

But Who Will Monitor the Monitor?^{†a}

By DAVID RAHMAN*

Suppose that providing incentives for a group of individuals in a strategic context requires a monitor to detect their deviations. What about the monitor's deviations? To address this question, I propose a contract that makes the monitor responsible for monitoring, and thereby provides incentives even when the monitor's observations are not only private, but costly, too. I also characterize exactly when such a contract can provide monitors with the right incentives to perform. In doing so, I emphasize virtual enforcement and suggest its implications for the theory of repeated games. (JEL C78, D23, D82, D86)

Ann owns a restaurant. She hires Bob to tally the till every night and report back any mismatch between the till and that night's bills. Ann is too busy to check the till herself and has to trust what Bob says. How can Ann provide Bob with appropriate incentives to exert the effort required to tally the till and report back the truth?

Ann's problem, basic as it is, seems to have eluded systematic analysis by economists. In studying incentives, most economists have focused on output-contingent contracts, such as bonuses for sales reps.¹ Thus, a great way of convincing a salesperson to exert effort is to promise him or her a greater reward the more he or she sells. This contract, however, gives Bob perverse incentives, since only he can know if there is a mismatch between the till and the bills. Hence, if Ann paid Bob a bonus for reporting a mismatch he would just report it without tallying the till, and similarly if the bonus was for reporting no mismatch. Some economists have emphasized auditing, perhaps at random to economize on its cost.² Unfortunately, Ann is just too busy to tally the till herself—this option is not credibly available to her. In any case, the solution I will propose economizes on auditing by avoiding

*Department of Economics, University of Minnesota, 4-101 Hanson Hall, 1925 Fourth Street South, Minneapolis, MN 55455 (e-mail: dmr@umn.edu). Financial support from the Spanish Ministry of Education's Research Grant No. SEJ 2004-07861 while at Universidad Carlos III de Madrid, the National Science Foundation's Grant No. SES 09-22253, and Grant-in-Aid No. 21600 from the Office of the Vice President of Research, University of Minnesota, are gratefully acknowledged. An early version of this paper was circulated under the title "Optimum Contracts with Public and Private Monitoring," which was based on Chapter 3 of my PhD dissertation at UCLA. I owe many thanks to Ben Bolitzer, Antonio Cabrales, V. V. Chari, Harold Demsetz, Andrew Dust (for excellent research assistance), Willie Fuchs, Larry Jones, Narayana Kocherlakota, David Levine, Roger Myerson, Ichiro Obara (whose collaboration on a related paper spilled over into this one), Joe Ostroy, Bill Zame, and numerous seminar audiences for insightful comments that helped me tremendously.

[†] To view additional materials, visit the article page at <http://dx.doi.org/10.1257/aer.102.6.2767>.

^a Alchian and Demsetz (1972, p. 782).

¹ Classic examples are Grossman and Hart (1983) and Holmström (1982), but Hermalin and Katz (1991); Legros and Matsushima (1991); Legros and Matthews (1993); and Strausz (1997) are also relevant.

² Starting from the costly state verification model of Townsend (1979), a vast literature includes Baron and Besanko (1984); Border and Sobel (1987); Mookherjee and Png (1989, 1992); and Kaplow and Shavell (1994). In this paper, I depart from this literature by allowing for the cost of state verification to be *infinite*.

it altogether, so it applies even if Ann's own tallying cost is exorbitant. Recently, incentives for truth-telling have been suggested,³ which in this setting boil down to paying Bob the same amount regardless of what he says, making him indifferent between honesty and any deception. This contract cannot help Ann either because then nothing would prevent Bob from neglecting to tally the till. Finally, mutual monitoring has also been studied in the literature.⁴ In these tough times, however, Ann cannot afford to hire anyone besides Bob.

I propose that Ann can solve her problem by sometimes secretly taking money from the till and offering Bob the following deal: if Ann took some money, she will pay Bob only when he reports a mismatch; if Ann did not take any money, she will pay Bob only when a mismatch is not reported. Bob's incentives are now aligned with Ann's. If Bob doesn't bother tallying the till, he won't know what to tell Ann in order to make sure he gets paid. On the other hand, if he does his job he'll discover whether or not there is a mismatch and deduce whether or not Ann took some money. Only then will Bob know what to tell Ann in order to get paid. By asking Bob a "trick question," Ann can now rest assured that he will acquire the requisite costly information and reveal it truthfully.

The insight behind Bob's contract has far-reaching consequences for the role of monitoring in organizations—exploring them is the purpose of this paper. Since Alchian and Demsetz (1972) posed the seminal question of how to remunerate monitors, it has generated much academic debate. Previously, a monitor's observations were assumed to be either verifiable (footnotes 1 and 2), costless (footnote 3), or mutual (footnote 4). I add to the debate by studying a theoretical model that accommodates unilateral, costly private monitoring. As I suggested already, existing solutions to apparently related problems cannot provide the right incentives in this richer environment. Nevertheless, using a version of Bob's contract I show how to make monitors responsible for monitoring. More broadly, I intuitively characterize the set of social outcomes that are attainable with such a contract.

I begin (Section I) by considering a firm with two agents: a worker and a monitor. I design a contract that constitutes a *communication equilibrium* (Forges 1986; Myerson 1986), with payments that crucially depend on both effort recommendations by the owner and the monitor's reports. Occasionally, the owner secretly asks the worker to shirk (or do something different), and rewards the monitor for "catching" these prompted deviations.⁵ Such a contract rewards the monitor for reporting accuracy in a way that the owner can confirm, thus respecting the ownership hierarchy. This calls into question Alchian and Demsetz's classic argument for making the monitor residual claimant (Section VIA).

³ See the literature on subjective evaluation, especially Prendergast (1999); Levin (2003); MacLeod (2003); and Fuchs (2007). In a principal-agent model where only the principal observes output, they make him indifferent over reports, so he tells the truth. But this contract breaks down if observing output is costly—no matter how small the cost. In statistics, this justifies proper scoring rules (Gneiting and Raftery 2007).

⁴ See Cremer and McLean (1985, 1988) and Mezzetti (2004, 2007) with exogenous types, as well as the mutual monitoring models of Ben-Porath and Kahneman (1996, 2003); Aoyagi (2005); and Miyagawa, Miyahara, and Sekiguchi (2008, 2009) in the context of a repeated game.

⁵ For related mechanisms, see footnote 4. In all these papers, however, agents' reports are cross-checked, whereas in my model the principal *tells* the worker his type. As a result, I can provide strict incentives, and thus avoid the critique of Bhaskar (2000), whereas these other papers cannot. Moreover, my results are robust to changes in information timing, and apply even if information is not nearly perfect. Section VI has details.

How these richer contracts enlarge the set of attainable social outcomes has important implications, both theoretically and for understanding real world institutions, as I argue next. Yet, with few exceptions, this question is not addressed in the literature (Section VI). The rest of the paper fills this gap. First I show that an outcome is enforceable; i.e., there is a single payment scheme that discourages every deviation, if and only if for every deviation there is a payment scheme that discourages it, where different schemes are allowed to discourage different deviations. In other words, discouraging deviations one by one is enough to discourage them simultaneously. Theorem 1 then argues that (statistically) detectable deviations don't matter, as they are easily discouraged, so an outcome is enforceable if and only if every undetectable deviation is unprofitable. Hence, if every deviation is detectable then the outcome is enforceable regardless of individual preferences (Theorem 2), as these determine the profitable deviations. This provides a robust benchmark for enforceability. I also show that with recommendation-contingent rewards different monitoring behavior may be used to deem different deviations detectable, whereas without such payments the same behavior must detect every deviation.⁶ Now more outcomes are enforceable because more deviations are detectable, which, as I suggested, don't determine enforceability.

To motivate the next results, recall that Bob's incentives break down if Ann never takes any money, and the less frequently she takes money, the larger must be Bob's contingent reward. This basic trade-off resonates with the model of crime and punishment of Becker (1968). Even if having neither crime nor enforcers is an impossible ideal, it may be approached with very little crime and very few enforcers (but large penalties to criminals) whenever crime is detectable, where fewer enforcers make detection less likely. Thus, although the ideal outcomes of neither crime nor enforcers and Ann never taking money are not enforceable, they are still *virtually enforceable*; i.e., there is an enforceable outcome arbitrarily close.

I derive two characterizations of virtual enforceability, in line with Theorems 1 and 2. The first completes my answer to the question of who will monitor the monitor. To gain intuition, suppose that providing incentives for a group of workers requires a monitor to detect their deviations. What about the monitor's deviations? Theorem 3 describes when the monitor's deviations can be discouraged robustly; i.e., regardless of preferences. By Theorem 1, discouraging deviations one by one is enough to discourage them simultaneously, and detectable deviations are easily discouraged, so the monitor's detectable deviations don't matter. Now only the monitor's undetectable deviations remain to be discouraged. By using recommendation-contingent rewards, a deviation is undetectable only if no matter what anybody does, it cannot be distinguished from honesty and obedience. Therefore, *it's still monitoring*; i.e., it still detects workers' deviations, so let the monitor play the deviation. Of course, this argument also applies to the monitor's deviations from this deviation, and so forth. I reconcile this potentially infinite regress⁷ (of the monitor playing a deviation of a deviation of ...) by showing that under reasonable assumptions (e.g., if agents

⁶This intuitive advantage distinguishes my results significantly from the literature (Section VI). Thus, my detectability requirement is much less demanding than the individual identifiability of Fudenberg, Levine, and Maskin (1994) as well as related generalizations. An exception is the independent work by Tomala (2009).

⁷A related regress on monitors for monitors is studied by Basu, Bhattacharya, and Mishra (1992).

have finitely many choices) not every behavior by the monitor can have a profitable, undetectable deviation. Thus, the principal monitors the monitor's detectable deviations (at the cost of occasional shirking by workers) and nobody needs to monitor the monitor's undetectable deviations. To motivate workers robustly, their deviations must be detectable with occasional monitoring, but deviations from monitoring itself need not be detectable.

For my last result, I fix both the information structure and individual preferences to show that an outcome is virtually enforceable if and only if profitable deviations are uniformly and credibly detectable (Theorem 4). I illustrate with examples that intuitively, this condition is roughly comparable to iterated elimination of weakly dominated undetectable strategies. This result has important consequences in the theory of repeated games, as I suggest towards the end of the paper. Specifically, it delineates the boundary of detectability for the Folk Theorem with mediated communication, as well as the fixed social cost of impatience, or as Radner, Myerson, and Maskin (1986) might put it, (lower hemi-) continuity of the equilibrium payoff correspondence of a repeated game with respect to the discount factor δ at $\delta = 1$.

Let me end this introduction with empirical motivation. First of all, Ann's problem is pervasive. For instance, consider airport security officials sitting behind an X-ray machine, watching suitcases pass them by. Their "output" is paying attention—only they can know if they are scrutinizing the baggage in front of them or just daydreaming. Of course, this problem is closely related to that of providing incentives for security guards and regulatory agencies, as well as maintaining police integrity. Without the right incentives, these agents might be tempted to shirk on their responsibilities or succumb to possibly unchecked corruption. Naturally, this problem appears in many other economic realms, such as management and supervision of workers, especially in service industries.

Moreover, contrived though it may seem, Bob's contract is ubiquitous. The Transportation Security Administration uses "covert testing" to evaluate airport inspectors (TSA 2004). Such testing ranges from superimposing images of bombs on computer screens to smuggling weapons. In 2005, the Government Accountability Office created a "Forensic Audits and Special Investigations Unit" (Walker 2007). This unit has undertaken several "red team" operations to expose vulnerabilities in government agencies, including the Federal Emergency Management Agency, the Nuclear Regulatory Commission, the Department of Defense, the Department of Transportation, Medicare, and the US Department of Labor, to name just a few. These operations have ranged from smuggling nuclear materials through the US border to test the effectiveness of border patrols, to making unfair labor practice claims to test the Wage and Hour Division's reporting standards, and applying for Medicare billing numbers without proper documentation to test due diligence in Medicare's protocol for granting billing rights (Kutz 2008a, b, c, d, 2009a, b).

