

# Robust Relational Contracts when Performance Evaluation is Subjective

V Bhaskar\*      Thomas Wiseman †

March 1, 2019

We study a repeated principal agent model with transferable utility, where the principal's evaluation of the agent's performance is subjective. Consequently, monitoring is noisy and private. We focus on purifiable equilibria, that are robust to small iid payoff shocks. Effort cannot be sustained in any finite memory purifiable equilibria; existing constructions fail to be purifiable. To address this problem, we allow the principal and agent to make simultaneous cheap talk announcements at the end of each period. This allows effort to be sustained with positive probability in every period, thereby we can approximate efficiency if the noise in monitoring is small. If we allow for a mediator, we can ensure payoffs arbitrarily close to full efficiency with non-vanishing noise, provided that the agents are sufficiently patient.

---

\*Department of Economics, University of Texas at Austin. [v.bhaskar@austin.utexas.edu](mailto:v.bhaskar@austin.utexas.edu).

†Department of Economics, University of Texas at Austin. [wiseman@austin.utexas.edu](mailto:wiseman@austin.utexas.edu).

In many organizations, the tasks that employees must perform lack an objective measure of performance. Thus performance evaluation is subjective, and the worker cannot observe the employer's evaluation of his own performance. When coupled with moral hazard, providing incentives becomes difficult. A series of papers such – see Levin (2003), MacLeod (2003) and Fuchs (2007)) – examine the use of *relational contracts* — i.e. repeated game mechanisms – in order to provide incentives. These contracts typically require the agent to exert effort, and for the principal to pay a bonus to the worker if and only if her evaluation of the worker's performance is good. However, in order to provide incentives for truth-telling, the employer is made indifferent between paying and not paying the bonus. To do this, the equilibrium must exhibit money-burning, for example by dissolving the relationship with some probability in the event that the bonus is not paid. Since the principal is indifferent between paying the bonus or not, she has incentives to truthfully disclose her subjective evaluation of the worker's performance. However, the equilibrium exhibits an inefficiency, since productive relationships must be dissolved with positive probability. Levin (2003) and Fuchs (2007) show that efficiency can be enhanced, by requiring the principal only to report on the agent's performance every  $T$  periods. As in Abreu, Milgrom, and Pearce (1991), the extent of inefficiency decreases with  $T$ , and when both players become arbitrarily patient, the one can approach full efficiency.

Our paper begins with the observation that the equilibria so constructed are fragile, and do not survive if principal and agent are subject to small iid payoff shocks. For example, the agent's cost of effort and his outside option in each period could be subject to a random shock, as could the principal's flow revenues. In this case, the equilibria where the principal is indifferent between paying the bonus or not do not survive, since the principal strictly prefers to pay the bonus when he learns that flow revenues in the next period are high, and strictly prefers not to pay it when he learns that they will be low. In consequence, she will condition her bonus payment on the shock, and not on the agent's performance. More generally, we show that in any *purifiable* finite memory equilibrium of the infinite horizon game played between the two players, the agent will not exert effort, rendering the relationship unprofitable. In particular, since all the above mentioned equilibria in the literature exhibit finite memory, none of them survive when there are payoff shocks.

This leads us to modify the game played between principal and agent, by allowing the players to make simultaneous cheap-talk announcements. We consider an equilibrium where the principal announces her private signal, while agent discloses his choice of effort (we require randomization on the part of the agent; otherwise, the agent's announcement would

be redundant). Incentives for truth-telling are provided by dissolving the relationship with positive probability whenever the announcements “differ” – i.e. when the principal announces a good signal and the agent announces that he shirked, or when the principal announces a bad signal and the agent worked. We are able to construct an equilibrium where the agent works with arbitrarily high probability, and this is close to being efficient if the noise in monitoring is small. Moreover, such an equilibrium is purifiable.

