

Renegotiation of long-term contracts as part of an implicit agreement

RUMEN KOSTADINOV

Department of Economics, McMaster University

I study a repeated principal–agent game with long-term output contracts that can be renegotiated at will. Actions are observable but not contractible, so they can only be incentivized through implicit agreements formed in equilibrium. I show that contract renegotiation is a powerful tool for incentive provision, despite the stationarity of the environment. Continuation contracts are designed to punish deviations in noncontractible behavior. If the equilibrium actions are observed, these contracts are renegotiated away. This form of anticipated renegotiation results in welfare improvements over outcomes attainable by one-period contracts or by long-term contracts that are not renegotiated. When the principal is not protected by limited liability, first-best outcomes are attainable regardless of the impatience of the players. Equilibrium strategies are shown to satisfy various concepts of renegotiation-proofness.

KEYWORDS. Long-term contracts, renegotiation, relational contracts.

JEL CLASSIFICATION. C73, C78, D86.

1. INTRODUCTION

Long-term relationships offer two channels for incentive provision: formal and informal contracting. Formal contracts are contingent on outcomes that can be enforced in a court of law. However, many aspects of behavior can be observable within the relationship, but difficult to verify by an outside party such as the court. They can be supported by informal contracting: implicit agreements sustained in equilibrium that reward and punish current actions through future behavior.

Previous work on the interplay between formal and informal contracting has restricted attention to formal contracts that either last for a single period or are stationary and cannot be renegotiated.¹ This paper studies long-term formal contracts that can be renegotiated with the consent of all parties. In this setting the renegotiation of a formal contract becomes part of informal contracting, and it has a powerful impact

Rumen Kostadinov: kostadir@mcmaster.ca

I thank David Pearce, Ennio Stacchetti, and Debraj Ray for their invaluable continued support. The three anonymous referees helped me significantly improve the paper. I also received valuable comments from Dilip Abreu, Sylvain Chassang, Marina Halac, Sam Kapon, Rohit Lamba, Alessandro Lizzeri, Erik Madsen, Trond Olsen, Francisco Roldán, Ariel Rubinstein, Alexis Akira Toda, Nikhil Vellodi, Joel Watson, and John Zhu, as well as seminar participants at Columbia, Iowa State, McGill, McMaster, NYU, Rice, Toronto, UCSD, and Western.

¹A notable exception is the simultaneous work of Watson et al. (2020) discussed extensively below.

on incentives. In efficient equilibria, the parties anticipate rewriting the future terms of their formal contract contingent on behavior observed in the interim. Such premeditated renegotiation creates a welfare improvement over outcomes attainable without renegotiation or long-term contracts.

These ideas are developed in a stylized repeated principal–agent model. A risk-neutral principal (she) hires a risk-averse agent (he) with unbounded utility. The parties can write long-term formal contracts contingent on the entire history of output. At the beginning of each period, the principal offers such a contract. The agent can accept, reject in favor of the existing contract, or take an outside option. Effort is observable but noncontractible. After observing effort and output, the principal pays the agent the salary specified in their current contract as well as a discretionary bonus. The model builds on [Pearce and Stacchetti \(1998\)](#), where output contracts specify wages for the current period only, leaving no scope for renegotiation.

The potential for future renegotiation combined with the unrestricted length and history dependence of the contracts make it challenging to identify equilibrium behavior. Signing a contract commits the principal to the salaries specified for the initial period, but the remaining terms can be discarded later on if the agent accepts a new offer. Despite this, the future terms of the contract are not rendered irrelevant by renegotiation. To the contrary, they *inform* renegotiation, as the agent can maintain them by refusing to renegotiate. The difficulty is that the agent's utility following no renegotiation is determined endogenously in equilibrium, and since contracts are incomplete, multiple equilibria are possible. This challenge is amplified by the potential complexity of the contracts the principal can offer and the contract currently in place.

The main result of this paper, Theorem 1, shows the attainability of a range of first-best outcomes for any parameters of the model, even when the players are arbitrarily impatient. On the equilibrium path, the players sign a two-period contract with heavily unbalanced compensation for the second period. These inefficient terms create punishments for (off-path) deviations. On the equilibrium path, however, the inefficient terms are never implemented, as the parties renegotiate them to the same two-period contract.

The crucial part of the construction lies in the punishment equilibria facilitated by the future terms. It may be intuitive to consider punishment strategies that maintain these inefficient terms following any deviation. However, the resulting payoffs for both players would be so low that the principal could offer a balanced contract that guarantees mutual improvement. Instead, the punishments exploit the multiplicity of equilibria in the subgame where the risky terms are not renegotiated. In one equilibrium, the agent receives no bonuses and his utility is low. In another, he receives a bonus when his contractual salary is low. This bonus has arbitrarily high marginal utility when the salaries are sufficiently unbalanced. Thus, the rejection of a future offer can be seen as an endogenous outside option for the agent, calibrated in equilibrium to punish either party for their past actions. If the agent deviates, the payoff he can secure by refusing to renegotiate is low. If the principal deviates, she receives a low payoff because the agent's utility following no renegotiation is high.

Signing contracts with the expectation of their future renegotiation is an essential feature of first-best equilibria. Proposition 2 shows that when the players are impatient, first-best outcomes with costly effort cannot be attained by one-period contracts as in [Pearce and Stacchetti \(1998\)](#) or without contractual rewriting. Thus, on-path renegotiation is strictly welfare-improving. This is distinct from the usual view of renegotiation as a response to changes in the underlying environment. Here, the production technology is static and there is no asymmetric information, so renegotiation occurs solely for the purpose of incentive provision.

While the first-best equilibria I construct use contracts with unfavorable future terms, they do not rely on inefficient punishment strategies that the players would find profitable to renegotiate. Proposition 3 shows that under the assumption of vanishing marginal utility, efficient equilibria can satisfy two notions of renegotiation-proofness: strong optimality ([Levin \(2003\)](#)) and contestable norms ([Safronov and Strulovici \(2018\)](#)).

I also explore an extension of the model where the principal can shut down the firm. Unlike in the baseline model, the first best is generally not attainable, since the worst punishments for both players are bounded below by their outside options. However, there exists a continuation contract that can hold either player down to their respective outside payoff in equilibrium. This makes the model equivalent to a one-period contracting problem with exogenously given worst punishments, akin to standard models in the relational contracting literature. Theorem 2 leverages this finding to provide a recursive characterization of efficient equilibrium payoffs.

To the best of my knowledge, the only other paper to consider long-term contract renegotiation in the presence of informal contracting is the independent and simultaneous work of [Watson et al. \(2020\)](#) (henceforth WMO). WMO's solution concept is contractual equilibrium ([Watson \(2013\)](#), [Miller and Watson \(2013\)](#)), which leverages risk neutrality to develop a cooperative approach to bargaining over contracts, transfers, and continuation strategies. In contrast, risk aversion is central to my analysis. My approach to bargaining also differs: formal contracts are renegotiated noncooperatively, and the renegotiation of strategies relies on equilibrium refinements distinct from contractual equilibrium.

The comparison to WMO is discussed in more detail in Sections 3.7 and 4.4. WMO's main result is that without loss of generality, formal contracts signed in equilibrium are semistationary: their terms for all non-initial periods are identical. Despite the modelling differences, this is also true in my setting as well as in the extended model of Section 4. In addition, I examine an equivalent definition of contractual equilibrium by [Miller and Watson \(2013\)](#) based on axiomatic equilibrium refinements. I show that the principal-optimal efficient equilibrium payoffs are attainable in an equilibrium that satisfies these refinements.

I proceed by presenting the baseline model in Section 2 and analyzing it in Section 3. The results of the extended model where the principal has limited liability are reported in Section 4. This is followed by a literature review and a short conclusion.

2. THE MODEL

2.1 Preferences and timing

A principal (she) and an agent (he) interact repeatedly at times $t = 1, 2, \dots$. They can sign long-term contracts on output, which takes values in $Y = \{h, l\}$ every period with $h > l \geq 0$.² A contract is a sequence $c = (c^t)_{t=1}^\infty$ of functions $c^t : Y^t \rightarrow [0, \bar{s}]$, where $\bar{s}$ is an exogenous upper bound needed to rule out Ponzi schemes (see Appendix A.1). While contracts are infinitely long, I interpret some of them as short-term contracts as follows: a contract c is a t -period contract if $c^\tau(y^\tau) = 0$ for all $\tau > t$ and $y^\tau \in Y^\tau$.³ Let C denote the space of all contracts.

NOTATION 1. For any contract $c \in C$, let $s_y := c^1(y)$ denote the contemporaneous salary for output y and let c_y denote the continuation contract following output y , i.e.,

$$c_y^t(y_1, \dots, y_t) = c^{t+1}(y, y_1, \dots, y_t) \quad \text{for all } t \in \mathbb{N}, (y_1, \dots, y_t) \in Y^t.$$

At the beginning of each period, the parties inherit a *residual contract* c^* that comprises the remaining terms of the last contract they signed. Thus, if they sign a contract c and output y is realized, the residual contract at the beginning of the next period is c_y . In the initial period the residual contract is the null contract $\mathbf{0}$ given by $\mathbf{0}^t(y^t) = 0$ for all $t \in \mathbb{N}$, $y^t \in Y^t$.

Each period the principal and the agent play a sequential game of perfect information. It begins with a *contract offer* c by the principal, which represents a proposal to renegotiate the residual contract c^* currently in place to c . The principal has an implicit option to refuse renegotiation by setting $c = c^*$. The agent's *contract response* is to accept c (A), reject c in favor of c^* (R) or take an outside option (O).⁴ The outside option terminates the relationship: the agent works for an outside employer at a constant wage $r \geq 0$ and the principal receives 0 in all subsequent periods.⁵

If the agent does not take the outside option, the game continues with his choice of effort $e \in [0, 1]$. His cost of effort is given by a convex, strictly increasing, and differentiable function $\psi : [0, 1] \rightarrow \mathbb{R}_+$ with $\psi(0) = 0$. Effort e generates output y with probability $p_y^e = ep_y^1 + (1-e)p_y^0$, where $0 < p_h^0 < p_h^1 < 1$, $0 < p_l^1 < p_l^0 < 1$, and $p_h^1 + p_l^1 = p_h^0 + p_l^0 = 1$. The agent is paid a wage from the contract currently in effect, amounting to s_y if he accepted the offer c and to s_y^* if he rejected. Finally, having observed the contract response, effort, and output, the principal makes a voluntary bonus payment $b \geq 0$ to the agent. Figure 1 summarizes this timing and shows the continuation contract following each stage-game history. This continuation contract becomes the residual contract in the next period.

²When output is not binary, the results remain unchanged as long as each effort level has full support over the realizations of output.

³It is possible to model short-term contracts explicitly by allowing them not to specify salaries past a certain period. This change in interpretation has no effect on the results.

⁴Alternatively, it is possible to treat the agent's outside option as a clause present in all contracts that allows him to terminate the relationship at any time during the period before output is realized. This has no effect on the results.

⁵The results remain unchanged if the players received their outside options for one period and continued interacting in the next period with residual contract $\mathbf{0}$.

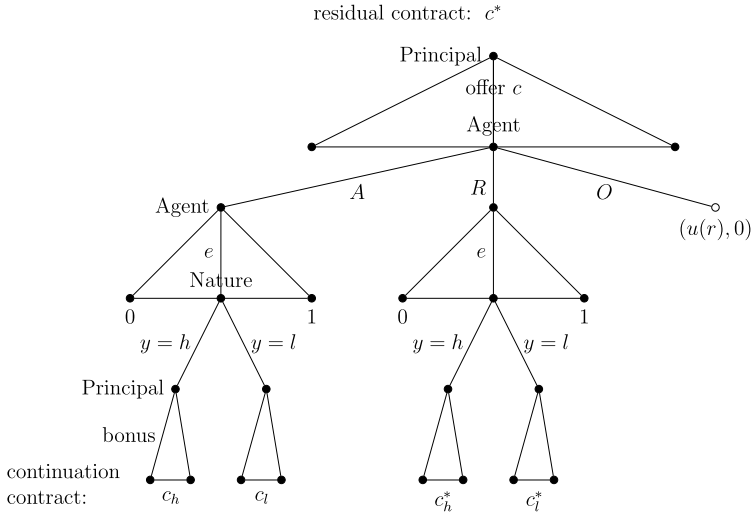


FIGURE 1. The stage game.

The per-period expected payoffs are given by

$$\text{agent: } \sum_y p_y^e u(s_y + b_y) - \psi(e)$$

$$\text{principal: } \sum_y p_y^e (y - s_y - b_y),$$

where e is the agent's effort, s_h and s_l are the salaries from the contract in effect after the negotiations, and b_h and b_l are the bonuses following high and low output. The agent's utility function $u : \mathbb{R}_+ \rightarrow \mathbb{R} \cup \{-\infty\}$ is assumed strictly concave, strictly increasing, and differentiable with $u(0) = -\infty$ and $\lim_{s \rightarrow \infty} u(s) = \infty$. Let u^{-1} denote the inverse of u .

The (expected) payoffs in the game are denoted as a pair (x, v) , which correspond to the payoffs of the agent and the principal, respectively. They are given by

$$x = (1 - \delta) \sum_{t=1}^{\infty} \delta^t x_t \quad \text{and} \quad v = (1 - \delta) \sum_{t=1}^{\infty} \delta^t v_t,$$

where (x_t, v_t) are the payoffs in period t and $\delta \in (0, 1)$ is a common discount factor.⁶

Let Θ be the space of parameters $\theta = (Y, p_h^0, p_l^0, p_h^1, p_l^1, r, u, \psi)$ that satisfy the assumptions of the model. In the results of the paper, θ is held fixed, while the remaining parameters δ and $\bar{s}$ may vary.

2.2 Strategies and equilibrium

Let $\Omega = C \times \{A, R\} \times [0, 1] \times Y \times \mathbb{R}_+$ be the space of all one-period histories where no separation has occurred, and let $\Omega_0 = \Omega \cup (C \times \{O\})$ be the space of all one-period histories.

⁶Equal patience is assumed for expositional purposes: it is not needed for any of the results.

A strategy for the agent is a pair of maps

$$\sigma_{A1} : \bigcup_{t=1}^{\infty} \Omega^{t-1} \times C \rightarrow \{A, R, O\} \quad \text{and} \quad \sigma_{A2} : \bigcup_{t=1}^{\infty} \Omega^{t-1} \times C \times \{A, R\} \rightarrow [0, 1],$$

which specify the agent's contract response and effort in each period t for each possible history. A strategy for the principal is a pair of maps

$$\sigma_{P1} : \bigcup_{t=1}^{\infty} \Omega^{t-1} \rightarrow C \quad \text{and} \quad \sigma_{P2} : \bigcup_{t=1}^{\infty} \Omega^{t-1} \times C \times \{A, R\} \times [0, 1] \rightarrow \mathbb{R}_+.$$

I consider the pure-strategy subgame-perfect equilibria of this game.

2.3 Recursive characterization

It is convenient to develop a recursive characterization of equilibrium payoffs by adapting the dynamic programming methods of [Abreu et al. \(1990\)](#), henceforth APS. One challenge is that long-term contracts break the stationarity of the game: subgames from the beginning of a period with different residual contracts may admit different equilibria. However, if the residual contracts are the same, the subgames are identical, because the history of play does not contain other payoff-relevant information. Hence, the model represents a stochastic game where the state variable is the residual contract and next period's state variable is the continuation contract, as shown in [Figure 1](#). Similarly, the prevailing contract in any subgame from the agent's effort determines the equilibrium payoff set in that subgame, as reflected in the following notation.

NOTATION 2. *The subgame from the beginning of a period with residual contract c^* is denoted subgame c^* . Hence, subgame $\mathbf{0}$ is the subgame from the initial period.*

The subgame starting from the agent's effort choice when contract c is in effect (either through acceptance of an offer c or rejection in favor of a residual contract equal to c) is denoted subgame (c, E) .

