duffing equation is not random at all. There is a very complicated order to it, and the strange attractor says something about this order. We cannot quite say what state the system will be in eventually, but given the fixed strobing frequency we narrow it down to the points on the attractor.

If we use a phase shift, for example $\pi/4$, and look at the times

$$\pi/4, \quad 2\pi + \pi/4, \quad 4\pi + \pi/4, \quad 6\pi + \pi/4, \quad \dots$$

we obtain a slightly different attractor. The picture is the right-hand side of Figure 8.14. It is as if we had rotated, moved, and slightly distorted the original. For each phase shift you can find the set of points towards which the system periodically keeps coming back to.

Study the pictures and notice especially the scales—where are these attractors located in the phase plane. Notice the regions where the strange attractor lives and compare it to the plot of the trajectories in Figure 8.12 on page 379.

Let us compare this section to the discussion in § 2.6 on forced oscillations. Consider

$$x'' + 2px' + \omega_0^2 x = \frac{F_0}{m} \cos(\omega t).$$

This is like the Duffing equation, but with no x^3 term. The steady periodic solution is of

the form

$$x = C \cos(\omega t + \gamma).$$

Strobing using the frequency ω , we obtain a single point in the phase space. The attractor in this setting is a single point—an expected result as the system is not chaotic. It was the opposite of chaotic: Any difference induced by the initial conditions dies away very quickly, and we settle into always the same steady periodic motion.

8.5.2 The Lorenz system

In two dimensions to find chaotic behavior, we must study forced, or non-autonomous, systems such as the Duffing equation. The Poincaré–Bendixson Theorem says that a solution to an autonomous two-dimensional system that exists for all time in the future either goes towards infinity, is periodic, approaches a loop, or perhaps goes towards a critical point. Hardly the chaotic behavior we are looking for.

In three dimensions, even autonomous systems can be chaotic. Let us very briefly return to the Lorenz system

$$x' = -10x + 10y, \quad y' = 28x - y - xz, \quad z' = -\frac{8}{3}z + xy.$$

The Lorenz system is an autonomous system in three dimensions exhibiting chaotic behavior. See [Figure 8.15](#) for a sample trajectory, which is now a curve in three-dimensional space.

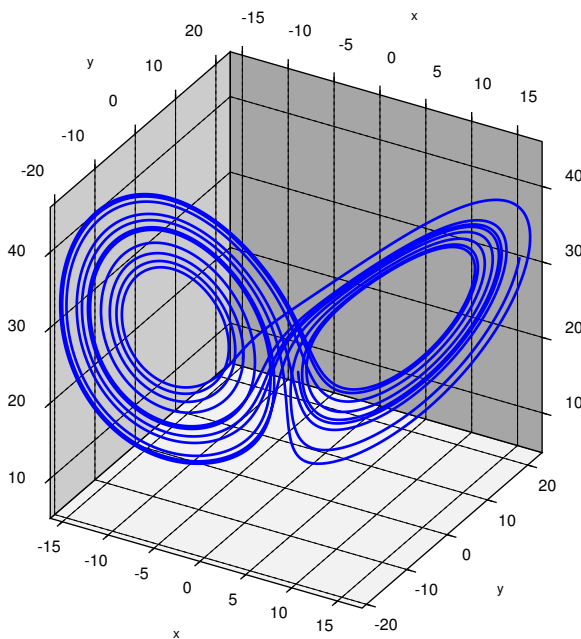


Figure 8.15: A trajectory in the Lorenz system.

The solutions tend to an *attractor* in space, the so-called *Lorenz attractor*. In this case, no strobing is necessary, the solution will tend towards the attractor set. Again we cannot quite see the attractor itself, but if we try to follow a solution for long enough, as in the figure, we get a pretty good picture of what the attractor looks like. The Lorenz attractor is also a strange attractor and has a complicated fractal structure. And, just as for the Duffing equation, what we want to draw is not the whole trajectory, but start drawing the trajectory after a while, once it is close to the attractor.

The path of the trajectory is not simply a repeating figure-eight. The trajectory spins some seemingly random number of times on the left, then spins a number of times on the right, and so on. As this system arose in weather prediction, one can perhaps imagine a few days of warm weather and then a few days of cold weather, where it is not easy to predict when the weather will change, just as it is not really easy to predict far in advance when the solution will jump onto the other side. See Figure 8.16 for a plot of the x component of the solution drawn above. A negative x corresponds to the left “loop” and a positive x corresponds to the right “loop”. On the other hand, while we cannot predict the weather, we *can* say something about the climate—the weather will be somewhere near the attractor.

Most of the mathematics we studied in this book is quite classical and well understood. On the other hand, chaos, including the Lorenz system, continues to be the subject of current research. Furthermore, chaos has found applications not just in the sciences, but also in art.

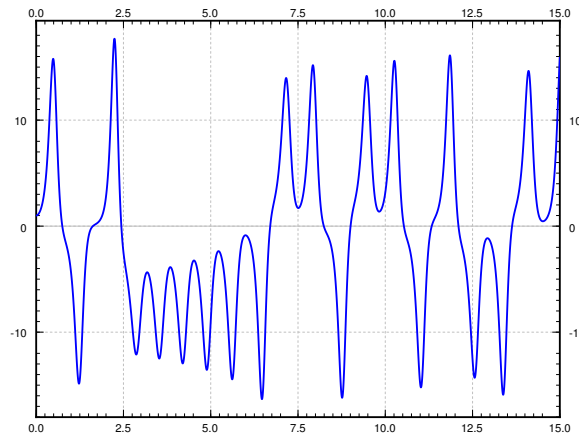


Figure 8.16: Graph of the $x(t)$ component of the solution.

8.5.3 Exercises

Exercise 8.5.1: For the non-chaotic equation $x'' + 2px' + \omega_0^2 x = \frac{F_0}{m} \cos(\omega t)$, suppose we strobe with frequency ω as we mentioned above. Use the known steady periodic solution to find precisely the point which is the attractor for the Poincaré section.

Exercise 8.5.2 (project): A simple fractal attractor can be drawn via the following chaos game. Draw the three vertices of a triangle and label them, say p_1 , p_2 and p_3 . Draw some random point p (it does not have to be one of the three points above). Roll a die to pick one of the p_1 , p_2 , or p_3 randomly (for example 1 and 4 mean p_1 , 2 and 5 mean p_2 , and 3 and 6 mean p_3). Suppose we picked p_2 . Let p_{new} be the point exactly halfway between p and p_2 . Draw this point and let p now refer to this new point p_{new} . Rinse, repeat. Try to be precise and draw as many iterations as possible. Your points will be attracted to the so-called Sierpinski triangle. A computer was used to run the game for 10,000 iterations to obtain the picture in [Figure 8.17](#).

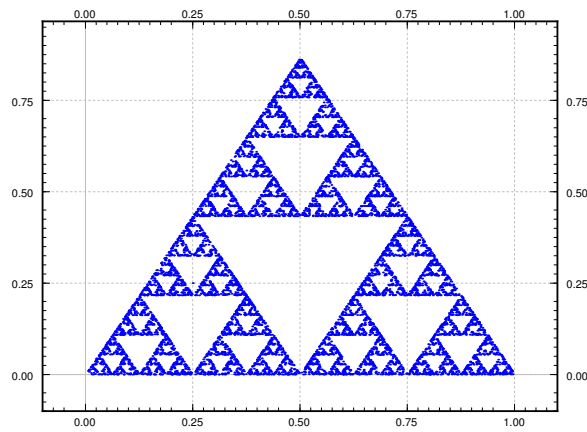


Figure 8.17: 10,000 iterations of the chaos game producing the Sierpinski triangle.

Exercise 8.5.3 (project): Construct the double pendulum described in the text with a string and two nuts (or heavy beads). Play around with the position of the middle nut, and perhaps use different weight nuts. Describe what you find.

Exercise 8.5.4 (computer project): Use computer software (such as Matlab, Octave, or perhaps even a spreadsheet) to plot the solution of the given forced Duffing equation with Euler's method. Plot the solution for t from 0 to 100 with several different (small) step sizes. Discuss.

Exercise 8.5.101: Find critical points of the Lorenz system and the associated linearizations.

Appendix A

Linear algebra

A.1 Vectors, mappings, and matrices

Note: 2 lectures

In real life, there is most often more than one variable. We wish to organize dealing with multiple variables in a consistent manner, and in particular organize dealing with linear equations and linear mappings, as those are both rather useful and rather easy to handle. Mathematicians joke that “to an engineer every problem is linear, and everything is a matrix.” And well, they (the engineers) are not wrong. Quite often, solving an engineering problem is figuring out the right finite-dimensional linear problem to solve, which is then solved with some matrix manipulation. Most importantly, linear problems are the ones that we know how to solve, and we have many tools to solve them. For engineers, mathematicians, physicists, and anybody else in a technical field, it is absolutely vital to learn linear algebra.

As motivation, suppose we wish to solve

$$\begin{aligned}x - y &= 2, \\ 2x + y &= 4,\end{aligned}$$

for x and y . That is, we desire numbers x and y such that the two equations are satisfied. Let us perhaps start by adding the equations together to find

$$x + 2x - y + y = 2 + 4, \quad \text{or} \quad 3x = 6.$$

In other words, $x = 2$. Once we have that, we plug $x = 2$ into the first equation to find $2 - y = 2$, so $y = 0$. OK, that was easy. What is all this fuss about linear equations? Well, try doing this if you have 5000 unknowns[†]. Also, we may have such equations not just of numbers, but of functions and derivatives of functions in differential equations. Clearly we need a systematic way of doing things. A nice consequence of making things systematic

[†]One of the downsides of making everything look like a linear problem is that the number of variables tends to become huge.

and simpler to write down is that it becomes easier to have computers do the work for us. Computers are very good at doing lots of repetitive tasks precisely, as long as we figure out a systematic way for them to perform the tasks.

A.1.1 Vectors and operations on vectors

Consider n real numbers as an n -tuple:

$$(x_1, x_2, \dots, x_n).$$

The set of such n -tuples is the so-called n -dimensional space, often denoted by $\mathbb{R}^n$. Sometimes we call this the n -dimensional *euclidean space*^{*}. In two dimensions, $\mathbb{R}^2$ is called the *cartesian plane*[†]. Each such n -tuple represents a point in the n -dimensional space. For example, the point $(1, 2)$ in the plane $\mathbb{R}^2$ is one unit to the right and two units up from the origin.

When we do algebra with these n -tuples of numbers, we call them *vectors*[‡]. Mathematicians are keen on separating what is a vector and what is a point of the space or in the plane, and it turns out to be an important distinction; however, for the purposes of linear algebra we can think of everything being represented by a vector. A way to think of a vector, which is especially useful in calculus and differential equations, is an arrow. It is an object that has a *direction* and a *magnitude*. For instance, the vector $(1, 2)$ is the arrow from the origin to the point $(1, 2)$ in the plane. The magnitude is the length of the arrow. See Figure A.1. If we think of vectors as arrows, the arrow does not always have to start at the origin. If we do move it around, however, it should always keep the same direction and the same magnitude.

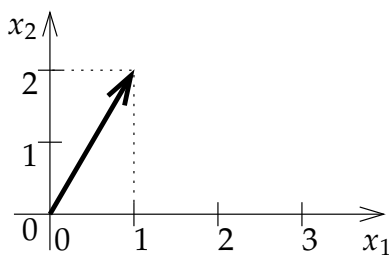


Figure A.1: The vector $(1, 2)$ drawn as an arrow from the origin to the point $(1, 2)$.

As vectors are arrows, when we want to give a name to a vector, we draw a little arrow above it:

$$\vec{x}$$

^{*}Named after the ancient Greek mathematician [Euclid of Alexandria](#) (around 300 BC), possibly the most famous of mathematicians; even small towns often have Euclid Street or Euclid Avenue.

[†]Named after the French mathematician [René Descartes](#) (1596–1650). It is “cartesian” as his name in Latin is Renatus Cartesius.

[‡]A common notation to distinguish vectors from points is to write $(1, 2)$ for the point and $\langle 1, 2 \rangle$ for the vector. We write both as $(1, 2)$.

Another popular notation is a bold $\mathbf{x}$, although we will use the little arrows. It may be easy to write a bold letter in a book, but it is not so easy to write it by hand on paper or on the board. Mathematicians often do not even write the arrows. A mathematician would write x and remember that x is a vector and not a number. Just like you remember that Bob is your uncle and you don't have to keep repeating "Uncle Bob". You can just say "Bob." In this book, however, we will call Bob "Uncle Bob" and write vectors with the little arrows.

The *magnitude* can be computed using the Pythagorean theorem. The vector $(1, 2)$ drawn in the figure has magnitude $\sqrt{1^2 + 2^2} = \sqrt{5}$. The magnitude is denoted by $\|\vec{x}\|$, and, in any number of dimensions, it can be computed in the same way:

$$\|\vec{x}\| = \|(x_1, x_2, \dots, x_n)\| = \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}.$$

For reasons that will become clear in the next section, we often write vectors as so-called *column vectors*:

$$\vec{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}.$$

Don't worry. It is just a different way of writing the same thing. For example, the vector $(1, 2)$ can be written as

$$\begin{bmatrix} 1 \\ 2 \end{bmatrix}.$$

The fact that we write arrows above vectors allows us to write several vectors $\vec{x}_1$, $\vec{x}_2$, etc., without confusing these with the components of some other vector $\vec{x}$.

So where is the *algebra* from *linear algebra*? Well, arrows can be added, subtracted, and multiplied by numbers. First we consider *addition*. If we have two arrows, we simply move along one, and then along the other. See [Figure A.2](#).

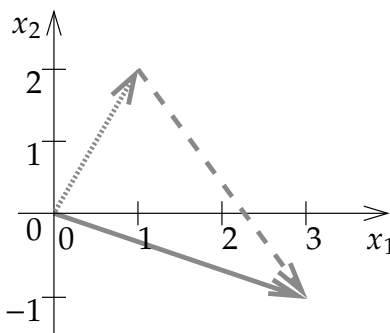


Figure A.2: Adding the vectors $(1, 2)$, drawn dotted, and $(2, -3)$, drawn dashed. The result, $(3, -1)$, is drawn as a solid arrow.

It is rather easy to see what it does to the numbers that represent the vectors. Suppose we want to add $(1, 2)$ to $(2, -3)$ as in the figure. We travel along $(1, 2)$ and then we travel along $(2, -3)$. What we did was travel one unit right, two units up, and then we traveled two units right, and three units down (the negative three). That means that we ended up at $(1 + 2, 2 + (-3)) = (3, -1)$. That is how addition always works:

$$\begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} + \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} = \begin{bmatrix} x_1 + y_1 \\ x_2 + y_2 \\ \vdots \\ x_n + y_n \end{bmatrix}.$$

Subtracting is similar. What $\vec{x} - \vec{y}$ means visually is that we first travel along $\vec{x}$, and then we travel backwards along $\vec{y}$. See Figure A.3. It is like adding $\vec{x} + (-\vec{y})$ where $-\vec{y}$ is the arrow we obtain by erasing the arrow head from one side and drawing it on the other side. That is, we reverse the direction. In terms of the numbers, we simply go backwards both horizontally and vertically, so we negate both numbers. For instance, if $\vec{y}$ is $(-2, 1)$, then $-\vec{y}$ is $(2, -1)$.

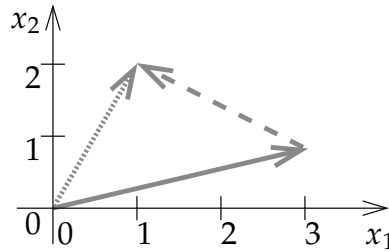


Figure A.3: Subtraction, the vector $(1, 2)$, drawn dotted, minus $(-2, 1)$, drawn dashed. The result, $(3, 1)$, is drawn as a solid arrow.

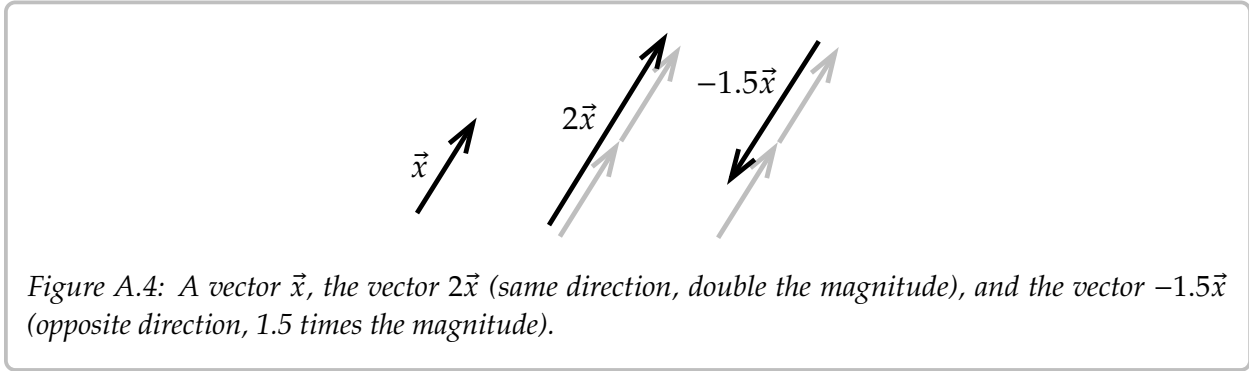
Another intuitive thing to do to a vector is to *scale* it. We represent this by multiplication of a number with a vector. Because of this, when we wish to distinguish between vectors and numbers, we call the numbers *scalars*. For example, suppose we want to travel three times further. If the vector is $(1, 2)$, travelling 3 times further means going 3 units to the right and 6 units up, so we get the vector $(3, 6)$. We just multiply each number in the vector by 3. If α is a number, then

$$\alpha \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} \alpha x_1 \\ \alpha x_2 \\ \vdots \\ \alpha x_n \end{bmatrix}.$$

Scaling (by a positive number) multiplies the magnitude and leaves direction untouched. The magnitude of $(1, 2)$ is $\sqrt{5}$. The magnitude of 3 times $(1, 2)$, that is, of $(3, 6)$, is $3\sqrt{5}$.

If we multiply a vector by a negative scalar, the vector is not only scaled, but it also switches direction. Multiplying $(1, 2)$ by -3 means we should go 3 times further but in the opposite direction, so 3 units to the left and 6 units down, or in other words, $(-3, -6)$. As we mentioned above, $-\vec{y}$ is the reverse of $\vec{y}$, and this is the same as $(-1)\vec{y}$.

In Figure A.4, you can see a couple of examples of what scaling a vector means visually.



We put all of these operations together to work out more complicated expressions. Let us compute a small example:

$$3 \begin{bmatrix} 1 \\ 2 \end{bmatrix} + 2 \begin{bmatrix} -4 \\ -1 \end{bmatrix} - 3 \begin{bmatrix} -2 \\ 2 \end{bmatrix} = \begin{bmatrix} 3(1) + 2(-4) - 3(-2) \\ 3(2) + 2(-1) - 3(2) \end{bmatrix} = \begin{bmatrix} 1 \\ -2 \end{bmatrix}.$$

We said a vector is a direction and a magnitude. Magnitude is easy to represent—it is just a number. The *direction* is usually given by a vector with magnitude one. We call such a vector a *unit vector*. That is, $\vec{u}$ is a unit vector when $\|\vec{u}\| = 1$. For instance, the vectors $(1, 0)$, $(1/\sqrt{2}, 1/\sqrt{2})$, and $(0, -1)$ are all unit vectors.

To represent the direction of a vector $\vec{x}$, we need to find the unit vector in the same direction. To do so, we simply rescale $\vec{x}$ by the reciprocal of the magnitude: $\frac{1}{\|\vec{x}\|}\vec{x}$, or more concisely, $\frac{\vec{x}}{\|\vec{x}\|}$.

As an example, the unit vector in the direction of $(1, 2)$ is the vector

$$\frac{1}{\sqrt{1^2 + 2^2}}(1, 2) = \left(\frac{1}{\sqrt{5}}, \frac{2}{\sqrt{5}} \right).$$

A.1.2 Linear mappings and matrices

A *vector-valued function* F is a rule that takes a vector $\vec{x}$ and returns another vector $\vec{y}$. For example, F could be a scaling that doubles the size of vectors:

$$F(\vec{x}) = 2\vec{x}.$$

Applied to say $(1, 3)$ we get

$$F\left(\begin{bmatrix} 1 \\ 3 \end{bmatrix}\right) = 2 \begin{bmatrix} 1 \\ 3 \end{bmatrix} = \begin{bmatrix} 2 \\ 6 \end{bmatrix}.$$

If F is a mapping that takes vectors in $\mathbb{R}^2$ to $\mathbb{R}^2$ (such as the above), we write

$$F: \mathbb{R}^2 \rightarrow \mathbb{R}^2.$$

The words *function* and *mapping* are used rather interchangeably, although more often than not, *mapping* is used when talking about a vector-valued function, and the word *function* is often used when the function is scalar-valued.

A beginning student of mathematics (and many a seasoned mathematician) who sees an expression such as

$$f(3x + 8y)$$

yearns to write

$$3f(x) + 8f(y).$$

After all, who has not wanted to write $\sqrt{x+y} = \sqrt{x} + \sqrt{y}$ or something like that at some point in their mathematical lives? Wouldn't life be simple if we could do that? Of course we cannot always do that (for example, not with the square roots!). But there are many other functions where we can do exactly the above. Such functions are called *linear*.

A mapping $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$ is called *linear* if

$$F(\vec{x} + \vec{y}) = F(\vec{x}) + F(\vec{y}),$$

for any vectors $\vec{x}$ and $\vec{y}$, and also

$$F(\alpha\vec{x}) = \alpha F(\vec{x}),$$

for any scalar α . The F we defined above that doubles the size of all vectors is linear. Let us check:

$$F(\vec{x} + \vec{y}) = 2(\vec{x} + \vec{y}) = 2\vec{x} + 2\vec{y} = F(\vec{x}) + F(\vec{y}),$$

and also

$$F(\alpha\vec{x}) = 2\alpha\vec{x} = \alpha 2\vec{x} = \alpha F(\vec{x}).$$

We also call a linear function a *linear transformation*. If you want to be really fancy and impress your friends, you can call it a *linear operator*. When a mapping is linear, we often do not write the parentheses. We write simply

$$F\vec{x}$$

instead of $F(\vec{x})$. We do this because linearity means that the mapping F behaves like multiplying $\vec{x}$ by "something." That something is a matrix.

A *matrix* is an $m \times n$ array of numbers (m rows and n columns). A 3×5 matrix is

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \end{bmatrix}.$$

The numbers a_{ij} are called *elements* or *entries*.

A column vector is simply an $m \times 1$ matrix. Similarly to a column vector, there is also a *row vector*, which is a $1 \times n$ matrix. If we have an $n \times n$ matrix, then we say that it is a *square matrix*.

How does a matrix A relate to a linear mapping? A matrix tells you where certain special vectors go. We give a name to those certain vectors. The *standard basis vectors* of $\mathbb{R}^n$ are

$$\vec{e}_1 = \begin{bmatrix} 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad \vec{e}_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad \vec{e}_3 = \begin{bmatrix} 0 \\ 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad \vec{e}_n = \begin{bmatrix} 0 \\ 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix}.$$

In $\mathbb{R}^3$, these vectors are

$$\vec{e}_1 = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \quad \vec{e}_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \quad \vec{e}_3 = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

You may recall from calculus of several variables that these are sometimes called $\vec{i}, \vec{j}, \vec{k}$. The reason these are called a *basis* is that every vector can be written as a unique *linear combination* of them. For example, in $\mathbb{R}^3$ the vector $(4, 5, 6)$ can be written as

$$4\vec{e}_1 + 5\vec{e}_2 + 6\vec{e}_3 = 4 \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + 5 \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + 6 \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 4 \\ 5 \\ 6 \end{bmatrix}.$$

So how does a matrix represent a linear mapping? Well, the columns of the matrix are the vectors where the matrix, as a linear mapping, takes $\vec{e}_1, \vec{e}_2$, etc. For instance, consider

$$M = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}.$$

As a linear mapping, $M: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ takes $\vec{e}_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ to $\begin{bmatrix} 1 \\ 3 \end{bmatrix}$ and $\vec{e}_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ to $\begin{bmatrix} 2 \\ 4 \end{bmatrix}$. In other words,

$$M\vec{e}_1 = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 3 \end{bmatrix}, \quad \text{and} \quad M\vec{e}_2 = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 2 \\ 4 \end{bmatrix}.$$

More generally, if we have an $n \times m$ matrix A , that is, we have n rows and m columns, then the mapping $A: \mathbb{R}^m \rightarrow \mathbb{R}^n$ takes $\vec{e}_j$ to the j^{th} column of A . For example,

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \end{bmatrix}$$

represents a mapping from $\mathbb{R}^5$ to $\mathbb{R}^3$ that does

$$A\vec{e}_1 = \begin{bmatrix} a_{11} \\ a_{21} \\ a_{31} \end{bmatrix}, \quad A\vec{e}_2 = \begin{bmatrix} a_{12} \\ a_{22} \\ a_{32} \end{bmatrix}, \quad A\vec{e}_3 = \begin{bmatrix} a_{13} \\ a_{23} \\ a_{33} \end{bmatrix}, \quad A\vec{e}_4 = \begin{bmatrix} a_{14} \\ a_{24} \\ a_{34} \end{bmatrix}, \quad A\vec{e}_5 = \begin{bmatrix} a_{15} \\ a_{25} \\ a_{35} \end{bmatrix}.$$

What about another vector $\vec{x}$ that is not in the standard basis? Where does it go? We use linearity. First, we write the vector as a linear combination of the standard basis vectors:

$$\vec{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{bmatrix} = x_1 \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} + x_2 \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \\ 0 \end{bmatrix} + x_3 \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \\ 0 \end{bmatrix} + x_4 \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix} + x_5 \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \end{bmatrix} = x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3 + x_4 \vec{e}_4 + x_5 \vec{e}_5.$$

Then

$$A\vec{x} = A(x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3 + x_4 \vec{e}_4 + x_5 \vec{e}_5) = x_1 A\vec{e}_1 + x_2 A\vec{e}_2 + x_3 A\vec{e}_3 + x_4 A\vec{e}_4 + x_5 A\vec{e}_5.$$

If we know where A takes all the basis vectors, we know where it takes all vectors.

Suppose M is the 2×2 matrix from above, then

$$M \begin{bmatrix} -2 \\ 0.1 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} -2 \\ 0.1 \end{bmatrix} = -2 \begin{bmatrix} 1 \\ 3 \end{bmatrix} + 0.1 \begin{bmatrix} 2 \\ 4 \end{bmatrix} = \begin{bmatrix} -1.8 \\ -5.6 \end{bmatrix}.$$

Every linear mapping from $\mathbb{R}^m$ to $\mathbb{R}^n$ can be represented by an $n \times m$ matrix. You just figure out where it takes the standard basis vectors. Conversely, every $n \times m$ matrix represents a linear mapping. Hence, we may think of matrices being linear mappings, and linear mappings being matrices.

Or can we? In this book we study mostly linear differential operators, and linear differential operators are linear mappings, although they are not acting on $\mathbb{R}^n$, but on an infinite-dimensional space of functions:

$$Lf = g.$$

For a function f we get a function g , and L is linear in the sense that

$$L(f + h) = Lf + Lh, \quad \text{and} \quad L(\alpha f) = \alpha Lf.$$

for any number (scalar) α and all functions f and h .

So the answer is not really. But if we consider vectors in finite-dimensional spaces $\mathbb{R}^n$, then yes, every linear mapping is a matrix. We have mentioned at the beginning of this section that we can “make everything a vector.” That’s not strictly true, but it is true approximately. Those “infinite-dimensional” spaces of functions can be approximated by a finite-dimensional space, and then linear operators are just matrices. So approximately, this is true. And as far as actual computations that we can do on a computer, we can work only with finitely many dimensions anyway. If you ask a computer or your calculator to plot a function, it samples the function at finitely many points and then connects the dots*. It does not actually give you infinitely many values. The way that you have been using the

*If you have ever used Matlab, you may have noticed that to plot a function, we take a vector of inputs, ask Matlab to compute the corresponding vector of values of the function, and then we ask it to plot the result.

computer or your calculator so far has already been a certain approximation of the space of functions by a finite-dimensional space.

To end the section, we notice how $A\vec{x}$ can be written more succinctly. Suppose

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \quad \text{and} \quad \vec{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}.$$

Then

$$A\vec{x} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 \end{bmatrix}.$$

For example,

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 2 \\ -1 \end{bmatrix} = \begin{bmatrix} 1 \cdot 2 + 2 \cdot (-1) \\ 3 \cdot 2 + 4 \cdot (-1) \end{bmatrix} = \begin{bmatrix} 0 \\ 2 \end{bmatrix}.$$

That is, you take the entries in a row of the matrix, you multiply them by the entries in your vector, you add things up, and that's the corresponding entry in the resulting vector.