When the noise in monitoring is large, we are lead to explore equilibria where announcements are made only every  $T$  periods. However, a difficulty arises, since any equilibrium with announcements requires the agent to shirk with positive probability in every period. We are unable to construct an equilibrium with cheap talk that achieves, and consequently, we consider the role of a mediator. The mediator makes recommendations to the agent, and receives reports from the principal every  $T$  periods. We show that as players become arbitrarily patient, there exists an equilibrium which approximates the fully efficient payoff.

## 1 The basic model

Time is discrete, the horizon infinite, with both players discounting payoffs at a common rate  $\delta$ . In each period, the agent chooses from  $\{E, S\}$ , where the cost of effort  $E$  is  $c > 0$ , while the cost of shirking is zero. Output  $y$  is stochastic, and takes values in  $\{G, B\}$ , with  $Pr(y = G|E) = 1 - \epsilon$ ,  $Pr(y = B|S) = 1 - \eta$ . Let  $\bar{y}$  denote the expected value of output when  $E$  is chosen, let  $-\ell$  denote expected output when  $S$  is chosen. Assume that  $\bar{y} - c > 0 > -\ell$ , and normalize the outside options of the two parties to zero. Thus it is efficient for the agent to be employed and to choose effort, but if the agent shirks, then it is preferable to dissolve the relationship. Both parties are risk-neutral and there do not face limited liability constraints. They maximize the discounted sum of payoffs, with common discount factor  $\delta$ .

In the interests of precision, let us consider the following stage game  $\Gamma$  that is played in every period, conditional on the relationship not having been terminated.

- The agent is paid a base wage  $w$  and chooses  $a \in \{E, S\}$ .
- The principal observes  $y \in \{G, B\}$  and decides whether to pay a bonus or not, over and above the base wage  $w$ .
- Principal and agent observe the realization of a public randomization device and simultaneously decide whether to terminate the relationship or not – the relationship continues to the next period if and only if both parties want to continue.

We denote the repeated game by  $\Gamma^\infty$ .

The fundamental problem is that monitoring is imperfect and private. The principal does not observe the agent's action, and the agent does not observe  $y$ .<sup>1</sup> To incentivize effort, the agent's bonus payments (or his continuation value) must depend upon the principal's observation of output. However, since this observation is private, the principal's continuation value, net of the cost of the bonus payment, cannot depend upon the signal that the principal observes. The solution to this problem, proposed by ? and Levin (2003), is to ensure that the principal is indifferent between paying the bonus, or not paying it. This can be achieved via a public randomization device that decrees that the relationship be dissolved with some probability, whenever the bonus is not paid. In other words, a part of the surplus from the relationship must be destroyed, since the agent cannot be punished while simultaneously rewarding the principal.

Two problems arise with this repeated game equilibrium/relational contract. First, it is inefficient, since some surplus is destroyed. Levin (2003) and Fuchs (2007) show that the inefficiency can be mitigated by if the players are patient, by dividing the interaction into blocks of  $T$  periods. The bonus is withheld only if the agent fails in every period in the block, and this reduces the loss in surplus. Second, and more importantly, the equilibrium relies on the principal's indifference between paying the bonus and not paying it, and on her breaking this indifference according to the private signal. Consequently it is fragile. In particular, if the value of the relationship to the principal is subject to small shocks that are privately observed by the principal, then she will condition her bonus payment on the realization of these shocks, and not upon output signals.

We now make this argument more precise. The perturbed version of the stage game,  $\Gamma(\xi)$ , is as follows:

- The agent observes a random shock  $z_1$  before he chooses his action, and his cost of effort is augmented by  $\xi z_1$ .
- The principal observes a shock  $z_2$  that affects her flow payoff from the relationship for the next period, before she makes the bonus decision, that is, the expected value of the output equals  $y + \xi z_2$ .
- A observes a random shock  $z_3$  that augments his outside option by  $\xi z_3$ , before the quitting decision.

---

<sup>1</sup>Later, we consider the possibility that the agent observes a different private signal that is correlated with  $y$ .

The shocks  $z_i$  are independently distributed and each has an atomless distribution. We denote the repeated perturbed game by  $\Gamma^\infty(\xi)$ .