Let $\mathcal{P}_0$ be the space of bounded correspondences $\Pi : C \rightrightarrows \mathbb{R}^2$. Let $\mathcal{E} \in \mathcal{P}_0$ be the correspondence describing the equilibrium payoff sets of all subgames from the start of a period. The recursive representation decomposes payoffs in these subgames into strategies for the initial period and continuation payoffs. This is accomplished through a class of auxiliary games with structure identical to the stage game in [Figure 1](#). Each auxiliary game is indexed by a residual contract c^* and a function $f : \Omega_0 \rightarrow \mathbb{R}^2$ that assigns continuation payoffs to each history ω such that $f(\omega) = (u(r), 0)$ for every $\omega \in \Omega_0 \setminus \Omega$. The payoffs of such a game (c^*, f) at any history ω are given by

$$(1 - \delta)(\text{first-period payoffs in subgame } c^* \text{ following history } \omega) + \delta f(\omega).$$

DEFINITION 1 (Admissibility). A profile (σ, f) of strategies σ in the stage game and a continuation payoff function $f : \Omega_0 \rightarrow \mathbb{R}^2$ is admissible with respect to $\Pi \in \mathcal{P}_0$ at c^* with value (x, v) if σ is an equilibrium of (c^*, f) with payoffs (x, v) and $f(\omega) \in \Pi(c')$ for any $\omega \in \Omega$ such that the continuation contract in subgame c^* following ω is c' .

Let (c^*, Π) denote the class of auxiliary games with residual contract c^* and continuation payoffs drawn from a payoff correspondence $\Pi \in \mathcal{P}_0$. An equilibrium of (c^*, Π) is defined as a strategy profile σ such that (σ, f) is admissible with respect to Π at c^* for some continuation payoff function f . In particular, an equilibrium of $(c^*, \mathcal{E})$ is an equilibrium of subgame c^* . Let $B\Pi(c^*)$ denote the set of values of profiles admissible with respect to Π at c^* , i.e., the equilibrium payoff set of (c^*, Π) .

DEFINITION 2. A correspondence $\Pi \in \mathcal{P}_0$ is called self-generating if $\Pi(c^*) \subseteq B\Pi(c^*)$ for all $c^* \in C$.

Any self-generating correspondence Π can be used to construct equilibrium payoffs as follows. Any payoff $(x, v) \in \Pi(\mathbf{0})$ is also in $B\Pi(\mathbf{0})$ so it can be decomposed into strategies for the initial period and continuation payoffs drawn from Π . But then the continuation payoffs following every first-period history are in $B\Pi$, so they can be decomposed into strategies for the second period and continuation payoffs. The equilibrium strategies can be constructed inductively in this manner, and incentive compatibility in each period is guaranteed by the admissibility condition. This is the self-generation result of APS stated below.⁷

PROPOSITION 1 (APS). *If Π is self-generating, then $B\Pi(c^*) \subseteq \mathcal{E}(c^*)$ for all $c^* \in C$.*

Proposition 1 is used throughout Section 3 to establish the main results of this paper. Unlike many applications of APS, the stage game played each period is sequential, which complicates working with the operator B . It is thus useful to consider an intermediate step in its computation: finding the payoffs in the subgame starting from the effort choice.

Let $\Omega_c = \{c\} \times \{A\} \times [0, 1] \times Y \times \mathbb{R}_+$ be the space of one-period histories where a contract c has been offered and accepted. Let (c, E, f_E) denote the subgame of (c^*, f) following the acceptance of a contract c , where f_E is the restriction of f to Ω_c . Since the residual contract becomes payoff-irrelevant after a new one is accepted, subgame (c, E, f_E) inherits no dependence on c^* .

DEFINITION 3 (E -admissibility). A profile (σ_E, f_E) of strategies σ_E in the stage game following acceptance and continuation payoffs $f_E : \Omega_c \rightarrow \mathbb{R}^2$ is E -admissible with respect to Π at c with value (x, v) if σ_E is an equilibrium of (c, E, f_E) with payoffs (x, v) and $f_E(\omega) \in \Pi(c_y)$ for any $\omega \in \Omega_c$, where the realized output is y .

An equilibrium of (c, E, Π) is defined as a strategy profile σ_E such that (σ_E, f_E) is E -admissible with respect to Π at c for some f_E . Let $B\Pi(c, E)$ denote the set of values of profiles E -admissible with respect to Π at c . The following lemma links the equilibrium payoffs of subgames from the beginning of a stage and from the agent's effort. Omitted proofs of all subsequent results are provided in the [Appendix](#).

⁷See Section 5.7.1 of [Mailath and Samuelson \(2006\)](#).

LEMMA 1. *Let $c^* \in C$, $\Pi \in \mathcal{P}_0$, and $(x^*, v^*) \in B\Pi(c^*)$. If $(x, v) \in B\Pi(c, E)$ for some $c \in C$ and $(x, v) \geq (x^*, v^*)$, then $(x, v) \in B\Pi(c^*)$.*

The proof of Lemma 1 alters the equilibrium of (c^*, Π) with payoffs (x^*, v^*) by having contract c offered and accepted, after which the equilibrium of (c, E, Π) with payoffs (x, v) is played.

2.4 First-best outcomes

For each level of agent utility x , there exists a corresponding first-best payoff for the principal, denoted by $v^{\text{FB}}(x)$. I assume that the agent never takes his outside option in a first-best outcome. Thus,

$$v^{\text{FB}}(x) = \max_e \sum_y p_y^e y - u^{-1}(x + \psi(e)).$$

Let $e^{\text{FB}}(x)$ denote the associated first-best effort, which is unique since u^{-1} and ψ are convex and strictly increasing, and u^{-1} is strictly convex. Let x^{FB} denote the inverse of v^{FB} .

3. RESULTS

The main result of this paper is the existence of first-best equilibria for any parameters of the model, including the patience of the players, provided that the contract space is sufficiently unrestricted.

THEOREM 1. *Fix any parameters $\theta \in \Theta$ and $\delta \in (0, 1)$. For any $\bar{s}$ sufficiently large, there exists $\bar{x} \geq u(r)$ such that any first-best outcome with agent utility $x \in [u(r), \bar{x}]$ is attainable in an equilibrium of subgame 0. Furthermore, $\bar{x} > u(r)$ if $e^{\text{FB}}(u(r)) > 0$.*

To attain the first best, the parties sign a contract c with an output-invariant continuation contract $c_h = c_l = c^*$. Only the first-period terms of c are realized on the equilibrium path, as the parties renegotiate c^* to c in each subsequent period. The contract c^* plays an instrumental role in enforcing the equilibrium actions: any deviation is punished by an adverse equilibrium for the corresponding player in the next-period subgame c^* . The heart of the proof of Theorem 1 lies in showing that subgame c^* admits punishment equilibria with sufficiently low payoffs for each player. Section 3.1 constructs equilibria that attain the agent's worst payoffs. The proof is completed in Section 3.2, where it is shown that subgame c^* admits an arbitrarily harsh punishment equilibrium for the principal whenever the upper bound on salaries $\bar{s}$ is sufficiently high.

3.1 Agent's worst equilibrium payoffs in key subgames

This section characterizes the agent's worst equilibrium payoffs in subgames from the start of a period and from his effort choice. Let $U(c^*)$ be the payoff the agent can guarantee in subgame c^* if he never accepts another contract offer. Thus, $U(c^*)$ is the highest

utility he can obtain by choosing each period whether to take his outside option or to exert effort, expecting the salaries from c^* and no bonuses. If effort is chosen and output y realizes, the agent receives s_y^* and faces the same problem in the subsequent period starting with residual contract c_y^* . Hence, U solves the Bellman equation

$$U(c^*) = \max\{U^E(c^*), u(r)\}$$

$$\text{where } U^E(c^*) = \max_e \sum_y p_y^e [(1 - \delta)(u(s_y^*) - \psi(e)) + \delta U(c_y^*)]. \quad (1)$$

Since salaries are upper bounded by $\bar{s}$ and the agent can guarantee $u(r)$, standard arguments establish that there exists a continuous solution for U , which is unique in the class of bounded functions on C .⁸ This uniquely determines the continuous function U^E .

It follows that $U(c^*)$ and $U^E(c^*)$ are lower bounds on the agent's equilibrium payoff in subgames c^* and (c^*, E) , respectively. Lemma 2 below states that these bounds can be attained in equilibrium.

LEMMA 2. *For any $c^* \in C$, there exist*

- *an equilibrium $\underline{\sigma}^A(c^*)$ of subgame c^* with agent payoff $U(c^*)$*
- *an equilibrium $\underline{\sigma}^A(c^*, E)$ of subgame (c^*, E) with agent payoff $U^E(c^*)$.*

The equilibrium $\underline{\sigma}^A(c^*)$ has payoffs $(U(c^*), V(c^*))$, where

$$V(c^*) = \begin{cases} \max_c V_c(c^*) & \text{if } U(c^*) > u(r) \\ \max\{\max_c V_c(c^*), 0\} & \text{if } U(c^*) = u(r) \end{cases}$$

$$V_c(c^*) = \max_e \sum_y p_y^e [(1 - \delta)(y - s_y) + \delta V(c_y)]$$

$$\text{such that } U^E(c) \geq U(c^*) \quad (2)$$

$$U^E(c) = \sum_y p_y^e [(1 - \delta)(u(s_y) - \psi(e)) + \delta U(c_y)] \quad (3)$$

and $V_c(c^*)$ is set to $-\infty$ whenever (2) does not hold. Note that the payoff functions U and V depend on all parameters of the model $\theta, \delta, \bar{s}$. This dependence is omitted from the notation for simplicity.

The equilibrium $\underline{\sigma}^A(c^*)$ is constructed by showing that the payoff correspondence Π with $\Pi(\hat{c}) = \{(U(\hat{c}), V(\hat{c}))\}$ for all $\hat{c} \in C$ is self-generating. Hence, the continuation equilibrium payoffs from the start of a period are fully determined by the residual contract. This makes it optimal for the principal to pay no bonuses. Since the agent's effort

⁸The contract space C is treated as the countable product $\times_{t=1}^{\infty} C^t$ endowed with the product topology, where C^t is the space of functions $c_t: Y^t \rightarrow [0, \bar{s}]$. This makes the functions $c^* \mapsto s_y^*$ and $c^* \mapsto c_y^*$ continuous for all y , so the Bellman operator preserves continuity.

does not affect his continuation payoff, he obtains $U^E(c^*)$ upon rejecting any contract offer. For the same reason, he obtains $U^E(c)$ if he accepts a contract c , reflected in (3). Given the agent's behavior, the principal has two options in subgame c^* . First, she can obtain $V_c(c^*)$ by offering a contract c that is acceptable to the agent, i.e., it satisfies (2). In this case, the agent receives $U(c^*)$, since (2) binds due to the lack of limited liability, reflected in the assumption $u(0) = -\infty$. Second, if the agent cannot guarantee more than his outside option by rejecting in favor of c^* , the principal can induce the outside option by offering an unattractive contract, e.g., c^* . In this case, the agent also receives his guaranteed payoff $U(c^*)$, which equals $u(r)$.

The equilibrium strategies depend on the residual contract c^* only through the agent's guaranteed payoff $U(c^*)$, so the following corollary is immediate.

COROLLARY 1. *For any $c, c^*, c^{**} \in C$, if $U(c^*) = U(c^{**})$, then $V(c^*) = V(c^{**})$ and $V_c(c^*) = V_c(c^{**})$.*

Consider any subgame c^* and suppose there exists an equilibrium in the further subgame (c^*, E) following rejection, where the agent receives $x^* \geq U(c^*)$. This can be used to modify the strategies $\underline{\sigma}^A(c^*)$ and obtain an equilibrium of subgame c^* with agent payoff x^* . The following lemma generalizes this result to any correspondence containing the payoffs of the equilibrium constructed to prove Lemma 2.

LEMMA 3. *Let $c^* \in C$ and $\Pi \in \mathcal{P}_0$ such that $(U(\hat{c}), V(\hat{c})) \in \Pi(\hat{c})$ for any $\hat{c} \in C$. If there exists an equilibrium of (c^*, E, Π) with agent payoff $x^* \geq U(c^*)$, then there exists an equilibrium of (c^*, Π) with agent payoff x^* .*

3.2 Proof of Theorem 1

Let $\theta \in \Theta$ and $\delta \in (0, 1)$. In what follows, θ and δ are held fixed, while $\bar{s}$ is varied. The statement of the theorem is shown for $\bar{x} = x^{\text{FB}}(V(\mathbf{0}))$.

Case 1: $e^{\text{FB}}(u(r)) = 0$. Consider a contract c with salaries $c^t(y^t) = r$ for any $t \in \mathbb{N}$, $y^t \in Y^t$. Since c exhibits constant salaries, it follows that $U^E(c) = u(r) = U(c)$. Then $V_c(c) = \sum_y p_y^0[(1 - \delta)(y - r) + \delta V(c)]$, since $c_h = c_l = c$ and the unique effort that satisfies (3) is $e = 0$. But $V(c) \geq V_c(c)$, so $V_c(c) \geq \sum_y p_y^0(y - r) = v^{\text{FB}}(u(r))$, since $e^{\text{FB}}(u(r)) = 0$. Moreover, $V(\mathbf{0}) \geq V_c(\mathbf{0}) = V_c(c)$, where the equality follows from Corollary 1 and $U(c) = U(\mathbf{0}) = u(r)$. Thus, $V(\mathbf{0}) = v^{\text{FB}}(u(r))$, so $\bar{x} = u(r)$. By Lemma 2, there exists an equilibrium of subgame $\mathbf{0}$ with payoffs $(U(\mathbf{0}), V(\mathbf{0})) = (u(r), v^{\text{FB}}(u(r)))$, as required.

Case 2: $e^{\text{FB}}(u(r)) > 0$. I begin by showing that $V(\mathbf{0})$ is bounded away from $v^{\text{FB}}(u(r))$ uniformly over all $\bar{s}$. To see this, consider the equilibrium $\underline{\sigma}^A(\mathbf{0})$ with principal payoff $V(\mathbf{0})$. Due to the agent's outside option, the principal's continuation payoff from the second period is bounded above by $v^{\text{FB}}(u(r))$. Thus, if $V(\mathbf{0})$ is arbitrarily close to $v^{\text{FB}}(u(r))$, then the principal's first-period payoff is arbitrarily close to $v^{\text{FB}}(u(r))$ and, consequently, the agent's first-period effort e is arbitrarily close to $e^{\text{FB}}(u(r))$. Then the strict concavity of u implies that the agent's first-period salaries must be arbitrarily close to $u^{-1}(u(r) + \psi(e))$, since $\underline{\sigma}^A(\mathbf{0})$ prescribes no bonuses. It follows that the agent can

deviate to zero effort to obtain approximately $(1 - \delta)u^{-1}(u(r) + \psi(e)) + \delta u(r)$, which exceeds his on-path payoff $u(r)$ when e is sufficiently close to $e^{\text{FB}}(u(r))$. Hence, there exists $\alpha > 0$ such that $v^{\text{FB}}(u(r)) - V(\mathbf{0}) > \alpha$ regardless of the value of $\bar{s}$. In particular, $\bar{x} > u(r)$, as required.

To show the existence of the desired first-best equilibria, it suffices, by Proposition 1, to show that the following correspondence Π is self-generating for all $\bar{s}$ sufficiently large:

$$\begin{aligned}\Pi(\mathbf{0}) &= \{(x, v^{\text{FB}}(x)) | u(r) \leq x \leq x^{\text{FB}}(V(\mathbf{0}))\} \cup \{(u(r), V(\mathbf{0}))\} \\ \Pi(c^*) &= \Pi(\mathbf{0}) \cup \{(x^*, v^p)\} \\ \Pi(\hat{c}) &= \{(U(\hat{c}), V(\hat{c}))\} \quad \text{for any } \hat{c} \in C \setminus \{\mathbf{0}, c^*\}.\end{aligned}$$

The idea of the construction is that in subgame $\mathbf{0}$, the parties sign a contract c with $s_h = s_l = 0$ and $c_h = c_l = c^*$, where c^* has contemporaneous salaries s_h^* and s_l^* given by

$$s_h^* = \bar{s} \quad \text{and} \quad \max_e \sum_y p_y^e u(s_y^*) - \psi(e) = u(r), \quad (4)$$

and continuation contracts $c_h^* = c_l^* = \mathbf{0}$. The unboundedness of the agent's utility implies that there exists $\bar{s}_1$ such that for all $\bar{s} \geq \bar{s}_1$, c^* is well defined with $e = 1$ solving the maximization problem in (4). Since $U(\mathbf{0}) = u(r)$, it follows that $U(c^*) = U^E(c^*) = u(r)$ and, by Corollary 1, $V(c^*) = V(\mathbf{0})$.