Similar arrangements for "policing the police" are also well documented. Internal Affairs Departments regularly use "integrity tests" to discourage police corruption.⁸

⁸Sherman (1978, p. 163–64) on integrity tests by police departments: "Both Oakland and New York constructed artificial situations giving police officers the opportunity to commit corrupt acts. The tests were designed to yield the evidence needed to arrest and convict an officer who failed the 'test.' [...] Some were random tests of conformity of procedures, such as the infamous 'wallet drop': wallets containing marked money were dropped by internal policing officers near randomly selected patrol officers to see if they would turn the wallet in to the police property clerk

Officers are infiltrated into police corruption rings to act as informants. In both cases, the mere possibility of monitoring can deter corruption.⁹ Even corrupt officers have been used to “test” and report on other officers, often relying on leniency to provide incentives.^{10, 11} Two important examples of the investigation and disciplining methods above according to the criminology literature are the Knapp Commission (Knapp 1972) and the Mollen Commission (Mollen 1994) investigations. See Marx (1992) for more interesting examples.

Similar contracts have also been used by managers. For instance, retailers routinely hire “mystery shoppers” (Ewoldt 2004) to secretly evaluate employees and provide feedback to managers. (Airlines call them “ghost riders”; see Dallos 1987.) The consulting branch of IBM offers “ethical hacking” services by “tiger teams” (Palmer 2001) that try to hack into clients’ IT network, to expose vulnerabilities. Amazon’s Mechanical Turk (Pontin 2007) decentralizes a wide variety of tasks to humans, such as verifying image quality. To provide workers with incentives, images whose quality is already known are occasionally included.

I. Robinson and Friday

Example 1: Consider a principal (Ann?) and two risk-neutral agents, Robinson (Bob?), the row player, and Friday, the column player, who interact with payoffs in the left bimatrix below. Intuitively, Friday is a worker and Robinson is a monitor. Each agent’s effort is costly—with cost normalized to unity—but unobservable.

	Work	Shirk		Work	Shirk
Monitor	0, 0	0, 1	Monitor	1, 0	0, 1
Rest	1, 0	1, 1	Rest	1/2, 1/2	1/2, 1/2
	Utility payoffs			Signal probabilities	

After actions have been taken, Robinson privately observes one of two possible signals, g and b . Their conditional probability (or the *monitoring technology*) appears in the right bimatrix above. In words, if Robinson monitors he observes Friday’s effort, but if he rests then his observation is completely uninformative.¹²

with the full amount of money. Other integrity tests had more specific targets. Money left in an illegally parked car was often used to test the integrity of certain police tow-truck drivers against whom allegations of stealing from towed cars had been made. Fake gambling operations were created to see if police officers tried to establish paid protection arrangements.”

⁹Sherman (1978, p. 156–57) on informants: “Under careful monitoring, honest police officers in New York were even assigned by internal investigators to join corruption conspiracies [...]. Quite apart from the value of the information these regular informants provided, the very fact that their existence was known to other police officers may have yielded a deterrent effect. Though the informants were few in number, no one was quite certain who could be trusted to keep silence.” See also Prenzler (2009) for several examples.

¹⁰Sherman (1978, p. 162): on incentives for informants: “[...] rewards were used to encourage informants to inform. These ranged from immunity from arrest for the informants’ own offenses to simple obligation for future considerations.” See also Skolnick (1966).

¹¹This is close to the Department of Justice’s Corporate Leniency Program, which, to discourage collusion, awards immunity to the first firm in a cartel to come forward with evidence of illegal activity (Harrington 2008; Miller 2009).

¹²Alternatively, we could assume that if Robinson rests he observes “no news.” The current assumption helps to compare with the literature that relies on publicly verifiable monitoring, such as Holmström (1982).

Finally, after Robinson observes the realized signal, he makes a verifiable report to the principal.

If monitoring were costless—following the subjective evaluation literature (footnote 3)—the principal could enforce the action profile (monitor, work) by paying Robinson a wage independent of his report. Robinson would be willing to monitor and report truthfully, and Friday could therefore be rewarded contingent on his effort via Robinson's report.

With costly monitoring, Robinson's effort becomes an issue. Suppose that the principal wants to enforce (rest, work) on the grounds that monitoring is unproductive. Unfortunately, this is impossible, since if Robinson rests then Friday's expected payment cannot depend on his own effort, so he will shirk. On the other hand, if Robinson's observations are publicly verifiable then not only can the principal enforce (monitor, work), but also *virtually enforce* (rest, work)—i.e., enforce an outcome arbitrarily close—using Holmström's group penalties: if news is good everyone gets paid and if news is bad nobody gets paid. Thus, the principal can induce Friday to always work and Robinson to secretly monitor with small but positive probability σ by paying Robinson \$2 and Friday $1/\sigma$ if g and both agents zero if b .

If Robinson's costly observations are unverifiable, Holmström's contracts break down, since Robinson will then just report g and rest, so Friday will shirk. Furthermore, although Robinson would happily tell the truth with a wage independent of his report, he would never monitor, so again Friday would shirk. This raises the question: how can we motivate Friday to work when Robinson's signal is both costly and private?

Having Friday always work is impossible, since then Robinson will never monitor, so Friday will shirk. The principal, however, can virtually enforce (rest, work) by asking Friday to shirk occasionally and correlating Robinson's payment with Friday's secret recommendation, thereby "monitoring the monitor." Indeed, the following contract is incentive compatible given $\mu \in (0, 1)$ and $\sigma \in (0, 1]$: (i) Robinson is asked to monitor with probability σ ; (ii) Friday is independently asked to work with probability μ ; and (iii) the principal pays Robinson and Friday, respectively, contingent on his recommendations and Robinson's report as follows.

	(monitor, work)	(monitor, shirk)	(rest, work)	(rest, shirk)
g	$1/\mu, 1/\sigma$	0, 0	0, 0	0, 0
b	0, 0	$1/(1 - \mu), 0$	0, 0	0, 0

Robinson and Friday's Recommendation- and Report-Contingent Payments

Friday is paid with Holmström's contract, whereas Robinson is paid $1/\mu$ if he reports g when (monitor, work) was recommended and $1/(1 - \mu)$ if he reports b when (monitor, shirk) was recommended. Robinson is not told Friday's recommendation—this he must discover by monitoring. Clearly, Friday is willing to obey the principal's recommendations if Robinson is honest and obedient. To see that Robinson will abide by the principal's requests, suppose that he was asked to monitor.

If he monitors, clearly it is optimal for him to also be honest, with expected payoff $\mu(1/\mu) + (1 - \mu)[1/(1 - \mu)] = 2$. If instead he rests, his expected payoff equals $1 + \mu(1/\mu) = 2$ if he reports g , and $1 + (1 - \mu)[1/(1 - \mu)] = 2$ if he reports b .

As $\sigma \rightarrow 0$ and $\mu \rightarrow 1$, Robinson and Friday's behavior tends to the profile (rest, work) with a contract that is incentive compatible along the way; i.e., (rest, work) is *virtually enforceable*. This requires arbitrarily large payments, yet in reality feasible payments may be bounded. (Section VB has more on this.) Nevertheless, virtual enforcement is a useful benchmark for otherwise approachable behavior. Interpreting payments as continuation values in a repeated game, virtual enforcement also describes asymptotic behavior as players become patient.

With verifiable monitoring, (rest, work) was virtually enforced by incurring the cost of monitoring Friday (Robinson's effort) with small probability. With private monitoring, an additional cost is incurred, also with small probability: the cost of monitoring Robinson. This cost is precisely the forgone productivity from Friday shirking. Such a loss may be avoided by asking Friday to take a costless action, like changing the color of his socks.

Robinson's contract pays him for matching his report to Friday's recommendation—he faces a “trick question” whose answer the principal already knows, just like Ann and Bob. Robinson is rewarded for reporting accuracy: he is responsible for monitoring through his ability to reproduce Friday's recommendation. As such, Robinson must not observe Friday's recommendation. Hence, the contract is not robust to “collusion”: both agents could avoid effort if Friday simply told Robinson his recommendation. This is cheap talk, however—not sharing the information is still an equilibrium. (Section VC has more on this.)

This example shows that a monitor's incentives can be aligned without having to become residual claimant, contrary to Alchian and Demsetz's claim. (Section VIA has more on this.) Moreover, if Robinson was residual claimant then he would never verify Friday's effort, as Friday's payment would come from his own pocket *after* effort had been exerted. Of course, this argument relies on Robinson and Friday not meeting in the future, so that Friday cannot threaten to quit if Robinson “cheats.”¹³ Notice that Robinson must not be telling people what to do (giving effort recommendations), since otherwise the above contracts would break down. Making Friday residual claimant would also create perverse incentives for two reasons. Firstly, if Robinson was the only one who could verify Friday's output at a cost then Friday would have to ask the trick questions to Robinson himself. In this case, it would be optimal for him to disobey his own recommendation to himself in order to save paying Robinson his wage. Secondly, it would be impossible to save on the costs of monitoring Friday by having Robinson monitor randomly if Friday was the one telling Robinson when to monitor.

A final comment: if recommendations are not verifiable, (rest, work) is still virtually enforceable without a third party by asking Friday if he worked. Section VIB has details.

¹³ See, e.g., Levin (2003) and Fuchs (2007). For any discount factor less than one, however, there is always some incentive for the principal to underreport effort, so some inefficiency remains.

II. Model

The previous example suggests that recommendation-contingent rewards open up a great deal of possibilities in terms of attainable social behavior. I now develop a general model to formalize this intuition by characterizing both enforceability and virtual enforceability.

Let $I = \{1, \dots, n\}$ be a finite set of risk-neutral agents, A_i a finite set of actions available to any agent $i \in I$, and $A = \prod_i A_i$ the space of action profiles. Let $v_i(a)$ denote the utility to agent i from action profile $a \in A$. A *correlated strategy* is a probability measure $\mu \in \Delta(A)$.¹⁴ Let S_i be a finite set of signals observable only by agent $i \in I$, S_0 a finite set of publicly verifiable signals and $S = \prod_{j=0}^n S_j$ be the space of signal profiles. A *monitoring technology* is a map $\Pr : A \rightarrow \Delta(S)$, where $\Pr(s|a)$ is the probability that $s \in S$ was observed when a was played. A *payment scheme* is a map $\zeta : I \times A \times S \rightarrow \mathbb{R}$ that assigns individual payments contingent on recommended actions and reported signals, each of which is assumed verifiable.

Time elapses as follows. First, I assume that the principal can and does commit to a *contract* (μ, ζ) , draws a profile of recommendations according to μ , and delivers them to the agents confidentially and verifiably.¹⁵ Agents then simultaneously take unobservable actions. Next, agents observe their private, unverifiable signals and submit a report to the principal before a public signal realizes (the order of signals is not essential, just simplifying). Finally, the agents pay the principal according to ζ contingent on recommendations and reports.

If all agents obey their recommendations and report truthfully, i 's expected utility equals

$$\sum_{a \in A} \mu(a) v_i(a) - \sum_{(a,s)} \mu(a) \zeta_i(a,s) \Pr(s|a).$$

Of course, i may disobey his recommendation and lie about his private signal. A *reporting strategy* is a map $\rho_i : S_i \rightarrow S_i$, where $\rho_i(s_i)$ is the reported signal when agent i observes s_i . Let R_i be the set of i 's reporting strategies. The *truthful* reporting strategy is the identity map $\tau_i : S_i \rightarrow S_i$ with $\tau_i(s_i) = s_i$. For every agent i and pair $(b_i, \rho_i) \in A_i \times R_i$, the probability that s is reported if everyone else is honest and plays a_{-i} equals

$$\Pr(s|a_{-i}, b_i, \rho_i) = \sum_{t_i \in \rho_i^{-1}(s_i)} \Pr(s_{-i}, t_i|a_{-i}, b_i).$$

A contract (μ, ζ) is *incentive compatible* if honesty and obedience is optimal, i.e., μ is a communication equilibrium (Myerson 1986; Forges 1986) of the game induced by ζ :

$$\begin{aligned} (1) \quad & \sum_{a_{-i}} \mu(a) [v_i(a_{-i}, b_i) - v_i(a)] \\ & \leq \sum_{(a_{-i}, s)} \mu(a) \zeta_i(a, s) [\Pr(s|a_{-i}, b_i, \rho_i) - \Pr(s|a)] \quad \forall (i, a_i, b_i, \rho_i). \end{aligned}$$

¹⁴ If X is a finite set, $\Delta(X) = \{\mu \in \mathbb{R}_+^X : \sum_x \mu(x) = 1\}$ is the set of probability vectors on X .