An equilibrium  $\sigma$  of the repeated game  $\Gamma^\infty$  is *purifiable* if for any sequence  $\xi \rightarrow 0$ , there exists a sequence of equilibria  $\sigma(\xi)$  of  $\Gamma^\infty(\xi)$  such that the associated behavior converges to  $\sigma$ .

**Proposition 1** *Let  $\sigma$  be a purifiable finite memory equilibrium of the unperturbed game. In any such equilibrium, the worker always shirks when hired; the principal never pays a bonus, and the agent always quits and the principal terminates the relation*

**Proof.** See appendix. ■

The idea of the proof is similar to that set out in Bhaskar, Mailath, and Morris (2013). It does not follow immediately from that proof for three reasons. First, in the present game, the principal has a continuous action space (of bonus payments), and the perturbations are lower dimensional. Second, the termination decisions are taken simultaneously, whereas the earlier result relies on sequential moves.

In particular, this proposition implies that the equilibria considered in work of MacLeod (2003), Levin (2003) and Fuchs (2007) are not purifiable, being finite memory equilibria.

## 2 The model with cheap talk

We modify the stage game  $\Gamma$  set out in the previous section, by allowing principal and agent to make simultaneous cheap talk announcements. Specifically, after the agent has chosen effort, and the principal has observed output, there is a cheap talk stage. In the cheap talk stage, the equilibrium requires that principal reports the signal she observed, i.e. either  $g$  or  $b$ , while the agent reports his action choice, i.e. he reports  $e$  or  $s$ . If the principal reports  $g$ , then the equilibrium requires that he pay a bonus  $B$ .

Our equilibrium requires the agent to choose both actions with positive probability, so that he chooses  $E$  with probability  $\pi$ . Equilibrium requires “truth-telling” at the cheap talk stage. If the reports “coincide”, i.e. are either  $(e, g)$  or  $(s, b)$ , then the relationship continues to the next period. If the reports differ, then the relationship is terminated with positive probability, depending upon the realization of a public randomization device. Termination occurs with probability  $q$  if the reports are  $(e, b)$  and with probability  $\kappa q$  if the reports are  $(s, g)$ .

We choose the bonus  $B$  so that it exactly incentivizes effort, i.e. it must satisfy:

$$c = ((1 - \eta) - \epsilon) B. \quad (1)$$

Given that the agent is compensated for effort via the bonus, he is indifferent between exerting effort or not when the probability of termination is equal after both actions, i.e.  $\kappa$  satisfies:

$$\kappa = \frac{\epsilon}{\eta}. \quad (2)$$

Given the expected bonus  $B$ , and the above value of  $\kappa$ , the agent is indifferent between  $E$  and  $S$  and thus willing to randomize. We turn to truth-telling incentives at the cheap talk stage.

If the agent has chosen  $E$ , then his payoff loss from announcing  $e$  arises when the principal sees a bad signal, and equals  $\epsilon qV^A$ . His payoff loss from announcing  $s$  is  $(1 - \epsilon)\kappa qV^A$ . Thus truth-telling after choosing  $E$  is optimal if

$$\kappa = \frac{\epsilon}{\eta} \geq \frac{\epsilon}{1 - \epsilon}, \quad (3)$$

which is satisfied since  $\eta > 1 - \epsilon$ . On the other hand, if the agent has chosen  $S$ , then his payoff loss from announcing  $s$  is  $\eta\kappa qV^A$ , and from announcing  $e$  is  $(1 - \eta)qV^A$ , and thus truth telling at  $S$  is optimal if

$$\kappa = \frac{\epsilon}{\eta} \leq \frac{1 - \eta}{\eta}, \quad (4)$$

which also follows from the same inequality  $1 - \eta > \epsilon$ .