After accepting c , the agent exerts the first-best effort, receiving zero salaries from the contract and a bonus independent of the output realization (but contingent on first-best effort). In each subsequent period, the residual contract is c^* , and the parties obtain the same first-best payoffs by renegotiating c^* to c and replicating their play from the initial period. Any deviation from the equilibrium path except for bonus payments is followed by the agent's worst equilibrium constructed in the proof of Lemma 2. If the principal reneges on the bonus, the equilibrium with payoffs (x^*, v^p) is played in the next-period subgame c^* , where

$$x^* = (1 - \delta) \sum_y p_y^1 (u(s_y^* + b_y^*) - \psi(1)) + \delta u(r) \quad (5)$$

$$b_h^* = 0, \quad \text{and} \quad b_l^* = \frac{\delta}{1 - \delta} (v^{\text{FB}}(u(r)) - V(\mathbf{0})). \quad (6)$$

As $\bar{s} \rightarrow \infty$, $u(s_h^*) \rightarrow \infty$, since u is unbounded above. Moreover, $u(s_l^* + b_l^*) \geq u(\alpha\delta/(1 - \delta))$. Thus, $x^* \rightarrow \infty$ as $\bar{s} \rightarrow \infty$. Since $v^p \leq v^{\text{FB}}(x^*)$, the principal's punishment for reneging on the bonus becomes arbitrarily harsh as $\bar{s} \rightarrow \infty$ in the sense that $v^p \rightarrow -\infty$.

In the remainder of the proof, I show that any payoffs in $\Pi(\hat{c})$ are also contained in $B\Pi(\hat{c})$ for any $\hat{c} \in C$. Since $(U(\hat{c}), V(\hat{c})) \in \Pi(\hat{c})$ for all $\hat{c}$, it follows from the proof of Lemma 2 that $(U(\hat{c}), V(\hat{c})) \in B\Pi(\hat{c})$ for all $\hat{c}$. Two steps remain. The first is to show that the desired first-best payoffs are in $B\Pi(\mathbf{0})$ and $B\Pi(c^*)$. The second step shows that $(x^*, v^p) \in B\Pi(c^*)$.

Step 1: First-best payoffs. Let $x \in [u(r), x^{\text{FB}}(V(\mathbf{0}))]$. To show that $(x, v^{\text{FB}}(x)) \in B\Pi(\mathbf{0}) \cap B\Pi(c^*)$, it suffices, by Lemma 1, to show that $(x, v^{\text{FB}}(x)) \in B\Pi(c, E)$. This is

because $U(\mathbf{0}) = U(c^*) = u(r) \leq x$ and $V(\mathbf{0}) = V(c^*) \leq v^{\text{FB}}(x)$, where the latter inequality follows from $x \leq x^{\text{FB}}(V(\mathbf{0}))$.

Consider the following strategies σ in (c, E, Π) . The agent exerts effort $e^{\text{FB}}(x)$. Following any output y , the principal pays a bonus $b = u^{-1}(x + \psi(e^{\text{FB}}(x)))$ if the agent's effort is $e^{\text{FB}}(x)$ and pays no bonus otherwise. The continuation payoffs f are given by

- $(x, v^{\text{FB}}(x))$ if no one deviated from σ
- $(u(r), V(\mathbf{0}))$ if the agent deviated from σ
- (x^*, v^p) if the principal deviated from σ but the agent did not.

Notice that the continuation contract following any history is c^* , and the above continuation payoffs are contained in $\Pi(c^*)$. Now consider the incentives to follow σ given the continuation payoffs in f . When the agent deviates, paying no bonus is optimal, as bonuses do not affect the principal's continuation payoff. Following effort $e^{\text{FB}}(x)$, paying the bonus b is incentive compatible only if

$$(1 - \delta)u^{-1}(x + \psi(e^{\text{FB}}(x))) \leq \delta(v^{\text{FB}}(x) - v^p).$$

Since $v^p \rightarrow -\infty$ as $\bar{s} \rightarrow \infty$, there exists $\bar{s}_2$ such that whenever $\bar{s} \geq \bar{s}_2$, the incentive constraint above is satisfied for any $x \in [u(r), x^{\text{FB}}(V(\mathbf{0}))]$. Moreover, the agent obtains $-\infty$ upon deviating, since following any output y , he receives salary $s_y = 0$ and no bonus. Thus, (σ, f) is E -admissible with respect to Π at c with value $(x, v^{\text{FB}}(x))$ whenever $\bar{s} \geq \bar{s}_2$.

Step 2: $(x^*, v^p) \in B\Pi(c^*)$ for some v^p . By Lemma 3, it suffices to show that there exists an equilibrium of (c^*, E, Π) with agent payoff x^* . To this end, consider the following strategies σ^* . The agent exerts effort $e = 1$. Following output y , the principal pays bonus b_y^* (defined in (6)) if $e = 1$ is observed, and pays no bonus otherwise. Continuation payoffs f^* are given by

- $(u(r), v^{\text{FB}}(u(r)))$ if no one deviated from σ^*
- $(u(r), V(\mathbf{0}))$ otherwise.

Notice that $c_h^* = c_l^* = \mathbf{0}$ and the continuation payoffs above are contained in $\Pi(\mathbf{0})$. Now consider the incentives to follow σ^* given the continuation payoffs in f^* . Paying no bonuses following $e \neq 1$ is optimal, since they do not affect the principal's continuation payoff. If $e = 1$ and output is y , the principal receives $v^{\text{FB}}(u(r))$ if she pays b_y^* and receives $V(\mathbf{0})$ otherwise, which makes b_y^* incentive compatible by (6). As for the agent, notice that he obtains x^* if he chooses $e = 1$. Alternatively, if he deviates in effort, he receives no more than

$$\max_e (1 - \delta) \sum_y p_y^e (u(s_y^*) - \psi(e)) + \delta u(r). \quad (7)$$

When $\bar{s} \geq \bar{s}_1$, effort 1 is optimal in (4) and, consequently, it is also optimal in problem (7). Moreover, setting $e = 1$ in (7) provides a lower bound on x^* by (5). Hence, the agent has no profitable deviation. It follows that (σ^*, f^*) is E -admissible with respect to Π at c^* with agent payoff x^* whenever $\bar{s} \geq \bar{s}_1$.

Thus, Π is self-generating when $\bar{s} \geq \max\{\bar{s}_1, \bar{s}_2\}$. This completes the proof.

3.3 Discussion of Theorem 1

The key part of the construction of the first-best equilibria is c^* —the next-period continuation contract at all on-path histories. Any deviation is followed by an adverse equilibrium for the corresponding player in the next-period subgame c^* . In particular, the agent can be held down to his outside option, while the principal's payoff can be made arbitrarily low when the contract space is large enough. Hence, any bonus can be incentive compatible, regardless of the principal's impatience. The construction takes advantage of this by making c , the contract the parties negotiate to every period, exhibit zero contemporaneous salaries. This trivializes effort incentives, since the agent's entire equilibrium compensation comprises effort-contingent bonuses.

The severity of the punishments is enabled by the unbounded risk aversion of the agent. The first-period compensation in c^* is heavily unbalanced: a high salary for high output and a low salary for low output.⁹ In the agent's punishment equilibrium $\underline{\sigma}^A(c^*)$ of subgame c^* , he is held down to the value he can guarantee by maintaining the terms of c^* in anticipation of no bonuses. Despite the high salary for high output, this guarantee equals his outside option due to his unbounded risk aversion ($u(0) = -\infty$) and his inability to save income in prior periods to insure against a low output realization. In contrast, the principal's punishment equilibrium in subgame c^* features a bonus for low output when the agent maintains c^* by rejecting the principal's offer. This bonus has arbitrarily high marginal utility, as the salary for low output approaches 0. Hence, the agent can obtain an arbitrarily large payoff by rejecting any contract offer, making the principal's payoff arbitrarily low. The size of the contract space is crucial: as the principal becomes impatient, the incentive compatibility of any positive bonus requires an arbitrarily harsh punishment, which can only be attained through a continuation contract with high salaries.

It is important to note that the punishments are not merely a consequence of the poor insurance c^* provides on paper. If maintaining the terms of c^* always resulted in low payoffs for both players, they would renegotiate c^* to a more balanced contract, regardless of any past deviations. Instead, the role of c^* is to admit a large multiplicity of equilibria, so that the value players attach to it can vary significantly based on the history, creating punishments for past transgressions. In this manner, c^* enables endogenous shifts in bargaining power. It can be seen as a source of strategic ambiguity, as described by [Bernheim et al. \(1998\)](#), who argue that environments where contracts are incomplete may make it optimal for parties to sign contracts that are even more incomplete. Here, c^* does not add strategic ambiguity by leaving out contractual terms directly, but it adds endogenous incompleteness through equilibrium multiplicity.

3.4 Necessity of contract renegotiation

A prominent feature of the first-best equilibria constructed above is that the players repeatedly sign a continuation contract c^* , only to renegotiate it in the following period.

⁹Even though c^* is a high-powered contract, this does not need to be the case. A contract with a high salary for low output and a low salary for high output can be used to the same effect.

They deliberately choose c^* for its potential to be extremely damaging to either of them, recognizing that it curbs their own opportunistic behavior. The terms of c^* are never meant to be realized: once it is observed that no one deviated from the equilibrium actions, c^* has served its purpose and is renegotiated to a contract the present terms of which are more conducive to attaining the first best. This is a robust feature: Proposition 2 below shows that whenever the players are sufficiently impatient, first-best outcomes with costly effort are not attainable in equilibria without on-path contract renegotiation.¹⁰ Hence, there are settings where even the combination of formal and informal contracting cannot attain maximum welfare unless the parties use renegotiation as part of their informal contract.

DEFINITION 4. An equilibrium without on-path renegotiation exhibits no acceptance of contracts different from the residual contract on the equilibrium path, except when the residual contract is $\mathbf{0}$.

Equilibria without on-path renegotiation require the parties to continue using their previously signed contract with a few exceptions. In the initial period, where the residual contract is $\mathbf{0}$, they can sign any contract in order to start their relationship. Recontracting is also permitted in any subsequent period where the residual contract is $\mathbf{0}$. One interpretation is that the parties can sign short-term contracts and renegotiate them after their expiry, signified by the null contract.

PROPOSITION 2. Fix any parameters $\theta \in \Theta$. For any effort level $e > 0$, there exists $\underline{\delta} > 0$ such that no first-best outcome with effort e can be attained in an equilibrium without on-path renegotiation of subgame $\mathbf{0}$ whenever $\delta < \underline{\delta}$, regardless of the value of $\bar{s}$.

PROOF. Consider any equilibrium without on-path renegotiation of subgame $\mathbf{0}$ that attains a first-best outcome with effort $e > 0$. Let x be the agent's utility so that $e = e^{\text{FB}}(x)$. The agent's total compensation, comprising a salary and a bonus, must equal $u^{-1}(x + \psi(e))$ at all on-path histories. Since bonuses are nonnegative, the lack of renegotiation implies that no contract signed on the equilibrium path can have a salary higher than $u^{-1}(x + \psi(e))$. Thus, at any on-path history up to the beginning of a period, the principal can guarantee $-u^{-1}(x + \psi(e))$ by offering the residual contract for the rest of the game and paying no bonuses (recall that output is nonnegative). Similarly, her equilibrium payoff in subgame $\mathbf{0}$ is bounded below by 0, which implies that $x \leq x^{\text{FB}}(\mathbf{0})$.

In what follows, I show that whenever δ is sufficiently small, at least one of the bonuses paid in the initial period on the equilibrium path must exceed

$$\underline{b} := \min_{\hat{x} \in [u(r), x^{\text{FB}}(\mathbf{0})]} \frac{u^{-1}(\hat{x} + \psi(e)) - u^{-1}(\hat{x})}{2} = \frac{u^{-1}(u(r) + \psi(e)) - r}{2}.$$

¹⁰If the players are patient enough, some first-best outcomes can be attained through informal contracting alone. For instance, it is possible to construct an equilibrium with payoffs $(u(r), v^{\text{FB}}(u(r)))$, where the null contract $\mathbf{0}$ is used on path and the agent's compensation consists entirely of bonuses. The continuation strategies following any deviation are taken from the agent's worst equilibrium constructed to prove Lemma 2. In particular, the principal has no incentive to deviate from the equilibrium bonus when δ is sufficiently high, as her continuation payoff $V(\mathbf{0})$ is lower than her on-path payoff $v^{\text{FB}}(u(r))$.

Let s_h and s_l be the first-period salaries for high and low output of the initial contract signed in equilibrium. Suppose $s_h \geq u^{-1}(x + \psi(e)) - \underline{b}$ and $s_l \geq u^{-1}(x + \psi(e)) - \underline{b}$ with at least one strict inequality. Then the concavity of u implies that

$$\begin{aligned} \sum_y p_y^0 u(s_y) &> u(u^{-1}(x + \psi(e)) - \underline{b}) \\ &\geq u\left(u^{-1}(x + \psi(e)) - \frac{u^{-1}(x + \psi(e)) - u^{-1}(x)}{2}\right) \\ &= u\left(\frac{u^{-1}(x + \psi(e)) + u^{-1}(x)}{2}\right) \geq x + \frac{\psi(e)}{2}. \end{aligned}$$

It follows that if the agent deviates to zero effort in the initial period, he can obtain a payoff greater than

$$(1 - \delta)\left(x + \frac{\psi(e)}{2}\right) + \delta u(r) = x + \frac{\psi(e)}{2} - \delta\left(x + \frac{\psi(e)}{2} - u(r)\right),$$

which exceeds his on-path payoff x whenever $\delta < \delta_1 := \psi(e)/(\psi(e) + 2(x^{\text{FB}}(0) - u(r)))$. Hence, one of the bonuses paid on the equilibrium path in the initial period, denoted by b , satisfies $b \geq \underline{b}$ whenever $\delta < \delta_1$.

On the equilibrium path, the principal receives a next-period continuation payoff $v^{\text{FB}}(x)$ after paying the bonus b . Hence, the lower bound on the principal's payoff derived earlier implies that b is incentive compatible only if

$$b \leq \frac{\delta}{1 - \delta} (v^{\text{FB}}(x) + u^{-1}(x + \psi(e))).$$

Let δ_2 be the unique discount factor that satisfies

$$\underline{b} = \frac{\delta_2}{1 - \delta_2} (v^{\text{FB}}(u(r)) + u^{-1}(x^{\text{FB}}(0) + \psi(e))).$$

It follows that whenever $\delta < \min\{\delta_1, \delta_2\}$, there exists no first-best equilibrium without on-path renegotiation of subgame **0** where the agent exerts effort e . $\square$

In a hypothetical first-best equilibrium without on-path renegotiation, the agent should not value any of the contractual salaries much higher than his equilibrium utility; otherwise, his compensation would be distorted. This limits the punishment that can be exerted on the principal, unlike in the equilibrium constructed to prove Theorem 1, where the punishment is unbounded. It follows that only vanishingly small bonus payments are credible, as the principal becomes arbitrarily impatient. However, when the first-best effort is costly, nontrivial bonuses are needed to provide incentives, rendering such an equilibrium impossible. Hence, by Theorem 1, when the parties are impatient, contract renegotiation on the equilibrium path provides a Pareto improvement as long as the contract space is large enough. Proposition 2 also shows that this welfare gap does not shrink as the contract space is expanded, i.e., $\bar{s}$ is increased.