A.1.3 Exercises

Exercise A.1.1: On a piece of graph paper draw the vectors:

$$a) \begin{bmatrix} 2 \\ 5 \end{bmatrix} \quad b) \begin{bmatrix} -2 \\ -4 \end{bmatrix} \quad c) (3, -4)$$

Exercise A.1.2: On a piece of graph paper draw the vector $(1, 2)$ starting at (based at) the given point:

$$a) \text{ based at } (0, 0) \quad b) \text{ based at } (1, 2) \quad c) \text{ based at } (0, -1)$$

Exercise A.1.3: On a piece of graph paper draw the following operations. Draw and label the vectors involved in the operations as well as the result:

$$a) \begin{bmatrix} 1 \\ -4 \end{bmatrix} + \begin{bmatrix} 2 \\ 3 \end{bmatrix} \quad b) \begin{bmatrix} -3 \\ 2 \end{bmatrix} - \begin{bmatrix} 1 \\ 3 \end{bmatrix} \quad c) 3 \begin{bmatrix} 2 \\ 1 \end{bmatrix}$$

Exercise A.1.4: Compute the magnitude of

$$a) \begin{bmatrix} 7 \\ 2 \end{bmatrix} \quad b) \begin{bmatrix} -2 \\ 3 \\ 1 \end{bmatrix} \quad c) (1, 3, -4)$$

Exercise A.1.5: Compute

$$a) \begin{bmatrix} 2 \\ 3 \end{bmatrix} + \begin{bmatrix} 7 \\ -8 \end{bmatrix}$$

$$b) \begin{bmatrix} -2 \\ 3 \end{bmatrix} - \begin{bmatrix} 6 \\ -4 \end{bmatrix}$$

$$c) -\begin{bmatrix} -3 \\ 2 \end{bmatrix}$$

$$d) 4 \begin{bmatrix} -1 \\ 5 \end{bmatrix}$$

$$e) 5 \begin{bmatrix} 1 \\ 0 \end{bmatrix} + 9 \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

$$f) 3 \begin{bmatrix} 1 \\ -8 \end{bmatrix} - 2 \begin{bmatrix} 3 \\ -1 \end{bmatrix}$$

Exercise A.1.6: Find the unit vector in the direction of the given vector

$$a) \begin{bmatrix} 1 \\ -3 \end{bmatrix}$$

$$b) \begin{bmatrix} 2 \\ 1 \\ -1 \end{bmatrix}$$

$$c) (3, 1, -2)$$

Exercise A.1.7: If $\vec{x} = (1, 2)$ and $\vec{y}$ are added together, we find $\vec{x} + \vec{y} = (0, 2)$. What is $\vec{y}$?

Exercise A.1.8: Write $(1, 2, 3)$ as a linear combination of the standard basis vectors $\vec{e}_1$, $\vec{e}_2$, and $\vec{e}_3$.

Exercise A.1.9: If the magnitude of $\vec{x}$ is 4, what is the magnitude of

$$a) 0\vec{x}$$

$$b) 3\vec{x}$$

$$c) -\vec{x}$$

$$d) -4\vec{x}$$

$$e) \vec{x} + \vec{x}$$

$$f) \vec{x} - \vec{x}$$

Exercise A.1.10: Suppose a linear mapping $F: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ takes $(1, 0)$ to $(2, -1)$ and it takes $(0, 1)$ to $(3, 3)$. Where does it take

$$a) (1, 1)$$

$$b) (2, 0)$$

$$c) (2, -1)$$

Exercise A.1.11: Suppose a linear mapping $F: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ takes $(1, 0, 0)$ to $(2, 1)$, it takes $(0, 1, 0)$ to $(3, 4)$, and it takes $(0, 0, 1)$ to $(5, 6)$. Write down the matrix representing the mapping F .

Exercise A.1.12: Suppose that a mapping $F: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ takes $(1, 0)$ to $(1, 2)$, $(0, 1)$ to $(3, 4)$, and $(1, 1)$ to $(0, -1)$. Explain why F is not linear.

Exercise A.1.13 (challenging): Let $\mathbb{R}^3$ represent the space of quadratic polynomials in t : A point (a_0, a_1, a_2) in $\mathbb{R}^3$ represents the polynomial $a_0 + a_1t + a_2t^2$. Consider the derivative $\frac{d}{dt}$ as a mapping of $\mathbb{R}^3$ to $\mathbb{R}^3$, and note that $\frac{d}{dt}$ is linear. Write down $\frac{d}{dt}$ as a 3×3 matrix.

Exercise A.1.101: Compute the magnitude of

$$a) \begin{bmatrix} 1 \\ 3 \end{bmatrix}$$

$$b) \begin{bmatrix} 2 \\ 3 \\ -1 \end{bmatrix}$$

$$c) (-2, 1, -2)$$

Exercise A.1.102: Find the unit vector in the direction of the given vector

$$a) \begin{bmatrix} -1 \\ 1 \end{bmatrix}$$

$$b) \begin{bmatrix} 1 \\ -1 \\ 2 \end{bmatrix}$$

$$c) (2, -5, 2)$$

Exercise A.1.103: Compute

a) $\begin{bmatrix} 3 \\ 1 \end{bmatrix} + \begin{bmatrix} 6 \\ -3 \end{bmatrix}$

b) $\begin{bmatrix} -1 \\ 2 \end{bmatrix} - \begin{bmatrix} 2 \\ -1 \end{bmatrix}$

c) $-\begin{bmatrix} -5 \\ 3 \end{bmatrix}$

d) $2 \begin{bmatrix} -2 \\ 4 \end{bmatrix}$

e) $3 \begin{bmatrix} 1 \\ 0 \end{bmatrix} + 7 \begin{bmatrix} 0 \\ 1 \end{bmatrix}$

f) $2 \begin{bmatrix} 2 \\ -3 \end{bmatrix} - 6 \begin{bmatrix} 2 \\ -1 \end{bmatrix}$

Exercise A.1.104: If the magnitude of $\vec{x}$ is 5, what is the magnitude of

a) $4\vec{x}$

b) $-2\vec{x}$

c) $-4\vec{x}$

Exercise A.1.105: Suppose a linear mapping $F: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ takes $(1, 0)$ to $(1, -1)$ and it takes $(0, 1)$ to $(2, 0)$. Where does it take

a) $(1, 1)$

b) $(0, 2)$

c) $(1, -1)$

A.2 Matrix algebra

Note: 2–3 lectures

A.2.1 One-by-one matrices

Let us motivate what we want to achieve with matrices. Real-valued linear mappings of the real line, linear functions that eat numbers and spit out numbers, are just multiplications by a number. Consider a mapping defined by multiplying by a number. Let's call this number α . The mapping then takes x to αx . We can *add* such mappings: If we have another mapping β , then

$$\alpha x + \beta x = (\alpha + \beta)x.$$

We get a new mapping $\alpha + \beta$ that multiplies x by, well, $\alpha + \beta$. If D is a mapping that doubles its input, $Dx = 2x$, and T is a mapping that triples, $Tx = 3x$, then $D + T$ is a mapping that multiplies by 5, $(D + T)x = 5x$.

Similarly we can *compose* such mappings. That is, we could apply one and then the other. We take x , we run it through the first mapping α to get α times x , then we run αx through the second mapping β . In other words,

$$\beta(\alpha x) = (\beta\alpha)x.$$

We just multiply those two numbers. Using our doubling and tripling mappings, if we double and then triple, that is, $T(Dx)$, then we obtain $3(2x) = 6x$. The composition TD is the mapping that multiplies by 6. For larger matrices, composition also ends up being a kind of multiplication.

A.2.2 Matrix addition and scalar multiplication

The mappings that multiply numbers by numbers are just 1×1 matrices. The number α above could be written as a matrix $[\alpha]$. Perhaps we would want to do to all matrices the same things that we did to those 1×1 matrices at the start of this section above. First, let us add matrices. If we have a matrix A and a matrix B that are of the same size, say $m \times n$, then they are mappings from $\mathbb{R}^n$ to $\mathbb{R}^m$. The mapping $A + B$ should also be a mapping from $\mathbb{R}^n$ to $\mathbb{R}^m$, and it should do the following to vectors:

$$(A + B)\vec{x} = A\vec{x} + B\vec{x}.$$

It turns out you just add the matrices element-wise: If the ij^{th} entry of A is a_{ij} , and the ij^{th} entry of B is b_{ij} , then the ij^{th} entry of $A + B$ is $a_{ij} + b_{ij}$. If

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \quad \text{and} \quad B = \begin{bmatrix} b_{11} & b_{12} & b_{13} \\ b_{21} & b_{22} & b_{23} \end{bmatrix},$$

then

$$A + B = \begin{bmatrix} a_{11} + b_{11} & a_{12} + b_{12} & a_{13} + b_{13} \\ a_{21} + b_{21} & a_{22} + b_{22} & a_{23} + b_{23} \end{bmatrix}.$$

We illustrate on a more concrete example:

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{bmatrix} + \begin{bmatrix} 7 & 8 \\ 9 & 10 \\ 11 & -1 \end{bmatrix} = \begin{bmatrix} 1+7 & 2+8 \\ 3+9 & 4+10 \\ 5+11 & 6-1 \end{bmatrix} = \begin{bmatrix} 8 & 10 \\ 12 & 14 \\ 16 & 5 \end{bmatrix}.$$

Let us check that this does the right thing to a vector. We use some of the vector algebra that we already know, and regroup things:

$$\begin{aligned} \begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{bmatrix} \begin{bmatrix} 2 \\ -1 \end{bmatrix} + \begin{bmatrix} 7 & 8 \\ 9 & 10 \\ 11 & -1 \end{bmatrix} \begin{bmatrix} 2 \\ -1 \end{bmatrix} &= \left(2 \begin{bmatrix} 1 \\ 3 \\ 5 \end{bmatrix} - \begin{bmatrix} 2 \\ 4 \\ 6 \end{bmatrix} \right) + \left(2 \begin{bmatrix} 7 \\ 9 \\ 11 \end{bmatrix} - \begin{bmatrix} 8 \\ 10 \\ -1 \end{bmatrix} \right) \\ &= 2 \left(\begin{bmatrix} 1 \\ 3 \\ 5 \end{bmatrix} + \begin{bmatrix} 7 \\ 9 \\ 11 \end{bmatrix} \right) - \left(\begin{bmatrix} 2 \\ 4 \\ 6 \end{bmatrix} + \begin{bmatrix} 8 \\ 10 \\ -1 \end{bmatrix} \right) \\ &= 2 \begin{bmatrix} 1+7 \\ 3+9 \\ 5+11 \end{bmatrix} - \begin{bmatrix} 2+8 \\ 4+10 \\ 6-1 \end{bmatrix} = 2 \begin{bmatrix} 8 \\ 12 \\ 16 \end{bmatrix} - \begin{bmatrix} 10 \\ 14 \\ 5 \end{bmatrix} \\ &= \begin{bmatrix} 8 & 10 \\ 12 & 14 \\ 16 & 5 \end{bmatrix} \begin{bmatrix} 2 \\ -1 \end{bmatrix} \quad \left(= \begin{bmatrix} 2(8) - 10 \\ 2(12) - 14 \\ 2(16) - 5 \end{bmatrix} = \begin{bmatrix} 6 \\ 10 \\ 27 \end{bmatrix} \right). \end{aligned}$$

If we replaced the numbers by letters, that would constitute a proof! Notice that we did not really have to compute what the result is to convince ourselves that the two expressions were equal.

If the sizes of the matrices do not match, then addition is undefined. If A is 3×2 and B is 2×5 , then we cannot add the matrices. We do not know what that could possibly mean.

It is also useful to have a matrix that when added to any other matrix does nothing. This is the zero matrix, the matrix of all zeros:

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} + \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}.$$

We often denote the zero matrix by 0 without specifying size. We would then just write $A + 0$, where we just assume that 0 is the zero matrix of the same size as A .

There are really two things we can multiply matrices by. We can multiply matrices by scalars or we can multiply by other matrices. Let us first consider multiplication by scalars. For a matrix A and a scalar α , we want αA to be the matrix that accomplishes

$$(\alpha A)\vec{x} = \alpha(A\vec{x}).$$

That is just scaling the result by α . If you think about it, scaling every term in A by α achieves just that: If

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix}, \quad \text{then} \quad \alpha A = \begin{bmatrix} \alpha a_{11} & \alpha a_{12} & \alpha a_{13} \\ \alpha a_{21} & \alpha a_{22} & \alpha a_{23} \end{bmatrix}.$$

For example,

$$2 \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} = \begin{bmatrix} 2 & 4 & 6 \\ 8 & 10 & 12 \end{bmatrix}.$$

Let us list some properties of matrix addition and scalar multiplication. Denote by 0 the zero matrix, by α, β scalars, and by A, B, C matrices. Then:

$$\begin{aligned} A + 0 &= A = 0 + A, \\ A + B &= B + A, \\ (A + B) + C &= A + (B + C), \\ \alpha(A + B) &= \alpha A + \alpha B, \\ (\alpha + \beta)A &= \alpha A + \beta A. \end{aligned}$$

These rules should look very familiar.

A.2.3 Matrix multiplication

As we mentioned above, composition of linear mappings is also a multiplication of matrices. Suppose A is an $m \times n$ matrix, that is, A takes $\mathbb{R}^n$ to $\mathbb{R}^m$, and B is an $n \times p$ matrix, that is, B takes $\mathbb{R}^p$ to $\mathbb{R}^n$. The composition AB should work as follows

$$AB\vec{x} = A(B\vec{x}).$$

First, a vector $\vec{x}$ in $\mathbb{R}^p$ gets taken to the vector $B\vec{x}$ in $\mathbb{R}^n$. Then the mapping A takes it to the vector $A(B\vec{x})$ in $\mathbb{R}^m$. In other words, the composition AB should be an $m \times p$ matrix. In terms of sizes we should have

$$“ \quad [m \times n][n \times p] = [m \times p]. \quad ”$$

Notice how the middle size must match.

OK, now we know what sizes of matrices we should be able to multiply and what the product should be. Let us see how to actually compute matrix multiplication. We start with the so-called *dot product* (or *inner product*) of two vectors. Usually, this is a row vector multiplied by a column vector of the same size. Dot product multiplies each pair of entries from the first and the second vector and sums these products. The result is a single number. For example,

$$\begin{bmatrix} a_1 & a_2 & a_3 \end{bmatrix} \cdot \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} = a_1b_1 + a_2b_2 + a_3b_3.$$

And similarly for larger (or smaller) vectors. A dot product is really a product of two matrices: a $1 \times n$ matrix and an $n \times 1$ matrix resulting in a 1×1 matrix—a number.

Armed with the dot product we define the *product of matrices*. We denote by $\text{row}_i(A)$ the i^{th} row of A and by $\text{column}_j(A)$ the j^{th} column of A . For an $m \times n$ matrix A and an $n \times p$ matrix B we can compute the product AB : The matrix AB is an $m \times p$ matrix whose ij^{th} entry is the dot product

$$\text{row}_i(A) \cdot \text{column}_j(B).$$

For example, given a 2×3 and a 3×2 matrix we should end up with a 2×2 matrix:

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{bmatrix} = \begin{bmatrix} a_{11}b_{11} + a_{12}b_{21} + a_{13}b_{31} & a_{11}b_{12} + a_{12}b_{22} + a_{13}b_{32} \\ a_{21}b_{11} + a_{22}b_{21} + a_{23}b_{31} & a_{21}b_{12} + a_{22}b_{22} + a_{23}b_{32} \end{bmatrix}, \quad (\text{A.1})$$

or with some numbers:

$$\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} \begin{bmatrix} -1 & 2 \\ -7 & 0 \\ 1 & -1 \end{bmatrix} = \begin{bmatrix} 1 \cdot (-1) + 2 \cdot (-7) + 3 \cdot 1 & 1 \cdot 2 + 2 \cdot 0 + 3 \cdot (-1) \\ 4 \cdot (-1) + 5 \cdot (-7) + 6 \cdot 1 & 4 \cdot 2 + 5 \cdot 0 + 6 \cdot (-1) \end{bmatrix} = \begin{bmatrix} -12 & -1 \\ -33 & 2 \end{bmatrix}.$$

A useful consequence of the definition is that the evaluation $A\vec{x}$ for a matrix A and a (column) vector $\vec{x}$ is also matrix multiplication. That is really why we think of vectors as column vectors, or $n \times 1$ matrices. For example,

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 2 \\ -1 \end{bmatrix} = \begin{bmatrix} 1 \cdot 2 + 2 \cdot (-1) \\ 3 \cdot 2 + 4 \cdot (-1) \end{bmatrix} = \begin{bmatrix} 0 \\ 2 \end{bmatrix}.$$

If you look at the last section, that is precisely the last example we gave.

You should stare at the computation of multiplication of matrices AB and the previous definition of $A\vec{y}$ as a mapping for a moment. What we are doing with matrix multiplication is applying the mapping A to the columns of B . This is usually written as follows. Suppose we write the $n \times p$ matrix $B = [\vec{b}_1 \ \vec{b}_2 \ \cdots \ \vec{b}_p]$, where $\vec{b}_1, \vec{b}_2, \dots, \vec{b}_p$ are the columns of B . Then for an $m \times n$ matrix A ,

$$AB = A[\vec{b}_1 \ \vec{b}_2 \ \cdots \ \vec{b}_p] = [A\vec{b}_1 \ A\vec{b}_2 \ \cdots \ A\vec{b}_p].$$

The columns of the $m \times p$ matrix AB are the vectors $A\vec{b}_1, A\vec{b}_2, \dots, A\vec{b}_p$. For example, in (A.1), the columns of

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{bmatrix}$$

are

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} b_{11} \\ b_{21} \\ b_{31} \end{bmatrix} \quad \text{and} \quad \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix} \begin{bmatrix} b_{12} \\ b_{22} \\ b_{32} \end{bmatrix}.$$

This is a very useful way to understand what matrix multiplication is. It should also make it easier to remember how to perform matrix multiplication.

A.2.4 Some rules of matrix algebra

For multiplication we want an analogue of a 1. That is, we desire a matrix that just leaves everything as it found it. This analogue is the so-called *identity matrix*. The identity matrix is a square matrix with 1s on the main diagonal and zeros everywhere else. It is usually denoted by I . For each size we have a different identity matrix and so sometimes we may denote the size as a subscript. For example, I_3 is the 3×3 identity matrix

$$I = I_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

Let us see how the matrix works on a smaller example,

$$\begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} a_{11} \cdot 1 + a_{12} \cdot 0 & a_{11} \cdot 0 + a_{12} \cdot 1 \\ a_{21} \cdot 1 + a_{22} \cdot 0 & a_{21} \cdot 0 + a_{22} \cdot 1 \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}.$$

Multiplication by the identity from the left looks similar, and also does not touch anything.

We have the following rules for matrix multiplication. Suppose that A, B, C are matrices of the correct sizes so that the following make sense. Let α denote a scalar (number). Then

$$\begin{aligned} A(BC) &= (AB)C && \text{(associative law),} \\ A(B + C) &= AB + AC && \text{(distributive law),} \\ (B + C)A &= BA + CA && \text{(distributive law),} \\ \alpha(AB) &= (\alpha A)B = A(\alpha B), && \\ IA = A = AI &&& \text{(identity).} \end{aligned}$$

Example A.2.1: Let us demonstrate a couple of these rules. For example, the associative law:

$$\underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\left(\underbrace{\begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}}_B \underbrace{\begin{bmatrix} -1 & 4 \\ 5 & 2 \end{bmatrix}}_C \right)}_{BC} = \underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\begin{bmatrix} 16 & 24 \\ -16 & -2 \end{bmatrix}}_{BC} = \underbrace{\begin{bmatrix} -96 & -78 \\ 64 & 52 \end{bmatrix}}_{A(BC)},$$

and

$$\underbrace{\left(\underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}}_B \right)}_{AB} \underbrace{\begin{bmatrix} -1 & 4 \\ 5 & 2 \end{bmatrix}}_C = \underbrace{\begin{bmatrix} -9 & -21 \\ 6 & 14 \end{bmatrix}}_{AB} \underbrace{\begin{bmatrix} -1 & 4 \\ 5 & 2 \end{bmatrix}}_C = \underbrace{\begin{bmatrix} -96 & -78 \\ 64 & 52 \end{bmatrix}}_{(AB)C}.$$

Or how about multiplication by scalars:

$$10 \left(\underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}}_B \right) = 10 \underbrace{\begin{bmatrix} -9 & -21 \\ 6 & 14 \end{bmatrix}}_{AB} = \underbrace{\begin{bmatrix} -90 & -210 \\ 60 & 140 \end{bmatrix}}_{10(AB)},$$

$$\underbrace{\left(10 \begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}\right)}_A \underbrace{\begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}}_B = \underbrace{\begin{bmatrix} -30 & 30 \\ 20 & -20 \end{bmatrix}}_{10A} \underbrace{\begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}}_B = \underbrace{\begin{bmatrix} -90 & -210 \\ 60 & 140 \end{bmatrix}}_{(10A)B},$$

and

$$\underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\left(10 \begin{bmatrix} 4 & 4 \\ 1 & -3 \end{bmatrix}\right)}_B = \underbrace{\begin{bmatrix} -3 & 3 \\ 2 & -2 \end{bmatrix}}_A \underbrace{\begin{bmatrix} 40 & 40 \\ 10 & -30 \end{bmatrix}}_{10B} = \underbrace{\begin{bmatrix} -90 & -210 \\ 60 & 140 \end{bmatrix}}_{A(10B)}.$$

A multiplication rule, one you have used since primary school on numbers, is quite conspicuously missing for matrices. That is, matrix multiplication is not commutative. Firstly, just because AB makes sense, it may be that BA is not even defined. For example, if A is 2×3 , and B is 3×4 , then we can multiply AB but not BA .

Even if AB and BA are both defined, it does not mean that they are equal. For example, take $A = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}$:

$$AB = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 1 & 2 \end{bmatrix} \quad \neq \quad \begin{bmatrix} 1 & 1 \\ 2 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} = BA.$$

A.2.5 Inverse

A couple of other algebra rules you know for numbers do not quite work on matrices:

- (i) $AB = AC$ does not necessarily imply $B = C$, even if A is not 0.
- (ii) $AB = 0$ does not necessarily mean that $A = 0$ or $B = 0$.

For example:

$$\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 0 & 2 \\ 0 & 0 \end{bmatrix}.$$

To make these rules hold, we do not just need one of the matrices to not be zero, we would need to “divide” by a matrix. This is where the *matrix inverse* comes in. Suppose that A and B are $n \times n$ matrices such that

$$AB = I = BA.$$

Then we call B the inverse of A and we denote B by A^{-1} . Perhaps not surprisingly, $(A^{-1})^{-1} = A$, since if the inverse of A is B , then the inverse of B is A . If the inverse of A exists, then we say A is *invertible*. If A is not invertible, we say A is *singular*.

If $A = [a]$ is a 1×1 matrix, then A^{-1} is $a^{-1} = \frac{1}{a}$. That is where the notation comes from. The computation is not nearly as simple when A is larger.

The proper formulation of the cancellation rule is:

$$\text{If } A \text{ is invertible, then } AB = AC \text{ implies } B = C.$$

The computation is what you would do in regular algebra with numbers, but you have to be careful never to commute matrices:

$$\begin{aligned} AB &= AC, \\ A^{-1}AB &= A^{-1}AC, \\ IB &= IC, \\ B &= C. \end{aligned}$$

And similarly for cancellation on the right:

$$\text{If } A \text{ is invertible, then } BA = CA \text{ implies } B = C.$$

The rule says, among other things, that the inverse of a matrix is unique if it exists: If $AB = I = AC$, then A is invertible and $B = C$.

We will see later how to compute an inverse of a matrix in general. For now, let us note that there is a simple formula for the inverse of a 2×2 matrix

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}.$$

For example:

$$\begin{bmatrix} 1 & 1 \\ 2 & 4 \end{bmatrix}^{-1} = \frac{1}{1 \cdot 4 - 1 \cdot 2} \begin{bmatrix} 4 & -1 \\ -2 & 1 \end{bmatrix} = \begin{bmatrix} 2 & -1/2 \\ -1 & 1/2 \end{bmatrix}.$$

Let's try it:

$$\begin{bmatrix} 1 & 1 \\ 2 & 4 \end{bmatrix} \begin{bmatrix} 2 & -1/2 \\ -1 & 1/2 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \quad \text{and} \quad \begin{bmatrix} 2 & -1/2 \\ -1 & 1/2 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 2 & 4 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}.$$

Just as we cannot divide by every number, not every matrix is invertible. In the case of matrices, however, we may have singular matrices that are not zero. For example,

$$\begin{bmatrix} 1 & 1 \\ 2 & 2 \end{bmatrix}$$

is a singular matrix. But didn't we just give a formula for an inverse? Let us try it:

$$\begin{bmatrix} 1 & 1 \\ 2 & 2 \end{bmatrix}^{-1} = \frac{1}{1 \cdot 2 - 1 \cdot 2} \begin{bmatrix} 2 & -1 \\ -2 & 1 \end{bmatrix} = ?$$

We get into a bit of trouble: We are trying to divide by zero.

So a 2×2 matrix A is invertible whenever

$$ad - bc \neq 0$$

and otherwise it is singular. The expression $ad - bc$ is called the *determinant* and we will look at it more carefully in a later section. There is a similar expression for a square matrix of any size.

A.2.6 Diagonal matrices

A simple (and surprisingly useful) type of a square matrix is a so-called *diagonal matrix*. It is a matrix whose entries are all zero except those on the main diagonal from top left to bottom right. For example a 4×4 diagonal matrix is of the form

$$\begin{bmatrix} d_1 & 0 & 0 & 0 \\ 0 & d_2 & 0 & 0 \\ 0 & 0 & d_3 & 0 \\ 0 & 0 & 0 & d_4 \end{bmatrix}.$$

Such matrices have nice properties when we multiply by them. If we multiply them by a vector, they multiply the k^{th} entry by d_k . For example,

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} 4 \\ 5 \\ 6 \end{bmatrix} = \begin{bmatrix} 1 \cdot 4 \\ 2 \cdot 5 \\ 3 \cdot 6 \end{bmatrix} = \begin{bmatrix} 4 \\ 10 \\ 18 \end{bmatrix}.$$

Similarly, when they multiply another matrix from the left, they multiply the k^{th} row by d_k . For example,

$$\begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix} = \begin{bmatrix} 2 & 2 & 2 \\ 3 & 3 & 3 \\ -1 & -1 & -1 \end{bmatrix}.$$

On the other hand, multiplying on the right, they multiply the columns:

$$\begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix} \begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & -1 \end{bmatrix} = \begin{bmatrix} 2 & 3 & -1 \\ 2 & 3 & -1 \\ 2 & 3 & -1 \end{bmatrix}.$$

And it is really easy to multiply two diagonal matrices together—we multiply the entries:

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & -1 \end{bmatrix} = \begin{bmatrix} 1 \cdot 2 & 0 & 0 \\ 0 & 2 \cdot 3 & 0 \\ 0 & 0 & 3 \cdot (-1) \end{bmatrix} = \begin{bmatrix} 2 & 0 & 0 \\ 0 & 6 & 0 \\ 0 & 0 & -3 \end{bmatrix}.$$

For this last reason, they are easy to invert—you simply invert each diagonal element:

$$\begin{bmatrix} d_1 & 0 & 0 \\ 0 & d_2 & 0 \\ 0 & 0 & d_3 \end{bmatrix}^{-1} = \begin{bmatrix} d_1^{-1} & 0 & 0 \\ 0 & d_2^{-1} & 0 \\ 0 & 0 & d_3^{-1} \end{bmatrix}.$$

Let us check an example

$$\underbrace{\begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 4 \end{bmatrix}^{-1}}_{A^{-1}} \underbrace{\begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 4 \end{bmatrix}}_A = \underbrace{\begin{bmatrix} \frac{1}{2} & 0 & 0 \\ 0 & \frac{1}{3} & 0 \\ 0 & 0 & \frac{1}{4} \end{bmatrix}}_{A^{-1}} \underbrace{\begin{bmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 4 \end{bmatrix}}_A = \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}}_I.$$

It is no wonder that the way we solve many problems in linear algebra (and in differential equations) is to try to reduce the problem to the case of diagonal matrices.

A.2.7 Transpose

Vectors do not always have to be column vectors. That is just a convention. From time to time, swapping rows and columns is needed. The operation that swaps rows and columns is the so-called *transpose*. The transpose of A is denoted by A^T . Example:

$$\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix}^T = \begin{bmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix}.$$

Transpose takes an $m \times n$ matrix to an $n \times m$ matrix.