We turn now to truth telling incentives for the principal. If the principal reports  $b$  at  $B$ , her payoff loss is

$$L(b|B) = \frac{\pi\epsilon}{\pi\epsilon + (1 - \pi)(1 - \eta)} qV^P. \quad (5)$$

Her expected loss from announcing  $g$  is

$$L(g|B) = (1 - \delta)B + \frac{(1 - \pi)(1 - \eta)}{\pi\epsilon + (1 - \pi)(1 - \eta)} \kappa qV^P. \quad (6)$$

Since the loss from the announcing  $b$  should be smaller, we get the condition:

$$\pi\epsilon[qV^P - (1 - \delta)B] \leq (1 - \pi)(1 - \eta)[\kappa qV^P + (1 - \delta)B]. \quad (7)$$

Thus the incentive constraint for the principal at  $B$  reduces to

$$\frac{\pi}{1-\pi} \leq \left( \frac{1-\eta}{\epsilon} \right) \frac{\kappa q V^P + (1-\delta)B}{q V^P - (1-\delta)B} \quad (8)$$

Second, the principal should be willing to announce  $g$  when she sees  $G$ . Her loss from announcing  $g$  at  $G$  is

$$L(g|G) = (1-\delta)B + \frac{(1-\pi)\eta}{\pi(1-\epsilon) + (1-\pi)\eta} \kappa q V^P. \quad (9)$$

Her loss from reporting  $b$  at  $G$  equals

$$L(b|G) = \frac{\pi(1-\epsilon)}{\pi(1-\epsilon) + (1-\pi)\eta} q V^P. \quad (10)$$

Thus the incentive constraint for the principal at  $G$  is:

$$\frac{\pi}{1-\pi} \geq \left( \frac{\eta}{1-\epsilon} \right) \frac{\kappa q V^P + (1-\delta)B}{q V^P - (1-\delta)B}. \quad (11)$$

An equilibrium consists of a triple  $(\pi, q, V^P)$  – a mixing probability for the agent, a termination probability after  $(e, s)$  and the principal's value – such that the principal's truth-telling constraints 11 and 28 are satisfied, and the  $V^P$  is indeed the value generated by the equilibrium:

$$V^P(q, \pi) = (1-\delta) (\pi(\bar{y} - c + \ell) - w - \ell) + \delta(1-\epsilon q) V^P. \quad (12)$$

Observe from that we must have

$$q V^P \geq (1-\delta) \geq \frac{(1-\delta)c}{1-\eta-\epsilon},$$

since the denominator the right-hand side must be positive. Furthermore, since  $V^P$  is increasing in  $\pi$ , and 2 must be satisfied, we have the following upper bound  $\bar{V}^P$  on the principal's value:

$$\bar{V}^P = (\bar{y} - c - w) - \delta \frac{\epsilon c}{1-\eta-\epsilon}. \quad (13)$$

Let the corresponding value of  $q$  be  $\bar{q}$ . Thus  $\bar{q} \bar{V}^P = (1-\delta)B$ .

**Proposition 2** *Let  $\Delta > 0$ . There exists a purifiable equilibrium with agent mixing probability  $\pi < 1$  and termination probability  $q > \bar{q}$  which generates a principal value  $V^P(\pi, q)$  such*

that  $|\bar{V}^P - V(\pi, q)| < \Delta$ .

**Proof.** Since  $V^P(q, \pi)$ , as defined by equation 12, is a differentiable function, it is straightforward to verify that it is increasing in  $\pi$  and decreasing in  $q$  (as long as  $V^P$  is positive). Furthermore,  $qV^P(q, \pi)$  is increasing in  $q$ , and thus at any  $(q, \pi)$  where the incentive constraint 2 binds, it is possible to increase  $q$  and make it hold strictly.