¹⁵ In Section VIB, I discuss relaxing the principal's full commitment power, as well as verifiability.

In words, the left-hand side reflects agent i 's utility gain from playing b_i even though a_i was recommended. (Since misreporting is costless, ρ_i is irrelevant.) The right-hand side reflects his monetary loss from deviating to (b_i, ρ_i) relative to playing a_i and reporting truthfully.

DEFINITION 1: A correlated strategy μ is called enforceable if there is a payment scheme ζ such that (μ, ζ) is incentive compatible, and virtually enforceable if there is a sequence $\{\mu^m\}$ of enforceable correlated strategies such that $\mu^m \rightarrow \mu$.

Thus, in Example 1, the profile (rest, work) is virtually enforceable but not enforceable. My goal is to understand enforceable and virtually enforceable outcomes, as well as the role played by recommendation-contingent rewards. To this end, I now introduce the notion of a detectable strategy, which will play a crucial role in all the results below.

A strategy for agent i is a map $\sigma_i: A_i \rightarrow \Delta(A_i \times R_i)$, where $\sigma_i(b_i, \rho_i | a_i)$ is the probability that i plays (b_i, ρ_i) after the recommendation a_i . Call σ_i a *deviation* if it ever differs from honesty and obedience; i.e., if $\sigma_i(b_i, \rho_i | a_i) > 0$ for some $(b_i, \rho_i) \neq (a_i, \tau_i)$. Thus, Robinson always shirking and reporting g is a deviation. Let $\Pr(\mu)$ be the vector of report probabilities if everyone is honest and obedient, defined pointwise by $\Pr(s | \mu) = \sum_a \mu(a) \Pr(s | a)$, and $\Pr(\mu, \sigma_i)$ the corresponding vector if i plays σ_i instead, defined similarly by

$$\Pr(s | \mu, \sigma_i) = \sum_{(a, b_i, \rho_i)} \mu(a) \Pr(s | a_{-i}, b_i, \rho_i) \sigma_i(b_i, \rho_i | a_i).$$

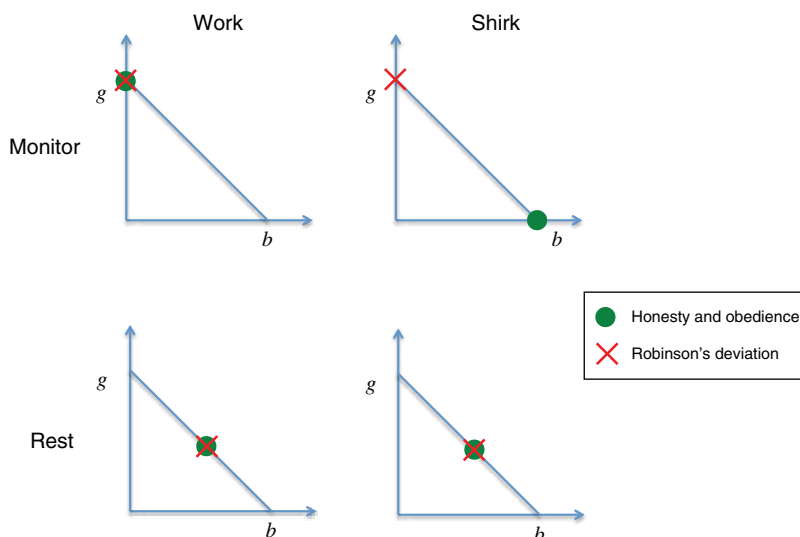
DEFINITION 2: Given any subset of action profiles $B \subset A$, a strategy σ_i is called B -detectable if $\Pr(s | a) \neq \Pr(s | a, \sigma_i)$ for some $a \in B$ and $s \in S$.¹⁶ Otherwise, σ_i is called B -undetectable. An A -detectable strategy is simply called detectable, etc.

A deviation is B -detectable if there is a recommendation profile in B such that the report probabilities induced by it differ from those due to honesty and obedience. In this weak sense, B -detectability means that the deviation can be statistically identified.

To illustrate, Figure 1 shows that, in Example 1, Robinson's strategy of always resting is $\{(\text{monitor}, \text{work})\}$ -undetectable as long as he reports g when asked to monitor. Similarly, always resting is $\{(\text{monitor}, \text{shirk})\}$ -undetectable if he reports b when asked to monitor. On the other hand, it is easy to see that always resting is still detectable regardless of Robinson's reporting strategy, yet always monitoring and mixing his reports evenly is undetectable.

I now present my four main results. I characterize enforceability and virtual enforceability in terms of the model's primitives, i.e., the monitoring technology $\Pr$ and the profile of utility functions v , as well as in terms of just the monitoring technology. This provides useful robust conditions for outcomes to be attainable regardless of preferences, as I argue later.

¹⁶I abuse notation slightly by identifying Dirac measure $[a] \in \Delta(A)$ with the action profile $a \in A$.

FIGURE 1. ROBINSON'S DEVIATION IS DETECTABLE BUT NOT $\{(MONITOR, WORK)\}$ -DETECTABLE

III. Enforceability

I begin with an intuitive characterization of enforceability. For any correlated strategy μ , consider the following zero-sum two-person game between the principal and a “surrogate” for the agents. The principal chooses a payment scheme ζ and the surrogate chooses a strategy σ_i for some agent i . (Each strategy set is clearly convex.) The principal pays the surrogate the expected deviation gains from i playing σ_i instead of honest and obediently,

$$\sum_{(a, b_i, \rho_i)} \mu(a) \sigma_i(b_i, \rho_i | a_i) [v_i(a_{-i}, b_i) - v_i(a)] \\ - \sum_{s \in S} \zeta_i(a, s) (\Pr(s | a_{-i}, b_i, \rho_i) - \Pr(s | a)).$$

The value of this game is at least zero for the surrogate, since he could always have his agents play honest and obediently. In fact, by construction, μ is enforceable if and only if this value equals zero, since then there is a payment scheme that discourages every deviation. As Figure 2 suggests, by the Minimax Theorem *it doesn't matter who moves first*. Letting the principal move second, μ is enforceable if and only if for every deviation there is a payment scheme that discourages it, where different schemes may be used to discourage different deviations. Intuitively, for enforceability it suffices to discourage deviations one by one.¹⁷

Now let us follow the logic of having the principal move second, so pick a strategy σ_i . If it is $\text{supp } \mu$ -detectable¹⁸ then there is an action profile $a \in \text{supp } \mu$ that

¹⁷This argument resonates with Hart and Schmeidler (1989). They fixed utilities and varied the correlated strategy, however, whereas I fix the correlated strategy and vary utilities via the payment scheme.

¹⁸Let $\text{supp } \mu = \{a \in A : \mu(a) > 0\}$ be the set of action profiles with positive probability under μ .

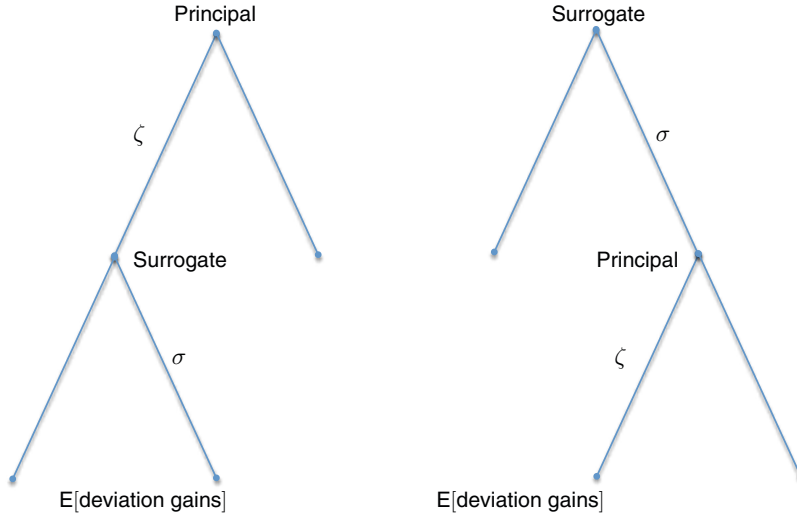


FIGURE 2. A ZERO-SUM GAME BETWEEN THE PRINCIPAL AND A SURROGATE FOR THE AGENTS

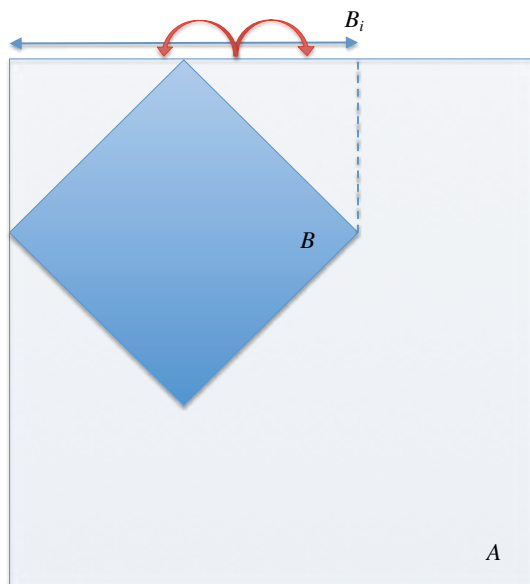
detects σ_i ; i.e., such that $\Pr(a) \neq \Pr(a, \sigma_i)$. Hence, there are some signals whose probability increases with σ_i (“bad” news) and others whose probability decreases (“good” news). The following payment scheme discourages σ_i : choose a sufficiently large wedge between good and bad news after a is recommended that the monetary loss outweighs the utility gain from playing σ_i instead of following the principal’s recommendations. This way, $\text{supp } \mu$ -detectable deviations are easily discouraged. On the other hand, if σ_i is $\text{supp } \mu$ -undetectable then the surrogate receives the same payoff regardless of the principal’s choice of payment scheme. If σ_i gives positive utility gain then there is nothing the principal can do to discourage it, so for μ to be enforceable, it better have been the case that σ_i was not more desirable than honesty and obedience. The next result formalizes this intuition. For all proofs, see Appendix A.

THEOREM 1 (Minimax Lemma): *A given correlated strategy μ is enforceable if and only if every $\text{supp } \mu$ -undetectable deviation σ_i is μ -unprofitable; i.e.,*

$$\Delta v_i(\mu, \sigma_i) = \sum_{(a, b_i, \rho_i)} \mu(a) \sigma_i(b_i, \rho_i | a_i) [v_i(a_{-i}, b_i) - v_i(a)] \leq 0.$$

Thus, in Example 1, the profile (rest, work) is not enforceable because Friday shirking is $\{(\text{rest}, \text{work})\}$ -undetectable and (rest, work)-profitable, yet any completely mixed correlated strategy is enforceable. To verify whether an outcome μ is enforceable we must in principle find a payment scheme that simultaneously discourages every deviation. By the Minimax Lemma, we may ignore payments and just verify that every $\text{supp } \mu$ -undetectable deviation is μ -unprofitable.¹⁹ If every

¹⁹The set of $\text{supp } \mu$ -undetectable strategies is easily seen to be a convex polyhedron. By linearity of $\Delta v_i(\mu, \sigma_i)$ with respect to σ_i , it is enough to check the polyhedron’s finitely many extreme points.

FIGURE 3. ILLUSTRATION OF A B -DISOBEDIENCE IN DEFINITION 3

relevant deviation is $\text{supp } \mu$ -detectable, then the consequent of the Minimax Lemma holds vacuously, therefore μ is enforceable regardless of the utility profile v . What makes a deviation relevant? Reports carry no utility cost, so only choosing an action that disobeys a recommendation in $\text{supp } \mu$ can be μ -profitable. This leads to my next result. First, I need a simple but important definition that will be used repeatedly.