A key feature of the transpose is that if the product AB makes sense, then $B^T A^T$ also makes sense, at least from the point of view of sizes. In fact, we get precisely the transpose of AB . That is:

$$(AB)^T = B^T A^T.$$

For example,

$$\left(\begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \\ 2 & -2 \end{bmatrix} \right)^T = \begin{bmatrix} 0 & 1 & 2 \\ 1 & 0 & -2 \end{bmatrix} \begin{bmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix}.$$

It is left to the reader to verify that computing the matrix product on the left and then transposing is the same as computing the matrix product on the right.

If we have a column vector $\vec{x}$ to which we apply a matrix A and we transpose the result, then the row vector $\vec{x}^T$ applies to A^T from the left:

$$(A\vec{x})^T = \vec{x}^T A^T.$$

Another place where transpose is useful is when we wish to apply the dot product* to two column vectors:

$$\vec{x} \cdot \vec{y} = \vec{y}^T \vec{x}.$$

That is the way that one often writes the dot product in software.

We say a matrix A is *symmetric* if $A = A^T$. For example,

$$\begin{bmatrix} 1 & 2 & 3 \\ 2 & 4 & 5 \\ 3 & 5 & 6 \end{bmatrix}$$

is a symmetric matrix. Notice that a symmetric matrix is always square, that is, $n \times n$. Symmetric matrices have many nice properties[†], and come up quite often in applications.

*As a side note, mathematicians write $\vec{y}^T \vec{x}$ and physicists write $\vec{x}^T \vec{y}$. Shhh. . . don't tell anyone, but the physicists are probably right on this.

[†]Although so far we have not learned enough about matrices to really appreciate them.

A.2.8 Exercises*Exercise A.2.1: Add the following matrices*

$$a) \begin{bmatrix} -1 & 2 & 2 \\ 5 & 8 & -1 \end{bmatrix} + \begin{bmatrix} 3 & 2 & 3 \\ 8 & 3 & 5 \end{bmatrix}$$

$$b) \begin{bmatrix} 1 & 2 & 4 \\ 2 & 3 & 1 \\ 0 & 5 & 1 \end{bmatrix} + \begin{bmatrix} 2 & -8 & -3 \\ 3 & 1 & 0 \\ 6 & -4 & 1 \end{bmatrix}$$

Exercise A.2.2: Compute

$$a) 3 \begin{bmatrix} 0 & 3 \\ -2 & 2 \end{bmatrix} + 6 \begin{bmatrix} 1 & 5 \\ -1 & 5 \end{bmatrix}$$

$$b) 2 \begin{bmatrix} -3 & 1 \\ 2 & 2 \end{bmatrix} - 3 \begin{bmatrix} 2 & -1 \\ 3 & 2 \end{bmatrix}$$

Exercise A.2.3: Multiply the following matrices

$$a) \begin{bmatrix} -1 & 2 \\ 3 & 1 \\ 5 & 8 \end{bmatrix} \begin{bmatrix} 3 & -1 & 3 & 1 \\ 8 & 3 & 2 & -3 \end{bmatrix}$$

$$b) \begin{bmatrix} 1 & 2 & 3 \\ 3 & 1 & 1 \\ 1 & 0 & 3 \end{bmatrix} \begin{bmatrix} 2 & 3 & 1 & 7 \\ 1 & 2 & 3 & -1 \\ 1 & -1 & 3 & 0 \end{bmatrix}$$

$$c) \begin{bmatrix} 4 & 1 & 6 & 3 \\ 5 & 6 & 5 & 0 \\ 4 & 6 & 6 & 0 \end{bmatrix} \begin{bmatrix} 2 & 5 \\ 1 & 2 \\ 3 & 5 \\ 5 & 6 \end{bmatrix}$$

$$d) \begin{bmatrix} 1 & 1 & 4 \\ 0 & 5 & 1 \end{bmatrix} \begin{bmatrix} 2 & 2 \\ 1 & 0 \\ 6 & 4 \end{bmatrix}$$

Exercise A.2.4: Compute the inverse of the given matrices

$$a) [-3]$$

$$b) \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 4 \\ 1 & 3 \end{bmatrix}$$

$$d) \begin{bmatrix} 2 & 2 \\ 1 & 4 \end{bmatrix}$$

Exercise A.2.5: Compute the inverse of the given matrices

$$a) \begin{bmatrix} -2 & 0 \\ 0 & 1 \end{bmatrix}$$

$$b) \begin{bmatrix} 3 & 0 & 0 \\ 0 & -2 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & 0.01 & 0 \\ 0 & 0 & 0 & -5 \end{bmatrix}$$

Exercise A.2.101: Add the following matrices

$$a) \begin{bmatrix} 2 & 1 & 0 \\ 1 & 1 & -1 \end{bmatrix} + \begin{bmatrix} 5 & 3 & 4 \\ 1 & 2 & 5 \end{bmatrix}$$

$$b) \begin{bmatrix} 6 & -2 & 3 \\ 7 & 3 & 3 \\ 8 & -1 & 2 \end{bmatrix} + \begin{bmatrix} -1 & -1 & -3 \\ 6 & 7 & 3 \\ -9 & 4 & -1 \end{bmatrix}$$

Exercise A.2.102: Compute

$$a) 2 \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} + 3 \begin{bmatrix} -1 & 3 \\ 1 & 2 \end{bmatrix}$$

$$b) 3 \begin{bmatrix} 2 & -1 \\ 1 & 3 \end{bmatrix} - 2 \begin{bmatrix} 2 & 1 \\ -1 & 2 \end{bmatrix}$$

Exercise A.2.103: Multiply the following matrices

$$a) \begin{bmatrix} 2 & 1 & 4 \\ 3 & 4 & 4 \end{bmatrix} \begin{bmatrix} 2 & 4 \\ 6 & 3 \\ 3 & 5 \end{bmatrix}$$

$$b) \begin{bmatrix} 0 & 3 & 3 \\ 2 & -2 & 1 \\ 3 & 5 & -2 \end{bmatrix} \begin{bmatrix} 6 & 6 & 2 \\ 4 & 6 & 0 \\ 2 & 0 & 4 \end{bmatrix}$$

$$c) \begin{bmatrix} 3 & 4 & 1 \\ 2 & -1 & 0 \\ 4 & -1 & 5 \end{bmatrix} \begin{bmatrix} 0 & 2 & 5 & 0 \\ 2 & 0 & 5 & 2 \\ 3 & 6 & 1 & 6 \end{bmatrix}$$

$$d) \begin{bmatrix} -2 & -2 \\ 5 & 3 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} 0 & 3 \\ 1 & 3 \end{bmatrix}$$

Exercise A.2.104: Compute the inverse of the given matrices

$$a) \begin{bmatrix} 2 \end{bmatrix}$$

$$b) \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 2 \\ 3 & 5 \end{bmatrix}$$

$$d) \begin{bmatrix} 4 & 2 \\ 4 & 4 \end{bmatrix}$$

Exercise A.2.105: Compute the inverse of the given matrices

$$a) \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix}$$

$$b) \begin{bmatrix} 4 & 0 & 0 \\ 0 & 5 & 0 \\ 0 & 0 & -1 \end{bmatrix}$$

$$c) \begin{bmatrix} -1 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0.1 \end{bmatrix}$$

A.3 Elimination

Note: 2–3 lectures

A.3.1 Linear systems of equations

One application of matrices is to solve systems of linear equations*. Consider the following system of linear equations

$$\begin{aligned} 2x_1 + 2x_2 + 2x_3 &= 2, \\ x_1 + x_2 + 3x_3 &= 5, \\ x_1 + 4x_2 + x_3 &= 10. \end{aligned} \tag{A.2}$$

There is a systematic procedure called *elimination* to solve such a system. In this procedure, we attempt to eliminate each variable from all but one equation. We want to end up with equations such as $x_3 = 2$, where we can just read off the answer.

We write a system of linear equations as a matrix equation:

$$A\vec{x} = \vec{b}.$$

The system (A.2) is written as

$$\underbrace{\begin{bmatrix} 2 & 2 & 2 \\ 1 & 1 & 3 \\ 1 & 4 & 1 \end{bmatrix}}_A \underbrace{\begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}}_{\vec{x}} = \underbrace{\begin{bmatrix} 2 \\ 5 \\ 10 \end{bmatrix}}_{\vec{b}}.$$

If we knew the inverse of A , then we would be done. We would simply solve the equation:

$$\vec{x} = A^{-1}A\vec{x} = A^{-1}\vec{b}.$$

Well, but that is part of the problem. We do not know how to compute the inverse for matrices bigger than 2×2 . We will see later that to compute the inverse, we are really solving $A\vec{x} = \vec{b}$ for several different $\vec{b}$. In other words, we will need to do elimination to find A^{-1} . In addition, we may wish to solve $A\vec{x} = \vec{b}$ if A is not invertible, or perhaps not even square.

Let us return to the equations themselves and see how we can manipulate them. There are a few operations we can perform on the equations that do not change the solution. First, perhaps an operation that may seem stupid, we can swap two equations in (A.2):

$$\begin{aligned} x_1 + x_2 + 3x_3 &= 5, \\ 2x_1 + 2x_2 + 2x_3 &= 2, \\ x_1 + 4x_2 + x_3 &= 10. \end{aligned}$$

*Although perhaps we have this backwards, quite often we solve a linear system of equations to find out something about matrices, rather than vice versa.

Clearly these new equations have the same solutions x_1, x_2, x_3 . A second operation is that we can multiply an equation by a nonzero number. For example, we multiply the third equation in (A.2) by 3:

$$\begin{aligned} 2x_1 + 2x_2 + 2x_3 &= 2, \\ x_1 + x_2 + 3x_3 &= 5, \\ 3x_1 + 12x_2 + 3x_3 &= 30. \end{aligned}$$

Finally, we can add a multiple of one equation to another equation. For instance, we add 3 times the third equation in (A.2) to the second equation:

$$\begin{aligned} 2x_1 + 2x_2 + 2x_3 &= 2, \\ (1 + 3)x_1 + (1 + 12)x_2 + (3 + 3)x_3 &= 5 + 30, \\ x_1 + 4x_2 + x_3 &= 10. \end{aligned}$$

The same x_1, x_2, x_3 should still be solutions to the new equations. These were just examples; we did not get any closer to the solution. We must do these three operations in some more logical manner, but it turns out these three operations suffice to solve every linear equation.

The first thing is to write the equations in a more compact manner. Given

$$A\vec{x} = \vec{b},$$

we write down the so-called *augmented matrix*

$$[A \mid \vec{b}],$$

where the vertical line is just a marker for us to know where the “right-hand side” of the equation starts. For the system (A.2) the augmented matrix is

$$\left[\begin{array}{ccc|c} 2 & 2 & 2 & 2 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right].$$

The entire process of elimination, which we will describe, is often applied to any sort of matrix, not just an augmented matrix. Simply think of the matrix as the 3×4 matrix

$$\left[\begin{array}{cccc} 2 & 2 & 2 & 2 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right].$$

A.3.2 Row echelon form and elementary operations

We apply the three operations above to the matrix. We call these the *elementary operations* or *elementary row operations*. Translating the operations to the matrix setting, the operations become:

- (i) Swap two rows.
- (ii) Multiply a row by a nonzero number.
- (iii) Add a multiple of one row to another row.

We run these operations until we get into a state where it is easy to read off the answer, or until we get into a contradiction indicating no solution.

More specifically, we run the operations until we obtain the so-called *row echelon form*. Let us call the first (from the left) nonzero entry in each row the *leading entry*. A matrix is in *row echelon form* if the following conditions are satisfied:

- (i) The leading entry in any row is strictly to the right of the leading entry of the row above.
- (ii) Any zero rows are below all the nonzero rows.
- (iii) All leading entries are 1.

A matrix is in *reduced row echelon form* if furthermore the following condition is satisfied.

- (iv) All the entries above a leading entry are zero.

Note that the definition applies to matrices of any size.

Example A.3.1: The following matrices are in row echelon form. The leading entries are marked:

$$\begin{bmatrix} \boxed{1} & 2 & 9 & 3 \\ 0 & 0 & \boxed{1} & 5 \\ 0 & 0 & 0 & \boxed{1} \end{bmatrix} \quad \begin{bmatrix} \boxed{1} & -1 & -3 \\ 0 & \boxed{1} & 5 \\ 0 & 0 & \boxed{1} \end{bmatrix} \quad \begin{bmatrix} \boxed{1} & 2 & 1 \\ 0 & \boxed{1} & 2 \\ 0 & 0 & 0 \end{bmatrix} \quad \begin{bmatrix} 0 & \boxed{1} & -5 & 2 \\ 0 & 0 & 0 & \boxed{1} \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

None of the matrices above are in *reduced* row echelon form. For example, in the first matrix none of the entries above the second and third leading entries are zero; they are 9, 3, and 5. The following matrices are in reduced row echelon form. The leading entries are marked:

$$\begin{bmatrix} \boxed{1} & 3 & 0 & 8 \\ 0 & 0 & \boxed{1} & 6 \\ 0 & 0 & 0 & 0 \end{bmatrix} \quad \begin{bmatrix} \boxed{1} & 0 & 2 & 0 \\ 0 & \boxed{1} & 3 & 0 \\ 0 & 0 & 0 & \boxed{1} \end{bmatrix} \quad \begin{bmatrix} \boxed{1} & 0 & 3 \\ 0 & \boxed{1} & -2 \\ 0 & 0 & 0 \end{bmatrix} \quad \begin{bmatrix} 0 & \boxed{1} & 2 & 0 \\ 0 & 0 & 0 & \boxed{1} \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

The procedure we will describe to find a reduced row echelon form of a matrix is called *Gauss–Jordan elimination*. The first part of it, which obtains a row echelon form, is called *Gaussian elimination* or *row reduction*. For some problems, a row echelon form is sufficient, and it is a bit less work to only do this first part.

To attain the row echelon form we work systematically. We go column by column, starting at the first column. We find the topmost entry in the first column that is not zero, and we call it the *pivot*. If there is no nonzero entry we move to the next column. We swap

rows to put the row with the pivot as the first row. We divide the first row by the pivot to make the pivot entry be a 1. Now look at all the rows below and subtract the correct multiple of the pivot row so that all the entries below the pivot become zero.

After this procedure we forget that we had a first row (it is now fixed), and we forget about the column with the pivot and all the preceding zero columns. Below the pivot row, all the entries in these columns are just zero. Then we focus on the smaller matrix and we repeat the steps above.

It is best shown by example, so let us go back to the example from the beginning of the section. We keep the vertical line in the matrix, even though the procedure works on any matrix, not just an augmented matrix. We start with the first column and we locate the pivot, in this case the first entry of the first column.

$$\left[\begin{array}{ccc|c} \boxed{2} & 2 & 2 & 2 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right]$$

We multiply the first row by $1/2$.

$$\left[\begin{array}{ccc|c} \boxed{1} & 1 & 1 & 1 \\ 1 & 1 & 3 & 5 \\ 1 & 4 & 1 & 10 \end{array} \right]$$

We subtract the first row from the second and third row (two elementary operations).

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 0 & 2 & 4 \\ 0 & 3 & 0 & 9 \end{array} \right]$$

We are done with the first column and the first row for now. We almost pretend the matrix does not have the first column and the first row.

$$\left[\begin{array}{ccc|c} * & * & * & * \\ * & 0 & 2 & 4 \\ * & 3 & 0 & 9 \end{array} \right]$$

OK, look at the second column, and notice that now the pivot is in the third row.

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 0 & 2 & 4 \\ 0 & \boxed{3} & 0 & 9 \end{array} \right]$$

We swap rows.

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & \boxed{3} & 0 & 9 \\ 0 & 0 & 2 & 4 \end{array} \right]$$

And we divide the pivot row by 3.

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & \boxed{1} & 0 & 3 \\ 0 & 0 & 2 & 4 \end{array} \right]$$

We do not need to subtract anything as everything below the pivot is already zero. We move on. We again start ignoring the second row and second column and focus on

$$\left[\begin{array}{ccc|c} * & * & * & * \\ * & * & * & * \\ * & * & 2 & 4 \end{array} \right].$$

We find the pivot, then divide that row by 2:

$$\left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & \boxed{2} & 4 \end{array} \right] \rightarrow \left[\begin{array}{ccc|c} 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 2 \end{array} \right].$$

The matrix is now in row echelon form.

The equation corresponding to the last row is $x_3 = 2$. We know x_3 and we could substitute it into the first two equations to get equations for x_1 and x_2 . Then we could do the same thing with x_2 , until we solve for all 3 variables. This procedure is called *backsubstitution* and we can achieve it via elementary operations. We start from the lowest pivot (leading entry in the row echelon form) and subtract the right multiple from the row above to make all the entries above this pivot zero. Then we move to the next pivot and so on. After we are done, we will have a matrix in reduced row echelon form.

We continue our example. Subtract the last row from the first to get

$$\left[\begin{array}{ccc|c} 1 & 1 & 0 & -1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 2 \end{array} \right].$$

The entry above the pivot in the second row is already zero. So we move onto the next pivot, the one in the second row. We subtract this row from the top row to get

$$\left[\begin{array}{ccc|c} 1 & 0 & 0 & -4 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 2 \end{array} \right].$$

The matrix is in reduced row echelon form.

If we now write down the equations for x_1, x_2, x_3 , we find

$$x_1 = -4, \quad x_2 = 3, \quad x_3 = 2.$$

In other words, we have solved the system.

A.3.3 Non-unique solutions and inconsistent systems

It is possible that the solution of a linear system of equations is not unique, or that no solution exists. Suppose for a moment that the row echelon form we found was

$$\left[\begin{array}{ccc|c} 1 & 2 & 3 & 4 \\ 0 & 0 & 1 & 3 \\ 0 & 0 & 0 & 1 \end{array} \right].$$

Then we have an equation $0 = 1$ coming from the last row. That is impossible, and the equations are what we call *inconsistent*. There is no solution to $A\vec{x} = \vec{b}$.

On the other hand, if we find a row echelon form

$$\left[\begin{array}{ccc|c} 1 & 2 & 3 & 4 \\ 0 & 0 & 1 & 3 \\ 0 & 0 & 0 & 0 \end{array} \right],$$

then there is no issue with finding solutions. In fact, we will find way too many. Let us continue with backsubstitution (subtracting 3 times the third row from the first) to find the reduced row echelon form and let's mark the pivots.

$$\left[\begin{array}{ccc|c} \boxed{1} & 2 & 0 & -5 \\ 0 & 0 & \boxed{1} & 3 \\ 0 & 0 & 0 & 0 \end{array} \right]$$

The last row is all zeros; it just says $0 = 0$ and we ignore it. The two remaining equations are

$$x_1 + 2x_2 = -5, \quad x_3 = 3.$$

Let us solve for the variables that corresponded to the pivots, that is, x_1 and x_3 as there was a pivot in the first column and in the third column:

$$\begin{aligned} x_1 &= -2x_2 - 5, \\ x_3 &= 3. \end{aligned}$$

The variable x_2 can be anything you wish, and we still get a solution. The x_2 is called a *free variable*. There are infinitely many solutions, one for every choice of x_2 . If we pick $x_2 = 0$, then $x_1 = -5$, and $x_3 = 3$ give a solution. But we also get a solution by picking say $x_2 = 1$, in which case $x_1 = -7$ and $x_3 = 3$, or by picking $x_2 = -5$ in which case $x_1 = 5$ and $x_3 = 3$.

The general idea is that if any row has all zeros in the columns corresponding to the variables, but a nonzero entry in the column corresponding to the right-hand side $\vec{b}$, then the system is inconsistent and has no solutions. In other words, the system is inconsistent if you find a pivot on the right side of the vertical line drawn in the augmented matrix. Otherwise, the system is *consistent*, and at least one solution exists.

Suppose the system is consistent (at least one solution exists):

- (i) If every column corresponding to a variable has a pivot element, then the solution is unique.
- (ii) If there are columns corresponding to variables with no pivot, then those are *free variables* that can be chosen arbitrarily, and there are infinitely many solutions.

When $\vec{b} = \vec{0}$, we have a so-called *homogeneous matrix equation*

$$A\vec{x} = \vec{0}.$$

There is no need to write an augmented matrix in this case. As the elementary operations do not do anything to a zero column, it always stays a zero column. Moreover, $A\vec{x} = \vec{0}$ always has at least one solution, namely $\vec{x} = \vec{0}$. Such a system is always consistent. It may have other solutions: If we find any free variables, then we get infinitely many solutions.

The set of solutions of $A\vec{x} = \vec{0}$ comes up quite often so people give it a name. It is called the *nullspace* or the *kernel* of A . One place where the kernel comes up is invertibility of a square matrix A . If the kernel of A contains a nonzero vector, then it contains infinitely many vectors (there was a free variable). But then it is impossible to invert $\vec{0}$, since infinitely many vectors go to $\vec{0}$, so there is no unique vector that A takes to $\vec{0}$. So if the kernel is nontrivial, that is, if there are any nonzero vectors in the kernel, in other words, if there are any free variables, or in yet other words, if the row echelon form of A has columns without pivots, then A is not invertible. We will return to this idea later.

A.3.4 Linear independence and rank

If rows of a matrix correspond to equations, it may be good to find out how many equations we really need to find the same set of solutions. Similarly, if we find a number of solutions to a linear equation $A\vec{x} = \vec{0}$, we may ask if we found enough so that all other solutions can be formed out of the given set. The concept we want is that of linear independence. That same concept is useful for differential equations, for example in [chapter 2](#).

Given row or column vectors $\vec{y}_1, \vec{y}_2, \dots, \vec{y}_n$, a *linear combination* is an expression of the form

$$\alpha_1\vec{y}_1 + \alpha_2\vec{y}_2 + \dots + \alpha_n\vec{y}_n,$$

where $\alpha_1, \alpha_2, \dots, \alpha_n$ are all scalars. For example, $3\vec{y}_1 + \vec{y}_2 - 5\vec{y}_3$ is a linear combination of $\vec{y}_1, \vec{y}_2$, and $\vec{y}_3$.

We have seen linear combinations before. The expression

$$A\vec{x}$$

is a linear combination of the columns of A , while

$$\vec{x}^T A = (A^T \vec{x})^T$$

is a linear combination of the rows of A .

The way linear combinations come up in our study of differential equations is similar to the following computation. Suppose that $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ are solutions to $A\vec{x}_1 = \vec{0}$, $A\vec{x}_2 = \vec{0}$, $\dots$, $A\vec{x}_n = \vec{0}$. Then the linear combination

$$\vec{y} = \alpha_1 \vec{x}_1 + \alpha_2 \vec{x}_2 + \dots + \alpha_n \vec{x}_n$$

is a solution to $A\vec{y} = \vec{0}$:

$$\begin{aligned} A\vec{y} &= A(\alpha_1 \vec{x}_1 + \alpha_2 \vec{x}_2 + \dots + \alpha_n \vec{x}_n) = \\ &= \alpha_1 A\vec{x}_1 + \alpha_2 A\vec{x}_2 + \dots + \alpha_n A\vec{x}_n = \alpha_1 \vec{0} + \alpha_2 \vec{0} + \dots + \alpha_n \vec{0} = \vec{0}. \end{aligned}$$

So if we find enough solutions, we have all solutions. The question is, when did we find enough of them?

We say the vectors $\vec{y}_1, \vec{y}_2, \dots, \vec{y}_n$ are *linearly independent* if the only solution to

$$\alpha_1 \vec{x}_1 + \alpha_2 \vec{x}_2 + \dots + \alpha_n \vec{x}_n = \vec{0}$$

is $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$. Otherwise, we say the vectors are *linearly dependent*.

For example, the vectors $\begin{bmatrix} 1 \\ 2 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$ are linearly independent. Let's try:

$$\alpha_1 \begin{bmatrix} 1 \\ 2 \end{bmatrix} + \alpha_2 \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} \alpha_1 \\ 2\alpha_1 + \alpha_2 \end{bmatrix} = \vec{0} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}.$$

So $\alpha_1 = 0$, and then it is clear that $\alpha_2 = 0$ as well. In other words, the two vectors are linearly independent.

If a set of vectors is linearly dependent, that is, some of the α_j s are nonzero, then we can solve for one vector in terms of the others. Suppose $\alpha_1 \neq 0$. Since $\alpha_1 \vec{x}_1 + \alpha_2 \vec{x}_2 + \dots + \alpha_n \vec{x}_n = \vec{0}$, then

$$\vec{x}_1 = \frac{-\alpha_2}{\alpha_1} \vec{x}_2 - \frac{-\alpha_3}{\alpha_1} \vec{x}_3 + \dots + \frac{-\alpha_n}{\alpha_1} \vec{x}_n.$$

For example,

$$2 \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix} - 4 \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} + 2 \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix},$$

and so

$$\begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix} = 2 \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} - \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix}.$$

You may have noticed that solving for those α_j s is just solving linear equations, and so you may not be surprised that to check if a set of vectors is linearly independent we use row reduction.

Given a set of vectors, we may not be interested in just finding if they are linearly independent or not, we may be interested in finding a linearly independent subset. Or

perhaps we may want to find some other vectors that give the same linear combinations and are linearly independent. The way to figure this out is to form a matrix out of our vectors. If we have row vectors we consider them as rows of a matrix. If we have column vectors we consider them columns of a matrix. The set of all linear combinations of a set of vectors is called their *span*.

$$\text{span}\{\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n\} = \{\text{Set of all linear combinations of } \vec{x}_1, \vec{x}_2, \dots, \vec{x}_n\}.$$

Given a matrix A , the maximal number of linearly independent rows is called the *rank* of A , and we write “rank A ” for the rank. For example,

$$\text{rank} \begin{bmatrix} 1 & 1 & 1 \\ 2 & 2 & 2 \\ -1 & -1 & -1 \end{bmatrix} = 1.$$

The second and third row are multiples of the first one. We cannot choose more than one row and still have a linearly independent set. But what is

$$\text{rank} \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} = ?$$

That seems to be a tougher question to answer. The first two rows are linearly independent (neither is a multiple of the other), so the rank is at least two. If we would set up the equations for the α_1 , α_2 , and α_3 , we would find a system with infinitely many solutions. One solution is

$$[1 \ 2 \ 3] - 2[4 \ 5 \ 6] + [7 \ 8 \ 9] = [0 \ 0 \ 0].$$

So the set of all three rows is linearly dependent, the rank cannot be 3. Therefore the rank is 2.

But how can we do this in a more systematic way? We find the row echelon form!

$$\text{Row echelon form of } \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} \text{ is } \begin{bmatrix} 1 & 2 & 3 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{bmatrix}.$$

The elementary row operations do not change the set of linear combinations of the rows (that was one of the main reasons for defining them as they were). In other words, the span of the rows of the A is the same as the span of the rows of the row echelon form of A . In particular, the number of linearly independent rows is the same. And in the row echelon form, all nonzero rows are linearly independent. This is not hard to see. Consider the two nonzero rows in the example above. Suppose we tried to solve for the α_1 and α_2 in

$$\alpha_1 [1 \ 2 \ 3] + \alpha_2 [0 \ 1 \ 2] = [0 \ 0 \ 0].$$

Since the first column of the row echelon matrix has zeros except in the first row, it follows that $\alpha_1 = 0$. Using that $\alpha_1 = 0$ and looking at the second column, we find $\alpha_2 = 0$. We only have two nonzero rows, and they are linearly independent, so the rank of the matrix is 2.

The span of the rows is called the *row space*. The row space of A and the row echelon form of A are the same. In the example,

$$\begin{aligned} \text{row space of } \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} &= \text{span} \{ [1 \ 2 \ 3], [4 \ 5 \ 6], [7 \ 8 \ 9] \} \\ &= \text{span} \{ [1 \ 2 \ 3], [0 \ 1 \ 2] \}. \end{aligned}$$

Similarly to row space, the span of columns is called the *column space*.

$$\text{column space of } \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} = \text{span} \left\{ \begin{bmatrix} 1 \\ 4 \\ 7 \end{bmatrix}, \begin{bmatrix} 2 \\ 5 \\ 8 \end{bmatrix}, \begin{bmatrix} 3 \\ 6 \\ 9 \end{bmatrix} \right\}.$$

So it may also be good to find the number of linearly independent columns of A . One way to do that is to find the number of linearly independent rows of A^T . It is a tremendously useful fact that the number of linearly independent columns is always the same as the number of linearly independent rows:

Theorem A.3.1. $\text{rank } A = \text{rank } A^T$

In particular, to find a set of linearly independent columns we need to look at where the pivots were. If you recall above, when solving $A\vec{x} = \vec{0}$ the key was finding the pivots. Any non-pivot columns corresponded to free variables. That means we can solve for the non-pivot columns in terms of the pivot columns. Let's see an example. First we reduce some random matrix:

$$\begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix}.$$

We find a pivot and reduce the rows below:

$$\begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix} \rightarrow \begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 0 & 0 & -1 & -2 \\ 3 & 6 & 7 & 8 \end{bmatrix} \rightarrow \begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 0 & 0 & -1 & -2 \\ 0 & 0 & -2 & -4 \end{bmatrix}.$$

We find the next pivot, make it one, and rinse and repeat:

$$\begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 0 & 0 & \boxed{-1} & -2 \\ 0 & 0 & -2 & -4 \end{bmatrix} \rightarrow \begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 0 & 0 & \boxed{1} & 2 \\ 0 & 0 & -2 & -4 \end{bmatrix} \rightarrow \begin{bmatrix} \boxed{1} & 2 & 3 & 4 \\ 0 & 0 & \boxed{1} & 2 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

The final matrix is the row echelon form of the matrix. Consider the pivots that we marked. The pivot columns are the first and the third column. All other columns correspond to free

variables when solving $A\vec{x} = \vec{0}$, so all other columns can be solved in terms of the first and the third column. In other words

$$\text{column space of } \begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix} = \text{span} \left\{ \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}, \begin{bmatrix} 2 \\ 4 \\ 6 \end{bmatrix}, \begin{bmatrix} 3 \\ 5 \\ 7 \end{bmatrix}, \begin{bmatrix} 4 \\ 6 \\ 8 \end{bmatrix} \right\} = \text{span} \left\{ \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}, \begin{bmatrix} 3 \\ 5 \\ 7 \end{bmatrix} \right\}.$$

We could perhaps use another pair of columns to get the same span, but the first and the third are guaranteed to work because they are pivot columns.