While Proposition 2 implies that renegotiation must occur at least once on the equilibrium path, it is silent on the minimal frequency of renegotiation that suffices to implement first-best outcomes. The equilibria constructed in the proof of Theorem 1 exhibit renegotiation after every on-path history. However, it is possible to reduce the renegotiation frequency by amending them as follows. Instead of signing a contract with zero salaries in the initial period and a continuation contract c^* , the parties sign a contract with little compensation in the first T periods, followed by c^* . If the principal deviates, the agent can reject her offers until c^* is reached, where he receives a high payoff, as in the proof of Theorem 1. When $\bar{s}$ is large enough, such a construction can support the first-best outcomes in Theorem 1 through on-path renegotiation that occurs every T periods for any finite T . However, obtaining less renegotiation in this manner necessitates a higher upper bound on salaries, as the agent needs to wait longer for the salaries in c^* .

3.5 Contract length

One implication of Proposition 2 is that one-period formal contracts as in Pearce and Stacchetti (1998) cannot replicate the first-best outcomes attained by long-term contracts with renegotiation. The advantage of longer contracts is that their future terms can create threats for bad behavior without being realized on the equilibrium path. This force is so powerful that even two-period contracts can attain the first-best outcomes in Theorem 1, as shown in the equilibrium construction. However, longer contracts can reduce the size of the contract space required to attain the first best.

This can be illustrated as follows. Recall the on-path continuation contract c^* with highly skewed first-period salaries $s_h^* \gg s_l^*$ and no further compensation. Consider a contract c^{**} that exhibits the same skewed compensation for two periods instead, i.e., $s_y^{**} = s_y^*$ and $c_y^{**} = c^*$ for all y . It can be shown that subgame c^{**} admits a worse equilibrium payoff for the principal in comparison to the equilibrium of subgame c^* with payoffs (x^*, v^p) constructed in Step 2 of the proof of Theorem 1. In the latter, the principal pays a bonus $b_l^* = (v^{FB}(u(r) - V(\mathbf{0}))\delta)/(1 - \delta)$ whenever the agent rejects in favor of c^* and output is low. But when the residual contract is c^{**} instead, this bonus can be increased to $(v^{FB}(u(r) - v^p)\delta)/(1 - \delta)$ by threatening the principal with continuation payoff v^p in the next-period subgame c^* whenever she reneges on the bonus.

Thus, the construction of first-best equilibria in the proof of Theorem 1 can be amended so that the contract c signed each period has a continuation contract c^{**} instead of c^* . In this new equilibrium, the principal has stronger incentives to pay the on-path bonuses due to the harsher punishment available in the continuation subgame c^{**} . Since this punishment payoff decreases in $\bar{s}$, it follows that equilibria with three-period contracts can attain the payoffs of the first-best equilibria with two-period contracts in the proof of Theorem 1 for lower values of $\bar{s}$.¹¹ It can be similarly shown that longer contracts decrease the required value of $\bar{s}$ further.

¹¹The improvement resulting from three-period contracts is not necessarily strict. When δ is high, there exist first-best equilibria where $\mathbf{0}$ is the only contract signed on the equilibrium path (see footnote 10). However, for sufficiently low δ , it is possible to show the following stronger result: there are values of $\bar{s}$ for which no equilibrium with two-period contracts can attain the first-best outcomes in Theorem 1, but an equilibrium with three-period contracts can.

It is interesting to combine this with the observation at the end of Section 3.4. Together, they suggest that, everything else equal, the frequency of renegotiation can be reduced by writing longer contracts.

3.6 Renegotiation proofness

In the baseline model, players are able to renegotiate contracts, but they cannot escape inefficient equilibria by renegotiating their strategies. To account for the latter, this section presents two renegotiation-proofness refinements.

The classic notions of renegotiation-proofness (Bernheim et al. (1989), Farrell and Maskin (1989), Pearce (1987)) were developed in the context of repeated simultaneous-move games. They do not apply directly to this model, because the stage game is sequential and varies with the long-term contract in place at the beginning of the period. Nevertheless, Levin (2003) considers a generalization to dynamic stage games, which can be seen as a stronger version of strong consistency (Bernheim et al. (1989)) and strong renegotiation-proofness (Farrell and Maskin (1989)). Levin's (2003) strong optimality essentially posits that the parties renegotiate to efficient equilibria on and off the equilibrium path, and this renegotiation occurs at the beginning of the period. Though this refinement is developed in the context of a stationary game, it is easily adapted to the stochastic game considered here.

DEFINITION 5. An equilibrium is strongly optimal if the continuation equilibrium at any history up to the beginning of a period is efficient in the sense that its payoffs are Pareto optimal in $\mathcal{E}(c^*)$, where c^* is the residual contract.

Safronov and Strulovici (2018) propose an alternative approach that models the negotiation of strategies explicitly. They consider an augmented game where at the end of each period, one of the players is randomly selected to propose a continuation equilibrium. This is similar to the noncooperative approach to contract renegotiation adopted here. Safronov and Strulovici (2018) define a contestable norm to be an equilibrium of the augmented game where within any subgame from the beginning of a period, any strategy profile that is accepted is played as long as no off-path proposals were accepted in the past.¹² All strongly optimal equilibrium payoffs are attainable by contestable norms, as there exist no Pareto improving proposals for continuation play from the start of a period.

The following proposition shows that under the assumption of vanishing marginal utility, all efficient equilibrium payoffs can be attained in strongly optimal equilibria, so they also satisfy the weaker refinement of contestable norms.

PROPOSITION 3. Fix any parameters $\theta \in \Theta$ and $\delta \in (0, 1)$ such that $\lim_{s \rightarrow \infty} u'(s) = 0$. For any $\bar{s}$ sufficiently large and any $c^* \in C$, the payoffs of any efficient equilibrium of subgame c^* are attainable in a strongly optimal equilibrium of subgame c^* .

¹²Safronov and Strulovici (2018) also consider the case of more than two players with threshold acceptance rules.

To see why Proposition 3 holds, consider any efficient equilibrium σ of any subgame c^* and any history up to the beginning of a stage. If this history is reached with positive probability on the equilibrium path, then continuation play must be efficient; otherwise, replacing the continuation strategies with a Pareto improving equilibrium would strengthen incentives at preceding histories and Pareto improve the payoffs from σ . In contrast, a Pareto improvement of continuation payoffs at off-path histories does not necessarily preserve incentives, as it may make deviations at preceding histories more attractive. However, since actions are sequential, incentives are unchanged if the Pareto improving equilibrium leaves the payoff of the *last* deviating player unchanged, improving only the other one.

It follows that Proposition 3 holds provided that for any player and any equilibrium σ of any subgame c^* , there exists an efficient equilibrium of subgame c^* with the same payoff for this player. If salaries were unrestricted, such an equilibrium can be constructed from another efficient equilibrium by adjusting the first-period salaries to increase one player's payoff at the expense of the other. This is the approach to proving strong optimality in the extended model of Section 4. However, in the baseline model considered here, the exogenous upper bound on salaries may make it impossible to decrease the principal's payoff by increasing salaries of equilibrium contracts. The proof of Proposition 3 overcomes this issue by showing that the principal's worst equilibrium payoff in any subgame c^* is lower bounded by $v^{\text{FB}}(u(\bar{s}))$, which is shown to be attainable in a first-best equilibrium with agent utility $u(\bar{s})$. This follows from the assumption of vanishing marginal utility, which ensures that the above first-best outcome involves zero effort. The proof offers a direct construction of all efficient equilibria, showing that they are first best.

3.7 Relation to WMO

Watson et al. (2020) consider a general setting where each period consists of a cooperative negotiation phase followed by a noncooperative action phase. Two risk-neutral players sign long-term formal contracts governing the structure of the action phase. In the bargaining phase, players negotiate a contract, transfers, and continuation strategies according to a contractual equilibrium (Watson (2013), Miller and Watson (2013)), where joint surplus is maximized and split according to the generalized Nash bargaining solution. The threat point for Nash bargaining is given by an equilibrium of the game from the action phase under the residual contract (i.e., when the formal contract has not been renegotiated), with next-period continuation payoffs drawn from the set of contractual equilibria. WMO find that, without loss of generality, the long-term contracts signed in any contractual equilibrium are semistationary, that is, they have identical terms for all periods except the initial one, regardless of the history of verifiable outcomes.

Contractual equilibrium is conceptually distinct from the noncooperative approach to contract negotiation and the notions of renegotiation-proofness used in this paper. Nevertheless, WMO's main result remains true. The construction from the proof of Theorem 1 can be amended so that the on-path continuation contract c^* has the same skewed compensation in every period (see Proposition 4 below). This makes the on-path contract c semistationary.

Another point of comparison is the attainability of the first-best outcomes from Theorem 1 in a contractual equilibrium. This cannot be accomplished directly, as the definition in WMO relies on transferable utility. However, WMO argue that an alternative definition of contractual equilibrium can be obtained by replacing their cooperative bargaining phase with a noncooperative bargaining protocol and applying an equilibrium refinement based on three axioms from Miller and Watson (2013).¹³ These axioms do not depend on risk preferences, so they can be interpreted in the context of my model. In general terms, the internal agreement consistency axiom (IAC) states that accepted proposals are implemented.¹⁴ The no-fault disagreement axiom (NFD) states that the outcome of disagreement over a proposal does not depend on the proposal. Instead of presenting the last axiom, Pareto external agreement consistency (PEAC), I consider a stronger version: payoffs following the acceptance of any proposal are Pareto optimal in the class of equilibria that satisfy IAC and NFD.

Proposition 4 below shows that the principal-optimal first-best outcome can be attained in an equilibrium of my model that satisfies IAC, NFD, and PEAC. A notable difference in the meaning of the axioms comes from the nature of proposals: in my model, the principal proposes contracts, whereas in WMO, players propose strategies as well. For example, IAC does not have bite here, since an accepted contract automatically comes into effect. However, Proposition 4 would continue to hold if the principal were able to propose strategies in addition to contracts. This is due to the nature of the equilibrium constructed in the proof, where the principal's continuation payoff following any history up to the beginning of a period equals her highest feasible (first-best) payoff given what the agent can obtain by rejection or his outside option. Notably, punishments are created by varying the agent's payoff following rejection based on prior deviations, but not on the current contract offer.

PROPOSITION 4. *Fix any parameters $\theta \in \Theta$ and $\delta \in (0, 1)$ such that $u'(s) \rightarrow 0$ as $s \rightarrow \infty$. For any $\bar{s}$ sufficiently large, the payoffs $(u(r), v^{FB}(u(r)))$ are attainable in a strongly optimal equilibrium with the following properties:*

- (i) *The continuation strategies at any history following the rejection of a contract do not depend on the contract that was offered.*
- (ii) *The continuation payoffs at any history following the acceptance of any contract c are Pareto optimal in the class of equilibria of subgame (c, E) that satisfy (i).*

In addition, all contracts signed on the path of this equilibrium are semistationary.

4. LIMITED LIABILITY FOR THE PRINCIPAL

The continuation contracts used to attain the first-best outcomes in Theorem 1 expose the principal to liability that may be orders of magnitude larger than the value of the firm

¹³See Appendix B.3 of WMO.

¹⁴This is a stronger version of the axiom than the one in Miller and Watson (2013), which considers only proposals to switch to the continuation equilibrium strategies following a different history.

when δ is small. Therefore, in this section I extend the model by allowing the principal to shut down the firm at the contract offer stage.¹⁵ As the exercise of the agent's outside option, this shutdown is permanent and results in payoffs $(u(r), 0)$. Contractual salaries are unrestricted: an upper bound is no longer needed as both parties can terminate the relationship. The model is otherwise identical to the game analyzed so far and retains the same notation. In addition, let $\underline{x}(c^*)$ and $\underline{v}(c^*)$ denote the worst equilibrium payoffs in subgame c^* for the agent and the principal, respectively.

For ease of exposition, I assume that there exists an equilibrium of subgame **0** such that the outside options are not taken in the initial period.

4.1 Equilibrium characterization

The main result of this section is a characterization of the frontier of efficient equilibrium payoffs using the recursive techniques of APS. Let $\mathcal{F}$ denote the space of all functions $f : [u(r), \bar{x}] \rightarrow [0, v^{FB}(u(r))]$, where $\bar{x} \geq u(r)$. For each $f \in \mathcal{F}$, let $\text{dom } f$ denote the domain of f and let

$$\beta f(x) = \max_{e, \{s_y, b_y, x_y\}_{y \in Y}} \sum_y p_y^e [(1 - \delta)(y - s_y - b_y) + \delta f(x_y)]$$

$$\text{such that } x = \sum_y p_y^e [(1 - \delta)(u(s_y + b_y) - \psi(e)) + \delta x_y]$$

$$x \geq \max_{e' \in [0, 1]} \sum_y p_y^{e'} [(1 - \delta)(u(s_y) - \psi(e')) + \delta u(r)] \quad (8)$$

$$0 \leq b_y \leq \frac{\delta}{1 - \delta} f(x_y) \quad \text{for all } y \in Y \quad (9)$$

$$e \in [0, 1], s_y \geq 0, x_y \in \text{dom } f \quad \text{for all } y \in Y.$$

Consider the operator $T : \mathcal{F} \rightarrow \mathcal{F}$ such that for any $f \in \mathcal{F}$, Tf is the restriction of βf to the domain $\{x \geq u(r) \mid \beta f(x) \geq 0\}$. Theorem 2 below shows that T can be used to characterize efficient equilibrium payoffs.

THEOREM 2. *Let $f_0 \in \mathcal{F}$ with $\text{dom } f_0 = [u(r), x^{FB}(0)]$ be given by $f_0(x) = v^{FB}(x)$. The sequence (f_n) given by $f_{n+1} = Tf_n$ converges pointwise to a function $f^* \in \mathcal{F}$ such that $f^* = Tf^*$.*

*The set of efficient equilibrium payoffs in subgame **0** consists of all payoffs $(x, f^*(x))$ in the graph of f^* such that $f^*(x) \geq \underline{v}(\mathbf{0})$.*

The characterization in Theorem 2 is based on the operator β , which obtains efficient equilibrium payoffs from the set of efficient next-period continuation payoffs. Strikingly, β contains no explicit reference to long-term contracts and renegotiation. The reason behind this simple characterization lies in the following result, which follows from the analysis in Appendix B.2.

¹⁵This timing of the outside option is prevalent in the relational contracting literature. If the principal could also shut down after the agent's contract response, the results would remain unchanged.

LEMMA 4. *There exists a contract c^* such that $\underline{x}(c^*) = u(r)$ and $\underline{v}(c^*) = 0$.*

Recall that the construction of first-best equilibria in the baseline model uses a punishment contract c^* such that subgame c^* admits equilibria with the harshest possible punishments for each player: an arbitrarily low payoff for the principal and $u(r)$ for the agent. Lemma 4 provides a similar result for the extended model, where neither player can be held below their outside utility in equilibrium. The punishment contract for the extended model has skewed compensation in each period, mirroring the one from the baseline model (see Lemma 10 and Corollary 2 in Appendix B.2). The reason is similar: by varying the bonus complementing the salary for low output, it is possible to construct equilibria of subgame (c^*, E) with any agent payoff greater than or equal to his outside option. In the agent's punishment equilibrium of subgame c^* , he receives $u(r)$ in subgame (c^*, E) following rejection. In the principal's punishment equilibrium the agent can guarantee such a high payoff from rejection that the principal prefers to shut down.¹⁶

Lemma 4 implies it is possible to restrict attention to equilibria where all contracts accepted on path have output-invariant continuations equal to c^* . Thus, the next-period worst punishments are exogenously given by the outside options, which is reflected in the incentive constraints (8) and (9). This is a tremendous simplification, even in comparison to the model of Pearce and Stacchetti (1998) with *one-period* contracts. In the latter, a similar characterization applies, but the next-period punishments are determined endogenously as the worst equilibria of subgame $\mathbf{0}$, since contracts do not specify future terms.