DEFINITION 3: Given $B \subset A$, a strategy σ_i is called a B -disobedience if $\sigma_i(b_i, \rho_i | a_i) > 0$ for some $a_i \in B_i$ and $b_i \neq a_i$, where $B_i = \{b_i \in A_i : \exists b_{-i} \in A_{-i} \text{ s.t. } b \in B\}$ is the projection of B on A_i . An A -disobedience is called simply a disobedience. (See Figure 3 for intuition.)

THEOREM 2: A given correlated strategy μ is enforceable for each prole of utility functions if and only if every $\text{supp } \mu$ -disobedience is $\text{supp } \mu$ -detectable.

Theorem 2 intuitively characterizes robust enforceability²⁰ of an outcome μ in terms of the monitoring technology: every disobedience must be detectable with behavior in the support of μ . Crucially, *different action profiles may be used to detect different disobediences*. Of course, this feature also applies to Theorem 1. Let me explain using Robinson and Friday. Suppose that Robinson deviates to rest instead of monitoring. If he plans to always report g then Friday can just shirk and render Robinson's deviation detectable, and if he always reports b then Friday can work instead. (If Robinson mixes his reports, Friday still can shirk.) On the other hand, it is easy to see that no single (pure or mixed) action for Friday can detect

²⁰By robust I mean "for all utility functions." For more on what this robustness means, see Section VC.

all of Robinson's deviations from monitoring. Thus, if Friday works then Robinson can pass undetected by reporting g , whereas if Friday shirks then Robinson can just report b instead. (If Friday mixes then Robinson can just mix his reports accordingly.)

This key feature renders the condition of Theorem 2 much weaker than those in the literature, such as individual full rank (IFR) by Fudenberg, Levine, and Maskin (1994).²¹ As I discuss in Section VA, this feature captures the enlargement of contractual opportunities from allowing for payment contingent on recommendations. To illustrate, consider an example where every disobedience is detectable, yet IFR fails.²²

Example 2: Two publicly verifiable signals and two agents, Ann and Bob. Ann has two choices, $\{U, D\}$, and Bob has three, $\{L, M, R\}$. The monitoring technology is given below.

	L	M	R
U	1, 0	0, 1	$1/2, 1/2$
D	1, 0	0, 1	$1/3, 2/3$

If Ann plays U then Bob playing R is indistinguishable from $\frac{1}{2}[L] + \frac{1}{2}[M]$. Similarly, if Ann plays D then Bob can deviate from R to $\frac{1}{3}[L] + \frac{2}{3}[M]$ without changing signal probabilities. Hence, IFR fails at any outcome that gives R positive probability. In Theorem 5, I show that it is therefore impossible to convince Bob to ever play R with a payment scheme that only depends on signals if Bob strictly prefers playing L and M . Every disobedience is detectable, however: for any deviation by Bob there is an action by Ann that detects it. By correlating Bob's payment with Ann's (recommended) action, the principal can keep Bob from knowing how he ought to mix between L and M for his payment to equal what he would obtain by playing R . This renders R enforceable by Theorem 2, although, just as with Robinson and Friday, recommendation-contingent rewards are required.

I end this section with a couple of quick corollaries. For the first one, notice that the detectability condition in Theorem 2 for enforceability of μ only depends on its support.

COROLLARY 1: *Every correlated strategy with support equal to B is enforceable for each profile of utility functions if and only if every B -disobedience is B -detectable.*

By Corollary 1, every completely mixed correlated strategy is enforceable for each profile of utility functions if and only if every disobedience is detectable. Approaching an arbitrary correlated strategy with completely mixed ones, it becomes virtually enforceable. Conversely, for every undetectable disobedience there is a utility profile that renders it profitable.

²¹ The spirit of IFR at some correlated strategy μ is that the same μ detects every deviation σ_i , i.e., $\Pr(\mu) \neq \Pr(\mu, \sigma_i)$. See footnote 25 for a definition and Section VA for further discussion.

²² It even fails *local* IFR of d'Aspremont and Gérard-Varet (1998), which requires that one correlated strategy per agent—possibly different for different agents—detect all of its agent's deviations.

COROLLARY 2: *Every correlated strategy is virtually enforceable for each profile of utility functions if and only if every disobedience is detectable.*

IV. Virtual Enforceability

Corollary 2 gives minimal conditions on the monitoring technology for every outcome to be virtually enforceable. Yet, it is easy to find examples where this condition fails. Perhaps most prominently, it fails for Robinson and Friday, as not every disobedience is detectable, e.g., Robinson monitoring and mixing his reports evenly when asked to rest. If this deviation were profitable, there would be no way to discourage it. This motivates seeking weaker results. First, I characterize virtual enforcement of a given correlated strategy μ , rather than every one, for each profile of utility functions. As I will argue soon, this completes my answer to the question of who will monitor the monitor. Afterwards, I'll drop the quantifier on utilities.

To introduce the first result, consider the following sufficient but unnecessary condition. It is easy to see that μ is virtually enforceable for each profile of utility functions if every C -disobedience is C -detectable for some $C \supset \text{supp } \mu$. Indeed, since C contains $\text{supp } \mu$, μ is approachable with correlated strategies whose support equals C (by a simple application of Corollary 2), which are enforceable by Corollary 1. However, μ can be virtually enforceable for each profile of utility functions even if this condition fails, as the next example shows.

Example 3: Ann and Bob, two public signals, and the following monitoring technology:

	L	M	R
U	1, 0	1, 0	1, 0
D	1, 0	0, 1	0, 1

Clearly, (U, L) is not enforceable for each profile of utility functions because Ann playing D if asked to play U is a $\{(U, L)\}$ -undetectable $\{(U, L)\}$ -disobedience. It is also easy to see that there is a C -undetectable C -disobedience for every $C \supset \{(U, L)\}$. Even though both M and R can detect Ann's deviation above, no single action can be used for each utility profile, since for some utility profiles M strictly dominates R , whereas for others it's the other way around, and M is completely indistinguishable from R . Nevertheless, (U, L) is still virtually enforceable by letting Bob choose between M or R to detect Ann's deviation. Notice that every $\{(U, L)\}$ -disobedience is detectable.

THEOREM 3: *A given correlated strategy μ is virtually enforceable for each profile of utility functions if and only if every $\text{supp } \mu$ -disobedience is detectable.*

Theorem 3 is one of the main results of the paper. It shows that μ is virtually enforceable for every utility profile as long as every disobedience from μ is detectable, perhaps with some occasional "monitoring" behavior. Crucially, there is no requirement on disobediences to behavior outside of μ ; i.e., deviations from monitoring need not be detectable.

To make intuitive sense of this result, let $B \subset A$ be the support of μ . Recall that by the Minimax Lemma, we may discourage disobediences one by one. Suppose that, to detect a disobedience $\sigma_i(a_i)$ away from $a_i \in B_i$, some $a_j \notin B_j$ must be played infrequently by $j \neq i$. Call this “monitoring.” What if a_j itself has a profitable deviation $\sigma_j(a_j)$? After all, the condition of Theorem 3 purposely says nothing about detection outside B .

If such $\sigma_j(a_j)$ is detectable then it is easily discouraged, as usual. If, on the other hand, $\sigma_j(a_j)$ is undetectable then playing $\sigma_j(a_j)$ instead of a_j still detects deviations from a_i by virtue of being undetectable, since this means that no matter what anybody else does, the deviation does not change any report probabilities. In other words, *it’s still monitoring*. Similarly, undetectable deviations from $\sigma_j(a_j)$ detect deviations from a_i , and so on. Now, since the game is finite there must exist a maximal deviation amongst the undetectable ones, which—by construction—has no profitable, undetectable deviation.

This argument completes my answer to the question of who will monitor the monitor. The principal monitors the monitor’s detectable deviations by occasionally asking his workers to secretly shirk, and nobody needs to monitor the monitor’s undetectable deviations.²³ How to monitor the monitor? By making the monitor responsible for monitoring with trick questions that follow Robinson’s contract in. When is it possible to monitor the monitor? Theorem 3 gives a minimal requirement in terms of just the monitoring technology.

I end the section by characterizing virtual enforceability of a given correlated strategy μ for a fixed profile of utility functions. To motivate the problem, notice that the condition in Theorem 3 that every supp μ -disobedience be detectable does not hold for (rest, work) in Example 1. This is because Robinson monitoring and mixing his report is not detectable. On the other hand, this deviation by Robinson is clearly strictly unprofitable.

Although enforceability has a simple characterization (Theorem 1), virtual enforceability does not. To see why, for some μ to be virtually enforceable not every supp μ -disobedience needs to be detectable: strictly μ -unprofitable supp μ -disobediences may be undetectable, for instance, without jeopardizing virtual enforceability. On the other hand, it is not enough that every profitable supp μ -disobedience be detectable, as the next example shows.

Example 4: Consider the following variation on Robinson and Friday (Example 1).

	Work	Shirk	Solitaire		Work	Shirk	Solitaire
Monitor	0, 0	0, 1	0, 1	Monitor	1, 0	0, 1	1, 0
Rest	0, 0	0, 1	0, 0	Rest	1/2, 1/2	1/2, 1/2	1/2, 1/2
Utility payoffs				Signal probabilities			

Assume that signals are *publicly verifiable* and Robinson’s utility is constant. Clearly, the profile (rest, work) is not enforceable, since Friday shirking is

²³ This argument relies crucially on recommendation-contingent rewards. Section VA has the details.

(rest, work)-profitable and $\{(\text{rest, work})\}$ -undetectable. Moreover, (rest, work) is *not virtually enforceable* either. Indeed, for Friday to ever work Robinson must monitor with positive probability. But then nothing can discourage Friday from playing solitaire when asked to work, since it is undetectable and weakly dominant. On the other hand, every (rest, work)-profitable disobedience is detectable: (rest, work)-profitability requires shirking with positive probability, which can be detected.

Detecting (rest, work)-profitable deviations is not enough here because solitaire weakly dominates work and is indistinguishable from it. Indeed, if solitaire strictly dominated work then there would exist a (rest, work)-profitable undetectable strategy that rendered (rest, work) virtually unenforceable. On the other hand, if Friday's payoff from (rest, solitaire) was negative instead of zero, so solitaire no longer weakly dominated working, then (rest, work) would be virtually enforceable, because playing solitaire when asked to work would be strictly unprofitable if Robinson monitored with sufficiently low probability.

So what *is* required beyond detecting profitable deviations? Below, I will argue that profitable deviations must be *uniformly and credibly* detectable. To illustrate, note that if solitaire is removed from Example 4 then (rest, work) is virtually enforceable, not just because every (rest, work)-profitable deviation is detectable (this is true with or without solitaire), but also because the utility gains from every (rest, work)-profitable deviation can be uniformly outweighed by monetary losses. To describe this "uniform detection" formally, let us introduce some notation. For any strategy σ_i and any correlated strategy μ , write

$$\|\Delta\text{Pr}(\mu, \sigma_i)\| = \sum_{s \in S} \left| \sum_{(a, b_i, \rho_i)} \mu(a) [\sigma_i(b_i, \rho_i | a_i) \text{Pr}(s | a_{-i}, b_i, \rho_i) - \text{Pr}(s | a)] \right|.$$

Intuitively, this norm describes the statistical difference between abiding by μ and deviating to σ_i . Thus, σ_i is *supp μ -undetectable* if and only if $\|\Delta\text{Pr}(\mu, \sigma_i)\| = 0$.

Say that *every μ -profitable deviation is uniformly detectable* if $z \geq 0$ exists such that for every μ -profitable deviation σ_i there is a correlated strategy η (possibly different for different σ_i) with σ_i being *supp η -detectable* and $\Delta v_i(\eta, \sigma_i) \leq z \sum_a \eta(a) \|\Delta\text{Pr}(a, \sigma_i)\|$. In other words, a bound $z \geq 0$ exists such that for every μ -profitable strategy σ_i there is a payment scheme ζ satisfying $-z \leq \zeta_i(a, s) \leq z$ that strictly discourages σ_i .²⁴ Intuitively, every μ -profitable deviation can be strictly discouraged with a uniformly bounded payment scheme. To see how uniform detectability fails in Example 4 but holds without solitaire, let α and β , respectively, be the probabilities that Friday shirks and plays solitaire after being asked to work. Clearly, (rest, work)-profitability requires $\alpha > 0$. To obtain uniform detectability we need z such that given $\alpha > 0$ and $\beta \geq 0$, a correlated strategy η exists with $(\alpha + \beta)\eta(\text{monitor, work}) + \alpha\eta(\text{rest, work}) \leq 2z\alpha\eta(\text{monitor, work})$. Therefore, $(\alpha + \beta)/\alpha \leq 2z$ is necessary for uniform detectability. No such z satisfies this for all relevant (α, β) , however. Removing solitaire restores uniform detectability: now $\beta = 0$, so any $z \geq 1/2$ works.