The discussion above could be expanded into a proof of the theorem if we wanted. As each nonzero row in the row echelon form contains a pivot, then the rank is the number of pivots, which is the same as the maximal number of linearly independent columns.

The idea also works in reverse. Suppose we have a bunch of column vectors and we just need to find a linearly independent set. For example, suppose we started with the vectors

$$\vec{v}_1 = \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}, \quad \vec{v}_2 = \begin{bmatrix} 2 \\ 4 \\ 6 \end{bmatrix}, \quad \vec{v}_3 = \begin{bmatrix} 3 \\ 5 \\ 7 \end{bmatrix}, \quad \vec{v}_4 = \begin{bmatrix} 4 \\ 6 \\ 8 \end{bmatrix}.$$

These vectors are not linearly independent as we saw above. In particular, the span of $\vec{v}_1$ and $\vec{v}_3$ is the same as the span of all four of the vectors. So $\vec{v}_2$ and $\vec{v}_4$ can both be written as linear combinations of $\vec{v}_1$ and $\vec{v}_3$. A common thing that comes up in practice is that one gets a set of vectors whose span is the set of solutions of some problem. But perhaps we get way too many vectors, and we want to simplify. For example above, all vectors in the span of $\vec{v}_1, \vec{v}_2, \vec{v}_3, \vec{v}_4$ can be written $\alpha_1\vec{v}_1 + \alpha_2\vec{v}_2 + \alpha_3\vec{v}_3 + \alpha_4\vec{v}_4$ for some numbers $\alpha_1, \alpha_2, \alpha_3, \alpha_4$. But it is also true that every such vector can be written as $a\vec{v}_1 + b\vec{v}_3$ for two numbers a and b . And one has to admit, that looks much simpler. Moreover, these numbers a and b are unique. More on that in the next section.

To find this linearly independent set, we simply take our vectors and form the matrix $[\vec{v}_1 \ \vec{v}_2 \ \vec{v}_3 \ \vec{v}_4]$, that is, the matrix

$$\begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix}.$$

We crank up the row-reduction machine, feed this matrix into it, find the pivot columns, and pick those. In this case, $\vec{v}_1$ and $\vec{v}_3$.

A.3.5 Computing the inverse

If the matrix A is square and there exists a unique solution $\vec{x}$ to $A\vec{x} = \vec{b}$ for any $\vec{b}$ (there are no free variables), then A is invertible. This is equivalent to the $n \times n$ matrix A being of rank n .

In particular, if $A\vec{x} = \vec{b}$ then $\vec{x} = A^{-1}\vec{b}$. Now we just need to compute what A^{-1} is. We can surely do elimination every time we want to find $A^{-1}\vec{b}$, but that would be ridiculous.

The mapping A^{-1} is linear and hence given by a matrix, and we have seen that to figure out the matrix we just need to find where A^{-1} takes the standard basis vectors $\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n$.

That is, to find the first column of A^{-1} , we solve $A\vec{x} = \vec{e}_1$, because then $A^{-1}\vec{e}_1 = \vec{x}$. To find the second column of A^{-1} , we solve $A\vec{x} = \vec{e}_2$. And so on. It is really just n eliminations that we need to do. But it gets even easier. If you think about it, the elimination is the same for everything on the left side of the augmented matrix. Doing n eliminations separately, we would redo most of the computations. Best is to do all at once.

Therefore, to find the inverse of A , we write an $n \times 2n$ augmented matrix $[A \mid I]$, where I is the identity matrix, whose columns are precisely the standard basis vectors. We then perform row reduction until we arrive at the reduced row echelon form. If A is invertible, then pivots can be found in every column of A , and so the reduced row echelon form of $[A \mid I]$ looks like $[I \mid A^{-1}]$. We then just read off the inverse A^{-1} . If you do not find a pivot in every one of the first n columns of the augmented matrix, then A is not invertible.

This is best seen by example. Suppose we wish to invert the matrix

$$\begin{bmatrix} 1 & 2 & 3 \\ 2 & 0 & 1 \\ 3 & 1 & 0 \end{bmatrix}.$$

We write the augmented matrix and we start reducing:

$$\begin{aligned} & \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 3 & 1 & 0 & 0 \\ 2 & 0 & 1 & 0 & 1 & 0 \\ 3 & 1 & 0 & 0 & 0 & 1 \end{array} \right] \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 3 & 1 & 0 & 0 \\ 0 & -4 & -5 & -2 & 1 & 0 \\ 0 & -5 & -9 & -3 & 0 & 1 \end{array} \right] \rightarrow \\ & \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 3 & 1 & 0 & 0 \\ 0 & \boxed{1} & 5/4 & 1/2 & -1/4 & 0 \\ 0 & -5 & -9 & -3 & 0 & 1 \end{array} \right] \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 3 & 1 & 0 & 0 \\ 0 & \boxed{1} & 5/4 & 1/2 & -1/4 & 0 \\ 0 & 0 & -11/4 & -1/2 & -5/4 & 1 \end{array} \right] \rightarrow \\ & \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 3 & 1 & 0 & 0 \\ 0 & \boxed{1} & 5/4 & 1/2 & -1/4 & 0 \\ 0 & 0 & \boxed{1} & 2/11 & 5/11 & -4/11 \end{array} \right] \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 2 & 0 & 5/11 & -15/11 & 12/11 \\ 0 & \boxed{1} & 0 & 3/11 & -9/11 & 5/11 \\ 0 & 0 & \boxed{1} & 2/11 & 5/11 & -4/11 \end{array} \right] \rightarrow \\ & \rightarrow \left[\begin{array}{ccc|ccc} \boxed{1} & 0 & 0 & -1/11 & 3/11 & 2/11 \\ 0 & \boxed{1} & 0 & 3/11 & -9/11 & 5/11 \\ 0 & 0 & \boxed{1} & 2/11 & 5/11 & -4/11 \end{array} \right]. \end{aligned}$$

So

$$\begin{bmatrix} 1 & 2 & 3 \\ 2 & 0 & 1 \\ 3 & 1 & 0 \end{bmatrix}^{-1} = \begin{bmatrix} -1/11 & 3/11 & 2/11 \\ 3/11 & -9/11 & 5/11 \\ 2/11 & 5/11 & -4/11 \end{bmatrix}.$$

Not too terrible, no? Perhaps harder than inverting a 2×2 matrix for which we had a simple formula, but not too bad. In practice, this is done efficiently by a computer.

A.3.6 Exercises

Exercise A.3.1: Compute the reduced row echelon form for the following matrices:

$$a) \begin{bmatrix} 1 & 3 & 1 \\ 0 & 1 & 1 \end{bmatrix}$$

$$b) \begin{bmatrix} 3 & 3 \\ 6 & -3 \end{bmatrix}$$

$$c) \begin{bmatrix} 3 & 6 \\ -2 & -3 \end{bmatrix}$$

$$d) \begin{bmatrix} 6 & 6 & 7 & 7 \\ 1 & 1 & 0 & 1 \end{bmatrix}$$

$$e) \begin{bmatrix} 9 & 3 & 0 & 2 \\ 8 & 6 & 3 & 6 \\ 7 & 9 & 7 & 9 \end{bmatrix}$$

$$f) \begin{bmatrix} 2 & 1 & 3 & -3 \\ 6 & 0 & 0 & -1 \\ -2 & 4 & 4 & 3 \end{bmatrix}$$

$$g) \begin{bmatrix} 6 & 6 & 5 \\ 0 & -2 & 2 \\ 6 & 5 & 6 \end{bmatrix}$$

$$h) \begin{bmatrix} 0 & 2 & 0 & -1 \\ 6 & 6 & -3 & 3 \\ 6 & 2 & -3 & 5 \end{bmatrix}$$

Exercise A.3.2: Compute the inverse of the given matrices

$$a) \begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix}$$

$$b) \begin{bmatrix} 1 & 1 & 1 \\ 0 & 2 & 1 \\ 0 & 0 & 1 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 2 & 3 \\ 2 & 0 & 1 \\ 0 & 2 & 1 \end{bmatrix}$$

Exercise A.3.3: Solve (find all solutions), or show no solution exists

$$a) \begin{aligned} 4x_1 + 3x_2 &= -2 \\ -x_1 + x_2 &= 4 \end{aligned}$$

$$x_1 + 5x_2 + 3x_3 = 7$$

$$b) 8x_1 + 7x_2 + 8x_3 = 8$$

$$4x_1 + 8x_2 + 6x_3 = 4$$

$$\begin{aligned} 4x_1 + 8x_2 + 2x_3 &= 3 \\ c) -x_1 - 2x_2 + 3x_3 &= 1 \\ 4x_1 + 8x_2 &= 2 \end{aligned}$$

$$x + 2y + 3z = 4$$

$$d) 2x - y + 3z = 1$$

$$3x + y + 6z = 6$$

Exercise A.3.4: By computing the inverse, solve the following systems for $\vec{x}$.

$$a) \begin{bmatrix} 4 & 1 \\ -1 & 3 \end{bmatrix} \vec{x} = \begin{bmatrix} 13 \\ 26 \end{bmatrix}$$

$$b) \begin{bmatrix} 3 & 3 \\ 3 & 4 \end{bmatrix} \vec{x} = \begin{bmatrix} 2 \\ -1 \end{bmatrix}$$

Exercise A.3.5: Compute the rank of the given matrices

$$a) \begin{bmatrix} 6 & 3 & 5 \\ 1 & 4 & 1 \\ 7 & 7 & 6 \end{bmatrix}$$

$$b) \begin{bmatrix} 5 & -2 & -1 \\ 3 & 0 & 6 \\ 2 & 4 & 5 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 2 & 3 \\ -1 & -2 & -3 \\ 2 & 4 & 6 \end{bmatrix}$$

Exercise A.3.6: For the matrices in [Exercise A.3.5](#), find a linearly independent set of row vectors that span the row space (they do not need to be rows of the matrix).

Exercise A.3.7: For the matrices in [Exercise A.3.5](#), find a linearly independent set of columns that span the column space. That is, find the pivot columns of the matrices.

Exercise A.3.8: Find a linearly independent subset of the following vectors that has the same span.

$$\begin{bmatrix} -1 \\ 1 \\ 2 \end{bmatrix}, \begin{bmatrix} 2 \\ -2 \\ -4 \end{bmatrix}, \begin{bmatrix} -2 \\ 4 \\ 1 \end{bmatrix}, \begin{bmatrix} -1 \\ 3 \\ -2 \end{bmatrix}$$

Exercise A.3.101: Compute the reduced row echelon form for the following matrices:

$$\begin{array}{llll}
 a) \begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix} & b) \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} & c) \begin{bmatrix} 1 & 1 \\ -2 & -2 \end{bmatrix} & d) \begin{bmatrix} 1 & -3 & 1 \\ 4 & 6 & -2 \\ -2 & 6 & -2 \end{bmatrix} \\
 e) \begin{bmatrix} 2 & 2 & 5 & 2 \\ 1 & -2 & 4 & -1 \\ 0 & 3 & 1 & -2 \end{bmatrix} & f) \begin{bmatrix} -2 & 6 & 4 & 3 \\ 6 & 0 & -3 & 0 \\ 4 & 2 & -1 & 1 \end{bmatrix} & g) \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} & h) \begin{bmatrix} 1 & 2 & 3 & 3 \\ 1 & 2 & 3 & 5 \end{bmatrix}
 \end{array}$$

Exercise A.3.102: Compute the inverse of the given matrices

$$\begin{array}{lll}
 a) \begin{bmatrix} 0 & 1 & 0 \\ -1 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix} & b) \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 0 \\ 1 & 0 & 0 \end{bmatrix} & c) \begin{bmatrix} 2 & 4 & 0 \\ 2 & 2 & 3 \\ 2 & 4 & 1 \end{bmatrix}
 \end{array}$$

Exercise A.3.103: Solve (find all solutions), or show no solution exists

$$\begin{array}{ll}
 a) \begin{array}{l} 4x_1 + 3x_2 = -1 \\ 5x_1 + 6x_2 = 4 \end{array} & \begin{array}{l} 5x + 6y + 5z = 7 \\ 6x + 8y + 6z = -1 \\ 5x + 2y + 5z = 2 \end{array} \\
 \begin{array}{l} a + b + c = -1 \\ a + 5b + 6c = -1 \\ -2a + 5b + 6c = 8 \end{array} & d) \begin{array}{l} -2x_1 + 2x_2 + 8x_3 = 6 \\ x_2 + x_3 = 2 \\ x_1 + 4x_2 + x_3 = 7 \end{array}
 \end{array}$$

Exercise A.3.104: By computing the inverse, solve the following systems for $\vec{x}$.

$$\begin{array}{ll}
 a) \begin{bmatrix} -1 & 1 \\ 3 & 3 \end{bmatrix} \vec{x} = \begin{bmatrix} 4 \\ 6 \end{bmatrix} & b) \begin{bmatrix} 2 & 7 \\ 1 & 6 \end{bmatrix} \vec{x} = \begin{bmatrix} 1 \\ 3 \end{bmatrix}
 \end{array}$$

Exercise A.3.105: Compute the rank of the given matrices

$$\begin{array}{lll}
 a) \begin{bmatrix} 7 & -1 & 6 \\ 7 & 7 & 7 \\ 7 & 6 & 2 \end{bmatrix} & b) \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 2 & 2 & 2 \end{bmatrix} & c) \begin{bmatrix} 0 & 3 & -1 \\ 6 & 3 & 1 \\ 4 & 7 & -1 \end{bmatrix}
 \end{array}$$

Exercise A.3.106: For the matrices in [Exercise A.3.105](#), find a linearly independent set of row vectors that span the row space (they do not need to be rows of the matrix).

Exercise A.3.107: For the matrices in [Exercise A.3.105](#), find a linearly independent set of columns that span the column space. That is, find the pivot columns of the matrices.

Exercise A.3.108: Find a linearly independent subset of the following vectors that has the same span.

$$\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, \quad \begin{bmatrix} 3 \\ 1 \\ -5 \end{bmatrix}, \quad \begin{bmatrix} 0 \\ 3 \\ -1 \end{bmatrix}, \quad \begin{bmatrix} -3 \\ 2 \\ 4 \end{bmatrix}$$

A.4 Subspaces, dimension, and the kernel

Note: 1 lecture

A.4.1 Subspaces, basis, and dimension

We often find ourselves looking at the set of solutions of a linear equation $L\vec{x} = \vec{0}$ for some matrix L . That is, we are interested in the kernel of L . The set of all such solutions has a nice structure: It looks and acts a lot like some euclidean space $\mathbb{R}^k$.

We say that a set S of vectors in $\mathbb{R}^n$ is a *subspace* if whenever $\vec{x}$ and $\vec{y}$ are members of S and α is a scalar, then

$$\vec{x} + \vec{y}, \quad \text{and} \quad \alpha\vec{x}$$

are also members of S . That is, we can add and multiply by scalars and we still land in S . So every linear combination of vectors of S is still in S . That is really what a subspace is. It is a subset where we can take linear combinations and still end up being in the subset. Consequently, the span of a number of vectors is automatically a subspace.

Example A.4.1:

- (i) If we let $S = \mathbb{R}^n$, then this S is a subspace of $\mathbb{R}^n$. Adding any two vectors in $\mathbb{R}^n$ gets a vector in $\mathbb{R}^n$, and so does multiplying by scalars.
- (ii) The set $S' = \{\vec{0}\}$, that is, the set of the zero vector by itself, is also a subspace of $\mathbb{R}^n$. There is only one vector in this subspace, so we only need to verify the definition for that one vector, and everything checks out: $\vec{0} + \vec{0} = \vec{0}$ and $\alpha\vec{0} = \vec{0}$.
- (iii) The set S'' of all the vectors of the form (a, a) for any real number a , such as $(1, 1)$, $(3, 3)$, or $(-0.5, -0.5)$, is a subspace of $\mathbb{R}^2$. Adding two such vectors, say $(1, 1) + (3, 3) = (4, 4)$ again gets a vector of the same form, and so does multiplying by a scalar, say $8(1, 1) = (8, 8)$.

If S is a subspace and we can find k linearly independent vectors in S

$$\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k,$$

such that every other vector in S is a linear combination of $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k$, then the set $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k\}$ is called a *basis* of S . In other words, S is the span of $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k\}$ and there is no smaller subset of these vectors that spans S . We say that S has dimension k , and we write

$$\dim S = k.$$

We have the following theorem.

Theorem A.4.1. *If $S \subset \mathbb{R}^n$ is a subspace and S is not the trivial subspace $\{\vec{0}\}$, then there exists a unique positive integer k (the dimension) and a (not unique) basis $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k\}$, such that every $\vec{w}$ in S can be uniquely represented by*

$$\vec{w} = \alpha_1 \vec{v}_1 + \alpha_2 \vec{v}_2 + \dots + \alpha_k \vec{v}_k,$$

for some scalars $\alpha_1, \alpha_2, \dots, \alpha_k$. By “uniquely represented” we mean that these scalars $\alpha_1, \alpha_2, \dots, \alpha_k$ are unique.

Just as a vector in $\mathbb{R}^k$ is represented by a k -tuple of numbers, so is a vector in a k -dimensional subspace of $\mathbb{R}^n$ represented by a unique k -tuple of numbers, at least once we have fixed a basis. A different basis would give a different k -tuple of numbers for the same vector.

We should reiterate that while k is unique (a subspace cannot have two different dimensions—every basis has the same number of vectors), the set of basis vectors is not at all unique. There are lots of different bases for any given subspace. Finding just the right basis for a subspace is a large part of what one does in linear algebra. In fact, that is what we spend a lot of time on in linear differential equations, although at first glance it may not seem like that is what we are doing.

Example A.4.2:

- (i) The standard basis

$$\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n,$$

is a basis of $\mathbb{R}^n$, (hence the name). So as expected

$$\dim \mathbb{R}^n = n.$$

- (ii) Both sets $\{(1, 0), (0, 1)\}$ and $\{(1, 1), (1, -1)\}$ are bases of $\mathbb{R}^2$. A vector $(3, 7)$ can be written as $(3, 7) = 3(1, 0) + 7(0, 1)$ (and in no other way with the first basis), and it can be written as $(3, 7) = 5(1, 1) - 2(1, -1)$ (and in no other way with the second basis).

- (iii) The subspace $\{\vec{0}\}$ is of dimension 0. A basis is simply an empty set of vectors.

- (iv) The subspace S'' from [Example A.4.1](#), that is, the set of vectors (a, a) , is of dimension 1. One possible basis is simply $\{(1, 1)\}$, the single vector $(1, 1)$: Every vector in S'' can be represented by $a(1, 1) = (a, a)$. Similarly, another possible basis would be $\{(-1, -1)\}$. Then the vector (a, a) would be represented as $(-a)(-1, -1)$.

- (v) The set $\{(1, 1), (-1, -1)\}$ does span S'' , but it is not a basis as it is not a linearly independent set. Consequently, a vector such as $(2, 2)$ can be written in multiple ways with this set of vectors: $(2, 2) = 2(1, 1) + 0(-1, -1)$, or $(2, 2) = 5(1, 1) + 3(-1, -1)$, etc.

Row and column spaces of a matrix are also examples of subspaces, as they are given as the span of vectors. We can use what we know about rank, row spaces, and column spaces from the previous section to find a basis.

Example A.4.3: In the last section, we considered the matrix

$$A = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix}.$$

Using row reduction to find the pivot columns, we found

$$\text{column space of } A \left(\begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 4 & 5 & 6 \\ 3 & 6 & 7 & 8 \end{bmatrix} \right) = \text{span} \left\{ \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}, \begin{bmatrix} 3 \\ 5 \\ 7 \end{bmatrix} \right\}.$$

What we did was we found a basis of the column space. This basis has two elements, and so the column space of A is two-dimensional. Notice that the rank of A is two.

We would have followed the same procedure if we wanted to find a basis of the subspace X spanned by

$$\begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}, \begin{bmatrix} 2 \\ 4 \\ 6 \end{bmatrix}, \begin{bmatrix} 3 \\ 5 \\ 7 \end{bmatrix}, \begin{bmatrix} 4 \\ 6 \\ 8 \end{bmatrix}.$$

We would have simply formed the matrix A with these vectors as columns and repeated the computation above. The subspace X is then the column space of A .

Example A.4.4: Consider the matrix

$$L = \begin{bmatrix} 1 & 2 & 0 & 0 & 3 \\ 0 & 0 & 1 & 0 & 4 \\ 0 & 0 & 0 & 1 & 5 \end{bmatrix}.$$

Conveniently, the matrix is in reduced row echelon form. The matrix is of rank 3. The column space is the span of the pivot columns. It is the 3-dimensional space

$$\text{column space of } L = \text{span} \left\{ \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} \right\} = \mathbb{R}^3.$$

The row space is the 3-dimensional space

$$\text{row space of } L = \text{span} \{ [1 \ 2 \ 0 \ 0 \ 3], [0 \ 0 \ 1 \ 0 \ 4], [0 \ 0 \ 0 \ 1 \ 5] \}.$$

As these vectors have 5 components, we think of the row space of L as a subspace of $\mathbb{R}^5$.

The way the dimensions worked out in the examples is not an accident. Since the number of vectors that we needed to take was always the same as the number of pivots, and the number of pivots is the rank, we get the following result.

Theorem A.4.2 (Rank). *The dimension of the column space and the dimension of the row space of a matrix A are both equal to the rank of A .*

A.4.2 Kernel

The set of solutions of a linear equation $L\vec{x} = \vec{0}$, the kernel of L , is a subspace: If $\vec{x}$ and $\vec{y}$ are solutions, then

$$L(\vec{x} + \vec{y}) = L\vec{x} + L\vec{y} = \vec{0} + \vec{0} = \vec{0}, \quad \text{and} \quad L(\alpha\vec{x}) = \alpha L\vec{x} = \alpha\vec{0} = \vec{0}.$$

So $\vec{x} + \vec{y}$ and $\alpha\vec{x}$ are solutions. The dimension of the kernel is called the *nullity* of the matrix.

The same sort of idea governs the solutions of linear differential equations. We try to describe the kernel of a linear differential operator, and as it is a subspace, we look for a basis of this kernel. Much of this book is dedicated to finding such bases.

The kernel of a matrix is the same as the kernel of its reduced row echelon form. For a matrix in reduced row echelon form, the kernel is rather easy to find. If we take a product of a matrix L and a vector $\vec{x}$, then each entry in $\vec{x}$ corresponds to a column of L , the column that the entry multiplies. To find the kernel, pick a non-pivot column and make a vector that has a -1 in the entry corresponding to this non-pivot column and zeros at all the other entries corresponding to the other non-pivot columns. Then for all the entries corresponding to pivot columns make it precisely the value in the corresponding row of the non-pivot column to make the vector be a solution to $L\vec{x} = \vec{0}$. This procedure is best understood by example.

Example A.4.5: Consider

$$L = \begin{bmatrix} \boxed{1} & 2 & 0 & 0 & 3 \\ 0 & 0 & \boxed{1} & 0 & 4 \\ 0 & 0 & 0 & \boxed{1} & 5 \end{bmatrix}.$$

This matrix is in reduced row echelon form, the pivots are marked. There are two non-pivot columns, so the kernel has dimension 2, that is, it is the span of 2 vectors. Let us find the first vector. We look at the first non-pivot column, the 2nd column, and we put a -1 in the 2nd entry of our vector. We put a 0 in the 5th entry as the 5th column is also a non-pivot column:

$$\begin{bmatrix} ? \\ -1 \\ ? \\ ? \\ 0 \end{bmatrix}.$$

Let us fill the rest. When this vector hits the first row of L , we get a -2 and 1 times whatever the first question mark is. We want to get zero when hitting the first row, so make the first question mark 2. To get zero when hitting the second and third rows, it is sufficient to make the remaining question marks 0. We are really filling in the non-pivot column into

the remaining entries. Let us check while marking which numbers went where:

$$\begin{bmatrix} 1 & \boxed{2} & 0 & 0 & 3 \\ 0 & \boxed{0} & 1 & 0 & 4 \\ 0 & \boxed{0} & 0 & 1 & 5 \end{bmatrix} \begin{bmatrix} \boxed{2} \\ -1 \\ \boxed{0} \\ \boxed{0} \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}.$$

Yay! How about the second vector? We start with

$$\begin{bmatrix} ? \\ 0 \\ ? \\ ? \\ -1 \end{bmatrix}$$

We set the first question mark to 3, the second to 4, and the third to 5. Let us check, marking things as previously,

$$\begin{bmatrix} 1 & 2 & 0 & 0 & \boxed{3} \\ 0 & 0 & 1 & 0 & \boxed{4} \\ 0 & 0 & 0 & 1 & \boxed{5} \end{bmatrix} \begin{bmatrix} \boxed{3} \\ 0 \\ \boxed{4} \\ \boxed{5} \\ -1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}.$$

There are two non-pivot columns, so we only need two vectors. We have found a basis of the kernel. So,

$$\text{kernel of } L = \text{span} \left\{ \begin{bmatrix} 2 \\ -1 \\ 0 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 3 \\ 0 \\ 4 \\ 5 \\ -1 \end{bmatrix} \right\}$$

What we did in finding a basis of the kernel is that we expressed all solutions of $L\vec{x} = \vec{0}$ as a linear combination of some given vectors.

The procedure to find a basis of the kernel of a matrix L :

- (i) Find the reduced row echelon form of L .
- (ii) Write down a basis of the kernel as above, one vector for each non-pivot column.

The rank of a matrix is the dimension of the column space, and that is the span of the pivot columns, while the kernel is the span of certain vectors, one for each non-pivot column. So the two numbers must add to the number of columns.

Theorem A.4.3 (Rank–Nullity). *If a matrix A has n columns, rank r , and nullity k (dimension of the kernel), then*

$$n = r + k.$$

The theorem is immensely useful in applications. It allows one to compute the rank r if one knows the nullity k and vice versa, without doing any extra work. An example application is a simple version of the so-called *Fredholm alternative*. A similar result is true for differential equations. Consider

$$A\vec{x} = \vec{b},$$

where A is a square $n \times n$ matrix. There are then two mutually exclusive possibilities:

- (i) A nonzero solution $\vec{x}$ to $A\vec{x} = \vec{0}$ exists.
- (ii) The equation $A\vec{x} = \vec{b}$ has a unique solution $\vec{x}$ for every $\vec{b}$.

How does the Rank–Nullity theorem come into the picture? Well, if A has a nonzero solution $\vec{x}$ to $A\vec{x} = \vec{0}$, then the nullity k is positive. But then the rank $r = n - k$ must be less than n . It means that the column space of A is of dimension less than n , so it is a subspace that does not include everything in $\mathbb{R}^n$. So $\mathbb{R}^n$ has to contain some vector $\vec{b}$ not in the column space of A . In fact, most vectors in $\mathbb{R}^n$ are not in the column space of A .