Many models of relational contracting exhibit exogenously given worst equilibrium threats. Typically, this follows from two assumptions: (i) no formal contracts are permitted and (ii) the agent's least costly action cannot generate payoffs that dominate the outside options. They imply the existence of an equilibrium where the relationship is terminated immediately; otherwise, the agent would exert no effort in anticipation of no bonuses. Neither of these assumptions holds in my model. Instead, the worst punishments arise from the availability of long-term contracts and the careful design of their future terms.

4.2 Necessity of contractual renegotiation

As in the baseline model, the renegotiation of contracts on the equilibrium path is necessary for efficiency in some environments. To demonstrate this in a simple example, consider any model where

$$\frac{\delta}{1-\delta} v^{\text{FB}}(u(r)) = u^{-1}(u(r) + \psi(e^{\text{FB}}(u(r)))) - r \quad (10)$$

so that equilibrium bonuses are upper bounded by the difference between the agent's compensation in the principal-optimal first-best outcome and his outside wage. Hence, in any first-best equilibrium with agent payoff $u(r)$, the agent's realized contractual

¹⁶It is also possible to construct efficient punishment equilibria (see Proposition 5 below).

salaries in any period are at least r . Given this, no salary can be strictly larger than r ; otherwise, the agent could profitably deviate to zero effort. Hence, the agent's on-path compensation in every period must consist of a salary r and a bonus given by (10).

Theorem 2 implies that the principal-optimal first-best outcome can be attained in an equilibrium of subgame **0** using the above compensation scheme. However, it is unattainable in an equilibrium without on-path renegotiation when $\sum_y p_y^0 y - r > 0$. To see this, consider any history up to the bonus payment in the first period on the path of such an equilibrium. Since the bonus is given by (10), a deviation must result in a second-period continuation payoff of 0 for the principal. However, in the second period, she can guarantee more by offering a contract c with contemporaneous salaries slightly higher than r and continuation contracts $c_h = c_l = \mathbf{0}$. The agent does not take his outside option in response to such an offer, as he can guarantee more by accepting it and exerting no effort. But regardless of whether he accepts or rejects, the principal receives at least $\sum_y p_y^0 y$ in expected output and pays close to or less than r to the agent. The latter holds because the contract signed in the initial period cannot contain salaries higher than r due to the lack of renegotiation. Thus, the principal has a profitable deviation.

4.3 Renegotiation-proofness

Unlike in the baseline model, the strong optimality of efficient equilibria can be shown without making additional assumptions. This is owed to the lack of upper bound on salaries, which makes it possible to continuously modify the payoffs of any equilibrium to either player's benefit by adjusting the underlying contract.

PROPOSITION 5. *Fix any contract c^* . The payoffs of any efficient equilibrium of subgame c^* are attainable in a strongly optimal equilibrium of subgame c^* .*

4.4 Relation to WMO

The main result of [Watson et al. \(2020\)](#) also holds in this model. As mentioned in the discussion following Lemma 4, it is without loss of generality that each contract offered in equilibrium has the same skewed compensation in all periods except possibly in the initial one. An analogue of Proposition 4 also holds: the principal-optimal efficient payoffs in subgame **0** can be attained in an equilibrium that satisfies the axioms underlying WMO's contractual equilibrium. I omit the argument, as it is similar to the construction used to prove Proposition 4.

These similarities suggest that WMO's results are not only robust to risk aversion, but also to the potential nonstationarity of efficient equilibrium strategies that arises from it, suggested by Theorem 2. In comparison, WMO show that risk neutrality implies that efficient equilibrium play is stationary without loss of generality, extending the result of [Levin \(2003\)](#).

5. LITERATURE REVIEW

This paper contributes to the literature on relational contracts, which studies informal contracting as an equilibrium of a repeated game ([Bull \(1987\)](#), [Thomas and Worrall](#)

(1988), MacLeod and Malcomson (1989), Levin (2003)). Many authors have considered environments with both formal and informal contracting (Baker et al. (1994), Schmidt and Schnitzer (1995), Bernheim et al. (1998), Pearce and Stacchetti (1998), Che and Yoo (2001), Battigalli and Maggi (2008), Kvaløy and Olsen (2009), Iossa and Spagnolo (2011), Hermalin et al. (2013), Itoh and Morita (2015)). These studies are limited to formal contracts that are either one-period long or are stationary but cannot be renegotiated. In comparison, I consider arbitrary long-term contracts that can be renegotiated (as previously discussed, this is also done in WMO). In addition, I contribute to the comparatively smaller literature on relational contracts with risk aversion (Thomas and Worrall (1988), Pearce and Stacchetti (1998), MacLeod (2003), Thomas and Worrall (2018)).

The incomplete contracts literature examines settings where some but not all observable outcomes are contractible, emphasizing the structure of formal contracts and their renegotiation. One strand, originating in Hart and Moore (1988), uses a mechanism design approach, later generalized by Maskin and Moore (1999) and Segal and Whinston (2002), where contracts are contingent on messages sent by the players after observing nonverifiable outcomes. Since messages identify the entire history of play and players can costlessly renegotiate their contract after sending them, the revelation principle implies that without loss of generality contracts are not renegotiated; otherwise, players could incorporate the renegotiated terms into their original contract (Brennan and Watson (2013)). Hence, under this approach, renegotiation is not an equilibrium phenomenon, but a restriction on the contract space.

In my paper, renegotiation is modelled strategically as part of a noncooperative game without making use of contracts contingent on messages. This approach is used in a variety of one-period settings, including holdup (Huberman and Kahn (1988b)), leveraged buyout (Huberman and Kahn (1988a)), and a principal–agent model (Hermalin and Katz (1991)). Guriev and Kvasov (2005) consider a dynamic holdup problem where a seller makes persistent investments that improve the value of trade with a buyer. They obtain conditions under which the first-best surplus can be attained in an equilibrium where the players repeatedly renegotiate to the same long-term contract. There is a unique equilibrium, so the players do not rely on implicit contracting to create incentives. Instead, the seller is punished for underinvesting because the persistence of his investment lowers the utility he can guarantee under the continuation contract, deteriorating his bargaining position at the next instant when the contract is renegotiated. In contrast, the actions unrelated to contract renegotiation have no persistent effect on my environment.

Rey and Salanie (1990) analyze a setting with moral hazard where the second best, attained by long-term contracts, can be replicated by overlapping two-period contracts renegotiated in equilibrium. In their setting every observable outcome is contractible, which makes optimal long-term contracts renegotiation-proof. Thus, renegotiation is only useful when contract length is exogenously limited.

6. CONCLUSION

This paper analyzes a novel framework combining noncontractible actions with long-term formal contracts that can be renegotiated at will. I show it is beneficial to sign

contracts in anticipation of their future renegotiation, despite the lack of persistence in the environment. This is a powerful form of implicit contracting: continuation contracts provide harsh punishments for deviations, but are renegotiated away on the path of efficient equilibria. The result is driven by the interaction of risk aversion and voluntary bonus payments, so it may apply to other environments with risk-sharing and informal contracting.

APPENDIX A: PROOFS FOR THE BASELINE MODEL

A.1 *Boundedness of contracts*

Contracts must be bounded to ensure equilibrium existence by preventing Ponzi schemes. To see this, suppose wages are unrestricted but an equilibrium exists. Fix any equilibrium of any subgame c^* with payoffs (x, v) . The maximum feasible payoff for the principal is $v^{\max} = \sum_y p_y^1 y$ since salaries and bonuses are nonnegative. If $v = v^{\max}$, the agent must receive zero salaries and bonuses, so he is better off deviating to his outside option. Hence, $v < v^{\max}$.

I now argue that the principal can secure a payoff arbitrarily close to $v^{\max}$ with a multistage deviation involving a Ponzi scheme. In the initial period, the principal offers a contract c with $c(y^t) = \underline{s}$ for all $t \in \mathbb{N}$, $y^t \in Y^t$ with the exception of $c(h, h)$. When $c(h, h)$ is sufficiently large, the agent can guarantee at least x following his acceptance of c , even when $\underline{s}$ is arbitrarily small (see Section 3.1). The principal's worst equilibrium payoff in subgame c_l is lower bounded by $-\underline{s}$, as she can offer the residual contract in each subsequent period and pay no bonuses. Thus, the highest bonus the agent can anticipate in equilibrium following the acceptance of c when output is low is $(v^{\max} + \underline{s})\delta/(1 - \delta)$. Hence, for $c(h, h)$ sufficiently large, the agent would accept a deviant offer c and exert effort 1 in the initial period regardless of the continuation equilibrium he anticipates.

Now suppose the principal makes the deviant offer c and consider the next period with residual contract c_h or c_l . Since $c(y^t) = \underline{s}$ for any $t > 2$, $y^t \in Y^t$, the bonus paid by the principal in the initial period of any equilibrium of subgame (c_h, E) or (c_l, E) is upper bounded by $(v^{\max} + \underline{s})\delta/(1 - \delta)$. Thus, the agent's equilibrium payoff in these subgames can be upper bounded. Hence, if the principal deviates in subgame c_h or c_l by offering a contract c' with the same structure as c , i.e., $c'(h, h)$ high and other salaries arbitrarily small, then the agent accepts in any equilibrium and exerts effort 1. By making the same deviation in each subsequent period, the principal can ensure that the agent exerts maximal effort for almost no pay, which gives her a payoff arbitrarily close to $v^{\max}$.

A.2 *Proof of Lemma 1*

Let (σ^*, f^*) be a profile admissible with respect to Π at c^* with value (x^*, v^*) . Let (σ_E, f_E) be a profile E -admissible with respect to Π at c with value $(x, v) \geq (x^*, v^*)$. Consider the strategies σ identical to σ^* except that the principal offers c , the agent accepts c , and the following changes after the agent's contract response to c . The continuation of σ following the acceptance of c equals σ_E . The continuation of σ following rejection of c is

given by the continuation strategies of σ^* following the rejection of the contract offered in σ^* .

Construct a function f identical to f^* except at histories where c is offered. The restriction of f to histories where c is accepted is identical to f_E . For any history ω where c is offered and rejected, let $f(\omega) = f^*(\omega')$, where ω' is identical to ω except that the principal's offer equals the one made in σ^* . It follows that (σ, f) is admissible with respect to Π at c^* with value (x, v) , as required.

A.3 Proof of Lemma 2

If V is continuous and bounded, then each V_c is a well defined continuous function due to the continuity of U , U^E , u , and ψ . Moreover, for any $c^* \in C$, the set of contracts c for which $V_c(c^*) > -\infty$ is nonempty (as it includes a contract with constant salaries equal to $u^{-1}(U(c^*))$ in every period) and compact (due to the continuity of U^E). It follows that $\max_c V_c$ is continuous. It is also bounded, since $U(c^*)$ is bounded below by the agent's outside option and bounded above due to the upper bound on salaries $\bar{s}$. Hence, standard dynamic programming arguments show that there is a unique bounded solution for V , which is continuous.

The strategies constructed below break indifferences in favor of the principal, so it is useful to denote by $e(c)$ one of the possible multitude of optimal effort levels in the definition of $V_c(c^*)$. Notice that $e(c)$ is independent of c^* , since c^* does not enter constraint (3).

By Proposition 1, an equilibrium where the agent receives $U(c^*)$ in any subgame c^* exists if the correspondence

$$\Pi(c^*) = \{(U(c^*), V(c^*))\} \quad \text{for all } c^*$$

is self-generating. Hence, it suffices to show that for any $c^* \in C$, there exists an equilibrium of (c^*, f) with payoffs $(U(c^*), V(c^*))$, where $f(\omega) = (U(c'), V(c'))$ for any history $\omega \in \Omega$ such that the continuation contract in subgame c^* following first-period history ω equals c' .

The strategy of the principal is to offer c^* if $U(c^*) = u(r)$ and $0 > \max_c V_c(c^*)$, and to offer a contract c that maximizes $V_c(c^*)$ otherwise. The agent takes his outside option if the principal offers c^* , $U(c^*) = u(r)$, and $0 > \max_c V_c(c^*)$. Otherwise, he accepts an offer c when $U^E(c) \geq U(c^*)$, rejects it when $U^E(c^*) > U(c)$, and takes his outside option in the remaining cases. The agent's effort is $e(c)$ if he accepted c and is $e(c^*)$ if he rejected.

I now verify that the strategies described above form an equilibrium of (c^*, f) . It is optimal for the principal to pay no bonuses, as they do not affect her continuation payoffs. The agent's effort choice following rejection is optimal since $e(c^*)$ solves (1) and the agent's continuation payoff following output y equals $U(c_y^*)$. Similarly, $e(c)$ is an optimal effort choice following acceptance of any contract c . Thus, accepting c gives the agent $U^E(c)$, while rejecting gives him $U^E(c^*)$, making his contract response optimal.

Finally, consider the principal's choice of contract offer c . Given the agent's strategy, the highest payoff the principal can obtain by offering a contract c that will be accepted by the agent is $V_c(c^*)$ and the highest payoff she can obtain by having her offer

rejected is $V_{c^*}(c^*)$. Thus, her payoff cannot exceed $\max\{\max_c V_c(c^*), 0\}$. Moreover, when $U(c^*) > u(r)$, the agent does not take his outside option, so the principal's payoff cannot exceed $\max_c V_c(c^*)$. In both cases, these upper bounds are attainable. If $U(c^*) = u(r)$ and $0 > \max_c V_c(c^*)$, she can offer c^* , which results in the agent taking his outside option. Otherwise, she can offer a contract c , which maximizes $V_c(c^*)$, resulting in the agent's acceptance and effort $e(c)$, which gives the desired payoff.

In both cases the agent's payoff equals $U(c^*)$: the agent takes the outside option only if $U(c^*) = u(r)$, and the contract c that maximizes $V_c(c^*)$ has $U^E(c) = U(c^*)$. To see the latter, notice that if c satisfies (2), then $s_h > 0$ and $s_l > 0$, as $u(0) = -\infty$. Hence, if $U^E(c) > U(c^*)$, a contract c' with $c'_y = c_y$ and $u(s'_y) = u(s_y) - \varepsilon$ for all y has $U^E(c') > U(c^*)$ when $\varepsilon > 0$ is small and $V_{c'}(c^*) > V_c(c^*)$ as (3) continues to hold at $e = e(c)$ when c is replaced by c' .

Thus, the strategies form an equilibrium of (c^*, f) with payoffs $(U(c^*), V(c^*))$, as required. Moreover, for any $c \in C$, the restriction of the strategies to the auxiliary game (c, E, f) is an equilibrium with agent payoff $U^E(c)$. This concludes the proof.

A.4 Proof of Lemma 3

Let c^* , x^* , and Π satisfy the hypothesis of the lemma. Then there exists a profile (σ^*, f^*) that is E -admissible with respect to Π at c^* with value (x^*, v^*) for some v^* . Consider any contract $\hat{c}$ with $U^E(\hat{c}) = x^*$. Let (σ, f) be the first-period strategies and second-period continuation payoffs in the equilibrium $\underline{\sigma}^A(\hat{c})$ of subgame $\hat{c}$ constructed to prove Lemma 2. By the hypothesis of the lemma, $f(\omega) = (U(c'), V(c')) \in \Pi(c')$ for any history $\omega \in \Omega$ where the continuation contract is given by c' .

Now modify (σ, f) as follows. First, following rejection of any contract offer, the strategies and continuation payoffs are given by σ^* and f^* , respectively. Second, set the principal's offer to $\mathbf{0}$ if $v^* > V(\hat{c})$. Thus, the principal either obtains v^* by offering $\mathbf{0}$, which will be rejected by the agent, or the game unfolds as in the agent's worst equilibrium $\underline{\sigma}^A(\hat{c})$, since the agent can guarantee $x^* = U(\hat{c})$ by rejecting.