²⁴ To find this payment scheme, invoking the Bang-Bang Principle, let $\zeta_i(a, s) = \pm z$ depending on the sign of the statistical change created by σ_i , namely $\sum_{(b_i, \rho_i)} \sigma_i(b_i, \rho_i | a_i) \text{Pr}(s | a_{-i}, b_i, \rho_i) - \text{Pr}(s | a)$.

Uniform detectability is still not enough for virtual enforcement, as Example 5 shows.

Example 5: Add a row to the table in Example 4; i.e., another action for Robinson, with utility payoff

-1,0	-1,1	-1,0
------	------	------

and signal probabilities

1,0	0,1	1,0
-----	-----	-----

In Example 5 every (rest, work)-profitable deviation is uniformly detectable when Robinson plays his new action, but this action is not credible because it is strictly dominated by and indistinguishable from monitoring. Hence, (rest, work) is not virtually enforceable.

DEFINITION 4: Say that every μ -profitable deviation is uniformly and credibly detectable if there exists $z \geq 0$ such that for every μ -profitable deviation σ_i , there exists a correlated strategy η satisfying (i) σ_i is $\text{supp } \eta$ -detectable, (ii) $\Delta v_i(\eta, \sigma_i) \leq z \sum_a \eta(a) \|\Delta \text{Pr}(a, \sigma_i)\|$, and (iii) $\Delta v_j(\eta, \sigma_j) \leq z \sum_a \eta(a) \|\Delta \text{Pr}(a, \sigma_j)\|$ for all other (j, σ_j) .

Intuitively, this means that we may still use different η to uniformly detect different σ_i , but these η must be credible in that incentives can be provided for η to be played.

THEOREM 4: A correlated strategy μ is virtually enforceable if and only if every μ -profitable deviation is uniformly and credibly detectable.

V. Discussion

In this section I begin by describing the value of recommendation-contingent payments and noting how Theorem 3 relies deeply on them. Next, in Section VB, I look at how adding realistic contractual restrictions might affect results. In Section VC, I discuss robustness to fundamentals; i.e., preferences and beliefs. I end the section with comments on collusion, multiplicity, and renegotiation.

A. The Margin of Recommendation-Contingent Rewards

All of this paper's results rely on the principal's ability to make payments that may vary with agents' recommendations. Examples 1 and 2 show that such schemes can yield a strict improvement for the principal relative to ones that just depend on reported signals. I now characterize this improvement in terms of detectability requirements for enforcement. Intuitively, I argue that with recommendation-contingent schemes, deviations may be detected effectively "after the fact" in the sense that different actions may be used to detect different deviations. Contrariwise,

the same behavior must detect every profitable deviation without these rewards. I then show how this relates to in important ways.

Given μ , call σ_i *detectable at μ* if $\Pr(s|\mu) \neq \Pr(s|\mu, \sigma_i)$ for some $s \in S$.²⁵

THEOREM 5: *A given correlated strategy μ is enforceable with payments that do not depend on recommendations if and only if every μ -profitable deviation is detectable at μ .*

Theorems 1 and 5 capture the margin of recommendation-contingent rewards. This is precisely the difference between *supp μ -detectability* and *detectability at μ* . The former allows for different actions to detect different deviations, whereas the latter does not.

To illustrate, recall Example 1. Suppose that Robinson is asked to monitor by a principal without access to recommendation-contingent rewards. By Theorem 5, the principal must ask Friday to work with some probability μ such that every deviation by Robinson is detectable at Friday's mixed strategy. Unfortunately, no such behavior exists. To see this, notice simply that if Friday works then Robinson resting and reporting g is undetectable at work, and if Friday shirks then Robinson resting and reporting b is undetectable at shirk. Finally, if Friday mixes, then, given Friday's behavior, all Robinson needs to do is rest and report g with probability μ , which produces exactly the same probability distribution over reported signals as if he had monitored and reported honestly. Therefore, Robinson's profitable deviation is undetectable at Friday's mixed strategy. By Theorem 5, it follows that Robinson monitoring is not enforceable without recommendation-contingent rewards.

On the other hand, if the principal has access to recommendation-contingent rewards then Robinson monitoring is enforceable, as Example 1 shows. In terms of the margin of recommendation-contingent rewards, now the principal can choose different actions to detect different deviations. If Robinson is asked to monitor but chooses to rest and report g , then the principal can react by asking Friday to shirk, which would have led to b had Robinson monitored and reported truthfully. Similarly, if Robinson rests and reports b then Friday can be asked to work instead, rendering Robinson's deviation detectable again. Finally, if Robinson rests and mixes his reports then Friday can just shirk so that Robinson should have reported b with probability one. (See also Example 2 for similar reasoning.)

Evidently, recommendation-contingent rewards cannot help to enforce a pure strategy profile a , even if they help to virtually enforce it. By enforcing, a requires that every a -profitable disobedience be $\{a\}$ -detectable. Of course, $\{a\}$ -detectability is the same as detectability at a . Since agents receive only one recommendation under a , there is no use in having payments depend on recommendations. For a correlated strategy with nonsingleton support, however, the two contract spaces differ, and as such so do the notions of detectability that characterize enforcement. In general, μ is enforceable (with or without recommendation-contingent rewards) if and only if every μ -profitable deviation is detectable (either *supp μ -* or *at μ*). Hence,

²⁵This definition is very close to (and implied by) IFR. Formally, IFR at μ means that every deviation σ_i is detectable at μ , where σ_i need not be nonnegative. Fudenberg, Levine, and Maskin (1994) just focused on mixed strategy profiles μ ; Kandori and Matsushima (1998) generalized IFR to nonnegative vectors σ_i .

different payment schemes may be used to discourage different detectable deviations. Furthermore, with recommendation-contingent payments, different actions (in the support of μ) may be used to detect different deviations, whereas without them μ itself must simultaneously detect every profitable deviation.

This shows how Theorem 3 relies deeply on recommendation-contingent rewards. To conclude that a monitor's undetectable deviations are irrelevant, we argued that any such deviation is just as good as monitoring because given any action profile, it does not change the distribution of reported signals. Moreover, the set of undetectable deviations is independent of the correlated strategy being enforced, and remains unchanged at every approximation stage of virtual enforcement. Without recommendation-contingent rewards, a deviation must be detected "at μ ." It does not follow that a deviation from monitoring, undetectable at μ , will still detect the original deviation for which monitoring was recommended *at the monitor's deviation* from μ . Moreover, the set of undetectable deviations depends on the correlated strategy, which changes at every approximation stage. Thus, a deviation from monitoring may be undetectable at one stage of the approximation but not at another. For these reasons, a comparable version of Theorem 3 without recommendation-contingent rewards must fail.

B. Contractual Restrictions

So far, I have ignored limited liability, budget balance, and individual rationality. Thus, Robinson and Friday's contracts in Example 1 and the characterizations of virtual enforcement in and use arbitrarily large payments. Even though many important contributions to contract theory rely on them (e.g., Becker 1968; Holmström 1982), it is also important to understand how these sometimes unrealistic assumptions affect results.

Imposing one-sided limited liability on agents does not change the paper's results simply because a constant can be added to any payment scheme without disrupting incentives. Therefore, an allocation is (virtually) enforceable if and only if it is so subject to agents' one-sided limited liability. On the other hand, two-sided limited liability may restrict the set of enforceable outcomes if payments cannot be made large enough to discourage profitable deviations. Nevertheless, the spirit of Theorems 1 and 2 remains: deviations can still be discouraged one by one and different actions can be used to detect different deviations, but now the amount of detection must outweigh the utility gained from each deviation.

THEOREM 6: *A correlated strategy μ is enforceable with payments bounded above and below by $\pm z$ if and only if $\Delta v_i(\mu, \sigma_i) \leq z \sum_a \mu(a) \|\Delta \Pr(a, \sigma_i)\|$ for each deviation σ_i .*

Interestingly, restricting payments in Theorem 6 to not depend on recommendations yields $\Delta v_i(\mu, \sigma_i) \leq z \|\Delta \Pr(\mu, \sigma_i)\|$ as a characterizing condition: μ must simultaneously detect every deviation by an amount that outweighs its utility gain. Mathematically, the set of virtually enforceable correlated strategies is the closure of the set of enforceable ones. With two-sided limited liability, the latter is already closed, so " z -constrained" virtual and exact enforceability coincide. Also, as liability

limits are relaxed (i.e., $z \rightarrow \infty$), “ z -constrained” enforceable outcomes converge to “unconstrained” virtually enforceable ones.

This motivates virtual enforceability as a benchmark for approachable outcomes before knowing agents’ liability limits. This benchmark helps to understand both the social cost of these limits and the gain from relaxing them. To illustrate, recall Becker’s basic trade-off between the probability of monitoring and the size of contingent penalties. Theorems 3 and 4 may be interpreted as yielding necessary (and sufficient) conditions for Becker’s contracts to provide the right incentives, and Theorem 6 as characterizing his trade-off in general.

As for limits on the principal’s liability, it is possible that, ex post, the principal cannot pay an aggregate amount larger than some upper bound. This constraint, together with one-sided limited liability on the agents, also affects the set of feasible outcomes. (On its own, it clearly does not affect the feasible set.) As in Theorem 6, the set of feasible payment schemes is compact, so virtual and exact enforcement coincide, with comparable results.

THEOREM 7: *A correlated strategy μ is enforceable with an incentive scheme ζ that satisfies (i) limited liability for the agents, i.e., $\zeta_i(a, s) \leq 0$, and (ii) limited liability for the principal, i.e., $\sum_i \zeta_i(a, s) \geq -z$ for some given $z \geq 0$, if and only if for each strategy profile $(\sigma_1, \dots, \sigma_n)$,*

$$(2) \quad \sum_{i \in I} \Delta v_i(\mu, \sigma_i) \leq z \sum_{(a, s)} \mu(a) \max_{i \in I} \{\Delta \Pr(s | a, \sigma_i)^-\},$$

where $\Delta \Pr(s | a, \sigma_i)^- = -\min\{\Delta \Pr(s | a, \sigma_i), 0\} \geq 0$ is the negative part of $\Delta \Pr(s | a, \sigma_i)$.

The proof of Theorem 7 is omitted, as it is similar to that for Theorem 6. Let us briefly interpret this result. Given a profile of strategies, the left-hand side of (2) stands for the sum of expected unilateral deviation gains across agents arising from this profile, whereas the right-hand side stands for the z -weighted expected maximal probability decrease across agents (think of this decrease as good news). The principal’s liability constraint relates each agent’s incentive scheme, so it is no longer without loss to provide incentives agent by agent. As in Theorem 1, however, it suffices to discourage strategy profiles one by one. To discourage any such profile, for each recommendation and signal profile (a, s) , simply reward with the maximum amount z whoever would have most decreased the probability of s given a with σ_i relative to honesty and obedience—everyone else is paid nothing. Heuristically, whoever is the least likely to have deviated gets rewarded.

To compare Theorem 6 with Theorem 7, first notice that letting σ consist of just one agent deviating, it is clear that every profitable deviation must be detectable. Moreover, $\Delta \Pr(s | a, \sigma_i)^- \leq \|\Delta \Pr(s | a, \sigma_i)\|$ implies that, given z , enforceability with two-sided limited liability follows from enforceability with one-sided limited liability. To economize on his liability, in the one-sided environment the principal optimally chooses who to reward, whereas in two-sided environment he can reward and punish agents independently.