A.4.3 Exercises

Exercise A.4.1: For the following sets of vectors, find a basis for the subspace spanned by the vectors, and find the dimension of the subspace.

a) $\begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}, \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix}$

b) $\begin{bmatrix} 1 \\ 0 \\ 5 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ -1 \\ 0 \end{bmatrix}$

c) $\begin{bmatrix} -4 \\ -3 \\ 5 \end{bmatrix}, \begin{bmatrix} 2 \\ 3 \\ 3 \end{bmatrix}, \begin{bmatrix} 2 \\ 0 \\ 2 \end{bmatrix}$

d) $\begin{bmatrix} 1 \\ 3 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 2 \\ 2 \end{bmatrix}, \begin{bmatrix} -1 \\ -1 \\ 2 \end{bmatrix}$

e) $\begin{bmatrix} 1 \\ 3 \\ 3 \end{bmatrix}, \begin{bmatrix} 0 \\ 2 \\ 2 \end{bmatrix}, \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix}$

f) $\begin{bmatrix} 3 \\ 1 \\ 3 \end{bmatrix}, \begin{bmatrix} 2 \\ 4 \\ -4 \end{bmatrix}, \begin{bmatrix} -5 \\ -5 \\ -2 \end{bmatrix}$

Exercise A.4.2: For the following matrices, find a basis for the kernel (nullspace).

a) $\begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 5 \\ 1 & 1 & -4 \end{bmatrix}$

b) $\begin{bmatrix} 2 & -1 & -3 \\ 4 & 0 & -4 \\ -1 & 1 & 2 \end{bmatrix}$

c) $\begin{bmatrix} -4 & 4 & 4 \\ -1 & 1 & 1 \\ -5 & 5 & 5 \end{bmatrix}$

d) $\begin{bmatrix} -2 & 1 & 1 & 1 \\ -4 & 2 & 2 & 2 \\ 1 & 0 & 4 & 3 \end{bmatrix}$

Exercise A.4.3: Suppose a 5×5 matrix A has rank 3. What is the nullity?

Exercise A.4.4: Suppose that X is the set of all the vectors of $\mathbb{R}^3$ whose third component is zero. Is X a subspace? And if so, find a basis and the dimension.

Exercise A.4.5: Consider a square matrix A , and suppose that $\vec{x}$ is a nonzero vector such that $A\vec{x} = \vec{0}$. What does the Fredholm alternative say about invertibility of A ?

Exercise A.4.6: Consider

$$M = \begin{bmatrix} 1 & 2 & 3 \\ 2 & ? & ? \\ -1 & ? & ? \end{bmatrix}.$$

If the nullity of this matrix is 2, fill in the question marks. Hint: What is the rank?

Exercise A.4.101: For the following sets of vectors, find a basis for the subspace spanned by the vectors, and find the dimension of the subspace.

$$\begin{array}{lll} a) \begin{bmatrix} 1 \\ 2 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \end{bmatrix} & b) \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}, \begin{bmatrix} 2 \\ 2 \\ 2 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix} & c) \begin{bmatrix} 5 \\ 3 \\ 1 \end{bmatrix}, \begin{bmatrix} 5 \\ -1 \\ 5 \end{bmatrix}, \begin{bmatrix} -1 \\ 3 \\ -4 \end{bmatrix} \\ d) \begin{bmatrix} 2 \\ 2 \\ 4 \end{bmatrix}, \begin{bmatrix} 2 \\ 2 \\ 3 \end{bmatrix}, \begin{bmatrix} 4 \\ 4 \\ -3 \end{bmatrix} & e) \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \begin{bmatrix} 2 \\ 0 \end{bmatrix}, \begin{bmatrix} 3 \\ 0 \end{bmatrix} & f) \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 2 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 2 \end{bmatrix} \end{array}$$

Exercise A.4.102: For the following matrices, find a basis for the kernel (nullspace).

$$\begin{array}{llll} a) \begin{bmatrix} 2 & 6 & 1 & 9 \\ 1 & 3 & 2 & 9 \\ 3 & 9 & 0 & 9 \end{bmatrix} & b) \begin{bmatrix} 2 & -2 & -5 \\ -1 & 1 & 5 \\ -5 & 5 & -3 \end{bmatrix} & c) \begin{bmatrix} 1 & -5 & -4 \\ 2 & 3 & 5 \\ -3 & 5 & 2 \end{bmatrix} & d) \begin{bmatrix} 0 & 4 & 4 \\ 0 & 1 & 1 \\ 0 & 5 & 5 \end{bmatrix} \end{array}$$

Exercise A.4.103: Suppose the column space of a 9×5 matrix A is of dimension 3. Find

- a) Rank of A .
- b) Nullity of A .
- c) Dimension of the row space of A .
- d) Dimension of the nullspace of A .
- e) Size of the maximum subset of linearly independent rows of A .

A.5 Inner product and projections

Note: 1–2 lectures

A.5.1 Inner product and orthogonality

To do basic geometry, we need length, and we need angles. We have already seen the euclidean length, so let us figure out angles. Mostly, we are worried about the right angle*.

Given two (column) vectors in $\mathbb{R}^n$, we define the (standard) *inner product* as the dot product:

$$\langle \vec{x}, \vec{y} \rangle = \vec{x} \cdot \vec{y} = \vec{y}^T \vec{x} = x_1 y_1 + x_2 y_2 + \cdots + x_n y_n = \sum_{i=1}^n x_i y_i.$$

Why do we seemingly give a new notation for the dot product ($\langle \vec{x}, \vec{y} \rangle$ instead of just $\vec{x} \cdot \vec{y}$)? Because there are other possible inner products, which are not the dot product, although we will not worry about others here. An inner product can even be defined on spaces of functions, as we do in [chapter 4](#):

$$\langle f(t), g(t) \rangle = \int_a^b f(t)g(t) dt.$$

But we digress.

The inner product satisfies the following rules:

- (i) $\langle \vec{x}, \vec{x} \rangle \geq 0$, and $\langle \vec{x}, \vec{x} \rangle = 0$ if and only if $\vec{x} = 0$,
- (ii) $\langle \vec{x}, \vec{y} \rangle = \langle \vec{y}, \vec{x} \rangle$,
- (iii) $\langle a\vec{x}, \vec{y} \rangle = \langle \vec{x}, a\vec{y} \rangle = a\langle \vec{x}, \vec{y} \rangle$,
- (iv) $\langle \vec{x} + \vec{y}, \vec{z} \rangle = \langle \vec{x}, \vec{z} \rangle + \langle \vec{y}, \vec{z} \rangle$ and $\langle \vec{x}, \vec{y} + \vec{z} \rangle = \langle \vec{x}, \vec{y} \rangle + \langle \vec{x}, \vec{z} \rangle$.

Anything that satisfies the properties above can be called an inner product, although in this section we are concerned with the standard inner product in $\mathbb{R}^n$.

The standard inner product gives the euclidean length:

$$\|\vec{x}\| = \sqrt{\langle \vec{x}, \vec{x} \rangle} = \sqrt{x_1^2 + x_2^2 + \cdots + x_n^2}.$$

How does it give angles? You may recall a formula for the standard inner product (the dot product) from multivariable calculus in two or three dimensions in terms of the angle θ between the vectors:

$$\langle \vec{x}, \vec{y} \rangle = \|\vec{x}\| \|\vec{y}\| \cos \theta.$$

That is, θ is the angle that $\vec{x}$ and $\vec{y}$ make when they are based at the same point.

*When Euclid defined angles in his *Elements*, the only angle he ever really defined was the right angle.

In $\mathbb{R}^n$ (any dimension), we are simply going to say that θ from the formula is what the angle is. This makes sense as any two vectors based at the origin lie in a 2-dimensional plane (subspace), and the formula works in 2 dimensions. In fact, one could even talk about angles between functions this way, and we do in [chapter 4](#), where we talk about orthogonal functions (functions at right angle to each other).

To compute the angle, we solve for $\cos \theta$:

$$\cos \theta = \frac{\langle \vec{x}, \vec{y} \rangle}{\|\vec{x}\| \|\vec{y}\|}.$$

Our angles are always in radians. We are computing the cosine of the angle, which is really the best we can do. Given two vectors at an angle θ , we can give the angle as $-\theta$, $2\pi - \theta$, etc. See [Figure A.5](#). Fortunately, $\cos \theta = \cos(-\theta) = \cos(2\pi - \theta)$. If we solve for θ using the inverse cosine $\cos^{-1}$, we can just decree that $0 \leq \theta \leq \pi$.

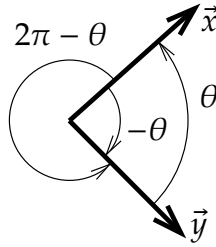


Figure A.5: Angle between vectors.

Example A.5.1: Let us compute the angle between the vectors $(3, 0)$ and $(1, 1)$ in the plane. Compute

$$\cos \theta = \frac{\langle (3, 0), (1, 1) \rangle}{\|(3, 0)\| \|(1, 1)\|} = \frac{3 + 0}{3\sqrt{2}} = \frac{1}{\sqrt{2}}.$$

Therefore $\theta = \pi/4$.

As we said, the most important angle is the right angle. A right angle is $\pi/2$ radians, and $\cos(\pi/2) = 0$, so the formula is particularly easy in this case. We say vectors $\vec{x}$ and $\vec{y}$ are *orthogonal* if they are at right angles, that is, if

$$\langle \vec{x}, \vec{y} \rangle = 0.$$

The vectors $(1, 0, 0, 1)$ and $(1, 2, 3, -1)$ are orthogonal. So are $(1, 1)$ and $(1, -1)$. However, $(1, 1)$ and $(1, 2)$ are not orthogonal as their inner product is 3 and not 0.

A.5.2 Orthogonal projection

A typical application of linear algebra is to take a difficult problem, write everything in the right basis, and in this new basis the problem becomes simple. A particularly useful

basis is an orthogonal basis, a basis where all the basis vectors are orthogonal. When we draw a coordinate system in two or three dimensions, we almost always draw our axes as orthogonal to each other.

Generalizing this concept to functions, it is particularly useful in [chapter 4](#) to express a function using a particular orthogonal basis, the Fourier series.

To express one vector in terms of an orthogonal basis, we need to first *project* one vector onto another. Given a nonzero vector $\vec{v}$, we define the *orthogonal projection* of $\vec{w}$ onto $\vec{v}$ as

$$\text{proj}_{\vec{v}}(\vec{w}) = \left(\frac{\langle \vec{w}, \vec{v} \rangle}{\langle \vec{v}, \vec{v} \rangle} \right) \vec{v}.$$

For the geometric idea, see [Figure A.6](#). That is, we find the “shadow of $\vec{w}$ ” on the line spanned by $\vec{v}$ if the direction of the sun’s rays were exactly perpendicular to the line. Another way of thinking about it is that the tip of the arrow of $\text{proj}_{\vec{v}}(\vec{w})$ is the closest point on the line spanned by $\vec{v}$ to the tip of the arrow of $\vec{w}$. In terms of euclidean distance, $\vec{u} = \text{proj}_{\vec{v}}(\vec{w})$ minimizes the distance $\|\vec{w} - \vec{u}\|$ among all vectors $\vec{u}$ that are multiples of $\vec{v}$. Because of this, this projection comes up often in applied mathematics in all sorts of contexts where we cannot solve a problem exactly: We cannot always solve “Find $\vec{w}$ as a multiple of $\vec{v}$,” but $\text{proj}_{\vec{v}}(\vec{w})$ is the best “solution.”

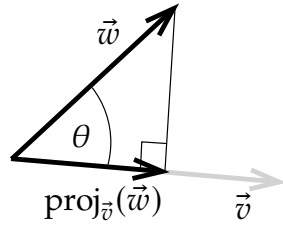


Figure A.6: Orthogonal projection.

The formula follows from basic trigonometry. If θ is acute, the length of $\text{proj}_{\vec{v}}(\vec{w})$ should be $\cos \theta$ times the length of $\vec{w}$, that is, $(\cos \theta)\|\vec{w}\|$. We take the unit vector in the direction of $\vec{v}$, that is, $\frac{\vec{v}}{\|\vec{v}\|}$ and we multiply it by the length of the projection. In other words,

$$\text{proj}_{\vec{v}}(\vec{w}) = (\cos \theta)\|\vec{w}\|\frac{\vec{v}}{\|\vec{v}\|} = \frac{(\cos \theta)\|\vec{w}\|\|\vec{v}\|}{\|\vec{v}\|^2}\vec{v} = \frac{\langle \vec{w}, \vec{v} \rangle}{\langle \vec{v}, \vec{v} \rangle}\vec{v}.$$

The computation works for obtuse θ ($\cos \theta < 0$) as we need to also switch direction.

Example A.5.2: Suppose we wish to project the vector $(3, 2, 1)$ onto the vector $(1, 2, 3)$. Compute

$$\begin{aligned} \text{proj}_{(1,2,3)}((3, 2, 1)) &= \frac{\langle (3, 2, 1), (1, 2, 3) \rangle}{\langle (1, 2, 3), (1, 2, 3) \rangle}(1, 2, 3) = \frac{3 \cdot 1 + 2 \cdot 2 + 1 \cdot 3}{1 \cdot 1 + 2 \cdot 2 + 3 \cdot 3}(1, 2, 3) \\ &= \frac{10}{14}(1, 2, 3) = \left(\frac{5}{7}, \frac{10}{7}, \frac{15}{7} \right). \end{aligned}$$

We double check that the projection is orthogonal. That is, $\vec{w} - \text{proj}_{\vec{v}}(\vec{w})$ ought to be orthogonal to $\vec{v}$. See the right angle in [Figure A.6](#) on the preceding page. In other words,

$$(3, 2, 1) - \text{proj}_{(1,2,3)}((3, 2, 1)) = \left(3 - \frac{5}{7}, 2 - \frac{10}{7}, 1 - \frac{15}{7}\right) = \left(\frac{16}{7}, \frac{4}{7}, \frac{-8}{7}\right)$$

ought to be orthogonal to $(1, 2, 3)$. The inner product had better be zero:

$$\left\langle \left(\frac{16}{7}, \frac{4}{7}, \frac{-8}{7}\right), (1, 2, 3) \right\rangle = \frac{16}{7} \cdot 1 + \frac{4}{7} \cdot 2 - \frac{8}{7} \cdot 3 = 0.$$

A.5.3 Orthogonal basis

As we said, a basis $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ is an *orthogonal basis* if all vectors in the basis are orthogonal to each other, that is, if

$$\langle \vec{v}_j, \vec{v}_k \rangle = 0$$

for all choices of j and k where $j \neq k$ (a nonzero vector cannot be orthogonal to itself). A basis is furthermore called an *orthonormal basis* if all the vectors in a basis are also unit vectors, that is, if all the vectors have magnitude 1. For example, the standard basis $\{(1, 0, 0), (0, 1, 0), (0, 0, 1)\}$ is an orthonormal basis of $\mathbb{R}^3$: Any pair is orthogonal, and each vector is of unit magnitude.

The reason why we are interested in orthogonal (or orthonormal) bases is that they make it really simple to represent a vector (or a projection onto a subspace) in the basis. The formula for the orthogonal projection onto a vector gives us the coefficients. In [chapter 4](#), we use the same idea by finding the correct orthogonal basis for the set of solutions of a differential equation. We are then able to find any particular solution by simply applying the orthogonal projection formula, which is just a couple of inner products.

Let us come back to linear algebra. Consider a subspace S and an orthogonal basis $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ for S . We wish to express some vector $\vec{x}$ in terms of this basis. If $\vec{x}$ is not in the span of this basis (when $\vec{x}$ is not in S), then of course it is not possible, but the following formula gives us at least the orthogonal projection onto the subspace S , or in other words, the best approximation to $\vec{x}$ in the subspace—the vector in S closest to $\vec{x}$.

First suppose that $\vec{x}$ is in the span. Then it is the sum of the orthogonal projections:

$$\vec{x} = \text{proj}_{\vec{v}_1}(\vec{x}) + \text{proj}_{\vec{v}_2}(\vec{x}) + \dots + \text{proj}_{\vec{v}_n}(\vec{x}) = \frac{\langle \vec{x}, \vec{v}_1 \rangle}{\langle \vec{v}_1, \vec{v}_1 \rangle} \vec{v}_1 + \frac{\langle \vec{x}, \vec{v}_2 \rangle}{\langle \vec{v}_2, \vec{v}_2 \rangle} \vec{v}_2 + \dots + \frac{\langle \vec{x}, \vec{v}_n \rangle}{\langle \vec{v}_n, \vec{v}_n \rangle} \vec{v}_n.$$

In other words, if we want to write $\vec{x} = a_1 \vec{v}_1 + a_2 \vec{v}_2 + \dots + a_n \vec{v}_n$, then

$$a_1 = \frac{\langle \vec{x}, \vec{v}_1 \rangle}{\langle \vec{v}_1, \vec{v}_1 \rangle}, \quad a_2 = \frac{\langle \vec{x}, \vec{v}_2 \rangle}{\langle \vec{v}_2, \vec{v}_2 \rangle}, \quad \dots, \quad a_n = \frac{\langle \vec{x}, \vec{v}_n \rangle}{\langle \vec{v}_n, \vec{v}_n \rangle}.$$

Another way to derive this formula is to work in reverse. Suppose that $\vec{x} = a_1 \vec{v}_1 + a_2 \vec{v}_2 + \dots + a_n \vec{v}_n$. Take an inner product with $\vec{v}_j$, and use the properties of the inner product:

$$\begin{aligned} \langle \vec{x}, \vec{v}_j \rangle &= \langle a_1 \vec{v}_1 + a_2 \vec{v}_2 + \dots + a_n \vec{v}_n, \vec{v}_j \rangle \\ &= a_1 \langle \vec{v}_1, \vec{v}_j \rangle + a_2 \langle \vec{v}_2, \vec{v}_j \rangle + \dots + a_n \langle \vec{v}_n, \vec{v}_j \rangle. \end{aligned}$$

As the basis is orthogonal, $\langle \vec{v}_k, \vec{v}_j \rangle = 0$ whenever $k \neq j$. That means that only one of the terms, the j^{th} one, on the right-hand side is nonzero and we get

$$\langle \vec{x}, \vec{v}_j \rangle = a_j \langle \vec{v}_j, \vec{v}_j \rangle.$$

Solving for a_j we find $a_j = \frac{\langle \vec{x}, \vec{v}_j \rangle}{\langle \vec{v}_j, \vec{v}_j \rangle}$ as before.

Example A.5.3: The vectors $(1, 1)$ and $(1, -1)$ form an orthogonal basis of $\mathbb{R}^2$. Suppose we wish to represent $(3, 4)$ in terms of this basis, that is, we wish to find a_1 and a_2 such that

$$(3, 4) = a_1(1, 1) + a_2(1, -1).$$

We compute:

$$a_1 = \frac{\langle (3, 4), (1, 1) \rangle}{\langle (1, 1), (1, 1) \rangle} = \frac{7}{2}, \quad a_2 = \frac{\langle (3, 4), (1, -1) \rangle}{\langle (1, -1), (1, -1) \rangle} = \frac{-1}{2}.$$

So

$$(3, 4) = \frac{7}{2}(1, 1) + \frac{-1}{2}(1, -1).$$

If the basis is orthonormal rather than orthogonal, then all the denominators are one. It is easy to make a basis orthonormal—divide all the vectors by their size. If you want to decompose many vectors, it may be better to find an orthonormal basis. In the example above, the orthonormal basis we would thus create is

$$\left(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right), \quad \left(\frac{1}{\sqrt{2}}, \frac{-1}{\sqrt{2}} \right).$$

Then the computation would have been

$$\begin{aligned} (3, 4) &= \left\langle (3, 4), \left(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right) \right\rangle \left(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right) + \left\langle (3, 4), \left(\frac{1}{\sqrt{2}}, \frac{-1}{\sqrt{2}} \right) \right\rangle \left(\frac{1}{\sqrt{2}}, \frac{-1}{\sqrt{2}} \right) \\ &= \frac{7}{\sqrt{2}} \left(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right) + \frac{-1}{\sqrt{2}} \left(\frac{1}{\sqrt{2}}, \frac{-1}{\sqrt{2}} \right). \end{aligned}$$

Maybe the example is not so awe inspiring, but given vectors in $\mathbb{R}^{20}$ rather than $\mathbb{R}^2$, then surely one would much rather do 20 inner products (or 40 if we did not have an orthonormal basis) rather than solving a system of twenty equations in twenty unknowns using row reduction of a 20×21 matrix.

As we said above, the formula still works even if $\vec{x}$ is not in the subspace, although then it does not get us the vector $\vec{x}$ but its projection. More concretely, suppose that S is a subspace that is the span of $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$ and $\vec{x}$ is any vector. Let $\text{proj}_S(\vec{x})$ be the vector in S that is the closest to $\vec{x}$. Then

$$\text{proj}_S(\vec{x}) = \frac{\langle \vec{x}, \vec{v}_1 \rangle}{\langle \vec{v}_1, \vec{v}_1 \rangle} \vec{v}_1 + \frac{\langle \vec{x}, \vec{v}_2 \rangle}{\langle \vec{v}_2, \vec{v}_2 \rangle} \vec{v}_2 + \dots + \frac{\langle \vec{x}, \vec{v}_n \rangle}{\langle \vec{v}_n, \vec{v}_n \rangle} \vec{v}_n.$$

Of course, if $\vec{x}$ is in S , then $\text{proj}_S(\vec{x}) = \vec{x}$, as the closest vector in S to $\vec{x}$ is $\vec{x}$ itself. But true utility is obtained when $\vec{x}$ is not in S . In much of applied mathematics, we cannot find an exact solution to a problem, but we try to find the best solution out of a small subset (subspace). The partial sums of Fourier series from [chapter 4](#) are one example. Another example is least square approximation to fit a curve to data. Yet another example is given by the most commonly used numerical methods to solve partial differential equations, the finite element methods.

Example A.5.4: The vectors $(1, 2, 3)$ and $(3, 0, -1)$ are orthogonal, and so they are an orthogonal basis of a subspace S :

$$S = \text{span}\{(1, 2, 3), (3, 0, -1)\}.$$

Let us find the vector in S that is closest to $(2, 1, 0)$. That is, let us find $\text{proj}_S((2, 1, 0))$.

$$\begin{aligned} \text{proj}_S((2, 1, 0)) &= \frac{\langle (2, 1, 0), (1, 2, 3) \rangle}{\langle (1, 2, 3), (1, 2, 3) \rangle} (1, 2, 3) + \frac{\langle (2, 1, 0), (3, 0, -1) \rangle}{\langle (3, 0, -1), (3, 0, -1) \rangle} (3, 0, -1) \\ &= \frac{2}{7} (1, 2, 3) + \frac{3}{5} (3, 0, -1) \\ &= \left(\frac{73}{35}, \frac{4}{7}, \frac{9}{35} \right). \end{aligned}$$

A.5.4 The Gram–Schmidt process

Before leaving orthogonal bases, let us note a procedure for manufacturing them out of any old basis. It may not be difficult to come up with an orthogonal basis for a 2-dimensional subspace, but for a 20-dimensional subspace, it seems a daunting task. Fortunately, the orthogonal projection can be used to “project away” the bits of the vectors that are making them not orthogonal. It is called the *Gram–Schmidt process*.

We start with a basis of vectors $\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n$. We construct an orthogonal basis $\vec{w}_1, \vec{w}_2, \dots, \vec{w}_n$ as follows.

$$\begin{aligned} \vec{w}_1 &= \vec{v}_1, \\ \vec{w}_2 &= \vec{v}_2 - \text{proj}_{\vec{w}_1}(\vec{v}_2), \\ \vec{w}_3 &= \vec{v}_3 - \text{proj}_{\vec{w}_1}(\vec{v}_3) - \text{proj}_{\vec{w}_2}(\vec{v}_3), \\ \vec{w}_4 &= \vec{v}_4 - \text{proj}_{\vec{w}_1}(\vec{v}_4) - \text{proj}_{\vec{w}_2}(\vec{v}_4) - \text{proj}_{\vec{w}_3}(\vec{v}_4), \\ &\vdots \\ \vec{w}_n &= \vec{v}_n - \text{proj}_{\vec{w}_1}(\vec{v}_n) - \text{proj}_{\vec{w}_2}(\vec{v}_n) - \dots - \text{proj}_{\vec{w}_{n-1}}(\vec{v}_n). \end{aligned}$$

What we do is at the k^{th} step, we take $\vec{v}_k$ and we subtract the projection of $\vec{v}_k$ to the subspace spanned by $\vec{w}_1, \vec{w}_2, \dots, \vec{w}_{k-1}$.

Example A.5.5: Consider the vectors $(1, 2, -1)$, and $(0, 5, -2)$ and call S the span of the two vectors. Let us find an orthogonal basis of S via the Gram–Schmidt process:

$$\begin{aligned}\vec{w}_1 &= (1, 2, -1), \\ \vec{w}_2 &= (0, 5, -2) - \text{proj}_{(1, 2, -1)}((0, 5, -2)) \\ &= (0, 5, -2) - \frac{\langle (0, 5, -2), (1, 2, -1) \rangle}{\langle (1, 2, -1), (1, 2, -1) \rangle} (1, 2, -1) = (0, 5, -2) - 2(1, 2, -1) = (-2, 1, 0).\end{aligned}$$

So $(1, 2, -1)$ and $(-2, 1, 0)$ span S and are orthogonal. Let us check: $\langle (1, 2, -1), (-2, 1, 0) \rangle = 0$.

Suppose we wish to find an orthonormal basis, not just an orthogonal one. Well, we simply make the vectors into unit vectors by dividing them by their magnitude. The two vectors making up the orthonormal basis of S are:

$$\frac{1}{\sqrt{6}}(1, 2, -1) = \left(\frac{1}{\sqrt{6}}, \frac{2}{\sqrt{6}}, \frac{-1}{\sqrt{6}} \right), \quad \frac{1}{\sqrt{5}}(-2, 1, 0) = \left(\frac{-2}{\sqrt{5}}, \frac{1}{\sqrt{5}}, 0 \right).$$

A.5.5 Exercises

Exercise A.5.1: Find the s that makes the following vectors orthogonal: $(1, 2, 3)$, $(1, 1, s)$.

Exercise A.5.2: Find the angle θ between $(1, 3, 1)$, $(2, 1, -1)$.

Exercise A.5.3: Given that $\langle \vec{v}, \vec{w} \rangle = 3$ and $\langle \vec{v}, \vec{u} \rangle = -1$ compute

$$\begin{array}{lll} a) \langle \vec{u}, 2\vec{v} \rangle & b) \langle \vec{v}, 2\vec{w} + 3\vec{u} \rangle & c) \langle \vec{w} + 3\vec{u}, \vec{v} \rangle\end{array}$$

Exercise A.5.4: Suppose $\vec{v} = (1, 1, -1)$. Find

$$\begin{array}{lll} a) \text{proj}_{\vec{v}}((1, 0, 0)) & b) \text{proj}_{\vec{v}}((1, 2, 3)) & c) \text{proj}_{\vec{v}}((1, -1, 0))\end{array}$$

Exercise A.5.5: Consider the vectors $(1, 2, 3)$, $(-3, 0, 1)$, $(1, -5, 3)$.

- Check that the vectors are linearly independent and so form a basis of $\mathbb{R}^3$.
- Check that the vectors are mutually orthogonal, and are therefore an orthogonal basis.
- Represent $(1, 1, 1)$ as a linear combination of this basis.
- Make the basis orthonormal.

Exercise A.5.6: Let S be the subspace spanned by $(1, 3, -1)$, $(1, 1, 1)$. Find an orthogonal basis of S by the Gram–Schmidt process.

Exercise A.5.7: Starting with $(1, 2, 3)$, $(1, 1, 1)$, $(2, 2, 0)$, follow the Gram–Schmidt process to find an orthogonal basis of $\mathbb{R}^3$.

Exercise A.5.8: Find an orthogonal basis of $\mathbb{R}^3$ such that $(3, 1, -2)$ is one of the vectors. Hint: First find two extra vectors to make a linearly independent set.

Exercise A.5.9: Using cosines and sines of θ , find a unit vector $\vec{u}$ in $\mathbb{R}^2$ that makes angle θ with $\vec{i} = (1, 0)$. What is $\langle \vec{i}, \vec{u} \rangle$?

Exercise A.5.101: Find the s that makes the following vectors orthogonal: $(1, 1, 1)$, $(1, s, 1)$.

Exercise A.5.102: Find the angle θ between $(1, 2, 3)$, $(1, 1, 1)$.

Exercise A.5.103: Given that $\langle \vec{v}, \vec{w} \rangle = 1$, $\langle \vec{v}, \vec{u} \rangle = -1$, and $\|\vec{v}\| = 3$, compute

a) $\langle 3\vec{u}, 5\vec{v} \rangle$

b) $\langle \vec{v}, 2\vec{w} + 3\vec{u} \rangle$

c) $\langle \vec{w} + 3\vec{v}, \vec{v} \rangle$

Exercise A.5.104: Suppose $\vec{v} = (1, 0, -1)$. Find

a) $\text{proj}_{\vec{v}}((0, 2, 1))$

b) $\text{proj}_{\vec{v}}((1, 0, 1))$

c) $\text{proj}_{\vec{v}}((4, -1, 0))$

Exercise A.5.105: The vectors $(1, 1, -1)$, $(2, -1, 1)$, $(0, 1, 1)$ form an orthogonal basis. Represent the following vectors in terms of this basis:

a) $(-1, 4, 0)$

b) $(4, -1, 3)$

c) $(-2, 6, -2)$

Exercise A.5.106: Let S be the subspace spanned by $(2, -1, 1)$, $(2, 2, 2)$. Find an orthogonal basis of S by the Gram–Schmidt process.