The modified profile is admissible with respect to Π at c^* with payoffs $(x^*, \hat{v})$, where $\hat{v} = \max\{v^*, V(\hat{c})\}$. To see this, first notice that incentives from the effort stage onward are satisfied by construction, since an equilibrium of (c^*, E, Π) is played following rejection and an equilibrium of (c, E, Π) is played following the acceptance of any contract c . The agent's contract response is also optimal, since he receives $x^* = U(\hat{c})$ upon rejection and $U^E(c)$ upon accepting any contract c . Finally, the incentives for the contract offer hold since the principal obtains $V(\hat{c})$ by making the best acceptable offer and obtains v^* by making an offer that will be rejected, e.g., $\mathbf{0}$.

A.5 Proof of Proposition 3

I begin by showing that the assumption of vanishing marginal utility implies that first-best outcomes with high agent utility involve no effort. Using this result, Lemma 6 obtains bounds on equilibrium payoffs.

LEMMA 5. Fix any parameters $\theta \in \Theta$ such that $u'(s) \rightarrow 0$ as $s \rightarrow \infty$. There exists $\bar{x}$ such that for any $\delta \in (0, 1)$ and $\bar{s} \geq 0$, $e^{\text{FB}}(x) = 0$ whenever $x \geq \bar{x}$.

PROOF. A sufficient condition for $e^{\text{FB}}(x) = 0$ is that

$$\sum_y (p_y^e - p_y^0)y < u^{-1}(x + \psi(e)) - u^{-1}(x) \quad (11)$$

holds for all $e \in (0, 1]$. Since u and ψ are differentiable, (11) can be rewritten as

$$e \sum_y (p_y^1 - p_y^0)y < e \frac{\psi'(0)}{u'(u^{-1}(x))} + o(e) \quad \text{as } e \rightarrow 0.$$

It follows that there exist $K_1 > 0$ and $\underline{e} > 0$ such that (11) holds for all $e \in (0, \underline{e})$ whenever $u'(u^{-1}(x)) < 1/K_1$. Let $K_2 = (h - l)/\psi(\underline{e})$ and $K = \max\{K_1, K_2\}$. Since u is unbounded and $u'(s) \rightarrow 0$ as $s \rightarrow \infty$, there exists $\bar{x}$ such that $u'(u^{-1}(\bar{x})) < 1/K$. Let $x \geq \bar{x}$. It follows from the monotonicity of u' and u^{-1} that (11) holds for all $e \in (0, \underline{e})$. Moreover, for any $e \in [\underline{e}, 1]$,

$$\sum_y (p_y^e - p_y^0)y < h - l \leq K\psi(\underline{e}) \leq \int_x^{x+\psi(e)} \frac{1}{u'(u^{-1}(z))} dz = u^{-1}(x + \psi(e)) - u^{-1}(x)$$

so (11) holds as well. Hence, $e^{\text{FB}}(x) = 0$ for all $x \geq \bar{x}$, as required. $\square$

LEMMA 6. Fix any parameters $\theta \in \Theta$ such that $u'(s) \rightarrow 0$ as $s \rightarrow \infty$ and let $\delta \in (0, 1)$. There exists $\bar{s}^{\min}$ such that $\mathcal{E}(c^*) \subseteq [u(r), u(\bar{s})] \times [v^{\text{FB}}(u(\bar{s})), v^{\text{FB}}(u(r))]$ for any $c^* \in C$ whenever $\bar{s} \geq \bar{s}^{\min}$.

PROOF. The lower bound on the agent's payoff and the upper bound on the principal's payoff are immediate. Suppose the principal offers a contract $\bar{c}$ with $\bar{c}^t(y^t) = \bar{s}$ for all t , y^t and pays no bonuses in all periods. When $\bar{s} > r$, a lower bound on her payoff from deviating to this strategy in any equilibrium is $\sum_y p_y^0 y - \bar{s}$, as the agent will not find it profitable to reject $\bar{c}$ in favor of his outside option. Moreover, Lemma 5 implies that the first-best outcome with agent utility $u(\bar{s})$ involves no effort when $\bar{s}$ is large, so $v^{\text{FB}}(u(\bar{s})) = \sum_y p_y^0 y - \bar{s}$. Hence, there exists $\bar{s}^{\min}$ such that the lower bound on the principal's payoff holds whenever $\bar{s} \geq \bar{s}^{\min}$. Finally, the upper bound on the agent's payoff is obtained from the first-best outcome where the principal's utility is at the lower bound. $\square$

To prove Proposition 3, consider the correspondence $\Pi \in \mathcal{P}_0$ such that

$$\Pi(c^*) = \bigcup_{(x,v) \in \mathcal{E}(c^*)} \{(x, v^{\text{FB}}(x)), (x^{\text{FB}}(v), v)\} \quad \text{for all } c^* \in C.$$

In what follows, I show that Π is self-generating. By Proposition 1, this implies that any efficient equilibrium payoffs in any subgame c^* are first best and can be attained in a strongly optimal equilibrium.

To this end, let $c^* \in C$. Consider any equilibrium of subgame c^* with payoffs (x, v) and denote its first-period strategies by σ . For any $\omega \in \Omega$, let c_ω^* and (x_ω, v_ω) be, respectively, the continuation contract and the continuation payoffs following history ω in subgame c . Since $(x_\omega, v_\omega) \in \mathcal{E}(c_\omega^*)$, it follows that $(x_\omega, v^{\text{FB}}(x_\omega)), (x^{\text{FB}}(v_\omega), v_\omega) \in \Pi(c_\omega^*)$.

I proceed to construct two equilibria of (c^*, Π) : one where the principal's payoff equals v and another where the agent's payoff equals x . Both of them are given by the strategies σ , though the continuation payoffs drawn from Π may differ. Toward the former, consider the continuation payoff function f such that for any $\omega \in \Omega$, $f(\omega)$ is given by

- $(x_\omega, v^{\text{FB}}(x_\omega))$ if ω shows that the agent deviated from σ but the principal paid the bonus specified by σ in the subsequent subgame
- $(x^{\text{FB}}(v_\omega), v_\omega)$ otherwise.

The restriction of σ to any subgame of (c^*, f) starting with the bonus payment is an equilibrium since the principal's continuation payoff following any history ω where she deviates from σ is unchanged relative to the equilibrium of subgame c^* , while her payoff from following σ resulting in some history ω is either unchanged (and equal to v_ω) or higher (equal to $v^{\text{FB}}(x_\omega)$). Given this, the restriction of σ to any subgame of (c^*, f) starting from the effort choice is an equilibrium since the agent's deviation payoff is unchanged, while his payoff from following σ is no smaller. Similarly, the agent's contract response and the principal's contract offer prescribed by σ are incentive compatible. Thus, (σ, f) is admissible with respect to Π at c with value (x', v) for some $x' \in [x, x^{\text{FB}}(v)]$.

Now consider the continuation payoff function $\hat{f}$ identical to f except that $\hat{f}(\omega) = (x_\omega, v^{\text{FB}}(x_\omega))$ at any history ω reached by σ with positive probability. Analogous arguments establish that $(\sigma, \hat{f})$ is admissible with respect to Π at c^* with value (x, v') for some $v' \in [v, v^{\text{FB}}(x)]$.

Since $(x, v'), (x', v) \in B\Pi(c^*)$, Lemma 1 implies that $(x, v^{\text{FB}}(x)), (x^{\text{FB}}(v), v) \in B\Pi(c^*)$ if there exist profiles that are E -admissible with respect to Π with these values. Hence, Lemma 6 implies that Proposition 3 is true if for any $x \in [u(r), u(\bar{s})]$, there exists an E -admissible profile with respect to Π at some c with value $(x, v^{\text{FB}}(x))$.

Let $\bar{x}$ satisfy the conditions of Lemma 5. Let $x \in [\bar{x}, u(\bar{s})]$ and consider the stationary contract c_x with $c_x^t(y^t) = u^{-1}(x)$ for all t, y^t . To see that $(x, v^{\text{FB}}(x)) \in B\Pi(c_x, E)$, consider a strategy profile where the agent exerts no effort and no bonuses are paid. Continuation payoffs after any history are given by $(x, v^{\text{FB}}(x))$. They are contained in $\Pi(c_x)$ by Lemma 2 and $U(c_x) = x$. This profile of strategies and continuation payoffs forms an equilibrium of (c_x, E, Π) with payoffs $(x, v^{\text{FB}}(x))$, since $x \geq \bar{x}$ implies $e^{\text{FB}}(x) = 0$.

Now suppose $x < \bar{x}$ and consider the contract c from the proof of Theorem 1, i.e., $s_h = s_l = 0$ and $c_h = c_l = c^*$, where s_h^* and s_l^* are defined in (4) and $c_h^* = c_l^* = \mathbf{0}$. Following Step 2 of the proof, whenever $\bar{s}$ is larger than some $\bar{s}^{\min}$, there exists an equilibrium of subgame (c^*, E) with agent payoff x^* , where x^* satisfies

$$u^{-1}(\bar{x} + \psi(1)) \leq \frac{\delta}{1 - \delta} (v^{\text{FB}}(\bar{x}) - v^{\text{FB}}(x^*)). \quad (12)$$

In this equilibrium, the principal pays a bonus $b_l^* = (v^{\text{FB}}(u(r)) - V(\mathbf{0}))\delta/(1-\delta)$ when the agent rejects any offer and output is low. Since reducing this bonus payment is incentive compatible, there exist equilibria of subgame (c^*, E) with any payoff smaller than x^* but no less than the agent's outside utility. It follows from Lemma 3 and Lemma 2 that there exists an equilibrium of subgame c^* with the same payoff for the agent. Hence, whenever $\bar{s} \geq \bar{s}^{\min}$, for any $\hat{x} \in [u(r), x^*]$, there exists $\hat{v}$ such that $(\hat{x}, \hat{v}) \in \mathcal{E}(c^*)$ and, consequently, $(\hat{x}, v^{\text{FB}}(\hat{x})) \in \Pi(c^*)$.

This can be used to construct the desired equilibrium of (c, E, Π) as follows. Consider the following strategies σ : the agent exerts effort $e^{\text{FB}}(x)$, the principal pays an output-independent bonus $u^{-1}(x + \psi(e^{\text{FB}}(x)))$ following effort $e^{\text{FB}}(x)$, and pays no bonuses otherwise. Continuation payoffs f are given by

- $(x, v^{\text{FB}}(x))$ if no player deviated from σ
- $(u(r), v^{\text{FB}}(u(r)))$ if the agent deviated from σ
- $(x^*, v^{\text{FB}}(x^*))$ if the principal deviated from σ but the agent did not.

By (12), the principal has no profitable deviation. Moreover, the agent's continuation payoff following any deviation equals $u(r) = U(c^*)$, so his deviation payoff cannot exceed $U(c) = u(r)$. Hence, (σ, f) is E -admissible with respect to Π at c with value $(x, v^{\text{FB}}(x))$ whenever $\bar{s} \geq \bar{s}^{\min}$, as required.

A.6 Proof of Proposition 4

Let c^* be a stationary contract with $c^{*,t}(y^{t-1}, y) = s_y^*$ for any $t \in \mathbb{N}$, $y \in Y$, $y^{t-1} \in Y^{t-1}$, where s_h^* and s_l^* are the salaries used in the proof of Theorem 1, defined in (4). Let $\bar{x}$ satisfy the statement of Lemma 5 and let x^* satisfy (12) from the proof of Proposition 3. Define the function $g : [x^*, u(\bar{s})] \rightarrow \mathbb{R}$ as

$$g(z) = (1 - \delta) \left(\sum_y p_y^1 u(s_y^* + b_y^*) \right) + \delta u(r),$$

$$\text{where } b_h^* = 0, b_l^* = \min \left\{ \bar{s} - s_l^*, \frac{\delta}{1 - \delta} (v^{\text{FB}}(u(r)) - v^{\text{FB}}(z)) \right\}. \quad (13)$$

Since $s_y^* + b_y^* \leq \bar{s}$ for all y , $g \leq u(\bar{s})$. Moreover, for high $\bar{s}$, effort $e = 1$ is optimal in (4), so $g(z) = u(r) + (1 - \delta) p_l^1 (u(s_y^* + b_y^*) - u(s_y^*))$. Since $s_l^* \rightarrow 0$ as $\bar{s} \rightarrow \infty$, it follows that $g \geq x^*$ for sufficiently high $\bar{s}$. Hence, by Brouwer's fixed point theorem, g has a fixed point z^* .

Consider the payoff correspondence $\Pi \in \mathcal{P}_0$:

$$\Pi(c^*) = \{(x, v^{\text{FB}}(x)) | u(r) \leq x \leq z^*\}$$

$$\Pi(\hat{c}) = \{(U(\hat{c}), v^{\text{FB}}(U(\hat{c})))\} \quad \text{for any } \hat{c} \in C \setminus \{c^*\}.$$

By Proposition 1 it suffices to show that for any $\hat{c} \in C$ and $(x, v^{\text{FB}}(x)) \in \Pi(\hat{c})$, there exist a semistationary contract c and the following equilibria:

(EQ1) An equilibrium of (c, E, Π) with payoffs $(x, v^{\text{FB}}(x))$.

(EQ2) An equilibrium of $(\hat{c}, E, \Pi)$ with agent payoff $\hat{x}$ such that $x = \max\{\hat{x}, u(r)\}$.

Then an equilibrium of $(\hat{c}, \Pi)$ with the desired properties can be constructed by playing (EQ1) following the acceptance of c and playing (EQ2) following the rejection of any offer. Given these strategies, the agent can guarantee exactly x by rejection or his outside option, so it is incentive compatible to accept c . Moreover, the principal cannot obtain more than $v^{\text{FB}}(x)$ given the agent's guarantee, so offering c is incentive compatible regardless of the outcomes following the acceptance of other contracts.¹⁷

In what follows I construct the contract c and equilibria (EQ1) and (EQ2) given any $\hat{c}$ and $(x, v^{\text{FB}}(x)) \in \Pi(\hat{c})$. I begin with the descriptions of c and (EQ1). If $x \geq \bar{x}$, let c equal the contract c_x from the proof of Proposition 3, and use the associated strategies and continuation payoffs. If $x < \bar{x}$, define c as $s_h = s_l = 0$ and $c_h = c_l = c^*$. Since c^* is stationary, c is semistationary. Consider the following strategies in (c, E, Π) : the agent exerts $e^{\text{FB}}(x)$, the principal rewards $e^{\text{FB}}(x)$ with an output-independent bonus $u^{-1}(x + \psi(e^{\text{FB}}(x)))$, and pays no bonuses otherwise. Continuation payoffs are given by $(z^*, v^{\text{FB}}(z^*))$ if the principal deviated from the prescribed bonus and by $(x, v^{\text{FB}}(x))$ otherwise. The continuation payoffs are contained in $\Pi(c_h^*) = \Pi(c_l^*) = \Pi(c^*)$, since $x < \bar{x} < x^*$ by (12). Since $z^* \geq x^*$, (12) implies that the bonus is incentive compatible. The agent's incentives hold since he receives no compensation if he deviates from $e^{\text{FB}}(x)$. Thus, the strategies form the required equilibrium of (c, E, Π) .

I proceed with the construction of (EQ2). If $x = U(\hat{c})$, use the first-period strategies of $\underline{\sigma}^A(\hat{c}, E)$ from Lemma 2, and continuation payoffs $(U(c'), v^{\text{FB}}(U(c')))$ at any history with continuation contract c' . The principal's incentives are trivial since she pays no bonuses, and the agent's incentives are unchanged relative to $\underline{\sigma}^A(\hat{c}, E)$ since he receives no bonuses and his continuation payoffs are unchanged. Hence, the above strategies form an equilibrium of $(\hat{c}, E, \Pi)$ with agent payoff $U^E(\hat{c})$. Since $U(\hat{c}) = \max\{U^E(\hat{c}), u(r)\}$, this forms the desired equilibrium (EQ2).