Consider  $q = \bar{q}, \pi = 1$ . Consider a small increase in  $q$  above  $\bar{q}$ . The right-hand sides of 11 and 28 are now finite and different from each other, so that it is possible to find a value of  $\pi < 1$  that satisfies these two incentive constraints. By choosing  $q$  sufficiently close to  $\bar{q}$ , the difference  $qV^P(1, q) - (1 - \delta)B$  can be made sufficiently small, so that  $\pi$  can be made as close to one as required, since the right-hand side of 28 converges to infinity as  $qV^P(1, q) - (1 - \delta)B \downarrow 0$ . Since  $V^P(\cdot)$  is continuous, this suffices to prove the theorem. ■

### 3 T-period equilibrium with cheap talk?

We now examine whether the basic construction, for the one-period case, can be extended to  $T$  periods, in the usual block manner.

Recall the one-period construction. Consider the non-purifiable pure strategy equilibrium, where the bonus compensates the worker for effort, and where the boss is made indifferent between paying the bonus and not, by choosing the termination probability  $q$ . We showed that with cheap talk, this equilibrium can be approximated by an equilibrium with the following features:

- The worker mixes between  $E$  and  $S$ .
- The bonus compensates the worker for the cost of effort.
- Both players have strict incentives at the cheap talk stage.

In the  $T$ -period construction, we still have a non-purifiable equilibrium with a block structure. However, the following problem arises. The critical incentive constraint for the worker arises in the first period of the block. Thus, if this incentive constraint is satisfied, then the worker strictly prefers to work in every subsequent period. However, this means that the worker is not indifferent between “always”  $E$  and “always”  $S$ , and strictly prefers the former.

Conversely, this means that if we have an equilibrium where the worker is indifferent between “always”  $E$  and “always”  $S$ , then he strictly prefers to choose  $E$  in the first period, and  $S$  in subsequent periods, to either of these alternatives.

We now examine in more detail the properties of the equilibrium. The proposed equilibrium is as follows. At the end of the block of  $T$  periods, there is a cheap talk stage. The bonus only depends upon the principal’s report. If the principal sees at least one  $G$ , he reports  $g$ ; otherwise he reports  $b$ . If he reports  $g$ , he pays a bonus  $B$  that incentivizes the agent to exert effort. Termination probabilities are as before; i.e. with probability  $q$  if the reports are  $(e, b)$  and with probability  $\kappa q$  if the reports are  $(s, g)$ .

For accounting purposes, we think of the bonus being paid at the beginning of the next block. We choose the expected bonus  $B$  so that it exactly incentivizes effort, i.e. it must satisfy:

$$c(1 - \delta^T) = \delta^T[(1 - \epsilon^T) - (1 - (1 - \eta)^T)]B = \delta^T[(1 - \eta)^T - \epsilon^T]B. \quad (14)$$

Given that the agent is compensated for effort via the bonus, he is indifferent between exerting effort or not when  $\kappa$  satisfies:

$$\kappa = \frac{\epsilon^T}{1 - (1 - \eta)^T}. \quad (15)$$

Suppose now that the cheap talk stage. If the principal reports  $b$  at all  $b$ , her payoff loss is

$$L(b|b) = \frac{\pi \epsilon^T}{\pi \epsilon^T + (1 - \pi)(1 - \eta)^T} q V^P. \quad (16)$$

Her expected loss from announcing  $g$  is

$$L(g|b) = (1 - \delta)B + \frac{(1 - \pi)(1 - \eta)^T}{\pi \epsilon^T + (1 - \pi)(1 - \eta)^T} \kappa q V^P. \quad (17)$$

Since the loss from the announcing  $b$  should be smaller, we get the condition:

$$\pi \epsilon^T [q V^P - (1 - \delta)B] \leq (1 - \pi)(1 - \eta)^T [\kappa q V^P + (1 - \delta)B]. \quad (18)$$

Thus the incentive constraint for the principal at “all  $b$ ” reduces to

$$\frac{\pi}{1 - \pi} \leq \left( \frac{1 - \eta}{\epsilon} \right)^T \frac{\kappa q V^P + (1 - \delta)B}{q V^P - (1 - \delta)B} \quad (19)$$