Exercise A.5.107: Starting with $(1, 1, -1)$, $(2, 3, -1)$, $(1, -1, 1)$, follow the Gram–Schmidt process to find an orthogonal basis of $\mathbb{R}^3$.

A.6 Determinant

Note: 1 lecture

For square matrices we define a useful quantity called the *determinant*. Define the determinant of a 1×1 matrix as the value of its only entry

$$\det([a]) \stackrel{\text{def}}{=} a.$$

For a 2×2 matrix, define

$$\det\left(\begin{bmatrix} a & b \\ c & d \end{bmatrix}\right) \stackrel{\text{def}}{=} ad - bc.$$

Before defining the determinant for larger matrices, we note the meaning of the determinant. An $n \times n$ matrix gives a mapping of the n -dimensional euclidean space $\mathbb{R}^n$ to itself. So a 2×2 matrix A is a mapping of the plane to itself. The determinant of A is the factor by which the area of objects changes. If we take the unit square (square of side 1) in the plane, then A takes the square to a parallelogram of area $|\det(A)|$. The sign of $\det(A)$ denotes a change of orientation (negative if the axes get flipped). For example, let

$$A = \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix}.$$

Then $\det(A) = 1 + 1 = 2$. Let us see where A sends the unit square—the square with vertices $(0, 0)$, $(1, 0)$, $(1, 1)$, and $(0, 1)$. The point $(0, 0)$ gets sent to $(0, 0)$.

$$\begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ -1 \end{bmatrix}, \quad \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 2 \\ 0 \end{bmatrix}, \quad \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}.$$

The image of the square is another square with vertices $(0, 0)$, $(1, -1)$, $(2, 0)$, and $(1, 1)$. The image square has a side of length $\sqrt{2}$, and it is therefore of area 2. See [Figure A.7](#).

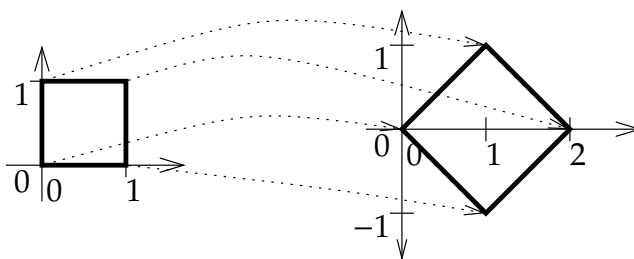


Figure A.7: Image of the unit square via the mapping A .

In general, the image of a square is a parallelogram. In high school geometry, you may have seen a formula for computing the area of a parallelogram with vertices $(0, 0)$, (a, c) ,

(b, d) and $(a + b, c + d)$. The area is

$$\left| \det \begin{pmatrix} a & b \\ c & d \end{pmatrix} \right| = |ad - bc|.$$

The vertical lines above mean absolute value. The matrix $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$ carries the unit square to the given parallelogram.

There are a number of ways to define the determinant for an $n \times n$ matrix. Let us use the so-called *cofactor expansion*. We define A_{ij} as the matrix A with the i^{th} row and the j^{th} column deleted. For example,

$$\text{if } A = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix}, \quad \text{then } A_{12} = \begin{bmatrix} 4 & 6 \\ 7 & 9 \end{bmatrix} \quad \text{and} \quad A_{23} = \begin{bmatrix} 1 & 2 \\ 7 & 8 \end{bmatrix}.$$

We now define the determinant recursively

$$\det(A) \stackrel{\text{def}}{=} \sum_{j=1}^n (-1)^{1+j} a_{1j} \det(A_{1j}),$$

or in other words

$$\det(A) = a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + a_{13} \det(A_{13}) - \cdots \begin{cases} +a_{1n} \det(A_{1n}) & \text{if } n \text{ is odd,} \\ -a_{1n} \det(A_{1n}) & \text{if } n \text{ is even.} \end{cases}$$

For a 3×3 matrix, we get $\det(A) = a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + a_{13} \det(A_{13})$. For example,

$$\begin{aligned} \det \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} &= 1 \cdot \det \begin{pmatrix} 5 & 6 \\ 8 & 9 \end{pmatrix} - 2 \cdot \det \begin{pmatrix} 4 & 6 \\ 7 & 9 \end{pmatrix} + 3 \cdot \det \begin{pmatrix} 4 & 5 \\ 7 & 8 \end{pmatrix} \\ &= 1(5 \cdot 9 - 6 \cdot 8) - 2(4 \cdot 9 - 6 \cdot 7) + 3(4 \cdot 8 - 5 \cdot 7) = 0. \end{aligned}$$

It turns out that we did not have to necessarily use the first row. That is, for any i ,

$$\det(A) = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det(A_{ij}).$$

It is sometimes useful to use a row other than the first. In the following example it is more convenient to expand along the second row. Notice that for the second row we are starting with a negative sign.

$$\begin{aligned} \det \begin{pmatrix} 1 & 2 & 3 \\ 0 & 5 & 0 \\ 7 & 8 & 9 \end{pmatrix} &= -0 \cdot \det \begin{pmatrix} 2 & 3 \\ 8 & 9 \end{pmatrix} + 5 \cdot \det \begin{pmatrix} 1 & 3 \\ 7 & 9 \end{pmatrix} - 0 \cdot \det \begin{pmatrix} 1 & 2 \\ 7 & 8 \end{pmatrix} \\ &= 0 + 5(1 \cdot 9 - 3 \cdot 7) + 0 = -60. \end{aligned}$$

Let us check if it is really the same as expanding along the first row,

$$\begin{aligned}\det\left(\begin{bmatrix} 1 & 2 & 3 \\ 0 & 5 & 0 \\ 7 & 8 & 9 \end{bmatrix}\right) &= 1 \cdot \det\left(\begin{bmatrix} 5 & 0 \\ 8 & 9 \end{bmatrix}\right) - 2 \cdot \det\left(\begin{bmatrix} 0 & 0 \\ 7 & 9 \end{bmatrix}\right) + 3 \cdot \det\left(\begin{bmatrix} 0 & 5 \\ 7 & 8 \end{bmatrix}\right) \\ &= 1(5 \cdot 9 - 0 \cdot 8) - 2(0 \cdot 9 - 0 \cdot 7) + 3(0 \cdot 8 - 5 \cdot 7) = -60.\end{aligned}$$

In computing the determinant, we alternately add and subtract the determinants of the submatrices A_{ij} multiplied by a_{ij} for a fixed i and all j . The numbers $(-1)^{i+j} \det(A_{ij})$ are called *cofactors* of the matrix. And that is why this method of computing the determinant is called the *cofactor expansion*.

Similarly, we do not need to expand along a row; we can expand along a column. For any j ,

$$\det(A) = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det(A_{ij}).$$

A related fact is that

$$\det(A) = \det(A^T).$$

A matrix is *upper triangular* if all elements below the main diagonal are 0. For example,

$$\begin{bmatrix} 1 & 2 & 3 \\ 0 & 5 & 6 \\ 0 & 0 & 9 \end{bmatrix}$$

is upper triangular. Similarly a *lower triangular* matrix is one where everything above the diagonal is zero. For example,

$$\begin{bmatrix} 1 & 0 & 0 \\ 4 & 5 & 0 \\ 7 & 8 & 9 \end{bmatrix}.$$

The determinant for triangular matrices is very simple to compute. Consider the lower triangular matrix. If we expand along the first row, we find that the determinant is 1 times the determinant of the lower triangular matrix $\begin{bmatrix} 5 & 0 \\ 8 & 9 \end{bmatrix}$. So the determinant is just the product of the diagonal entries:

$$\det\left(\begin{bmatrix} 1 & 0 & 0 \\ 4 & 5 & 0 \\ 7 & 8 & 9 \end{bmatrix}\right) = 1 \cdot 5 \cdot 9 = 45.$$

Similarly for upper triangular matrices

$$\det\left(\begin{bmatrix} 1 & 2 & 3 \\ 0 & 5 & 6 \\ 0 & 0 & 9 \end{bmatrix}\right) = 1 \cdot 5 \cdot 9 = 45.$$

In general, if A is triangular, then

$$\det(A) = a_{11}a_{22} \cdots a_{nn}.$$

If A is diagonal, then it is also triangular (upper and lower), so the same formula applies. For example,

$$\det \begin{pmatrix} 2 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 5 \end{pmatrix} = 2 \cdot 3 \cdot 5 = 30.$$

In particular, the identity matrix I is diagonal, and the diagonal entries are all 1. Thus,

$$\det(I) = 1.$$

The determinant is telling you how geometric objects scale. If B doubles the sizes of geometric objects and A triples them, then AB (which applies B to an object and then it applies A) should make size go up by a factor of 6. This is true in general:

Theorem A.6.1.

$$\det(AB) = \det(A) \det(B).$$

This property is one of the most useful, and it is employed often to actually compute determinants. A particularly interesting consequence is to note what it means for the existence of inverses. Take A and B to be inverses, that is, $AB = I$. Then

$$\det(A) \det(B) = \det(AB) = \det(I) = 1.$$

Neither $\det(A)$ nor $\det(B)$ can be zero. This fact is an extremely useful property of the determinant, and one which is used often in this book:

Theorem A.6.2. *An $n \times n$ matrix A is invertible if and only if $\det(A) \neq 0$.*

In fact, $\det(A^{-1}) \det(A) = 1$ says that

$$\det(A^{-1}) = \frac{1}{\det(A)}.$$

So we know what the determinant of A^{-1} is without computing A^{-1} .

Let us return to the formula for the inverse of a 2×2 matrix:

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}.$$

Notice the determinant of the matrix $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$ in the denominator of the fraction. The formula only works if the determinant is nonzero. Otherwise, we are dividing by zero.

A common notation for the determinant is a pair of vertical lines:

$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} = \det \left(\begin{bmatrix} a & b \\ c & d \end{bmatrix} \right).$$

Personally, I find this notation confusing as vertical lines usually mean a positive quantity, while determinants can be negative. Also think about how to write the absolute value of a determinant. This notation is not used in this book.

A.6.1 Exercises

Exercise A.6.1: Compute the determinant of the following matrices:

$$\begin{array}{llll}
 a) [3] & b) \begin{bmatrix} 1 & 3 \\ 2 & 1 \end{bmatrix} & c) \begin{bmatrix} 2 & 1 \\ 4 & 2 \end{bmatrix} & d) \begin{bmatrix} 1 & 2 & 3 \\ 0 & 4 & 5 \\ 0 & 0 & 6 \end{bmatrix} \\
 e) \begin{bmatrix} 2 & 1 & 0 \\ -2 & 7 & -3 \\ 0 & 2 & 0 \end{bmatrix} & f) \begin{bmatrix} 2 & 1 & 3 \\ 8 & 6 & 3 \\ 7 & 9 & 7 \end{bmatrix} & g) \begin{bmatrix} 0 & 2 & 5 & 7 \\ 0 & 0 & 2 & -3 \\ 3 & 4 & 5 & 7 \\ 0 & 0 & 2 & 4 \end{bmatrix} & h) \begin{bmatrix} 0 & 1 & 2 & 0 \\ 1 & 1 & -1 & 2 \\ 1 & 1 & 2 & 1 \\ 2 & -1 & -2 & 3 \end{bmatrix}
 \end{array}$$

Exercise A.6.2: For which x are the following matrices singular (not invertible).

$$\begin{array}{llll}
 a) \begin{bmatrix} 2 & 3 \\ 2 & x \end{bmatrix} & b) \begin{bmatrix} 2 & x \\ 1 & 2 \end{bmatrix} & c) \begin{bmatrix} x & 1 \\ 4 & x \end{bmatrix} & d) \begin{bmatrix} x & 0 & 1 \\ 1 & 4 & 2 \\ 1 & 6 & 2 \end{bmatrix}
 \end{array}$$

Exercise A.6.3: Compute

$$\det \left(\begin{bmatrix} 2 & 1 & 2 & 3 \\ 0 & 8 & 6 & 5 \\ 0 & 0 & 3 & 9 \\ 0 & 0 & 0 & 1 \end{bmatrix}^{-1} \right)$$

without computing the inverse.

Exercise A.6.4: Suppose

$$L = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 2 & 1 & 0 & 0 \\ 7 & \pi & 1 & 0 \\ 2^8 & 5 & -99 & 1 \end{bmatrix} \quad \text{and} \quad U = \begin{bmatrix} 5 & 9 & 1 & -\sin(1) \\ 0 & 1 & 88 & -1 \\ 0 & 0 & 1 & 3 \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

Let $A = LU$. Compute $\det(A)$ in a simple way, without computing what is A . Hint: First read off $\det(L)$ and $\det(U)$.

Exercise A.6.5: Consider the linear mapping from $\mathbb{R}^2$ to $\mathbb{R}^2$ given by the matrix $A = \begin{bmatrix} 1 & x \\ 2 & 1 \end{bmatrix}$ for some number x . You wish to make A such that it doubles the area of every geometric figure. What are the possibilities for x (there are two answers).

Exercise A.6.6: Suppose A and S are $n \times n$ matrices, and S is invertible. Suppose that $\det(A) = 3$. Compute $\det(S^{-1}AS)$ and $\det(SAS^{-1})$. Justify your answer using the theorems in this section.

Exercise A.6.7: Let A be an $n \times n$ matrix such that $\det(A) = 1$. Compute $\det(xA)$ given a number x . Hint: First try computing $\det(xI)$, then note that $xA = (xI)A$.

Exercise A.6.101: Compute the determinant of the following matrices:

$$\begin{array}{llll}
 a) \begin{bmatrix} -2 \end{bmatrix} & b) \begin{bmatrix} 2 & -2 \\ 1 & 3 \end{bmatrix} & c) \begin{bmatrix} 2 & 2 \\ 2 & 2 \end{bmatrix} & d) \begin{bmatrix} 2 & 9 & -11 \\ 0 & -1 & 5 \\ 0 & 0 & 3 \end{bmatrix} \\
 e) \begin{bmatrix} 2 & 1 & 0 \\ -2 & 7 & 3 \\ 1 & 1 & 0 \end{bmatrix} & f) \begin{bmatrix} 5 & 1 & 3 \\ 4 & 1 & 1 \\ 4 & 5 & 1 \end{bmatrix} & g) \begin{bmatrix} 3 & 2 & 5 & 7 \\ 0 & 0 & 2 & 0 \\ 0 & 4 & 5 & 0 \\ 2 & 1 & 2 & 4 \end{bmatrix} & h) \begin{bmatrix} 0 & 2 & 1 & 0 \\ 1 & 2 & -3 & 4 \\ 5 & 6 & -7 & 8 \\ 1 & 2 & 3 & -2 \end{bmatrix}
 \end{array}$$

Exercise A.6.102: For which x are the following matrices singular (not invertible):

$$\begin{array}{llll}
 a) \begin{bmatrix} 1 & 3 \\ 1 & x \end{bmatrix} & b) \begin{bmatrix} 3 & x \\ 1 & 3 \end{bmatrix} & c) \begin{bmatrix} x & 3 \\ 3 & x \end{bmatrix} & d) \begin{bmatrix} x & 1 & 0 \\ 1 & 4 & 0 \\ 1 & 6 & 2 \end{bmatrix}
 \end{array}$$

Exercise A.6.103: Compute

$$\det \left(\begin{bmatrix} 3 & 4 & 7 & 12 \\ 0 & -1 & 9 & -8 \\ 0 & 0 & -2 & 4 \\ 0 & 0 & 0 & 2 \end{bmatrix}^{-1} \right)$$

without computing the inverse.

Exercise A.6.104 (challenging): Find all the x that make the matrix inverse

$$\begin{bmatrix} 1 & 2 \\ 1 & x \end{bmatrix}^{-1}$$

have only integer entries (no fractions). Note that there are two answers.

Appendix B

Table of Laplace Transforms

The function u is the Heaviside function, δ is the Dirac delta function, and

$$\Gamma(t) = \int_0^\infty e^{-\tau} \tau^{t-1} d\tau, \quad \operatorname{erf}(t) = \frac{2}{\sqrt{\pi}} \int_0^t e^{-\tau^2} d\tau, \quad \operatorname{erfc}(t) = 1 - \operatorname{erf}(t).$$

$f(t)$	$F(s) = \mathcal{L}\{f(t)\} = \int_0^\infty e^{-st} f(t) dt$
C	$\frac{C}{s}$
t	$\frac{1}{s^2}$
t^n	$\frac{n!}{s^{n+1}}$
$t^p \quad (p > -1)$	$\frac{\Gamma(p+1)}{s^{p+1}}$
e^{-at}	$\frac{1}{s+a}$
$\sin(\omega t)$	$\frac{\omega}{s^2+\omega^2}$
$\cos(\omega t)$	$\frac{s}{s^2+\omega^2}$
$\sinh(\omega t)$	$\frac{\omega}{s^2-\omega^2}$
$\cosh(\omega t)$	$\frac{s}{s^2-\omega^2}$
$u(t-a) \quad (a \geq 0)$	$\frac{e^{-as}}{s}$
$\delta(t)$	1
$\delta(t-a) \quad (a \geq 0)$	e^{-as}
$\operatorname{erf}\left(\frac{t}{2a}\right)$	$\frac{1}{s} e^{(as)^2} \operatorname{erfc}(as)$
$t \sin(\omega t)$	$\frac{2\omega s}{(s^2+\omega^2)^2}$
$t \cos(\omega t)$	$\frac{s^2-\omega^2}{(s^2+\omega^2)^2}$
$e^{-at} \sin(\omega t)$	$\frac{\omega}{(s+a)^2+\omega^2}$
$e^{-at} \cos(\omega t)$	$\frac{s+a}{(s+a)^2+\omega^2}$

$f(t)$	$F(s) = \mathcal{L}\{f(t)\} = \int_0^\infty e^{-st} f(t) dt$
$e^{-at} \sinh(\omega t)$	$\frac{\omega}{(s+a)^2 - \omega^2}$
$e^{-at} \cosh(\omega t)$	$\frac{s+a}{(s+a)^2 - \omega^2}$
$\frac{1}{\sqrt{\pi t}} \exp\left(\frac{-a^2}{4t}\right) \quad (a > 0)$	$\frac{e^{-a\sqrt{s}}}{\sqrt{s}}$
$\frac{a}{\sqrt{4\pi t^3}} \exp\left(\frac{-a^2}{4t}\right) \quad (a > 0)$	$e^{-a\sqrt{s}}$
$\frac{1}{\sqrt{\pi t}} - ae^{a^2 t} \operatorname{erfc}(a\sqrt{t}) \quad (a > 0)$	$\frac{1}{\sqrt{s+a}}$
$\operatorname{erfc}\left(\frac{a}{2\sqrt{t}}\right) \quad (a > 0)$	$\frac{e^{-a\sqrt{s}}}{s}$
$\frac{1}{2\sqrt{\pi t^3}}(e^{bt} - e^{at})$	$\sqrt{s-a} - \sqrt{s-b}$
$\frac{1}{t}(e^{bt} - e^{at})$	$\ln \frac{s-a}{s-b}$
$J_0(at)$	$\frac{1}{\sqrt{s^2+a^2}}$
$J_0(a\sqrt{t})$	$\frac{\exp\left(\frac{-a^2}{4s}\right)}{s}$
$\frac{1}{\sqrt{\pi t}} \cos(a\sqrt{t})$	$\frac{\exp\left(\frac{-a^2}{4s}\right)}{\sqrt{s}}$
$\frac{2}{\sqrt{\pi}} \sinh(a\sqrt{t})$	$\frac{a \exp\left(\frac{a^2}{4s}\right)}{s^{3/2}}$
$\frac{1}{\sqrt{\pi t}} e^{at}(1 + 2at)$	$\frac{s}{(s-a)^{3/2}}$
$a f(t) + b g(t)$	$aF(s) + bG(s)$
$f(at) \quad (a > 0)$	$\frac{1}{a} F\left(\frac{s}{a}\right)$
$f(t-a)u(t-a) \quad (a \geq 0)$	$e^{-as}F(s)$
$e^{-at} f(t)$	$F(s+a)$
$g'(t)$	$sG(s) - g(0)$
$g''(t)$	$s^2G(s) - sg(0) - g'(0)$
$g^{(n)}(t)$	$s^n G(s) - s^{n-1}g(0) - \dots - g^{(n-1)}(0)$
$(f * g)(t) = \int_0^t f(\tau)g(t-\tau) d\tau$	$F(s)G(s)$
$t f(t)$	$-F'(s)$
$t^n f(t)$	$(-1)^n F^{(n)}(s)$
$\int_0^t f(\tau) d\tau$	$\frac{1}{s} F(s)$
$\frac{f(t)}{t}$	$\int_s^\infty F(\sigma) d\sigma$
$f(t)$ periodic with period P	$\frac{1}{1-e^{-Ps}} \int_0^P e^{-st} f(t) dt$

Further Reading

- [BM] Paul W. Berg and James L. McGregor, *Elementary Partial Differential Equations*, Holden-Day, San Francisco, CA, 1966.
- [BD] William E. Boyce, Richard C. DiPrima, and Douglas B. Meade, *Elementary Differential Equations and Boundary Value Problems*, 12th edition, John Wiley & Sons Inc., New York, NY, 2021.
- [EP] C.H. Edwards and D.E. Penney, *Differential Equations and Boundary Value Problems: Computing and Modeling*, 6th edition, Pearson, 2022.
- [F] Stanley J. Farlow, *An Introduction to Differential Equations and Their Applications*, McGraw-Hill, Inc., Princeton, NJ, 1994. (Published also by Dover Publications, 2006.)
- [I] E.L. Ince, *Ordinary Differential Equations*, Dover Publications, Inc., New York, NY, 1956.
- [J] Thomas W. Judson, *The Ordinary Differential Equations Project*, <https://judsonbooks.org/the-ordinary-differential-equations-project/>.
- [T] William F. Trench, *Elementary Differential Equations with Boundary Value Problems*. Books and Monographs. Book 9. 2013. <https://digitalcommons.trinity.edu/mono/9>

Solutions to Selected Exercises

0.2.101: Compute $x' = -2e^{-2t}$ and $x'' = 4e^{-2t}$. Then $(4e^{-2t}) + 4(-2e^{-2t}) + 4(e^{-2t}) = 0$.

0.2.102: Yes.

0.2.103: $y = x^r$ is a solution for $r = 0$ and $r = 2$.

0.2.104: $C_1 = 100, C_2 = -90$

0.2.105: $\varphi = -9e^{8s}$

0.2.106: a) $x = 9e^{-4t}$ b) $x = \cos(2t) + \sin(2t)$ c) $p = 4e^{3q}$ d) $T = 3 \sinh(2x)$

0.3.101: a) PDE, equation, 2nd-order, linear, nonhomogeneous, constant-coefficient.

b) ODE, equation, 1st-order, linear, nonhomogeneous, not constant-coefficient, not autonomous.

c) ODE, equation, 7th-order, linear, homogeneous, constant-coefficient, autonomous.

d) ODE, equation, 2nd-order, linear, nonhomogeneous, constant-coefficient, autonomous.

e) ODE, system, 2nd-order, nonlinear.

f) PDE, equation, 2nd-order, nonlinear.

0.3.102: equation: $a(x)y = b(x)$, solution: $y = \frac{b(x)}{a(x)}$.

0.3.103: $k = 0$ or $k = 1$

1.1.101: $y = e^x + \frac{x^2}{2} + 9$

1.1.102: $x = (3t - 2)^{1/3}$

1.1.103: $x = \sin^{-1}(t + 1/\sqrt{2})$

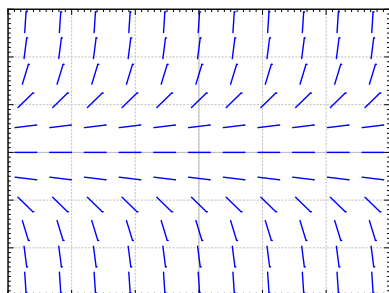
1.1.104: 170

1.1.105: If $n \neq 1$, then $y = ((1 - n)x + 1)^{1/(1-n)}$. If $n = 1$, then $y = e^x$.

1.1.106: The equation is $r' = -C$ for some constant C . The snowball will be completely melted in 25 minutes from time $t = 0$.

1.1.107: $y = Ax^3 + Bx^2 + Cx + D$, so 4 constants.

1.2.101:



$y = 0$ is a solution such that $y(0) = 0$.

1.2.102: Yes a solution exists. The equation is $y' = f(x, y)$ where $f(x, y) = xy$. The function $f(x, y)$ is continuous and $\frac{\partial f}{\partial y} = x$, which is also continuous near $(0, 0)$. So a solution exists and is unique. (In fact, $y = 0$ is the solution.)

1.2.103: No, the equation is not defined at $(x, y) = (1, 0)$.

1.2.104: a) $y' = \cos y$, b) $y' = y \cos(x)$, c) $y' = \sin x$. Justification left to reader.

1.2.105: Picard does not apply as f is not continuous at $y = 0$. The equation does not have a continuously differentiable solution. Suppose it did. Notice that $y'(0) = 1$. By the first derivative test, $y(x) > 0$ for small positive x . But then for those x , we have $y'(x) = f(y(x)) = 0$. It is not possible for y' to be continuous, $y'(0) = 1$ and $y'(x) = 0$ for arbitrarily small positive x .

1.2.106: The solution is $y(x) = \int_{x_0}^x f(s) ds + y_0$, and this does indeed exist for every x .

1.3.101: $y = Ce^{x^2}$

1.3.102: $x = e^{t^3} + 1$

1.3.103: $x^3 + x = t + 2$

1.3.104: $y = \frac{1}{1 - \ln x}$

1.3.105: $\sin(y) = -\cos(x) + C$

1.3.106: The range is approximately 7.45 to 12.15 minutes.

1.3.107: a) $x = \frac{1000e^t}{e^t + 24}$. b) 102 rabbits after one month, 861 after 5 months, 999 after 10 months, 1000 after 15 months.

1.4.101: $y = Ce^{-x^3} + 1/3$

1.4.102: $y = 2e^{\cos(2x)+1} + 1$

1.4.103: 250 grams

1.4.104: $P(5) = 1000e^{2 \times 5 - 0.05 \times 5^2} = 1000e^{8.75} \approx 6.31 \times 10^6$

1.4.105: $Ah' = I - kh$, where k is a constant with units m^2/s .

1.5.101: $y = \frac{2}{3x-2}$

1.5.102: $y = \frac{3-x^2}{2x}$

1.5.103: $y = (7e^{3x} + 3x + 1)^{1/3}$

1.5.104: $y = \pm \sqrt{x^2 - \ln(C - x)}$

1.6.101: a) 0, 1, 2 are critical points. b) $x = 0$ is unstable (semistable), $x = 1$ is stable, and $x = 2$ is unstable. c) 1

1.6.102: a) There are no critical points. b) ∞

1.6.103: a) $\frac{dx}{dt} = kx(M - x) + A$ b) $\frac{kM + \sqrt{(kM)^2 + 4Ak}}{2k}$

1.6.104: a) α is a stable critical point, β is an unstable one. b) α , c) α , d) ∞ or DNE.

1.7.101: Approximately: 1.0000, 1.2397, 1.3829

1.7.102: a) 0, 8, 12 b) $x(4) = 16$, so errors are: 16, 8, 4. c) Factors are 0.5 and 0.5.

1.7.103: a) 0, 0, 0 b) $x = 0$ is a solution so errors are: 0, 0, 0.

1.7.104: a) Improved Euler: $y(1) \approx 3.3897$ for $h = 1/4$, $y(1) \approx 3.4237$ for $h = 1/8$, b) Standard Euler: $y(1) \approx 2.8828$ for $h = 1/4$, $y(1) \approx 3.1316$ for $h = 1/8$, c) $y = 2e^x - x - 1$, so $y(1)$ is approximately 3.4366. d) Approximate errors for improved Euler: 0.046852 for $h = 1/4$ and 0.012881 for $h = 1/8$. For standard Euler: 0.55375 for $h = 1/4$ and 0.30499 for $h = 1/8$. Factor is approximately 0.27 for improved Euler and 0.55 for standard Euler.

1.8.101: a) $e^{xy} + \sin(x) = C$ b) $x^2 + xy - 2y^2 = C$ c) $e^x + e^y = C$ d) $x^3 + 3xy + y^3 = C$

1.8.102: a) Integrating factor is y , equation becomes $dx + 3y^2 dy = 0$. b) Integrating factor is e^x , equation becomes $e^x dx - e^{-y} dy = 0$. c) Integrating factor is y^2 , equation becomes $(\cos(x) + y) dx + x dy = 0$. d) Integrating factor is x , equation becomes $(2xy + y^2) dx + (x^2 + 2xy) dy = 0$.