Regarding budget balance, first notice that ex ante budget constraints do not affect any results: just add a constant to any payments so that they hold. Ex post budget

balance, however—i.e., the sum of payments across agents is always the same—does matter. Thus, in Example 1, the profile (rest, work) is not virtually enforceable with budget balance. Rahman and Obara (2010) characterize enforceability subject to such budget constraints, and a similar general theme prevails. As they show, in addition to detecting disobediences, identifying obedient agents is necessary (and sufficient) for budget-balanced enforcement. The paper fails to characterize virtual enforceability, however, with or without budget balance.

Finally, participation constraints are easily satisfied alone without disrupting incentives, again by shifting enforcing payments. Together with other constraints they generally bind, though. Rahman and Obara (2010) study how they interact with budget balance.

C. Robustness

I will now discuss the model's robustness to preferences, agents' beliefs, and the monitoring technology. Of course, Theorem 1 cannot be robust to the set of profitable, detectable deviations because its purpose is to characterize enforceability: the theorem sacrifices robustness for fine-tuning. Still, other important results are robust in some sense, as I argue next.

COROLLARY 3: *Fix any correlated strategy μ . There is a payment scheme ζ such that*

$$\mathbf{1}_{\{a_i \neq b_i\}} \leq \sum_{(a_{-i}, s)} \mu(a) \zeta_i(a, s) [\Pr(s | a_{-i}, b_i, \rho_i) - \Pr(s | a)] \quad \forall (i, a_i, b_i, \rho_i)$$

if and only if every supp μ -disobedience is supp μ -detectable. Therefore, the same payment scheme ζ implements μ for any utility profile v such that $\sup_{(i, a, b_i)} \{v_i(b_i, a_{-i}) - v_i(a)\} \leq 1$.

Corollary 3 follows from Theorem 2 and implies that *there is a single payment scheme that simultaneously discourages every disobedience regardless of preferences v* if deviation gains are bounded and this bound is known. Corollary 3 says more: there is a payment scheme that makes honesty and obedience a “strict equilibrium”; i.e., all the disobedience constraints above hold with strict inequality.²⁶ This yields robustness to just about everything. Thus, by Corollary 3, there is a payment scheme that enforces a given allocation even if agents' interim beliefs about others' actions or the monitoring technology are slightly perturbed.

Another robustness measure is genericity. As long as there are enough action-signal pairs for every agent's opponents, I argue next that every disobedience is detectable generically on the set of monitoring technologies; i.e., except for those in a set of Lebesgue measure zero.

²⁶To see this, scale the scheme in Corollary 3 enough to strictly outweigh the gains from any disobedience. To also strictly discourage dishonesty, just replace “disobedience” with “deviation” above.

THEOREM 8: *Every disobedience is detectable generically if for every agent i ,*

- (i) $|A_i| - 1 \leq |A_{-i}| (|S_{-i}| - 1)$ *when* $|S_{-i}| > 1$,²⁷ *and*
- (ii) $|A_i| (|S_i| - 1) \leq ||A_{-i}| - 1$ *when* $|S_{-i}| = 1$.

Intuitively, genericity holds even if $|S| = 2$, as long as agents have enough actions. Hence, a group of agents may overcome their incentive constraints generically even if only one individual can make substantive observations and these observations are just a binary bit of information. If others' action spaces are large enough and their actions have generic effects on the bit's probability, this uniquely informed individual may still be controlled by testing him with unpredictable combinations of others' actions.²⁸ Thus, if the monitoring technology were chosen uniformly at random then almost surely every disobedience would be detectable. By Theorem 2, every correlated strategy would be virtually enforceable.

D. Collusion, Multiple Equilibria, and Renegotiation

Sections VB and VC discussed extensions that demonstrate the central conclusions of the paper hold in more general environments. In particular, analogues to Theorems 1 and 2 can be obtained when payments are bounded or subject to budget constraints.

On the other hand, the paper's results would not hold if I permitted collusion. Thus, in Example 1, Robinson and Friday could communicate "extracontractually" and break down incentives,²⁹ or Robinson could buy Friday's recommendation.³⁰ At the same time, collusion is a problem in general.³¹ For instance, the surplus-extracting scheme of Cremer and McLean (1988) is not collusion-proof for similar reasons. In this paper I have not tried to overcome collusion because I view it as an additional layer, though it is possible to find conditions under which collusion among agents may be overcome. Thus, one may restrict attention to constant-surplus contracts, as in Che and Kim (2006). Generally, though, I take the view that in some settings collusion is more of a problem than in others.

The contracts of this paper rely on communication, so they exhibit multiple equilibria: there are "babbling" equilibria where everyone ignores recommendations. This precludes "full implementation." (See also Kar, Ray, and Serrano 2010.) Multiplicity brings with it the usual "baggage" of selection, but this baggage may not be all bad: it is an equilibrium for everyone to ignore extracontractual messages, which restores some resilience to collusion.

Finally, renegotiation may also be included as an additional layer to the problem outlined in this paper, to obtain conditions for mediated contracts that are robust in this sense. Some forms of renegotiation have been likened in the literature to efficiency and budget balance (e.g., Neeman and Pavlov 2010, and references therein).

²⁷ In comparison, genericity for IFR with public monitoring requires much more: $|A_i| \leq |S| - 1$ for every i .

²⁸ I thank Roger Myerson for urging me to emphasize this point.

²⁹ With more contrived contracts, such communication can be deterred. Details are available on request.

³⁰ If there are more workers, some forms of collusion may be averted, with budget balanced contracts.

³¹ The likelihood of collusion depends on agents' ability to communicate, so it need not always be a problem.

Again, Rahman and Obara (2010) characterize allocations that are enforceable with budget balanced payments.

VI. Literature

In this section I compare the results of this paper with the relevant literature, specifically partnerships, mechanism design, and the theory of repeated games.

A. The Partnership Problem

Alchian and Demsetz's partnership problem may be described intuitively as follows. Consider two people working together in an enterprise that involves mutual effort. The efficient amount of effort would align each party's marginal effort cost with its marginal social benefit, which in a competitive economy coincides with the firm's profit. Each individual has the incentive, however, to align his marginal effort cost with just his share of the marginal benefit, rather than the entire marginal benefit. This inevitably leads to shirking. One way to solve, or at least mitigate, this shirking problem would be for the firm to hire a monitor in order to contract directly for the workers' effort. But then who will monitor the monitor?

According to Alchian and Demsetz (1972, p. 778), "[t]wo key demands are placed on an economic organization—metering input productivity and metering rewards."³² At the heart of their "metering problem" lies the question of how to give incentives to monitors, which they answered by making the monitor residual claimant. This can leave the monitor with incentives to misreport input productivity, however, if his report influences input rewards, like workers' wages, since—given efforts—paying workers hurts him directly.³³ Hence, making the monitor residual claimant, or principal, fails to provide the right incentives.

On the other hand, Holmström (1982, p. 325) argues that "...the principal's role is not essentially one of monitoring ... the principal's primary role is to break the budget-balance constraint." He shows that if output is publicly verifiable then the principal can provide the right incentives to agents with "group penalties" that reward all agents when output is good and punish them all when it is bad. Where Alchian and Demsetz seem to overemphasize the role of monitoring in organizations, Holmström seems to underemphasize it. By assuming that output is publicly verifiable, he finds little role for monitoring, and as such Holmström (1982, p. 339) concludes, wondering "...how should output be shared so as to provide all members of the organization (including monitors) with the best incentives to perform?"

In this paper I accommodate costly private monitoring and find a contract that gives both workers and monitors the right incentives to perform. It also addresses the partnership problem. Although I had just one worker (Friday), it is easy to add more and find a similar contract that gives everyone incentives without the principal having

³² "Meter means to measure and also to apportion. One can meter (measure) output and one can also meter (control) the output. We use the word to denote both; the context should indicate which." (Alchian and Demsetz 1972, p. 778).

³³ Similarly, Strausz (1997) observed that delegated monitoring dominates monitoring by a principal who cannot commit to verifying privately observed effort. Strausz assumes, however, that monitoring signals are so-called hard evidence, so a monitor cannot misreport his information. I allow for soft evidence.

to spend any money. (Details are available on request.) Since the principal observes reports and makes recommendations at no cost, he would happily disclose them even if they were not verifiable. Thus, *nobody needs to monitor the principal*. Rahman and Obara (2010) have related work.³⁴

Legros and Matthews (1993) explored virtual enforcement before, using mixed strategies. They found sufficient conditions for virtually enforcing an efficient outcome with ex post budget balance, so-called “nearly efficient” partnerships.³⁵ They only considered output-contingent payments, however. See Rahman and Obara (2010) for a detailed comparison.

As for rich contracts spaces, “random” contracts have been used in the literature to relax incentive constraints, either by exploiting risk aversion (Rasmusen 1987; Arnott and Stiglitz 1988; Cole 1989; Bennardo and Chiappori 2003; Rahman 2005; Strausz 2010), or option values (e.g., Ederer, Holden, and Meyer 2009, in the context of multitasking).

Partnership dynamics are interesting (see, e.g., Radner, Myerson, and Maskin 1986; Levin 2003; Fuchs 2007) because they add useful specificity: now the principal has dynamic instruments, such as timing of dissolution. Reinterpreting continuation values as payments, I abstract from such important details but acknowledge that they are implicit in the paper.

There is also an important literature on market-based incentives, such as MacLeod and Malcomson (1998); Prendergast (1999); Tadelis (2002); and others. Although this model is not market-based, market incentives may be incorporated via participation constraints.

Let us turn to mechanism design theory. The seminal work on surplus extraction by Cremer and McLean (1985, 1988) relies on exogenous correlation in agents’ private information to discipline reporting and extract agents’ surplus. There are similarities with the contracts of this paper, but also important differences. First, types are correlated endogenously in my model: the principal allocates private information to provide incentives. Second, I don’t always need every agent’s report to provide incentives. Thus, in Example 1 the principal told Friday his type, whereas Cremer and McLean solicit information from every agent. Third, I focus on enforceability rather than surplus extraction, which clearly yields less restrictive conditions. Finally, since their contracts extract all surplus, participation constraints bind, so they are vulnerable to even small perturbations in fundamentals.

The work of Mezzetti (2004, 2007) is also related. He observes that when agents’ values are interdependent, their realized utilities are correlated conditional on the public outcome (here the outcome is the action profile). Hence, it is possible to discipline agents further by conditioning their payments on reported utility profiles. In a sense, this paper may be viewed as generalizing his results by viewing utility realizations as monitoring signals.

³⁴They find sufficient but not necessary conditions for virtual enforcement with ex post budget balance.

³⁵Miller (1997) enriches the model of Legros and Matthews by adding costless private monitoring.

B. Detection and Enforcement

The duality between detection and enforcement is a classic theme in the design of incentives. Early papers to point this out are Abreu, Milgrom, and Pearce (1991) and Fudenberg, Levine, and Maskin (1994), in the context of repeated games. Together with the literature on partnerships, such as Hermalin and Katz (1991); Legros and Matsushima (1991); d'Aspremont and Gérard-Varet (1998); and Legros and Matthews (1993), these papers focus on public monitoring. With private monitoring, Compte (1998) and Kandori and Matsushima (1998) derived Folk Theorems with public communication, whereas Kandori and Obara (2006); Ely, Hörner, and Olszewski (2005); and Kandori (2011) only permitted communication through actions. None of the papers above fully exploits mediation: payments/continuation values cannot depend on one's own intended/recommended action *and* reported signal. Thus, none can virtually enforce (rest, work) in Example 1.³⁶

Some recent papers have studied richer contracts in specific settings, such as Kandori (2003) and its private monitoring version by Obara (2008), Aoyagi (2005), and Tomala (2009). Aoyagi uses dynamic mediated strategies that rely on " ε -perfect" monitoring, but fail if monitoring is costly or one-sided. Tomala studies recursive communication equilibria and independently uses recommendation-contingent continuation values to prove a folk theorem. He derives a version of the Minimax Lemma, but does not study virtual enforcement.

Kandori has agents play mixed strategies and report the realization of their mixtures; payments may depend on these reports. Thus, in Example 1 Robinson can be "monitored" by having Friday mix and report what he did.³⁷ With limited liability, the principal pays more with Kandori's contracts. Also, if Robinson and Friday could commit to destroy or sell value, they could write a contract by themselves, without the principal. But more importantly, since agents are asked what they did, they require incentives to tell the truth. Such reporting constraints do not exist if instead the principal tells agents what to do, so recommendation-contingent rewards generally dominate Kandori's, as in the next example.