A similar inequality for the incentive constraint when the principal sees exactly one  $G$ ,

that she should be willing to report  $g$ . We do not set this out here, but it requires as a necessary condition:

$$(1 - \delta)B = \frac{(1 - \delta)c}{\delta^T[(1 - \eta)^T - \epsilon^T]} \leq qV^P. \quad (20)$$

$$V^P = (\pi(y + \ell) - w - \ell) - \frac{\delta^T \epsilon^T q V^P}{1 - \delta^T}. \quad (21)$$

$$V^P \leq ((y + \ell - w) - \frac{\epsilon^T c}{(1 + \delta + \dots + \delta^{T-1})[(1 - \eta)^T - \epsilon^T]}). \quad (22)$$

From this, we can conclude that there is no problem with providing incentives for the principal, and further, the efficiency loss due to this vanishes as  $T$  becomes large.

So let us turn to the agent's side, where things appear more difficult.

Recall that the agent's value function is given by

$$\begin{aligned} V^A &= w(1 - \delta^T) + \delta^T(1 - \epsilon^T q)V^A \\ &= w - \frac{\epsilon^T \delta^T q V^A}{1 - \delta^T}. \end{aligned} \quad (23)$$

The critical incentive constraint is

$$(1 - \delta)c \leq (1 - \eta - \epsilon)\epsilon^{T-1}\delta^T\{(1 - \delta)B + qV^A\}. \quad (24)$$

Using the expression for  $B$  in 20, this reduces to

$$(1 - \delta)c \leq (1 - \eta - \epsilon)\epsilon^{T-1}\left[\frac{(1 - \delta)c}{(1 - \delta^T)(1 - \eta)^T - \epsilon^T} + qV^A\right]. \quad (25)$$

We cannot conclude that this incentive constraint holds, especially when  $V^A$  is small. To see this, consider the case where the  $V^A$  is close to zero. Then the bonus must prevent a deviation where the agent shirks only in the first period. However, we know that if the bonus is such that the agent is indifferent between shirking and working in the first period, then he strictly prefers to work in every subsequent period. That is, he strictly prefers effort in every period to shirking in every period. Thus, when he is indifferent between always effort and always shirking, then he strictly prefers to shirk in the first period, and work in the remaining periods, to either of these options.

Note that we have not established that there is no efficient equilibrium with cheap talk, only that it appears to be difficult to construct one.

## 4 Introducing a mediator

We now allow for a mediator, who recommends actions to the agent, and collects reports from the principal.<sup>2</sup> Specifically:

- At the beginning of the  $T$ -period block, the mediator recommends that the agent plays  $E$  or  $S$ , where the agent is required to play the same action for the entire block.
- At the end of the block, the principal reports whether every signal was bad, i.e. report  $b$  or at least one signal was good, report  $g$ .
- If the principal reports  $g$ , and the agent is recommended  $E$ , the principal pays a bonus  $B$  to the agent; otherwise, he does not.
- The relationship is terminated with some probability after:  $(E, b)$ ,  $(S, g)$ , and continues for sure after  $(E, g)$  or  $(S, b)$ .

The critical difference as compared to cheap talk is that the agent does not have to be indifferent between always playing  $E$  and always playing  $S$ . Indeed, the efficient equilibrium has his incentive constraint after the recommendation to play  $E$  holding with equality, so that he is indifferent between working and shirking in the first period, but strictly prefers to work thereafter. Consequently, it should be possible to approximate full efficiency as  $\delta \rightarrow 1$ , by choosing  $T$  sufficiently large.

Let us first verify the agent's incentives. If the mediator recommends  $S$ , then choosing effort in any period does not result in any bonus payment, and only increases the probability of termination, by increasing the probability that at least one  $G$  is observed. To verify incentives after the recommendation  $E$ , let us set the bonus so that it exactly compensates the agent for his first period effort, given that he has been recommended to play  $E$ . That is the discounted expected loss of bonus equals the cost of the effort in the first period of the block.

$$B = \frac{c}{(1 - \eta - \epsilon)\epsilon^{T-1}\delta^T}. \quad (26)$$

---

<sup>2</sup>See Rahman (2012) for previous examination of the role of a mediator.