1.8.103: a) The equation is $-f(x) dx + \frac{1}{g(y)} dy = 0$, and this is exact because $M = -f(x)$, $N = \frac{1}{g(y)}$, so $M_y = 0 = N_x$. b) $-x dx + \frac{1}{y} dy = 0$, leads to potential function $F(x, y) = -\frac{x^2}{2} + \ln|y|$, solving $F(x, y) = C$ leads to the same solution as the example.

1.9.101: a) $u = \frac{1}{1+(x+5t)^2}$ b) $u = \cos(x - 2t)$

1.9.102: $u = \cos(x - t)e^{-t^2/2}$

1.9.103: $u = x + 4t$

2.1.101: Yes. To justify try to find a constant A such that $\sin(x) = Ae^x$ for all x .

2.1.102: No. $e^{x+2} = e^2 e^x$.

2.1.103: $y = 5$

2.1.104: $y = C_1 \ln(x) + C_2$

2.1.105: $y'' - 3y' + 2y = 0$

2.2.101: $y = C_1 e^{(-2+\sqrt{2})x} + C_2 e^{(-2-\sqrt{2})x}$

2.2.102: $y = C_1 e^{3x} + C_2 x e^{3x}$

2.2.103: $y = e^{-x/4} \cos((\sqrt{7}/4)x) - \sqrt{7} e^{-x/4} \sin((\sqrt{7}/4)x)$

2.2.104: $y = \frac{2(a-b)}{5} e^{-3x/2} + \frac{3a+2b}{5} e^x$

2.2.105: $z(t) = 2e^{-t} \cos(t)$

2.2.106: $y = \frac{a\beta-b}{\beta-\alpha} e^{\alpha x} + \frac{b-a\alpha}{\beta-\alpha} e^{\beta x}$

2.2.107: $y'' - y' - 6y = 0$

2.3.101: $y = C_1 e^x + C_2 x^3 + C_3 x^2 + C_4 x + C_5$

2.3.102: a) $r^3 - 3r^2 + 4r - 12 = 0$ b) $y''' - 3y'' + 4y' - 12y = 0$ c) $y = C_1 e^{3x} + C_2 \sin(2x) + C_3 \cos(2x)$

2.3.103: $y = 0$

2.3.104: No. $e^1 e^x - e^{x+1} = 0$.

2.3.105: Yes. (Hint: First note that $\sin(x)$ is bounded. Then note that x and $x \sin(x)$ cannot be multiples of each other.)

2.3.106: $y''' - y'' + y' - y = 0$

2.4.101: $k = 8/9 \text{ N/m}$ (and larger)

2.4.102: a) $0.05I'' + 0.1I' + (1/5)I = 0$ b) $I = Ce^{-t} \cos(\sqrt{3}t - \gamma)$ c) $I = 10e^{-t} \cos(\sqrt{3}t) + \frac{10}{\sqrt{3}}e^{-t} \sin(\sqrt{3}t)$

2.4.103: a) $k = 500000 \text{ N/m}$ b) $\frac{1}{5\sqrt{2}} \approx 0.141 \text{ m}$ c) 45000 kg d) 11250 kg

2.4.104: $m_0 = 1/3 \text{ kg}$. If $m < m_0$, then the system is overdamped and will not oscillate.

2.5.101: $y = \frac{-16 \sin(3x) + 6 \cos(3x)}{73}$

2.5.102: a) $y = \frac{2e^x + 3x^3 - 9x}{6}$ b) $y = C_1 \cos(\sqrt{2}x) + C_2 \sin(\sqrt{2}x) + \frac{2e^x + 3x^3 - 9x}{6}$

2.5.103: $y(x) = x^2 - 4x + 6 + e^{-x}(x - 5)$

2.5.104: $y = \frac{2xe^x - (e^x + e^{-x}) \ln(e^{2x} + 1)}{4}$

2.5.105: $y = \frac{-\sin(x+c)}{3} + C_1 e^{\sqrt{2}x} + C_2 e^{-\sqrt{2}x}$

2.5.106: $y = x^2 + 2x + 3$

2.6.101: $\omega = \frac{\sqrt{31}}{4\sqrt{2}} \approx 0.984 \text{ rad/s}$ $C(\omega) = \frac{16}{3\sqrt{7}} \approx 2.016 \text{ m}$

2.6.102: $x_{sp} = \frac{(\omega_0^2 - \omega^2)F_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2} \cos(\omega t) + \frac{2\omega p F_0}{m(2\omega p)^2 + m(\omega_0^2 - \omega^2)^2} \sin(\omega t) + \frac{A}{k}$, where $p = \frac{c}{2m}$ and $\omega_0 = \sqrt{\frac{k}{m}}$.

2.6.103: a) $\omega = 2 \text{ rad/s}$ b) 25 m

3.1.101: $y_1 = C_1 e^{3x}$, $y_2 = C_2 e^x + \frac{C_1}{2} e^{3x}$, $y_3 = C_3 e^x + \frac{C_1}{2} e^{3x}$

3.1.102: $x = \frac{5}{3}e^{2t} - \frac{2}{3}e^{-t}$, $y = \frac{5}{3}e^{2t} + \frac{4}{3}e^{-t}$

3.1.103: $u'_1 = u_2$, $u'_2 = u_3$, $u'_3 = u_1 + t$, (here $u_1 = x$)

3.1.104: $y'_1 = y_3$, $y'_2 = y_4$, $y'_3 = -y_1 - y_2 + t$, $y'_4 = -y_1 + y_2 + t^2$

3.1.105: $x_1 = x_2 = at$. Explanation of the intuition is left to reader.

3.1.106: a) Left to reader. b) $x'_1 = \frac{r-s}{V}x_2 - \frac{r}{V}x_1$, $x'_2 = \frac{r}{V}(x_1 - x_2)$. c) As t goes to infinity, both x_1 and x_2 go to zero, explanation is left to reader.

3.2.101: -15

3.2.102: -2

3.2.103: $\vec{x} = \begin{bmatrix} 15 \\ -5 \end{bmatrix}$

3.2.104: a) $\begin{bmatrix} 1/a & 0 \\ 0 & 1/b \end{bmatrix}$ b) $\begin{bmatrix} 1/a & 0 & 0 \\ 0 & 1/b & 0 \\ 0 & 0 & 1/c \end{bmatrix}$

3.3.101: Yes.

3.3.102: No. $2 \begin{bmatrix} \cosh(t) \\ 1 \end{bmatrix} - \begin{bmatrix} e^t \\ 1 \end{bmatrix} - \begin{bmatrix} e^{-t} \\ 1 \end{bmatrix} = \vec{0}$

$$3.3.103: \begin{bmatrix} x \\ y \end{bmatrix}' = \begin{bmatrix} 3 & -1 \\ t & 0 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} e^t \\ 0 \end{bmatrix}$$

$$3.3.104: \text{ a) } \vec{x}' = \begin{bmatrix} 0 & 2t \\ 0 & 2t \end{bmatrix} \vec{x} \quad \text{ b) } \vec{x} = \begin{bmatrix} C_2 e^{t^2} + C_1 \\ C_2 e^{t^2} \end{bmatrix}$$

$$3.4.101: \text{ a) Eigenvalues: } 4, 0, -1 \quad \text{ Eigenvectors: } \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \begin{bmatrix} 3 \\ 5 \\ -2 \end{bmatrix}$$

$$\text{ b) } \vec{x} = C_1 \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} e^{4t} + C_2 \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + C_3 \begin{bmatrix} 3 \\ 5 \\ -2 \end{bmatrix} e^{-t}$$

$$3.4.102: \text{ a) Eigenvalues: } \frac{1+\sqrt{3}i}{2}, \frac{1-\sqrt{3}i}{2}, \quad \text{ Eigenvectors: } \begin{bmatrix} -2 \\ 1-\sqrt{3}i \end{bmatrix}, \begin{bmatrix} -2 \\ 1+\sqrt{3}i \end{bmatrix}$$

$$\text{ b) } \vec{x} = C_1 e^{t/2} \begin{bmatrix} -2 \cos(\frac{\sqrt{3}t}{2}) \\ \cos(\frac{\sqrt{3}t}{2}) + \sqrt{3} \sin(\frac{\sqrt{3}t}{2}) \end{bmatrix} + C_2 e^{t/2} \begin{bmatrix} -2 \sin(\frac{\sqrt{3}t}{2}) \\ \sin(\frac{\sqrt{3}t}{2}) - \sqrt{3} \cos(\frac{\sqrt{3}t}{2}) \end{bmatrix}$$

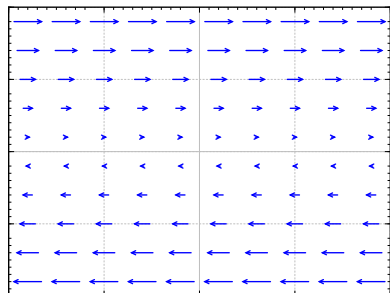
$$3.4.103: \vec{x} = C_1 \begin{bmatrix} 1 \\ 1 \end{bmatrix} e^t + C_2 \begin{bmatrix} 1 \\ -1 \end{bmatrix} e^{-t}$$

$$3.4.104: \vec{x} = C_1 \begin{bmatrix} \cos(t) \\ -\sin(t) \end{bmatrix} + C_2 \begin{bmatrix} \sin(t) \\ \cos(t) \end{bmatrix}$$

3.5.101: a) Two eigenvalues: $\pm\sqrt{2}$ so the behavior is a saddle. b) Two eigenvalues: 1 and 2, so the behavior is a source. c) Two eigenvalues: $\pm 2i$, so the behavior is a center (ellipses). d) Two eigenvalues: -1 and -2, so the behavior is a sink. e) Two eigenvalues: 5 and -3, so the behavior is a saddle.

3.5.102: Spiral source.

3.5.103:



The solution does not move anywhere if $y = 0$. When y is positive, the solution moves (with constant speed) in the positive x direction. When y is negative, the solution moves (with constant speed) in the negative x direction. It is not one of the behaviors we saw.

Note that the matrix has a double eigenvalue 0 and the general solution is $x = C_1 t + C_2$ and $y = C_1$, which agrees with the description above.

$$3.6.101: \vec{x} = \begin{bmatrix} 1 \\ -1 \end{bmatrix} (a_1 \cos(\sqrt{3}t) + b_1 \sin(\sqrt{3}t)) + \begin{bmatrix} 0 \\ 1 \\ -2 \end{bmatrix} (a_2 \cos(\sqrt{2}t) + b_2 \sin(\sqrt{2}t)) + \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} (a_3 \cos(t) + b_3 \sin(t)) + \begin{bmatrix} -1 \\ 1/2 \\ -1/3 \end{bmatrix} \cos(2t)$$

$$3.6.102: \begin{bmatrix} m & 0 & 0 \\ 0 & m & 0 \\ 0 & 0 & m \end{bmatrix} \vec{x}'' = \begin{bmatrix} -k & k & 0 \\ k & -2k & k \\ 0 & k & -k \end{bmatrix} \vec{x}. \text{ Solution: } \vec{x} = \begin{bmatrix} 1 \\ -2 \\ 1 \end{bmatrix} (a_1 \cos(\sqrt{3k/m}t) + b_1 \sin(\sqrt{3k/m}t)) + \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} (a_2 \cos(\sqrt{k/m}t) + b_2 \sin(\sqrt{k/m}t)) + \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} (a_3 t + b_3).$$

$$3.6.103: x_2 = (2/5) \cos(\sqrt{1/6}t) - (2/5) \cos(t)$$

$$3.7.101: \text{ a) } 3, 0, 0 \quad \text{ b) No defects.} \quad \text{ c) } \vec{x} = C_1 \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} e^{3t} + C_2 \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} + C_3 \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}$$

3.7.102: a) 1, 1, 2

b) Eigenvalue 1 has a defect of 1

$$c) \vec{x} = C_1 \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} e^t + C_2 \left(\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + t \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} \right) e^t + C_3 \begin{bmatrix} 3 \\ 3 \\ -2 \end{bmatrix} e^{2t}$$

3.7.103: a) 2, 2, 2

b) Eigenvalue 2 has a defect of 2

$$c) \vec{x} = C_1 \begin{bmatrix} 0 \\ 3 \\ 1 \end{bmatrix} e^{2t} + C_2 \left(\begin{bmatrix} 0 \\ -1 \\ 0 \end{bmatrix} + t \begin{bmatrix} 0 \\ 3 \\ 1 \end{bmatrix} \right) e^{2t} + C_3 \left(\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + t \begin{bmatrix} 0 \\ -1 \\ 0 \end{bmatrix} + \frac{t^2}{2} \begin{bmatrix} 0 \\ 3 \\ 1 \end{bmatrix} \right) e^{2t}$$

$$\mathbf{3.7.104:} \quad A = \begin{bmatrix} 5 & 5 \\ 0 & 5 \end{bmatrix}$$

$$\mathbf{3.8.101:} \quad e^{tA} = \begin{bmatrix} \frac{e^{3t} + e^{-t}}{2} & \frac{e^{-t} - e^{3t}}{2} \\ \frac{e^{-t} - e^{3t}}{2} & \frac{e^{3t} + e^{-t}}{2} \end{bmatrix}$$

$$\mathbf{3.8.102:} \quad e^{tA} = \begin{bmatrix} 2e^{3t} - 4e^{2t} + 3e^t & \frac{3e^t}{2} - \frac{3e^{3t}}{2} & -e^{3t} + 4e^{2t} - 3e^t \\ 2e^t - 2e^{2t} & e^t & 2e^{2t} - 2e^t \\ 2e^{3t} - 5e^{2t} + 3e^t & \frac{3e^t}{2} - \frac{3e^{3t}}{2} & -e^{3t} + 5e^{2t} - 3e^t \end{bmatrix}$$

$$\mathbf{3.8.103:} \quad a) e^{tA} = \begin{bmatrix} (t+1)e^{2t} & -te^{2t} \\ te^{2t} & (1-t)e^{2t} \end{bmatrix} \quad b) \vec{x} = \begin{bmatrix} (1-t)e^{2t} \\ (2-t)e^{2t} \end{bmatrix}$$

$$\mathbf{3.8.104:} \quad \begin{bmatrix} 1+2t+5t^2 & 3t+6t^2 \\ 2t+4t^2 & 1+2t+5t^2 \end{bmatrix} e^{0.1A} \approx \begin{bmatrix} 1.25 & 0.36 \\ 0.24 & 1.25 \end{bmatrix}$$

$$\mathbf{3.8.105:} \quad a) \begin{bmatrix} 5(3^n) - 2^{n+2} & 4(3^n) - 2^{n+2} \\ 5(2^n) - 5(3^n) & 5(2^n) - 4(3^n) \end{bmatrix} \quad b) \begin{bmatrix} 3 - 2(3^n) & 2(3^n) - 2 \\ 3 - 3^{n+1} & 3^{n+1} - 2 \end{bmatrix}$$

$$c) \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \text{ if } n \text{ is even, and } \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \text{ if } n \text{ is odd.}$$

3.9.101: The general solution is (particular solutions should agree with one of these):

$$x(t) = C_1 e^{9t} + 4C_2 e^{4t} - t/3 - 5/54, \quad y(t) = C_1 e^{9t} - C_2 e^{4t} + t/6 + 7/216$$

3.9.102: The general solution is (particular solutions should agree with one of these):

$$x(t) = C_1 e^t + C_2 e^{-t} + te^t, \quad y(t) = C_1 e^t - C_2 e^{-t} + te^t$$

$$\mathbf{3.9.103:} \quad \vec{x} = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \left(\frac{5}{2} e^t - t - 1 \right) + \begin{bmatrix} 1 \\ -1 \end{bmatrix} \frac{-1}{2} e^{-t}$$

$$\mathbf{3.9.104:} \quad \vec{x} = \begin{bmatrix} 1 \\ 9 \end{bmatrix} \left(\left(\frac{1}{140} + \frac{1}{120\sqrt{6}} \right) e^{\sqrt{6}t} + \left(\frac{1}{140} - \frac{1}{120\sqrt{6}} \right) e^{-\sqrt{6}t} - \frac{t}{60} - \frac{\cos(t)}{70} \right) \\ + \begin{bmatrix} 1 \\ -1 \end{bmatrix} \left(\frac{-9}{80} \sin(2t) + \frac{1}{30} \cos(2t) + \frac{9t}{40} - \frac{\cos(t)}{30} \right)$$

$$\mathbf{4.1.101:} \quad \omega = \pi \sqrt{\frac{15}{2}}$$

$$\mathbf{4.1.102:} \quad \lambda_k = 4k^2\pi^2 \text{ for } k = 1, 2, 3, \dots \quad x_k = \cos(2k\pi t) + B \sin(2k\pi t) \quad (\text{for any } B)$$

$$\mathbf{4.1.103:} \quad x(t) = -\sin(t)$$

4.1.104: General solution is $x = Ce^{-\lambda t}$. Since $x(0) = 0$ then $C = 0$, and so $x(t) = 0$. Therefore, the solution is always identically zero. One condition is always enough to guarantee a unique solution for a first-order equation.

$$\mathbf{4.1.105:} \quad \frac{\sqrt{3}}{3} e^{\frac{-3}{2}\sqrt[3]{\lambda}} - \frac{\sqrt{3}}{3} \cos\left(\frac{\sqrt{3}}{2}\sqrt[3]{\lambda}\right) + \sin\left(\frac{\sqrt{3}}{2}\sqrt[3]{\lambda}\right) = 0$$

$$\mathbf{4.2.101:} \quad \sin(t)$$

$$4.2.102: \sum_{n=1}^{\infty} \frac{2n(-1)^n \sin(\pi^2)}{\pi(\pi^2 - n^2)} \sin(nt)$$

$$4.2.103: \frac{1}{2} - \frac{1}{2} \cos(2t)$$

$$4.2.104: \frac{\pi^4}{5} + \sum_{n=1}^{\infty} \frac{(-1)^n (8\pi^2 n^2 - 48)}{n^4} \cos(nt)$$

$$4.3.101: \text{a) } \frac{8}{6} + \sum_{n=1}^{\infty} \frac{16(-1)^n}{\pi^2 n^2} \cos\left(\frac{n\pi}{2}t\right) \quad \text{b) } \frac{8}{6} - \frac{16}{\pi^2} \cos\left(\frac{\pi}{2}t\right) + \frac{4}{\pi^2} \cos(\pi t) - \frac{16}{9\pi^2} \cos\left(\frac{3\pi}{2}t\right) + \dots$$

$$4.3.102: \text{a) } \sum_{n=1}^{\infty} \frac{(-1)^{n+1} 2\lambda}{n\pi} \sin\left(\frac{n\pi}{\lambda}t\right) \quad \text{b) } \frac{2\lambda}{\pi} \sin\left(\frac{\pi}{\lambda}t\right) - \frac{\lambda}{\pi} \sin\left(\frac{2\pi}{\lambda}t\right) + \frac{2\lambda}{3\pi} \sin\left(\frac{3\pi}{\lambda}t\right) - \dots$$

$$4.3.103: f'(t) = \sum_{n=1}^{\infty} \frac{\pi}{n^2+1} \cos(n\pi t)$$

$$4.3.104: \text{a) } F(t) = \frac{t}{2} + C + \sum_{n=1}^{\infty} \frac{1}{n^4} \sin(nt) \quad \text{b) no}$$

$$4.3.105: \text{a) } \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} \sin(nt) \quad \text{b) } f \text{ is continuous at } t = \pi/2 \text{ so the Fourier series converges}$$

to $f(\pi/2) = \pi/4$. Obtain $\pi/4 = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{2n-1} = 1 - 1/3 + 1/5 - 1/7 + \dots$. c) Using the first 4 terms get $76/105 \approx 0.72$ (quite a bad approximation, you would have to take about 50 terms to start to get to within 0.01 of $\pi/4$).

$$4.3.106: \text{a) } F(0) = 1, \text{ b) } F(-1) = 0, \text{ c) } F(1) = 2, \text{ d) } F(-2) = 1, \text{ e) } F(4) = 1, \text{ f) } F(-9) = 0$$

$$4.4.101: \text{a) } 1/2 + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{\pi^2 n^2} \cos\left(\frac{n\pi}{3}t\right) \quad \text{b) } \sum_{n=1}^{\infty} \frac{2(-1)^{n+1}}{\pi n} \sin\left(\frac{n\pi}{3}t\right)$$

$$4.4.102: \text{a) } \cos(2t) \quad \text{b) } \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{4n}{\pi n^2 - 4\pi} \sin(nt)$$

$$4.4.103: \text{a) } f(t) \quad \text{b) } 0$$

$$4.4.104: \sum_{n=1}^{\infty} \frac{-1}{n^2(1+n^2)} \sin(nt)$$

$$4.4.105: \frac{t}{\pi} + \sum_{n=1}^{\infty} \frac{1}{2^n(\pi - n^2)} \sin(nt)$$

$$4.5.101: x = \frac{1}{\sqrt{2-4\pi^2}} \sin(2\pi t) + \frac{0.1}{\sqrt{2-100\pi^2}} \cos(10\pi t)$$

$$4.5.102: x = \sum_{n=1}^{\infty} \frac{e^{-n}}{3-(2n)^2} \cos(2nt)$$

$$4.5.103: x = \frac{1}{2\sqrt{3}} + \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{n^2\pi^2(\sqrt{3}-n^2\pi^2)} \cos(n\pi t)$$

$$4.5.104: x = \frac{1}{2\pi^2} - \frac{2}{\pi^3} t \sin(\pi t) + \sum_{\substack{n=3 \\ n \text{ odd}}}^{\infty} \frac{-4}{n^2\pi^4(1-n^2)} \cos(n\pi t)$$

$$4.6.101: u(x, t) = 5 \sin(x) e^{-3t} + 2 \sin(5x) e^{-75t}$$

$$4.6.102: u(x, t) = 1 + 2 \cos(x) e^{-0.1t}$$

$$4.6.103: u(x, t) = e^{\lambda t} e^{\lambda x} \text{ for some } \lambda$$

$$4.6.104: u(x, t) = A e^x + B e^t$$

$$4.6.105: \text{a) } 0, \quad \text{b) minimum } -100, \text{ maximum } 100, \quad \text{c) } t = \frac{\ln 2}{4\pi^2}.$$

$$4.7.101: y(x, t) = \sin(x)(\sin(t) + \cos(t))$$

$$4.7.102: y(x, t) = \frac{1}{5\pi} \sin(\pi x) \sin(5\pi t) + \frac{1}{100\pi} \sin(2\pi x) \sin(10\pi t)$$

$$4.7.103: y(x, t) = \sum_{n=1}^{\infty} \frac{2(-1)^{n+1}}{n} \sin(nx) \cos(n\sqrt{2}t)$$

$$4.7.104: y(x, t) = \sin(2x) + t \sin(x)$$

$$4.8.101: y(x, t) = \frac{\sin(2\pi(x-3t)) + \sin(2\pi(3t+x))}{2} + \frac{\cos(3\pi(x-3t)) - \cos(3\pi(3t+x))}{18\pi}$$

$$4.8.102: \text{a) } y(x, 0.1) = \begin{cases} x - x^2 - 0.04 & \text{if } 0.2 \leq x \leq 0.8 \\ 0.6x & \text{if } x \leq 0.2 \\ 0.6 - 0.6x & \text{if } x \geq 0.8 \end{cases}$$

$$\text{b) } y(x, 1/2) = -x + x^2 \quad \text{c) } y(x, 1) = x - x^2$$

$$4.8.103: \text{a) } y(1, 1) = -1/2 \quad \text{b) } y(4, 3) = 0 \quad \text{c) } y(3, 9) = 1/2$$

$$4.9.101: u(x, y) = \sum_{n=1}^{\infty} \frac{1}{n^2} \sin(n\pi x) \left(\frac{\sinh(n\pi(1-y))}{\sinh(n\pi)} \right)$$

$$4.9.102: u(x, y) = 0.1 \sin(\pi x) \left(\frac{\sinh(\pi(2-y))}{\sinh(2\pi)} \right)$$

$$4.10.101: u = 1 + \sum_{n=1}^{\infty} \frac{1}{n^2} r^n \sin(n\theta)$$

$$4.10.102: u = 1 - x$$

$$4.10.103: \text{a) } u = \frac{-1}{4} r^2 + \frac{1}{4} \quad \text{b) } u = \frac{-1}{4} r^2 + \frac{1}{4} + r^2 \sin(2\theta)$$

$$4.10.104: u(r, \theta) = \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{\rho^2 - r^2}{\rho^2 - 2r\rho \cos(\theta - \alpha) + r^2} g(\alpha) d\alpha$$

$$5.1.101: \lambda_n = \frac{n^2\pi^2}{4}, n = 1, 2, 3, \dots \text{ For odd } n, y_n = \cos\left(\frac{n\pi}{2}x\right), \text{ for even } n, y_n = \sin\left(\frac{n\pi}{2}x\right).$$

5.1.102: a) $p(x) = 1, q(x) = 0, r(x) = \frac{1}{x}, \alpha_1 = 1, \alpha_2 = 0, \beta_1 = 1, \beta_2 = 0$. The problem is not regular. b) $p(x) = 1 + x^2, q(x) = x^2, r(x) = 1, \alpha_1 = 1, \alpha_2 = 0, \beta_1 = 1, \beta_2 = 1$. The problem is regular.

$$5.2.101: y(x, t) = \sin(\pi x) \cos(2\pi^2 t)$$

$$5.2.102: 9y_{xxxx} + y_{tt} = 0 \quad (0 < x < 10, t > 0), \quad y(0, t) = y_x(0, t) = 0, \quad y(10, t) = y_x(10, t) = 0, \quad y(x, 0) = \sin^2(\pi x), \quad y_t(x, 0) = x(10 - x).$$

$$5.3.101: y_p(x, t) = \sum_{\substack{n=1 \\ n \text{ odd}}}^{\infty} \frac{-4}{\pi n^4} \left(\cos(nx) - \frac{\cos(n)-1}{\sin(n)} \sin(nx) - 1 \right) \cos(nt).$$

5.3.102: Approximately 1991 centimeters

6.1.101: $\frac{8}{s^3} + \frac{8}{s^2} + \frac{4}{s}$

6.1.102: $2t^2 - 2t + 1 - e^{-2t}$

6.1.103: $\frac{1}{(s+1)^2}$

6.1.104: $\frac{1}{s^2+2s+2}$

6.2.101: $f(t) = (t-1)(u(t-1) - u(t-2)) + u(t-2)$

6.2.102: $x(t) = (2e^{t-1} - t^2 - 1)u(t-1) - \frac{1}{2}e^{-t} + \frac{3}{2}e^t$

6.2.103: $H(s) = \frac{1}{s+1}$

6.2.104: $F(s) = \frac{4}{s}, X(s) = \frac{1}{s} - \frac{s}{s^2+4}, H(s) = \frac{1}{s^2+4}.$

6.3.101: $\frac{1}{2}(\cos t + \sin t - e^{-t})$

6.3.102: $5t - 5 \sin t$

6.3.103: $\frac{1}{2}(\sin t - t \cos t)$

6.3.104: $\int_0^t f(\tau)(1 - \cos(t - \tau)) d\tau$

6.4.101: $x(t) = t$

6.4.102: $x(t) = e^{-at}$

6.4.103: $x(t) = (e^t) * (e^t \sin(t)) = e^t(1 - \cos(t))$

6.4.104: $\delta(t) - \sin(t)$

6.4.105: $3\delta(t-1) + 2t$

6.5.101: $y = (x-t)u(t-x) + t$

6.5.102: $y = e^{-cx} \sin(t-x)u(t-x)$

6.5.103: $s^2Y(x) - s(1-x^2) + 3Y''(x) + Y(x) = \frac{x}{s} + \frac{1}{s^2}, \quad Y(-1) = 0, \quad Y(1) = 0.$

6.5.104: $y = \sin(x) \cos(t)$

6.5.105: $y(x, t) = \int_0^t f(\tau) \frac{x}{2\sqrt{\pi(t-\tau)^3}} e^{\frac{-x^2}{4(t-\tau)}} d\tau$

7.1.101: Yes. Radius of convergence is 10.

7.1.102: Yes. Radius of convergence is e .

7.1.103: $\frac{1}{1-x} = -\frac{1}{1-(2-x)}$ so $\frac{1}{1-x} = \sum_{n=0}^{\infty} (-1)^{n+1} (x-2)^n$, which converges for $1 < x < 3$.

7.1.104: $\sum_{n=7}^{\infty} \frac{1}{(n-7)!} x^n$

7.1.105: $f(x) - g(x)$ is a polynomial. Hint: Use Taylor series.