Example 6: One agent, three actions (L, M, R), two publicly verifiable signals (g, b).

L	M	R	L	M	R
0	2	0	1, 0	1/2, 1/2	0, 1
Utility payoffs			Signal probabilities		

The mixed strategy $\sigma = \frac{1}{2}[L] + \frac{1}{2}[R]$ is enforceable, but not with Kandori's contracts. Indeed, offering \$1 for g if asking to play L and \$1 for b if asking to play R makes σ enforceable. With Kandori's contracts, the agent supposedly plays σ and is then asked what he played before getting paid. He gains two "utils" by playing

³⁶ A similar comment applies to Phelan and Skrzypacz (2008) and Kandori and Obara (2010).

³⁷ Let σ and μ be the probability that Robinson monitors and Friday (independently) works. If Robinson says he monitored, he gets $\$1/\mu$ and Friday $\$1/\sigma$ if both say Friday worked, but $\$1/(1-\mu)$ and Friday \$0 if both say Friday shirked. Otherwise, both get \$0. If Robinson says he rested, he gets \$1 and Friday \$0.

M instead and reporting $L(R)$ if the realized signal is $g(b)$, with the same expected monetary payoff.

If agents see nothing before reporting the realization of their mixed strategy, since they must be indifferent over what to play, they will also be indifferent over what action to report. Therefore, if agents can secretly report their actions before observing anything relevant, then Kandori's contracts induce the same enforceable mixed strategy profiles as with recommendation-contingent payments. This suggests another improvement: if possible, have agents report their *intended* action (as in pool tables everywhere) before playing it.

Kandori's "unmediated" contracts are not without fault. First, agents may not be able to commit to report intended actions before observing any information. Second, without recommendations, only mixed strategy profiles are enforceable. This may be undesirable: in the classic chicken game (see, e.g., Aumann 1974), maximizing welfare involves correlated equilibria that are not a public randomization over Nash equilibria. Third, when agents mix they must be indifferent, whereas with recommendation-contingent payments they may be given strict incentives. Such robustness—as described in the critique of Bhaskar (2000) (see also Bhaskar, Mailath, and Morris 2008)—fails elsewhere, but holds in this paper. Finally, virtually enforcing even a pure strategy profile may require mediation, as in the next example.

Example 7: There are three agents: Rowena picks a row, Colin a column, and Matt a matrix. Rowena and Colin are indifferent over everything. Here is Matt's utility function.

	L	R		L	R		L	R
U	1	2	U	0	0	U	-1	2
D	2	-1	D	0	0	D	2	1
	A			B			C	

There are two publicly verifiable signals. The monitoring technology is below.

	L	R		L	R		L	R
U	$\frac{1}{2}, \frac{1}{2}$	1, 0	U	$\frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{2}$	U	$\frac{1}{2}, \frac{1}{2}$	0, 1
D	1, 0	1, 0	D	$\frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{2}$	D	0, 1	0, 1
	A			B			C	

Clearly, the action profile (U, L, B) is not enforceable, since playing A instead of B is a (U, L, B) -profitable, (U, L, B) -undetectable deviation. To virtually enforce (U, L, B) , Rowena and Colin cannot play just (U, L) . But then playing $\frac{1}{2}[A] + \frac{1}{2}[C]$ instead of B is Matt's only undetectable deviation. Call this deviation σ . If only Rowena mixes and plays U with probability $0 < p < 1$ then σ is profitable: Matt's profit equals $2(1 - p) > 0$. Similarly, if only Colin mixes between L and R then σ is profitable. If Rowena and Colin mix independently, with $p = \Pr(U)$ and $q = \Pr(L)$, Matt still profits from σ : he gets $2p(1 - q) + 2(1 - p)q > 0$. On the other hand, the correlated strategy $r[(U, L, B)] + (1 - r)[(D, R, B)]$ for $0 < r < 1$ renders σ

unprofitable, which is still Matt's only undetectable deviation. Therefore, (U, L, B) is virtually enforceable (by letting $r \rightarrow 1$), although only with correlated behavior.

Lastly, the work of Lehrer (1992) is especially noteworthy. He characterizes the set of uniform equilibrium payoffs (heuristically, discount factor equals one) in a two-player repeated game with imperfect monitoring in terms of sustainability. A payoff profile is *sustainable* if there is a mixed strategy profile $\mu = (\mu_1, \mu_2)$ that attains it and every μ -profitable deviation is detectable. Lehrer has players monitor each other with rapidly decreasing probability so that monitoring costs are negligible in the long run. With discounted utility, his argument fails. To describe equilibrium payoffs as the discount factor δ tends to one rather than at the limit of $\delta = 1$, virtual enforceability is the appropriate notion, not sustainability. Thus, the profile (rest, work) of Example 4 is clearly sustainable but not virtually enforceable. Using Theorem 4 to help describe limiting equilibrium payoffs in a repeated game as $\delta \rightarrow 1$ and any discontinuity in the equilibrium correspondence with respect to δ at $\delta = 1$ (e.g., Radner, Myerson, and Maskin 1986) is the purpose of future research.

VII. Conclusion

In this paper, I offer a new answer to Alchian and Demsetz's classic question of who will monitor the monitor: the principal "monitors" the monitor's detectable deviations by having his workers shirk occasionally, and nobody needs to monitor the monitor's undetectable deviations (Theorem 3). How to monitor the monitor? With "trick questions," as in Robinson's contract (Example 1). This contract aligns incentives by making the monitor responsible for monitoring. When is this outcome (virtually) enforceable? When every deviation from the desired outcome is detectable—even if the detecting behavior is undesirable.

Alchian and Demsetz argued that the monitor must be made residual claimant for his incentives to be aligned. In a sense, they "elevated" the role of monitoring in organizations. On the other hand, I have argued for "demoting" their monitor to a security guard—low down in the ownership hierarchy. As such, the question remains: what is the economic role of residual claimant? Answering this classic question is the purpose of future research.

APPENDIX: PROOFS

First, I state the Alternative Theorem (Rockafellar 1970, p. 198): *Let $A \in \mathbb{R}^{\ell \times m}$ and $b \in \mathbb{R}^m$. There exists $x \in \mathbb{R}^m$ such that $Ax \leq b$ if and only if for every $\lambda \in \mathbb{R}_+^m$, $\lambda A = 0$ implies $\lambda \cdot b \geq 0$.*

PROOF OF THEOREM 1:

By the Alternative Theorem, μ is *not* enforceable if and only if

$$\sum_{(b_i, \rho_i)} \mu(a) \lambda_i(a_i, b_i, \rho_i) [\Pr(s|a_{-i}, b_i, \rho_i) - \Pr(s|a)] = 0 \quad \forall (a, s)$$

and $\Delta v_i(\mu, \lambda_i) > 0$ for some i and vector $\lambda_i \geq 0$. Such λ_i exists if and only if σ_i , defined by

$$\sigma_i(b_i, \rho_i | a_i) := \begin{cases} \lambda_i(a_i, b_i, \rho_i) / \sum_{(b'_i, \rho'_i)} \lambda_i(a_i, b'_i, \rho'_i) & \text{if } \sum_{(b'_i, \rho'_i)} \lambda_i(a_i, b'_i, \rho'_i) > 0, \text{ and} \\ [(a_i, \tau_i)](b_i, \rho_i) & \text{otherwise (where } [\cdot] \text{ denotes Dirac measure),} \end{cases}$$

is μ -profitable and $\text{supp } \mu$ -undetectable.

PROOF OF THEOREM 2:

Let $B = \text{supp } \mu$. By the Alternative Theorem, every B -disobedience is B -detectable if and only if a scheme ξ exists such that $\xi_i(a, s) = 0$ if $a \notin B$ and

$$0 \leq \sum_{(a_{-i}, s)} \xi_i(a, s) (\Pr(s | a_{-i}, b_i, \rho_i) - \Pr(s | a))$$

$$\forall i \in I, a_i \in B_i, b_i \in A_i, \rho_i \in R_i,$$

with a strict inequality whenever $a_i \neq b_i$, where $B_i = \{a_i \in A_i : \exists a_{-i} \in A_{-i} \text{ s.t. } a \in B\}$. Replacing $\xi_i(a, s) = \mu(a) \zeta_i(a, s)$ for any μ with $\text{supp } \mu = B$, this is equivalent to there being, for every v , an appropriate rescaling of ζ that satisfies the incentive constraints (1).

PROOF OF THEOREM 3:

Let $B = \text{supp } \mu$. For necessity, suppose there is a B -disobedient, undetectable disobedience σ_i , so $\sigma_i(b_i, \rho_i | a_i) > 0$ for some $a_i \in B_i$, $b_i \neq a_i$, and $\rho_i \in R_i$. Letting $v_i(a_{-i}, b_i) < v_i(a)$ for every a_{-i} , clearly no correlated strategy with positive probability on a_i is virtually enforceable. Sufficiency follows by Lemmata B.3, B.4, and B.10 of online Appendix B.

PROOF OF THEOREM 4:

See the end of online Appendix B.

PROOF OF THEOREM 5:

Fix any $\mu \in \Delta(A)$. By the Alternative Theorem, every μ -profitable deviation is detectable at μ if and only if a scheme $\zeta : I \times S \rightarrow \mathbb{R}$ exists such that for all (i, a_i, b_i, ρ_i) , $\sum_{a_{-i}} \mu(a) [v_i(a_{-i}, b_i) - v_i(a)] \leq \sum_{(a_{-i}, s)} \mu(a) \zeta_i(s) [\Pr(s | a_{-i}, b_i, \rho_i) - \Pr(s | a)]$, as required.

PROOF OF THEOREM 6:

Follows from Lemma B.2(i) of online Appendix B.

PROOF OF THEOREM 7:

See online Appendix B.