Since the agent also loses  $qV^A > 0$  after  $(e, b)$ , this implies that his first period incentive constraint holds strictly, for any  $V^A > 0$ . Furthermore, by the standard Abreu, Milgrom, and Pearce (1991) argument, if the agent intends to announce  $e$ , then he does not have an incentive to shirk in any other period of the block, independent of his own actions.

We also choose  $\kappa$  as before, so that the probability of termination is the same after either recommended action of the mediator – this is no longer necessary, but simplifies calculations.

For the principal, we have the necessary condition:

$$qV^P \geq (1 - \delta)B = \frac{(1 - \delta)c}{(1 - \eta - \epsilon)\epsilon^{T-1}\delta^T}, \quad (27)$$

which is satisfied, for any given  $T$ , if  $\delta$  is large enough. The remaining condition is verifying truth-telling for the principal. Let  $\pi$  denote the probability that the mediator recommends  $E$  to the agent.

The incentive constraint for the principal at “all  $B$ ” reduces to

$$\frac{\pi}{1 - \pi} \leq \left( \frac{1 - \eta}{\epsilon} \right)^T \frac{\kappa q V^P}{q V^P - (1 - \delta)B}. \quad (28)$$

The incentive constraint for the principal when she observes a single  $G$ , that she is willing to announce  $g$ , is

$$\frac{\pi}{1 - \pi} \geq \left( \frac{(1 - \eta)}{\epsilon} \right)^{T-1} \left( \frac{\eta}{1 - \epsilon} \right) \frac{\kappa q V^P}{q V^P - (1 - \delta)B} \quad (29)$$

Since the right-hand side of 28 is strictly greater than that of 29, it is possible to find  $\pi$  such that both incentive constraints are satisfied. Further, by taking  $q$  sufficiently small, we can ensure that the right hand side of 28 can be made arbitrarily large, and so  $\pi$  can be chosen close to one.

The principal’s payoff, when  $q$  is chosen to be minimal, is given by:

$$V^P \leq (\pi(y + \ell - c) - w - \ell) - \frac{\epsilon c}{(1 - \eta - \epsilon)(1 + \delta + \dots + \delta^{T-1})}. \quad (30)$$

Since  $\pi$  can be made arbitrarily close to one for any  $T$ , and since  $T$  can be chosen to be arbitrarily large when  $\delta$  is large enough, we can approximate the efficient payoff as  $\delta \rightarrow 1$ .

We therefore have the following proposition:

**Proposition 3** *In the repeated game with a mediator, there exist purifiable equilibria that can approximate the fully efficient payoff provided that  $\delta$  is sufficiently close to one.*

## 5 Concluding Comments

We have claimed, without proof, that the equilibria we have constructed are purifiable. This remains to be done, but we do not anticipate any problem.

We have focused on extending the basic interaction between principal and agent, either by allowing for cheap talk or a mediator. Alternatively, if we allow for the agent's effort to have persistent effects on output, then one can construct purifiable equilibria.

## References

- ABREU, D., P. MILGROM, AND D. PEARCE (1991): "Information and Timing in Repeated Partnerships," *Econometrica*, 59(6), 1713–1733.
- BHASKAR, V., G. J. MAILATH, AND S. MORRIS (2013): "A Foundation for Markov Equilibria in Sequential Games with Finite Social Memory," *Review of Economic Studies*, 80(3), 925–948.
- FUCHS, W. (2007): "Contracting with repeated moral hazard and private evaluations," *American Economic Review*, 97(4), 1432–1448.
- LEVIN, J. (2003): "Relational Incentive Contracts," *American Economic Review*, 93(3), 835–857.
- MACLEOD, W. B. (2003): "Optimal contracting with subjective evaluation," *American Economic Review*, 93(1), 216–240.
- RAHMAN, D. (2012): "But who will monitor the monitor?," *American Economic Review*, 102(6), 2767–97.