7.1.106: a) $\sum_{k=6}^{\infty} (k-3)(k-4)x^k$ b) $\sum_{k=0}^{\infty} (k+2)x^k$ c) $\sum_{k=3}^{\infty} 2(k+2)x^k$

7.2.101: $a_2 = 0, a_3 = 0, a_4 = 0$, recurrence relation (for $k \geq 5$): $a_k = \frac{-2a_{k-5}}{k(k-1)}$, so:

$$y(x) = a_0 + a_1x - \frac{a_0}{10}x^5 - \frac{a_1}{15}x^6 + \frac{a_0}{450}x^{10} + \frac{a_1}{825}x^{11} - \frac{a_0}{47250}x^{15} - \frac{a_1}{99000}x^{16} + \dots$$

7.2.102: a) $a_2 = \frac{1}{2}$, and for $k \geq 3$ we have $a_k = \frac{a_{k-3}+1}{k(k-1)}$, so

$$y(x) = a_0 + a_1x + \frac{1}{2}x^2 + \frac{a_0+1}{6}x^3 + \frac{a_1+1}{12}x^4 + \frac{3}{40}x^5 + \frac{a_0+7}{180}x^6 + \frac{a_1+13}{504}x^7 + \frac{43}{2240}x^8 + \frac{a_0+187}{12960}x^9 + \frac{a_1+517}{45360}x^{10} + \dots$$

$$\text{b) } y(x) = \frac{1}{2}x^2 + \frac{1}{6}x^3 + \frac{1}{12}x^4 + \frac{3}{40}x^5 + \frac{7}{180}x^6 + \frac{13}{504}x^7 + \frac{43}{2240}x^8 + \frac{187}{12960}x^9 + \frac{517}{45360}x^{10} + \dots$$

7.2.103: Applying the method of this section directly we obtain $a_k = 0$ for all k and so $y(x) = 0$ is the only solution we find.

7.3.101: a) ordinary, b) singular but not regular singular, c) regular singular, d) regular singular, e) ordinary.

$$\text{7.3.102: } y = Ax^{\frac{1+\sqrt{5}}{2}} + Bx^{\frac{1-\sqrt{5}}{2}}$$

$$\text{7.3.103: } y = x^{3/2} \sum_{k=0}^{\infty} \frac{(-1)^k}{k!(k+2)!} x^k \quad (\text{Note that for convenience we did not pick } a_0 = 1.)$$

$$\text{7.3.104: } y = Ax + Bx \ln(x)$$

8.1.101: a) Critical points $(1, 0)$ and $(1, 1)$. At $(1, 0)$ using $u = x - 1, v = y$ the linearization is $u' = \pi v, v' = -v$. At $(1, 1)$ using $u = x - 1, v = y - 1$ the linearization is $u' = -\pi v, v' = v$.

b) Critical points $(0, 0)$ and $(0, -1)$. At $(0, 0)$ using $u = x, v = y$ the linearization is $u' = u + v, v' = u$. At $(0, -1)$ using $u = x, v = y + 1$ the linearization is $u' = u - v, v' = u$.

c) Critical point $(1/2, -1/4)$. Using $u = x - 1/2, v = y + 1/4$ the linearization is $u' = -u + v, v' = u + v$.

8.1.102: 1) is c), 2) is a), 3) is b)

8.1.103: Critical points are $(0, 0, 0)$, and $(-1, 1, -1)$. The linearization at the origin using variables $u = x, v = y, w = z$ is $u' = u, v' = -v, w' = w$. The linearization at the point $(-1, 1, -1)$ using variables $u = x + 1, v = y - 1, w = z + 1$ is $u' = u - 2w, v' = -v - 2w, w' = w - 2u$.

$$\text{8.1.104: } u' = f(u, v, w), v' = g(u, v, w), w' = 1.$$

8.2.101: a) $(0, 0)$: saddle (unstable), $(1, 0)$: source (unstable), b) $(0, 0)$: spiral sink (asymptotically stable), $(0, 1)$: saddle (unstable), c) $(1, 0)$: saddle (unstable), $(0, 1)$: source (unstable)

8.2.102: a) $\frac{1}{2}y^2 + \frac{1}{3}x^3 - 4x = C$, critical points: $(-2, 0)$, an unstable saddle, and $(2, 0)$, a stable center. b) $\frac{1}{2}y^2 + e^x = C$, no critical points. c) $\frac{1}{2}y^2 + xe^x = C$, critical point at $(-1, 0)$ is a stable center.

8.2.103: Critical point at $(0, 0)$. Trajectories are $y = \pm\sqrt{2C - (1/2)x^4}$, for $C > 0$, these give closed curves around the origin, so the critical point is a stable center.

8.2.104: A critical point x_0 is stable if $f'(x_0) < 0$ and unstable when $f'(x_0) > 0$.

8.3.101: a) Critical points are $\omega = 0, \theta = k\pi$ for any integer k . When k is odd, we have a saddle point. When k is even we get a sink. b) The findings mean the pendulum will simply go to one of the sinks, for example $(0, 0)$ and it will not swing back and forth. The friction is too high for it to oscillate, just like an overdamped mass-spring system.

8.3.102: a) Solving for the critical points we get $(0, -h/d)$ and $(\frac{bh+ad}{ac}, \frac{a}{b})$. The Jacobian matrix at $(0, -h/d)$ is $\begin{bmatrix} a+bh/d & 0 \\ -ch/d & -d \end{bmatrix}$ whose eigenvalues are $a + bh/d$ and $-d$. The eigenvalues are real of opposite signs and we get a saddle. (In the application, however, we are only looking at the positive quadrant so this critical point is irrelevant.) At $(\frac{bh+ad}{ac}, \frac{a}{b})$ we get Jacobian matrix $\begin{bmatrix} 0 & -\frac{b(bh+ad)}{ac} \\ \frac{ac}{b} & \frac{bh+ad}{a} - d \end{bmatrix}$. b) For the specific numbers given, the second critical point is $(\frac{550}{3}, 40)$ the matrix is $\begin{bmatrix} 0 & -11/6 \\ 3/25 & 1/4 \end{bmatrix}$, which has eigenvalues $\frac{5 \pm i\sqrt{327}}{40}$. Therefore there is a spiral source; the solution spirals outwards. The solution eventually hits one of the axes, $x = 0$ or $y = 0$, so something will die out in the forest.

8.3.103: The critical points are on the line $x = 0$. In the positive quadrant the y' is always positive and so the fox population always grows. The constant of motion is $C = y^a e^{-cx-by}$. For any C , this curve must hit the y -axis (why?), so the trajectory will simply approach a point on the y axis somewhere and the number of hares will go to zero.

8.4.101: Use Bendixson–Dulac Theorem. a) $f_x + g_y = 1 + 1 > 0$, so no closed trajectories. b) $f_x + g_y = -\sin^2(y) + 0 < 0$ for all x, y except the lines given by $y = k\pi$ (where we get zero), so no closed trajectories. c) $f_x + g_y = y^2 + 0 > 0$ for all x, y except the line given by $y = 0$ (where we get zero), so no closed trajectories.

8.4.102: Using Poincaré–Bendixson Theorem, the system has a limit cycle, which is the unit circle centered at the origin, as $x = \cos(t)(1 + e^{-t})$, $y = \sin(t)(1 + e^{-t})$ gets closer and closer to the unit circle. Thus $x = \cos(t)$, $y = \sin(t)$ is the periodic solution.

8.4.103: $f(x, y) = y$, $g(x, y) = \mu(1 - x^2)y - x$. So $f_x + g_y = \mu(1 - x^2)$. The Bendixson–Dulac Theorem says there is no closed trajectory lying entirely in the set $x^2 < 1$.

8.4.104: The closed trajectories are those where $\sin(r) = 0$, therefore, all the circles centered at the origin with radius that is a positive multiple of π are closed trajectories.

8.5.101: Critical points: $(0, 0, 0)$, $(6\sqrt{2}, 6\sqrt{2}, 27)$, $(-6\sqrt{2}, -6\sqrt{2}, 27)$. Linearization at $(0, 0, 0)$ using $u = x$, $v = y$, $w = z$ is $u' = -10u + 10v$, $v' = 28u - v$, $w' = -(8/3)w$. Linearization at $(6\sqrt{2}, 6\sqrt{2}, 27)$ using $u = x - 6\sqrt{2}$, $v = y - 6\sqrt{2}$, $w = z - 27$ is $u' = -10u + 10v$, $v' = u - v - 6\sqrt{2}w$, $w' = 6\sqrt{2}u + 6\sqrt{2}v - (8/3)w$. Linearization at $(-6\sqrt{2}, -6\sqrt{2}, 27)$ using $u = x + 6\sqrt{2}$, $v = y + 6\sqrt{2}$, $w = z - 27$ is $u' = -10u + 10v$, $v' = u - v + 6\sqrt{2}w$, $w' = -6\sqrt{2}u - 6\sqrt{2}v - (8/3)w$.

A.1.101: a) $\sqrt{10}$ b) $\sqrt{14}$ c) 3

A.1.102: a) $\begin{bmatrix} -1 \\ \sqrt{2} \\ 1 \\ \sqrt{2} \end{bmatrix}$ b) $\begin{bmatrix} \frac{1}{\sqrt{6}} \\ \frac{-1}{\sqrt{6}} \\ \frac{2}{\sqrt{6}} \end{bmatrix}$ c) $(\frac{2}{\sqrt{33}}, \frac{-5}{\sqrt{33}}, \frac{2}{\sqrt{33}})$

A.1.103: a) $\begin{bmatrix} 9 \\ -2 \end{bmatrix}$ b) $\begin{bmatrix} -3 \\ 3 \end{bmatrix}$ c) $\begin{bmatrix} 5 \\ -3 \end{bmatrix}$ d) $\begin{bmatrix} -4 \\ 8 \end{bmatrix}$ e) $\begin{bmatrix} 3 \\ 7 \end{bmatrix}$ f) $\begin{bmatrix} -8 \\ 0 \end{bmatrix}$

A.1.104: a) 20 b) 10 c) 20

A.1.105: a) $(3, -1)$ b) $(4, 0)$ c) $(-1, -1)$

A.2.101: a) $\begin{bmatrix} 7 & 4 & 4 \\ 2 & 3 & 4 \end{bmatrix}$ b) $\begin{bmatrix} 5 & -3 & 0 \\ 13 & 10 & 6 \\ -1 & 3 & 1 \end{bmatrix}$

A.2.102: a) $\begin{bmatrix} -1 & 13 \\ 9 & 14 \end{bmatrix}$ b) $\begin{bmatrix} 2 & -5 \\ 5 & 5 \end{bmatrix}$

A.2.103: a) $\begin{bmatrix} 22 & 31 \\ 42 & 44 \end{bmatrix}$ b) $\begin{bmatrix} 18 & 18 & 12 \\ 6 & 0 & 8 \\ 34 & 48 & -2 \end{bmatrix}$ c) $\begin{bmatrix} 11 & 12 & 36 & 14 \\ -2 & 4 & 5 & -2 \\ 13 & 38 & 20 & 28 \end{bmatrix}$ d) $\begin{bmatrix} -2 & -12 \\ 3 & 24 \\ 1 & 9 \end{bmatrix}$

A.2.104: a) $[1/2]$ b) $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$ c) $\begin{bmatrix} -5 & 2 \\ 3 & -1 \end{bmatrix}$ d) $\begin{bmatrix} 1/2 & -1/4 \\ -1/2 & 1/2 \end{bmatrix}$

A.2.105: a) $\begin{bmatrix} 1/2 & 0 \\ 0 & 1/3 \end{bmatrix}$ b) $\begin{bmatrix} 1/4 & 0 & 0 \\ 0 & 1/5 & 0 \\ 0 & 0 & -1 \end{bmatrix}$ c) $\begin{bmatrix} -1 & 0 & 0 & 0 \\ 0 & 1/2 & 0 & 0 \\ 0 & 0 & 1/3 & 0 \\ 0 & 0 & 0 & 10 \end{bmatrix}$

A.3.101: a) $\begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix}$ b) $\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$ c) $\begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix}$ d) $\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & -1/3 \\ 0 & 0 & 0 \end{bmatrix}$ e) $\begin{bmatrix} 1 & 0 & 0 & 77/15 \\ 0 & 1 & 0 & -2/15 \\ 0 & 0 & 1 & -8/5 \end{bmatrix}$
 f) $\begin{bmatrix} 1 & 0 & -1/2 & 0 \\ 0 & 1 & 1/2 & 1/2 \\ 0 & 0 & 0 & 0 \end{bmatrix}$ g) $\begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$ h) $\begin{bmatrix} 1 & 2 & 3 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$

A.3.102: a) $\begin{bmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix}$ b) $\begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & -1 \\ 1 & -1 & 0 \end{bmatrix}$ c) $\begin{bmatrix} 5/2 & 1 & -3 \\ -1 & -1/2 & 3/2 \\ -1 & 0 & 1 \end{bmatrix}$

A.3.103: a) $x_1 = -2, x_2 = 7/3$ b) no solution c) $a = -3, b = 10, c = -8$ d) x_3 is free, $x_1 = -1 + 3x_3, x_2 = 2 - x_3$

A.3.104: a) $\begin{bmatrix} -1 \\ 3 \end{bmatrix}$ b) $\begin{bmatrix} -3 \\ 1 \end{bmatrix}$

A.3.105: a) 3 b) 1 c) 2

A.3.106: a) $[1 \ 0 \ 0], [0 \ 1 \ 0], [0 \ 0 \ 1]$ b) $[1 \ 1 \ 1]$ c) $[1 \ 0 \ 1/3], [0 \ 1 \ -1/3]$

A.3.107: a) $\begin{bmatrix} 7 \\ 7 \\ 7 \end{bmatrix}, \begin{bmatrix} -1 \\ 7 \\ 6 \end{bmatrix}, \begin{bmatrix} 7 \\ 6 \\ 2 \end{bmatrix}$ b) $\begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix}$ c) $\begin{bmatrix} 0 \\ 6 \\ 4 \end{bmatrix}, \begin{bmatrix} 3 \\ 3 \\ 7 \end{bmatrix}$

A.3.108: $\begin{bmatrix} 3 \\ 1 \\ -5 \end{bmatrix}, \begin{bmatrix} 0 \\ 3 \\ -1 \end{bmatrix}$

A.4.101: a) $\begin{bmatrix} 1 \\ 2 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \end{bmatrix}$ dimension 2, b) $\begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix}$ dimension 2, c) $\begin{bmatrix} 5 \\ 3 \\ 1 \end{bmatrix}, \begin{bmatrix} 5 \\ -1 \\ 5 \end{bmatrix}, \begin{bmatrix} -1 \\ 3 \\ -4 \end{bmatrix}$ dimension 3, d) $\begin{bmatrix} 2 \\ 2 \\ 4 \end{bmatrix}, \begin{bmatrix} 2 \\ 2 \\ 3 \end{bmatrix}$ dimension 2, e) $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ dimension 1, f) $\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ 2 \end{bmatrix}$ dimension 2

A.4.102: a) $\begin{bmatrix} 3 \\ -1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 3 \\ 0 \\ 3 \\ -1 \end{bmatrix}$ b) $\begin{bmatrix} -1 \\ -1 \\ 0 \end{bmatrix}$ c) $\begin{bmatrix} 1 \\ 1 \\ -1 \end{bmatrix}$ d) $\begin{bmatrix} -1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}$

A.4.103: a) 3 b) 2 c) 3 d) 2 e) 3

A.5.101: $s = -2$

A.5.102: $\theta \approx 0.3876$

A.5.103: a) -15 b) -1 c) 28

A.5.104: a) $(-1/2, 0, 1/2)$ b) $(0, 0, 0)$ c) $(2, 0, -2)$

A.5.105: a) $(1, 1, -1) - (2, -1, 1) + 2(0, 1, 1)$ b) $2(2, -1, 1) + (0, 1, 1)$ c) $2(1, 1, -1) - 2(2, -1, 1) + 2(0, 1, 1)$

A.5.106: $(2, -1, 1), (2/3, 8/3, 4/3)$

A.5.107: $(1, 1, -1), (0, 1, 1), (4/3, -2/3, 2/3)$

A.6.101: a) -2 b) 8 c) 0 d) -6 e) -3 f) 28 g) 16 h) -24

A.6.102: a) 3 b) 9 c) ± 3 d) $1/4$

A.6.103: $1/12$

A.6.104: 1 and 3

Index

- absolute convergence, 328
- acceleration, 24
- adding vectors, 387
- addition of matrices, 127
- Airy's equation, 336
- algebraic multiplicity, 162
- almost linear, 357
- amplitude, 98
- analytic functions, 330
- angular frequency, 98
- ansatz, 84
- antiderivative, 22
- antidifferentiate, 22
- associated homogeneous equation, 104, 138
- associative law, 400
- asymptotically stable, 358
- asymptotically stable limit cycle, 374
- atan, 99
- atan2, 99
- attractor, 381
- augmented matrix, 132, 408
- autonomous, 19
- autonomous equation, 51
- autonomous system, 124

- backsubstitution, 411
- basis, 391, 421
- beating, 112
- Bendixson–Dulac Theorem, 376
- Bernoulli equation, 47
- Bessel function of second kind, 348
- Bessel function of the first kind, 348
- Bessel's equation, 347
- boundary conditions for a PDE, 232
- boundary value problem, 189
- Burgers' equation, 18

- cartesian plane, 386
- catenary, 14
- Cauchy–Euler equation, 82, 89
- center, 149, 359
- cgs units, 291, 292
- chaotic systems, 378
- characteristic coordinates, 72, 253
- characteristic curve, 73
- characteristic equation, 85
- characteristics, 72
- Chebyshev's equation of order p , 340
- Chebyshev's equation of order 1, 83
- clamped end of beam, 284
- closed curves, 362
- closed form, 22
- closed trajectory, 373
- coefficients, 18
- cofactor, 130, 438
- cofactor expansion, 130, 437
- column space, 416
- column vector, 127, 387
- commute, 129
- complementary solution, 104
- complete eigenvalue, 162
- complex conjugate, 143
- complex number, 86
- complex roots, 87
- conservative equation, 361
- conservative vector field, 64
- consistent, 412
- constant of motion, 370
- constant-coefficient, 19, 84, 137
- convection, 73
- convection equation, 320
- convergence of a power series, 328

- convergent power series, 328
- converges absolutely, 328
- convolution, 308
- corresponding eigenfunction, 190
- cosine series, 220
- critical point, 51, 352
- critically damped, 101
- d'Alembert solution to the wave equation, 252
- damped, 99
- damped motion, 96
- damped nonlinear pendulum equation, 371
- defect, 163
- defective eigenvalue, 163
- deficient matrix, 163
- delta function, 314
- dependent variable, 10
- determinant, 129, 436
- diagonal matrix, 153, 403
 - matrix exponential of, 170
- diagonalization, 171
- differential equation, 10
- Dirac delta function, 314
- direction, 386, 389
- direction field, 124
- Dirichlet boundary conditions, 221, 273
- Dirichlet problem, 259
- displacement vector, 153
- distance, 24
- distributive law, 400
- divergence, 376
- divergent power series, 328
- dot product, 128, 199, 398
- Duffing equation, 379
- dynamic damping, 161
- echelon form, 409
- eigenfunction, 190, 274
- eigenfunction decomposition, 273, 278
- eigenvalue, 140, 274
- eigenvalue of a boundary value problem, 190
- eigenvector, 140
 - eigenvector decomposition, 179, 186
- element of a matrix, 127, 390
- elementary operations, 408
- elementary row operations, 132, 408
- elimination, 407
- ellipses (vector field), 149
- elliptic PDE, 232
- endpoint problem, 189
- entry of a matrix, 127, 390
- envelope curves, 101
- equilibrium, 51, 353
- equilibrium solution, 51, 353
- equipotential curves, 63
- euclidean space, 386
- Euler's equation, 82
- Euler's formula, 87
- Euler's method, 57
- Euler–Bernoulli equation, 317
- even function, 202, 218
- even periodic extension, 218
- exact equation, 63
- existence and uniqueness, 30, 80, 90
- existence and uniqueness for systems, 124
- exponential growth, 17
- exponential growth model, 12
- exponential of a matrix, 169
- exponential order, 296
- extend periodically, 198
- first shifting property, 298
- first-order differential equation, 10
- first-order linear equation, 40
- first-order linear system of ODEs, 136
- first-order method, 58
- first-order system, 119
- fixed end of beam, 284
- forced motion, 96
 - systems, 159
- four fundamental equations, 13
- Fourier series, 200
- fourth-order method, 60
- Fredholm alternative, 426
 - simple case, 194

- Sturm–Liouville problems, 278
- free end of beam, 284
- free motion, 96
- free variable, 133, 412
- frequency, 98
- Frobenius method, 345
- Frobenius-type solution, 345
- fundamental frequency, 205, 248
- fundamental matrix, 137
- fundamental matrix solution, 137, 170
- Gauss–Jordan elimination, 409
- Gaussian elimination, 409
- general solution, 13
- general solution to a system, 120
- generalized eigenvectors, 163, 166
- generalized function, 314
- Genius software, 9
- geometric multiplicity, 162
- geometric series, 270, 333
- Gibbs phenomenon, 205
- Gram–Schmidt process, 433
- half period, 208
- Hamiltonian, 361
- harmonic, 200
- harmonic conjugate, 71
- harmonic function, 71, 258
- harvesting, 53
- heat equation, 17, 232
- Heaviside function, 294
- Hermite’s equation of order n , 339
- Hermite’s equation of order 2, 83
- hinged end of beam, 284
- homogeneous, 19
- homogeneous equation, 48
- homogeneous linear equation, 79
- homogeneous matrix equation, 413
- homogeneous side conditions, 233
- homogeneous system, 137
- Hooke’s law, 96, 152
- hyperbolic cosine, 14
- hyperbolic PDE, 232
- hyperbolic sine, 14
- identity matrix, 128, 400
- ill-posed, 323
- imaginary part, 87
- implicit solution, 35
- Improved Euler’s method, 62
- impulse response, 313, 315
- inconsistent, 412
- inconsistent system, 133
- indefinite integral, 22
- independent variable, 10
- indicial equation, 344
- inhomogeneous, 19
- initial condition, 13
- initial condition for a PDE, 72
- initial condition for a system, 120
- initial conditions for a PDE, 232
- inner product, 128, 398, 428
- inner product of functions, 200, 277
- integral equation, 305, 311
- integrate, 22
- integrating factor, 40
- integrating factor method, 40
 - for systems, 177
- inverse Laplace transform, 297
- invertible matrix, 129, 401
- IODE software, 9
- isolated critical point, 357
- Jacobian matrix, 354
- kernel, 413
- la vie, 106
- Laplace equation, 232, 258
- Laplace transform, 293
- Laplacian, 258
- Laplacian in polar coordinates, 265
- leading entry, 133, 409
- Leibniz notation, 23, 33
- limit cycle, 373
- linear combination, 80, 90, 391, 413
- linear equation, 18, 40, 79

- linear first-order system, 120
- linear mapping, 390
- linear operator, 80, 104
- linear PDE, 232
- linear systems, 120
- linearity of the Laplace transform, 295
- linearization, 354
- linearly dependent, 90, 414
- linearly independent, 81, 90, 414
 - for vector-valued functions, 137
- logistic equation, 51
 - with harvesting, 53
- Lorenz attractor, 383
- Lorenz system, 382
- Lotka–Volterra, 368
- lower triangular, 438
- magnitude, 387
- mass matrix, 153
- mathematical model, 12
- mathematical solution, 12
- matrix, 127, 390
- matrix exponential, 169
- matrix inverse, 129, 401
- matrix product, 399
- matrix-valued function, 136
- Maxwell's equations, 17
- mechanical vibrations, 17
- method of characteristics, 72
- Method of Frobenius, 345
- method of partial fractions, 297
- Mixed boundary conditions, 273
- mks units, 99, 102, 227
- multiplication of complex numbers, 86
- multiplicity, 93
- multiplicity of an eigenvalue, 162
- n -dimensional space, 386
- natural (angular) frequency, 98
- natural frequency, 111, 155
- natural mode of oscillation, 155
- Neumann boundary conditions, 222, 273
- Newton's law of cooling, 17, 37, 44, 51
- Newton's second law, 96, 97, 122, 152
- nilpotent, 171
- nonhomogeneous, 19
- nonlinear equation, 18
- normal mode of oscillation, 155
- nullity, 424
- nullspace, 413
- odd function, 201, 218
- odd periodic extension, 218
- ODE, 11, 17
- one-dimensional heat equation, 232
- one-dimensional wave equation, 243
- operator, 80, 390
- order, 17
- Ordinary differential equation, 17
- ordinary differential equation, 11
- ordinary point, 335
- orthogonal, 429
 - functions, 193, 200
 - vectors, 198
 - with respect to a weight, 276
- orthogonal basis, 431
- orthogonal projection, 430
- orthogonality, 193
- orthonormal basis, 431
- overdamped, 100
- overtones, 205, 248
- parabolic PDE, 232
- parallelogram, 130, 436
- Partial differential equation, 17
- partial differential equation, 11, 232
- partial sum, 327
- particular solution, 13, 104
- PDE, 11, 17, 232
- period, 98
- periodic, 198
- periodic extension, 198
- periodic orbit, 374
- phase diagram, 52, 352
- phase plane portrait, 124
- phase portrait, 52, 124, 352

- phase shift, 98
- Picard's theorem, 30, 124
- piecewise continuous, 211
- piecewise smooth, 211
- pivot, 409
- Poincaré Lemma, 65
- Poincaré section, 380
- Poincaré–Bendixson Theorem, 374
- Poisson kernel, 269
- potential function, 63
- power series, 327
- practical resonance, 115
- practical resonance amplitude, 115
- practical resonance frequency, 115
- practical resonance,, 231
- predator-prey, 368
- product of matrices, 128, 399
- projection, 200
 - orthogonal, 430
- proper rational function, 298
- pseudo-frequency, 102
- pulse, 313
- pure resonance, 113, 229
- quadratic formula, 85
- radius of convergence, 328
- rank, 415
- ratio test for series, 329
- real part, 87
- real-world problem, 12
- rectangular pulse, 313
- recurrence relation, 336
- reduced row echelon form, 133, 409
- reduction of order method, 81
- regular singular point, 345
- regular Sturm–Liouville problem, 275
- reindexing the series, 332
- relaxation oscillation, 373
- repeated roots, 92
- resonance, 113, 160, 229, 310
- RLC circuit, 96
- Robin boundary conditions, 273
- root test for series, 329
- row echelon form, 409
- row operations, 132, 408
- row reduction, 409
- row space, 416
- row vector, 127, 391
- Runge–Kutta method, 61
- saddle point, 148
- sawtooth, 201, 306
- scalar, 127, 388
- scalar multiplication, 127
- scale a vector, 388
- second shifting property, 302
- second-order differential equation, 14
- second-order linear differential equation, 79
- second-order method, 58
- second-order system, 119
- semistable critical point, 53
- separable, 33
- separation of variables, 234
- shifting property, 298, 302
- side conditions for a PDE, 232
- Sierpinski triangle, 384
- simple harmonic motion, 98
- simply connected region, 375
- sine series, 220
- singular matrix, 129, 401
- singular point, 335
- singular solution, 35
- sink, 147
- slope field, 27
- solution, 10
- solution curve, 124
- solution to a system, 119
- source, 147
- span, 415
- spectrum, 205, 248
- spiral sink, 150
- spiral source, 149
- square matrix, 127, 391
- square wave, 116, 203
- stable center, 359

- stable critical point, 51, 357
- stable node, 147
- standard basis vectors, 391
- standard inner product, 428
- steady periodic solution, 114, 226
- steady-state temperature, 241, 258
- step size, 57
- stiff problem, 60
- stiffness matrix, 153
- strange attractor, 380
- Sturm–Liouville problem, 274
 - regular, 275
- subspace, 421
- subtracting vectors, 388
- superposition, 79, 90, 137, 233
- symmetric matrix, 193, 198, 404
- system of differential equations, 17, 119

- Taylor series, 330
- tedious, 106, 108, 114, 182, 269, 378
- thermal diffusivity, 232
- three mass system, 152
- three-point beam bending, 317
- timbre, 248, 284
- total derivative, 63
- trajectory, 124
- transfer function, 304
- transformation, 390
- transient solution, 114
- transport equation, 17, 72, 320
- transpose, 127, 404
- transversal vibrations, 282
- trigonometric series, 200

- undamped, 98
- undamped motion, 96
 - systems, 152
- underdamped, 101
- undetermined coefficients, 105
 - for second-order systems, 159, 185
 - for systems, 182
- unforced motion, 96
- unit step function, 294
- unit vector, 389
- unstable critical point, 51, 357
- unstable node, 147
- upper triangular, 438
- upper triangular matrix, 164

- Van der Pol oscillator, 373
- variable-coefficient, 19
- variation of parameters, 108
 - for systems, 184
- vector, 127, 386
- vector field, 64, 124, 352
- vector-valued function, 136, 389
- velocity, 24
- Volterra integral equation, 311

- wave equation, 232, 243, 252
- wave equation in 2 dimensions, 17
- wavefronts, 256
- weight function, 276