REFERENCES

- Abreu, Dilip, Paul Milgrom, and David Pearce. 1991. "Information and Timing in Repeated Partnerships." *Econometrica* 59 (6): 1713–33.
- Alchian, Armen A., and Harold Demsetz. 1972. "Production, Information Costs, and Economic Organization." *American Economic Review* 62 (5): 777–95.
- Aoyagi, Masaki. 2005. "Collusion through Mediated Communication in Repeated Games with Imperfect Private Monitoring." *Economic Theory* 25 (2): 455–75.
- Arnott, Richard, and Joseph E. Stiglitz. 1988. "Randomization with Asymmetric Information." *RAND Journal of Economics* 19 (3): 344–62.
- Aumann, Robert J. 1974. "Subjectivity and Correlation in Randomized Strategies." *Journal of Mathematical Economics* 1 (1): 67–96.
- Baron, David P., and David Besanko. 1984. "Regulation, Asymmetric Information, and Auditing." *RAND Journal of Economics* 15 (4): 447–70.
- Basu, Kaushik, Sudipto Bhattacharya, and Ajit Mishra. 1992. "Notes on Bribery and the Control of Corruption." *Journal of Public Economics* 48 (3): 349–59.
- Becker, G. S. 1968. "Crime and Punishment: An Economic Approach." *Journal of Political Economy* 76 (2): 169–217.
- Bennardo, Alberto, and Pierre-Andre Chiappori. 2003. "Bertrand and Walras Equilibria under Moral Hazard." *Journal of Political Economy* 111 (4): 785–817.
- Ben-Porath, Elchanan, and Michael Kahneman. 1996. "Communication in Repeated Games with Private Monitoring." *Journal of Economic Theory* 70 (2): 281–97.
- Ben-Porath, Elchanan, and Michael Kahneman. 2003. "Communication in Repeated Games with Costly Monitoring." *Games and Economic Behavior* 44 (2): 227–50.
- Bhaskar, V. 2000. "The Robustness of Repeated Game Equilibria to Incomplete Payoff Information." Unpublished.
- Bhaskar, V., George J. Mailath, and Stephen Morris. 2008. "Purification in the Infinitely-Repeated Prisoners' Dilemma." *Review of Economic Dynamics* 11 (3): 515–28.
- Border, Kim C., and Joel Sobel. 1987. "Samurai Accountant: A Theory of Auditing and Plunder." *Review of Economic Studies* 54 (4): 525–40.
- Che, Yeon-Koo, and Jinwoo Kim. 2006. "Robustly Collusion-Proof Implementation." *Econometrica* 74 (4): 1063–1107.
- Cole, Harold Linh. 1989. "General Competitive Analysis in an Economy with Private Information: Comment." *International Economic Review* 30 (1): 249–52.
- Compte, Olivier. 1998. "Communication in Repeated Games with Imperfect Private Monitoring." *Econometrica* 66 (3): 597–626.
- Cremer, Jacques, and Richard P. McLean. 1985. "Optimal Selling Strategies under Uncertainty for a Discriminating Monopolist When Demands Are Interdependent." *Econometrica* 53 (2): 345–61.
- Cremer, Jacques, and Richard P. McLean. 1988. "Full Extraction of the Surplus in Bayesian and Dominant Strategy Auctions." *Econometrica* 56 (6): 1247–57.
- Dallos, Robert E. 1987. "'Ghost Riders': Airlines Spy on Selves in Service War." *Los Angeles Times*, July 21.
- d'Aspremont, Claude, and Louis-Andre Gérard-Varet. 1998. "Linear Inequality Methods to Enforce Partnerships under Uncertainty: An Overview." *Games and Economic Behavior* 25 (2): 311–36.
- Ederer, Florian, Richard Holden, and Margaret Meyer. 2009. "Gaming and Strategic Ambiguity in Incentive Provision." Unpublished.
- Ely, Jeffrey C., Johannes Hörner, and Wojciech Olszewski. 2005. "Belief-Free Equilibria in Repeated Games." *Econometrica* 73 (2): 377–415.
- Ewoldt, John. 2004. "Dollars and Sense: Undercover Shoppers." *Star Tribune*, October 27.
- Forges, Françoise M. 1986. "An Approach to Communication Equilibria." *Econometrica* 54 (6): 1375–85.
- Fuchs, William. 2007. "Contracting with Repeated Moral Hazard and Private Evaluations." *American Economic Review* 97 (4): 1432–48.
- Fudenberg, Drew, David I. Levine, and Eric Maskin. 1994. "The Folk Theorem with Imperfect Public Information." *Econometrica* 62 (5): 997–1039.
- Gneiting, Tilmann, and Adrian E. Raftery. 2007. "Strictly Proper Scoring Rules, Prediction, and Estimation." *Journal of the American Statistical Association* 102 (477): 359–78.
- Grossman, Sanford J., and Oliver D. Hart. 1983. "An Analysis of the Principal-Agent Problem." *Econometrica* 51 (1): 7–45.
- Harrington, Joseph E., Jr. 2008. "Optimal Corporate Leniency Programs." *Journal of Industrial Economics* 56 (2): 215–46.

- Hart, Sergiu, and David Schmeidler.** 1989. "Existence of Correlated Equilibria." *Mathematics of Operations Research* 14 (1): 18–25.
- Hermalin, Benjamin E., and Michael L. Katz.** 1991. "Moral Hazard and Verifiability: The Effects of Renegotiation in Agency." *Econometrica* 59 (6): 1735–53.
- Holmström, Bengt.** 1982. "Moral Hazard in Teams." *Bell Journal of Economics* 13 (2): 324–40.
- Kandori, Michihiro.** 2003. "Randomization, Communication, and Efficiency in Repeated Games with Imperfect Public Monitoring." *Econometrica* 71 (1): 345–53.
- Kandori, Michihiro.** 2011. "Weakly Belief-Free Equilibria in Repeated Games with Private Monitoring." *Econometrica* 79 (3): 877–92.
- Kandori, Michihiro, and Hitoshi Matsushima.** 1998. "Private Observation, Communication and Collusion." *Econometrica* 66 (3): 627–52.
- Kandori, Michihiro, and Ichiro Obara.** 2006. "Efficiency in Repeated Games Revisited: The Role of Private Strategies." *Econometrica* 74 (2): 499–519.
- Kandori, Michihiro, and Ichiro Obara.** 2010. "Towards a Belief-Based Theory of Repeated Games with Private Monitoring: An Application of POMDP." Unpublished.
- Kaplow, Louis, and Steven Shavell.** 1994. "Optimal Law Enforcement with Self-Reporting of Behavior." *Journal of Political Economy* 102 (3): 583–606.
- Kar, Anirban, Indrajit Ray, and Roberto Serrano.** 2010. "A Difficulty in Implementing Correlated Equilibrium Distributions." *Games and Economic Behavior* 69 (1): 189–93.
- Knapp, W.** 1972. *Report of the Commission to Investigate Alleged Police Corruption*. New York: George Braziller, Inc.
- Kutz, Gregory D.** 2008a. "Border Security: Summary of Covert Tests and Security Assessments for the Senate Committee on Finance." Government Accountability Office Report GAO-08-757.
- Kutz, Gregory D.** 2008b. "Investigative Operations: Use of Covert Testing to Identify Security Vulnerabilities and Fraud, Waste, and Abuse." Government Accountability Office Testimony GAO-08-286T.
- Kutz, Gregory D.** 2008c. "Medicare: Covert Testing Exposes Weaknesses in the Durable Medical Equipment Supplier Screening Process." Government Accountability Office Report GAO-08-955.
- Kutz, Gregory D.** 2008d. "Drug Testing: Undercover Tests Reveal Significant Vulnerabilities in DOT's Drug Testing Program." Government Accountability Office Testimony GAO-08-225T.
- Kutz, Gregory D.** 2009a. "Military and Dual-Use Technology: Covert Testing Shows Continuing Vulnerabilities of Domestic Sales for Illegal Export." Government Accountability Office Testimony GAO-09-725T.
- Kutz, Gregory D.** 2009b. "Department of Labor: Wage and Hour Division Needs Improved Investigative Processes and Ability to Suspend Statute of Limitations to Better Protect Workers Against Wage Theft." Government Accountability Office Report GAO-09-629.
- Legros, Patrick, and Hitoshi Matsushima.** 1991. "Efficiency in Partnerships." *Journal of Economic Theory* 55 (2): 296–322.
- Legros, Patrick, and Steven A. Matthews.** 1993. "Efficient and Nearly-Efficient Partnerships." *Review of Economic Studies* 60 (3): 599–611.
- Lehrer, Ehud.** 1992. "On the Equilibrium Payoffs Set of Two Player Repeated Games with Imperfect Monitoring." *International Journal of Game Theory* 20 (3): 211–26.
- Levin, Jonathan.** 2003. "Relational Incentive Contracts." *American Economic Review* 93 (3): 835–57.
- MacLeod, W. Bentley.** 2003. "Optimal Contracting with Subjective Evaluation." *American Economic Review* 93 (1): 216–40.
- MacLeod, W. Bentley, and James M. Malcomson.** 1998. "Motivation and Markets." *American Economic Review* 88 (3): 388–411.
- Marx, Gary T.** 1992. "When the Guards Guard Themselves: Undercover Tactics Turned Inward." *Policing and Society* 2 (3): 151–72.
- Mezzetti, Claudio.** 2004. "Mechanism Design with Interdependent Valuations: Efficiency." *Econometrica* 72 (5): 1617–26.
- Mezzetti, Claudio.** 2007. "Mechanism Design with Interdependent Valuations: Surplus Extraction." *Economic Theory* 31 (3): 473–88.
- Miller, Nolan H.** 1997. "Efficiency in Partnerships with Joint Monitoring." *Journal of Economic Theory* 77 (2): 285–99.
- Miller, Nathan H.** 2009. "Strategic Leniency and Cartel Enforcement." *American Economic Review* 99 (3): 750–68.
- Miyagawa, Eiichi, Yasuyuki Miyahara, and Tadashi Sekiguchi.** 2008. "The Folk Theorem for Repeated Games with Observation Costs." *Journal of Economic Theory* 139 (1): 192–221.
- Miyagawa, Eiichi, Yasuyuki Miyahara, and Tadashi Sekiguchi.** 2009. "Repeated Games with Costly Imperfect Monitoring." Unpublished.

- Mollen, M.** 1994. *The City of New York Commission to Investigate Allegations of Police Corruption and the Anti-Corruption Procedures of the Police Department: Commission Report*. New York: New York City Police Department.
- Mookherjee, Dilip, and Ivan Png.** 1989. "Optimal Auditing, Insurance, and Redistribution." *Quarterly Journal of Economics* 104 (2): 399–415.
- Mookherjee, Dilip, and Ivan Png.** 1992. "Monitoring vis-a-vis Investigation in Enforcement of Law." *American Economic Review* 82 (3): 556–65.
- Myerson, Roger B.** 1986. "Multistage Games with Communication." *Econometrica* 54 (2): 323–58.
- Myerson, Roger B.** 1997. "Dual Reduction and Elementary Games." *Games and Economic Behavior* 21 (1–2): 183–202.
- Nau, Robert F., and Kevin F. McCardle.** 1990. "Coherent Behavior in Noncooperative Games." *Journal of Economic Theory* 50 (2): 424–44.
- Neeman, Zvika, and Gregory Pavlov.** 2010. "Renegotiation-proof Mechanism Design." University of Western Ontario Department of Economics Working Paper 20101.
- Obara, Ichiro.** 2008. "The Full Surplus Extraction Theorem with Hidden Actions." *The B.E. Journal of Theoretical Economics* 8 (1): 8.
- Palmer, C. C.** 2001. "Ethical Hacking." *IBM Systems Journal* 40 (3): 769–780.
- Phelan, Christopher, and Andrzej Skrzypacz.** 2008. "Beliefs and Private Monitoring." Unpublished.
- Pontin, Jason.** 2007. "Artificial Intelligence, with Help from the Humans." *New York Times*, March 25.
- Prendergast, Canice.** 1999. "The Provision of Incentives in Firms." *Journal of Economic Literature* 37 (1): 7–63.
- Prenzler, Tim.** 2009. *Police Corruption: Preventing Misconduct and Maintaining Integrity*. Boca Raton, FL: Taylor & Francis Group, CRC Press.
- Radner, Roy, Roger Myerson, and Eric Maskin.** 1986. "An Example of a Repeated Partnership Game with Discounting and with Uniformly Inefficient Equilibria." *Review of Economic Studies* 53 (1): 59–69.
- Rahman, David M.** 2005. "Team Formation and Organization." PhD Diss., University of California, Los Angeles. ProQuest (AAT 3175198).
- Rahman, David, and Ichiro Obara.** 2010. "Mediated Partnerships." *Econometrica* 78 (1): 285–308.
- Rasmusen, Eric.** 1987. "Moral Hazard in Risk-Averse Teams." *RAND Journal of Economics* 18 (3): 428–35.
- Rockafellar, R. T.** 1970. *Convex Analysis*. Princeton, NJ: Princeton University Press.
- Sherman, Lawrence W.** 1978. *Scandal and Reform: Controlling Police Corruption*. Berkeley, CA: University of California Press.
- Skolnick, J. H.** 1966. *Justice Without Trial: Law Enforcement in Democratic Society*. New York: John Wiley & Sons.
- Strausz, Roland.** 1997. "Delegation of Monitoring in a Principal-Agent Relationship." *Review of Economic Studies* 64 (3): 337–57.
- Strausz, Roland.** 2010. "Mediated Contracts and Mechanism Design." Unpublished.
- Tadelis, Steven.** 2002. "The Market for Reputations as an Incentive Mechanism." *Journal of Political Economy* 110 (4): 854–82.
- Tomala, Tristan.** 2009. "Perfect Communication Equilibria in Repeated Games with Imperfect Monitoring." *Games and Economic Behavior* 67 (2): 682–94.
- Townsend, Robert M.** 1979. "Optimal Contracts and Competitive Markets with Costly State Verification." *Journal of Economic Theory* 21 (2): 265–93.
- Transportation Security Administration.** 2004. Guidance on Screening Partnership Program. www.tsa.gov/assets/pdf/SPP_OptOut_Guidance_6.21.04.pdf.
- Walker, David M.** 2007. "GAO Strategic Plan 2007–2012." Government Accountability Office Report GAO-07-1SP.