If $x \neq U(\hat{c})$, then $\hat{c} = c^*$ and $x \in (u(r), z^*]$. Then the desired equilibrium of (c^*, E, Π) can be constructed as follows. The agent exerts effort $e = 1$, the principal pays a bonus $b_l \in [0, b_l^*]$ if effort is 1 and output is low, and pays no bonus otherwise. Continuation payoffs are $(z^*, v^{\text{FB}}(z^*)) \in \Pi(c^*)$ if the principal reneges on the bonus, and $(u(r), v^{\text{FB}}(u(r))) \in \Pi(c^*)$ otherwise. Hence, the agent receives $u(r)$ if $b_l = 0$ and z^* if $b_l = b_l^*$, since z^* is a fixed point of g . Thus, $b_l \in [0, b_l^*]$ can be chosen to make the agent's payoff equal to x . The incentive compatibility of the bonuses follows from (12). Recall that for high $\bar{s}$, effort 1 is optimal in (4), so the agent obtains at most $u(r)$ by deviating. Thus, the above strategies form an equilibrium of (c^*, E, Π) with agent payoff x , as required.

¹⁷Strictly speaking, condition (ii) in the statement of Proposition 4 must be verified by establishing the existence of Pareto optimal equilibrium payoffs in any subgame (c', E) in the class of equilibria satisfying (i). This class is nonempty, as it contains $\underline{\sigma}^A(c', E)$. Furthermore, it can be shown by adapting arguments from APS that the set of payoffs of equilibria in subgame (c', E) satisfying (i) is compact.

APPENDIX B: PROOFS FOR THE EXTENDED MODEL WITH LIMITED LIABILITY FOR THE PRINCIPAL

Let C be the space of all contracts and let $\mathcal{P}_0$ be the space of correspondences $\Pi: C \rightrightarrows F$, where $F = [u(r), x^{\text{FB}}(0)] \times [0, v^{\text{FB}}(u(r))]$ is the set of feasible and individually rational payoffs. For any $\Pi \in \mathcal{P}_0$, define the auxiliary games (c^*, Π) and (c, E, Π) as well as their payoff sets $B\Pi(c^*)$ and $B\Pi(c, E)$ analogously to Section 2.3. As in the baseline model, $\mathcal{E} \in \mathcal{P}_0$ denotes the equilibrium payoff correspondence, and a correspondence Π is self-generating if $\Pi \subseteq B\Pi$.¹⁸ The following results can be adapted from APS.

PROPOSITION 6 (Self-generation). *If Π is self-generating, then $B\Pi \subseteq \mathcal{E}$.*

PROPOSITION 7 (Factorization). *We have $B\mathcal{E} = \mathcal{E}$.*

Let $\mathcal{P}$ be the set of nonempty, compact-valued, and upper hemicontinuous payoff correspondences Π in $\mathcal{P}_0$. For any $c^* \in C$, let $\underline{x}^\Pi(c^*)$ denote the agent's worst payoff in $\Pi(c^*)$, and let $\underline{v}^\Pi(c^*)$ and $\bar{v}^\Pi(c^*)$ denote the principal's worst and best payoffs in $\Pi(c^*)$, respectively. The following sections characterize $B\Pi$ and $B\Pi(\cdot, E)$ for any payoff correspondence $\Pi \in \mathcal{P}$.

Equilibrium set of (c, E, Π)

For any $c \in C$ and $\Pi \in \mathcal{P}$, the payoff set $B\Pi(c, E)$ consists of the values of all profiles that are E -admissible with respect to Π at c , as defined below.

DEFINITION 6 (E -admissibility). Let $c \in C$ and $\Pi \in \mathcal{P}$. A profile of actions (e, b_h, b_l) and continuation payoffs $(x_h, v_h) \in \Pi(c_h)$, $(x_l, v_l) \in \Pi(c_l)$ is E -admissible with respect to Π at c with value (x, v) if

$$\begin{aligned} x &= \sum_y p_y^e [(1 - \delta)(u(s_y + b_y) - \psi(e)) + \delta x_y] \\ v &= \sum_y p_y^e [(1 - \delta)(y - s_y - b_y) + \delta v_y] \\ x &\geq \max_{e' \in [0, 1]} \sum_y p_y^{e'} [(1 - \delta)(u(s_y) - \psi(e')) + \delta \underline{x}^\Pi(c_y)] \end{aligned} \tag{14}$$

$$0 \leq b_y \leq \frac{\delta}{1 - \delta} (v_y - \underline{v}^\Pi(c_y)) \quad \text{for all } y. \tag{15}$$

To see this, notice that in any equilibrium of (c, E, Π) , it is without loss of generality that each deviation is punished with the worst continuation equilibrium for the deviator (Abreu (1988)). Hence, the principal receives her worst continuation payoff $\underline{v}^\Pi(c_y)$ when she reneges on the equilibrium bonus b_y for effort e following output y . If the agent deviates to effort e' , he receives no bonuses and his worst continuation payoff

¹⁸Set operations on correspondences are defined pointwise.

$\underline{x}^\Pi(c_y)$ following any output y . Hence, the incentive constraints for effort and bonuses are given by (14) and (15), respectively. Notice that $B\Pi(\cdot, E)$ is compact-valued and upper hemicontinuous. Let $\underline{x}(c, E, \Pi)$ and $\underline{v}(c, E, \Pi)$ denote, respectively, the agent's and the principal's worst payoffs in $B\Pi(c, E)$.

Equilibrium set of (c^, Π)*

Let $\Pi \in \mathcal{P}$ and $c^*, c \in C$. Let (c^*, c, Π) denote the subgame of (c^*, Π) following an offer c by the principal. The principal's worst payoff in this subgame is obtained in one of three types of equilibria. In the first type, the agent accepts c , so without loss of generality he receives his worst equilibrium payoff $\underline{x}(c^*, E, \Pi)$ following rejection. The lowest payoff the principal can receive in such an equilibrium equals

$$\underline{v}(c^*, c, \Pi; A) = \inf\{v \mid (x, v) \in B\Pi(c, E), x \geq \max\{u(r), \underline{x}(c^*, E, \Pi)\}\},$$

where $\inf \emptyset = \infty$; otherwise the infimum is attained due to the compactness of $B\Pi(c, E)$.

In the second type of equilibrium, the agent rejects c . The worst payoff for the principal obtained in this manner equals $\underline{v}(c, c^*, \Pi; A)$ by symmetry of the subgames following acceptance and rejection.

Finally, consider an equilibrium where the agent takes his outside option. Such an equilibrium exists if and only if there are equilibria following both acceptance and rejection where the agent's payoff is no larger than his outside option. Thus, the principal's worst equilibrium payoff in subgame (c^*, c, Π) is given by

$$\underline{v}(c^*, c, \Pi) = \begin{cases} \min\{0, \underline{v}(c^*, c, \Pi; A), \underline{v}(c, c^*, \Pi; A)\} \\ \quad \text{if } u(r) \geq \max\{\underline{x}(c, E, \Pi), \underline{x}(c^*, E, \Pi)\} \\ \min\{\underline{v}(c^*, c, \Pi; A), \underline{v}(c, c^*, \Pi; A)\} \\ \quad \text{otherwise.} \end{cases}$$

In the principal's worst equilibrium in (c^*, Π) , she receives her lowest payoff in subgame (c^*, c, Π) following any offer $c \in C$. Hence,

$$\underline{v}^{B\Pi}(c^*) = \max\left\{0, \sup_{c \in C} \underline{v}(c^*, c, \Pi)\right\}. \quad (16)$$

The payoff set $B\Pi(c^*)$ consists of the values of all profiles admissible with respect to Π at c^* , as defined below.

DEFINITION 7 (Admissibility). Let $c^* \in C$ and $\Pi \in \mathcal{P}$. A profile “out” with value $(u(r), 0)$ is admissible with respect to Π at c^* if $\underline{v}^{B\Pi}(c^*) = 0$. A profile $c \in C$ with value $(x, v) \in B\Pi(c, E)$ is admissible with respect to Π at c^* if

$$v \geq \underline{v}^{B\Pi}(c^*) \quad (17)$$

$$\text{and } x \geq \max\{u(r), \underline{x}(c^*, E, \Pi)\}. \quad (18)$$

To see this, consider any equilibrium of (c^*, Π) where the outside options are taken by either player. It follows that $\underline{v}^{B\Pi}(c^*) = 0$; otherwise, there exists a contract the principal can offer to guarantee a positive payoff. Moreover, if $\underline{v}^{B\Pi}(c^*) = 0$, there is an equilibrium where the principal shuts down, because she receives her worst continuation payoff following any contract offer.

Now consider an equilibrium of (c^*, Π) with payoffs (x, v) , where the outside options are not taken. It is without loss of generality that the principal's contract offer c is accepted by the agent, as any equilibrium with rejection can be replicated with an offer equal to c^* . Thus, $(x, v) \in B\Pi(c, E)$. Since $\underline{v}^{B\Pi}(c^*)$ equals the principal's worst payoff from deviation at the contract offer stage, her incentive compatibility is equivalent to (17). As for the agent, accepting c is incentive compatible if and only if (18) holds, as he receives his worst payoff following rejection.

B.1 Proof of Proposition 5

The key intermediate result in the proof of Proposition 5 is Lemma 7 below, which establishes that equilibrium payoffs in subgames from the agent's effort can be modified continuously to the benefit of one of the players.

LEMMA 7. *Let $\Pi \in \mathcal{P}$. Consider any contract c and any equilibrium of subgame (c, E, Π) with payoffs (x, v) such that $x \neq -\infty$. For any $\varepsilon > 0$, the following situations exist:*

- *A contract c' and an equilibrium of (c', E, Π) with payoffs (x', v') such that $x < x' < x + \varepsilon$ and $v - \varepsilon < v' < v + \varepsilon$.*
- *A contract c' and an equilibrium of (c', E, Π) with payoffs (x', v') such that $x - \varepsilon < x' < x + \varepsilon$ and $v < v' < v + \varepsilon$.*

PROOF. Consider a profile σ of strategies (e, b_h, b_l) and continuation payoffs $(x_y, v_y)_y$ that are E -admissible with respect to Π at c .

For any e', s'_h, s'_l , let

$$\begin{aligned} A(e', s'_h, s'_l) &= \sum_y p_y^{e'} [(1 - \delta)(u(s'_y + b_y) - \psi(e)) + \delta x_y] \\ \underline{A}(e', s'_h, s'_l) &= \sum_y p_y^{e'} [(1 - \delta)(u(s'_y) - \psi(e)) + \delta \underline{x}^\Pi(c_y)] \\ P(e', s'_h, s'_l) &= \sum_y p_y^{e'} [(1 - \delta)(y - s'_y - b_y) + \delta v_y]. \end{aligned}$$

In what follows, I obtain a contract c' from c by changing the first-period salaries to s'_h and s'_l . Then the profile of strategies (e', b_h, b_l) and continuation payoffs $(x_y, v_y)_y$ is E -admissible with respect to Π at c' with value $(A(e', s'_h, s'_l), P(e', s'_h, s'_l))$ if and only if $A(e', s'_h, s'_l) \geq \max_{\hat{e}} \underline{A}(\hat{e}, s'_h, s'_l)$. In particular, when $c' = c$ the E -admissibility of σ implies that

$$A(e, s_h, s_l) \geq \max_{\hat{e}} \underline{A}(\hat{e}, s_h, s_l). \quad (19)$$

Let $\varepsilon > 0$. I begin by constructing the equilibrium that improves the principal's payoff. Suppose $s_h > 0$ and $s_l > 0$. Consider $s'_h \in (0, s_h)$ and $s'_l \in (0, s_l)$ such that

$$u(s_h) - u(s'_h) = u(s_l) - u(s'_l).$$

The profile $(e, b_h, b_l), (x_y, v_y)_y$ is E -admissible with respect to Π at c' , since

$$\begin{aligned} \max_{\hat{e} \in [0,1]} \underline{A}(\hat{e}, s'_h, s'_l) &= \max_{\hat{e} \in [0,1]} \underline{A}(\hat{e}, s_h, s_l) - (1 - \delta)(u(s_h) - u(s'_h)) \\ &\leq A(e, s_h, s_l) - (1 - \delta) \sum_y p_y^e (u(s_y + b_y) - u(s'_y + b_y)) \\ &= A(e, s'_h, s'_l), \end{aligned}$$

where the inequality follows from (19) and the concavity of u . Thus, there is an equilibrium of (c', E, Π) with payoffs $(A(e, s'_h, s'_l), P(e, s'_h, s'_l))$. Since A and P are continuous, these payoffs are within ε of (x, v) whenever (s'_h, s'_l) is sufficiently close to (s_h, s_l) . Moreover, $P(e, \cdot, \cdot)$ is strictly decreasing in both arguments, so $P(e, s'_h, s'_l) > v$, as required.

Now suppose $s_y = 0$ for some y . Without loss of generality, let $y = h$. Since $u(0) = -\infty$ and $x \neq -\infty$, it follows that $b_h > 0$. Let $b'_h < b_h$ and $b'_l = b_l$. The profile $(e, b'_h, b'_l), (x_y, v_y)_y$ is E -admissible with respect to Π at c since lower bonuses are incentive compatible and the right-hand side of (14) equals $-\infty$ due to $u(s_h) = -\infty$. Hence, there exists an equilibrium of subgame (c, E, Π) where the only difference in on-path behavior relative to the equilibrium with payoffs (x, v) lies in the smaller bonus for output h . Thus, the principal's payoff is improved at the expense of the agent. The equilibrium payoffs are within ε of (x, v) when b'_h is sufficiently close to b_h .

It remains to construct an equilibrium that improves the agent's payoff. Suppose (19) holds at equality, $b_h = b_l = 0$, and $x_y = \underline{x}^\Pi(c_y)$ for all y . Thus, $A = \underline{A}$. Let $s'_h > s_h$ and $s'_l = s_l$. Let e' be a maximizer of $\underline{A}(\cdot, s'_h, s_l)$. The profile $(e', b_h, b_l), (x_y, v_y)_y$ is E -admissible with respect to Π at c' , since

$$A(e', s'_h, s_l) = \underline{A}(e', s'_h, s_l) = \max_{\hat{e} \in [0,1]} \underline{A}(\hat{e}, s'_h, s_l).$$

Hence, there exists an equilibrium of (c', E, Π) with payoffs $(A(e', s'_h, s_l), P(e', s'_h, s_l))$. Since $s'_h > s_h$, it follows that $A(e', s'_h, s_l) = \underline{A}(e', s'_h, s_l) > \underline{A}(e, s_h, s_l) = A(e, s_h, s_l)$. Moreover, A and P are continuous, so the equilibrium payoffs are within ε of (x, v) when s'_h is sufficiently close to s_h .

Now suppose (19) holds at equality and there exists y such that at least one of $b_y > 0$ and $x_y > \underline{x}^\Pi(c_y)$ holds. Thus, $A > \underline{A}$ and

$$\max_{e' \in [0,1]} A(e', s_h, s_l) > \max_{\hat{e} \in [0,1]} \underline{A}(\hat{e}, s_h, s_l) = A(e, s_h, s_l).$$

Since A and P are continuous, there exists e' such that

$$\begin{aligned} x &= A(e, s_h, s_l) < A(e', s_h, s_l) < x + \varepsilon \\ v - \varepsilon &< P(e', s_h, s_l) < v + \varepsilon. \end{aligned}$$

It follows from (19) that $A(e', s_h, s_l) > \max_{\hat{e} \in [0,1]} A(\hat{e}, s_h, s_l)$. Thus, $(e', b_h, b_l), (x_y, v_y)_y$ is E -admissible with respect to Π at c with the desired payoffs.

Finally, suppose that (19) holds at a strict inequality. Let $s'_h > s_h$ and $s'_l = s_l$. By the maximum theorem, $A(e, \cdot, s_l)$ and $\max_{\hat{e} \in [0,1]} A(\hat{e}, \cdot, s_l)$ are continuous. It follows that when s'_h is sufficiently close to s_h , the profile $(e, b_h, b_l), (x_y, v_y)_y$ is E -admissible with respect to Π at c' . Hence, there exists an equilibrium of (c', E, Π) with payoffs within ε of (x, v) such that the agent's payoff is higher than x . $\square$

Recall from the discussion following Proposition 3 that showing the attainability of efficient equilibrium payoffs in a strongly optimal equilibrium amounts to the statement of Lemma 8 below for $\Pi = \mathcal{E}$. Thus, Proposition 5 follows from $\mathcal{E} \in \mathcal{P}$, which is established in the proof of Proposition 8 in Appendix B.2.

LEMMA 8. *Let $\Pi \in \mathcal{P}$ and $c^* \in C$. Consider any equilibrium of (c^*, Π) with payoffs (x^*, v^*) . There exist an efficient equilibrium of (c^*, Π) where the agent's payoff is x^* and an efficient equilibrium of (c^*, Π) where the principal's payoff is v^* .*

PROOF. Consider any equilibrium of (c^*, Π) with payoffs (x^*, v^*) . Let

$$\bar{v} = \max\{v \mid (x, v) \in B\Pi(c^*), x \geq x^*\}$$

$$\bar{x} = \max\{x \mid (x, \bar{v}) \in B\Pi(c^*)\},$$

which are well defined by Lemma 12 in Appendix B.2. An efficient equilibrium of (c^*, Π) with agent payoff x^* exists if and only if $\bar{x} = x^*$.

Toward a contradiction, suppose $\bar{x} > x^*$. It follows that $(\bar{x}, \bar{v}) \neq (u(r), 0)$, so there exists a profile c with value $(\bar{x}, \bar{v})$ admissible with respect to Π at c^* . By Lemma 7, there exists an equilibrium of (c, E, Π) with payoffs (x, v) such that $x > x^*$ and $v > \bar{v} \geq v^*$. Since $(x^*, v^*) \in B\Pi(c^*)$, $x > x^*$, and $v > v^*$, it follows from (17) and (18) that $(x, v) \in B\Pi(c^*)$, contradicting the maximality of $\bar{v}$. Hence, there exists an efficient equilibrium of (c^*, Π) where the agent receives x^* . A symmetric argument for the principal completes the proof. $\square$

The results of this section imply that any correspondence containing the equilibrium payoffs generates a nontrivial Pareto frontier of payoffs from the agent's effort including the outside utilities of each player. This is shown in Lemma 9 below, which is used to prove Theorem 2 in the following section.

LEMMA 9. *Let $\Pi \in \mathcal{P}$ such that $\Pi \supseteq \mathcal{E}$. Then there exists $\bar{x} > u(r)$ and a function $f : [u(r), \bar{x}] \rightarrow [0, v^{FB}(u(r))]$ such that $f(\bar{x}) = 0$, and $(x, f(x))$ is Pareto optimal in $\cup_{c \in C} B\Pi(c, E) \cap F$ for any $x \in [u(r), \bar{x}]$.*

PROOF. Let $\Pi \in \mathcal{P}$ such that $\Pi \supseteq \mathcal{E}$. The set $G_\Pi := \cup_{c \in C} B\Pi(c, E) \cap F$ is compact, since $B\Pi(\cdot, E)$ is compact-valued and upper hemicontinuous. By assumption, there exists $(x, v) \in \mathcal{E}(\mathbf{0})$ with $(x, v) \neq (u(r), 0)$. By Proposition 7 and admissibility, $(x, v) \in B\mathcal{E}(c, E) \cap F$ for some $c \in C$. Hence, the monotonicity of $B(\cdot, E)$ implies that $(x, v) \in G_\Pi$. The result then follows from Lemma 7 and Lemma 8. $\square$

B.2 Proof of Theorem 2

The proof of Theorem 2 is based on the algorithmic result of APS stated below.

PROPOSITION 8 (APS algorithm). *Let $\Pi_1 = F$ and $\Pi_{n+1} = B\Pi_n$ for all $n \in \mathbb{N}$. Then $\Pi_{n+1} \supseteq \Pi_n$ for all n and $\bigcap_n \Pi_n = \mathcal{E}$.*

In stationary repeated games, this result is shown in two steps. First, the operator B maps compact payoff sets (instead of correspondences) to compact sets, so the algorithm yields a decreasing sequence of compact sets. Second, this sequence converges to a compact self-generating set. In the stochastic game considered here, the second step requires the upper hemicontinuity of each correspondence Π_n . Thus, Lemma 12, the analogue to the first step, shows that B preserves compact-valuedness and upper hemicontinuity, as well as the properties of the following class $\mathcal{P}^+ \subseteq \mathcal{P}$. A payoff correspondence $\Pi \in \mathcal{P}$ is in $\mathcal{P}^+$ if $\Pi \supseteq \mathcal{E}$ and there exists an optimal punishment contract for Π in the following sense.

DEFINITION 8. A contract $c^\Pi \in C$ is an optimal punishment contract for $\Pi \in \mathcal{P}$ if $\underline{v}^{B\Pi}(c^\Pi) = 0$ and $\underline{x}(c^\Pi, E, \Pi) \leq u(r)$.

By admissibility, an optimal punishment contract c^Π satisfies $B\Pi(c^\Pi) \supseteq B\Pi(c^*)$ for all $c^* \in C$. Lemma 10 is used to show that each Π_n has an optimal punishment contract c^{Π_n} . This is essential to establishing the compact-valuedness of each Π_n and their limit $\mathcal{E}$. For the latter, one needs to obtain a convergent subsequence of a sequence of contracts (c^n) such that c^n is admissible with respect to Π_n . This is not immediate due to the unboundedness of C , so the convergence proceeds in the following two steps. First, the continuation contracts of each c^{n+1} are equal to the optimal punishment contract c^{Π_n} without loss of generality. The sequence (c^{Π_n}) is shown to converge to a stationary optimal punishment contract for $\mathcal{E}$ (Corollary 2). Second, the first-period salaries of each c^n are uniformly bounded, as shown in Lemma 11.

LEMMA 10. *Let $\Pi \in \mathcal{P}$ and let $\varepsilon > 0$ such that $\bar{v}^\Pi(c^*) - \underline{v}^\Pi(c^*) \geq \varepsilon$ for some $c^* \in C$. Then there exists an optimal punishment contract c^Π for Π . In addition, if $\Pi \supseteq \hat{\Pi} \supseteq \mathcal{E}$ for some $\hat{\Pi} \in \mathcal{P}$, there exists $\eta(\hat{\Pi}) > 0$ independent of Π and ε such that $\bar{v}^{B\Pi}(c^\Pi) - \underline{v}^{B\Pi}(c^\Pi) \geq \eta(\hat{\Pi})$.*

PROOF. Fix any Π , c^* , and ε as in the statement of the lemma. Let $(\bar{x}, \bar{v})$, $(\underline{x}, \underline{v}) \in \Pi(c^*)$ such that $\bar{v} = \bar{v}^\Pi(c^*)$ and $\underline{v} = \underline{v}^\Pi(c^*)$. Let c^Π be a contract with continuation contracts $c_h^\Pi = c_l^\Pi = c^*$ and salaries s_h^Π and s_l^Π such that

$$(1 - \delta) \sum_y p_y^1(u(s_y^\Pi) - \psi(1)) + \delta x^{\text{FB}}(0) = u(r). \quad (20)$$

Since u is unbounded, there exist multiple salary pairs (s_h^Π, s_l^Π) satisfying (20), where s_h^Π can be arbitrarily large and $s_l^\Pi \rightarrow 0$ as $s_h^\Pi \rightarrow \infty$.

Consider an equilibrium of (c^Π, E, Π) where no bonuses are paid and continuation payoffs are constant, the latter being possible due to $c_h^\Pi = c_l^\Pi$. The agent's continuation

payoffs are upper bounded by $x^{\text{FB}}(0)$, his highest payoff in F , since $\Pi \in \mathcal{P}$. For s_h^Π sufficiently large, s_l^Π is small, so effort 1 is optimal. It follows from (20) that $\underline{x}(c^\Pi, E, \Pi) \leq u(r)$.

Now consider the following equilibrium of (c^Π, E, Π) . The agent exerts effort 1 and the principal pays no bonuses unless effort equals 1 and output is low; in this case, she pays $b = \varepsilon\delta/(1 - \delta)$. Continuation payoffs are given by $(\bar{x}, \bar{v})$ if the principal did not deviate and by $(\underline{x}, \underline{v})$ if she did. Since $\varepsilon \leq \bar{v} - \underline{v}$, the principal's incentives hold. The agent's payoff is at least

$$(1 - \delta)(p_h^1 u(s_h^\Pi) + p_l^1 u(s_l^\Pi + b) - \psi(1)) + \delta u(r).$$

When s_h^Π is sufficiently high, s_l^Π is small, so the marginal utility of the bonus b becomes arbitrarily high. It follows from (20) that there exists an equilibrium of (c^Π, E, Π) with an arbitrarily high payoff for the agent. In particular, if the agent receives more than $x^{\text{FB}}(0)$, then $\underline{v}(c, c^\Pi, \Pi; A) \leq 0$ for any $c \in C$. Thus, $\underline{v}^{B\Pi}(c^\Pi) = 0$, as required.

Consider any $\hat{\Pi} \in \mathcal{P}$ such that $\Pi \supseteq \hat{\Pi} \supseteq \mathcal{E}$. By Lemma 9, there exist $c, c' \in C$, $(x, v) \in B\hat{\Pi}(c, E) \cap F$, and $(x', v') \in B\hat{\Pi}(c', E) \cap F$ such that $v' > v$. Since $\Pi \supseteq \hat{\Pi}$, it follows from the properties of c^Π and admissibility that $(x, v), (x', v') \in B\Pi(c^\Pi)$. The proof is then completed by setting $\eta(\hat{\Pi}) = v' - v$. $\square$

LEMMA 11. *Let $\Pi \in \mathcal{P}$ and $c \in C$. If a profile c with some value (x, v) is admissible with respect to Π at some $c^* \in C$, then s_h and s_l are upper bounded by*

$$\bar{s} = \frac{h + \frac{\delta}{1 - \delta} v^{\text{FB}}(u(r))}{\min\{p_l^1, p_h^0\}}.$$

PROOF. A profile c with value $(x, v) \in B\Pi(c, E)$ is admissible with respect to $\Pi \in \mathcal{P}$ at c^* only if $v \geq \underline{v}^{B\Pi}(c^*) \geq 0$. However, if $\max\{s_h, s_l\} > \bar{s}$, every equilibrium of (c, E, Π) has a negative payoff for the principal. To see this, notice that her continuation payoff is at most $v^{\text{FB}}(u(r))$, since $\Pi \in \mathcal{P}$. Hence, her payoff is upper bounded by

$$\begin{aligned} & (1 - \delta) \left(\max_e \sum_y p_y^e y - \min_e \sum_y p_y^e s_y \right) + \delta v^{\text{FB}}(u(r)) \\ & < (1 - \delta) (h - \min\{p_l^1, p_h^0\} \bar{s}) + \delta v^{\text{FB}}(u(r)) = 0. \end{aligned} \quad \square$$

LEMMA 12. *Let $\Pi \in \mathcal{P}^+$ such that $\Pi \supseteq B\Pi$. Then $B\Pi \in \mathcal{P}^+$.*

PROOF. Since $\Pi \supseteq \mathcal{E}$, Proposition 7 and the monotonicity of B imply that $B\Pi \supseteq \mathcal{E}$. Hence, it follows from the assumption of equilibrium existence that $B\Pi$ is nonempty.¹⁹ Moreover, $\Pi \supseteq B\Pi \supseteq \mathcal{E}$ and Lemma 10 imply that an optimal punishment contract for $B\Pi$ exists if $B\Pi \in \mathcal{P}$, which is demonstrated below.

To see that $B\Pi$ is compact-valued, consider any $c^* \in C$ and a sequence of profiles (σ_n) with value (x_n, v_n) admissible with respect to Π at c^* . Since $(x_n, v_n) \in F$ for all n ,

¹⁹Existence can also be established directly by a construction similar to Lemma 2.

there exists a subsequence along which the values converge to some $(x, v) \in F$. It suffices to show that (x, v) is the value of a profile admissible with respect to Π at c^* . If there exists a (further) subsequence along which $\sigma_n = \text{out}$, this is immediate. Otherwise, there exists a subsequence along which $\sigma_n = c^n$ for some sequence of contracts (c^n) . Since $\Pi \in \mathcal{P}^+$, there exists c_Π^* such that $\Pi(c_\Pi^*) \supseteq \Pi(c^n)$ for all n . Thus, it is without loss of generality that $c_h^n = c_l^n = c_\Pi^*$ for all n . Moreover, by Lemma 11, $s_y^n \leq \bar{s}$ for all n, y . It follows that c^n converges to some contract c , taking a subsequence if necessary. By the upper hemicontinuity of $B\Pi(\cdot, E)$, a profile c with value (x, v) is admissible with respect to Π at c^* , as required.

The upper hemicontinuity of $B\Pi$ obtains if for any c^* and any sequence of profiles c^n with value (x_n, v_n) admissible with respect to Π at c^* such that $c^n \rightarrow c$ and $(x_n, v_n) \rightarrow (x, v)$, a profile c with value (x, v) is admissible with respect to Π at c^* . It suffices to show that $\underline{x}(\cdot, E, \Pi)$ and $\underline{v}^{B\Pi}$ are lower semicontinuous so that (17) and (18) from the definition of admissibility hold in the limit as $c^n \rightarrow c$. The lower semicontinuity of $\underline{x}(\cdot, E, \Pi)$ follows from the upper hemicontinuity of $B\Pi(\cdot, E)$. Since the pointwise supremum of lower semicontinuous functions is lower semicontinuous, (16) implies that $\underline{v}^{B\Pi}$ is lower semicontinuous if $\underline{v}(\cdot, c, \Pi; A)$ and $\underline{v}(c, \cdot, \Pi; A)$ are lower semicontinuous for any c . This follows from the lower semicontinuity of $\underline{x}(\cdot, E, \Pi)$ and the upper hemicontinuity of $B\Pi(\cdot, E)$. $\square$

PROOF OF PROPOSITION 8. Since $\Pi_1 = F$, it is immediate that $\Pi_1 \in \mathcal{P}^+$ and $\Pi_1 \supseteq \Pi_2$. It follows from Lemma 12 and the monotonicity of B that $\Pi_n \in \mathcal{P}^+$ and $\Pi_n \supseteq \Pi_{n+1}$ for all $n \in \mathbb{N}$. Let $\Pi := \cap_{n \in \mathbb{N}} \Pi_n$. As the intersection of nested compact sets, $\Pi(c^*)$ is compact for any c^* by the finite intersection property. Similar arguments establish that Π is upper hemicontinuous. Thus, $\Pi \in \mathcal{P}$, so the definitions of admissibility and E -admissibility can be used to characterize $B\Pi$ and $B\Pi(\cdot, E)$, respectively. Since $\Pi \supseteq \mathcal{E}$, Proposition 6 implies that $\Pi = \mathcal{E}$ obtains if Π is self-generating. Toward this goal, consider any $c^* \in C$ and any $(x, v) \in \Pi(c^*)$.

Case 1: $(x, v) = (u(r), 0)$.

Then $\underline{v}^{B\Pi_n}(c^*) = 0$ for all n since $(x, v) \in \Pi(c^*) \subseteq \Pi_n(c^*)$. Hence, $(u(r), 0) \in B\Pi(c^*)$ obtains if $\lim_{n \rightarrow \infty} \underline{v}^{B\Pi_n}(c^*) = \underline{v}^{B\Pi}(c^*)$. This is shown in the sequel by a backward induction argument through the stage game.

Step 1: $\underline{x}^{\Pi_n}(c^*) \nearrow \underline{x}^\Pi(c^*)$ and $\underline{v}^{\Pi_n}(c^*) \nearrow \underline{v}^\Pi(c^*)$ for any $c^* \in C$. By the monotonicity of (Π_n) (i.e., $\Pi_{n+1} \supseteq \Pi_n$), $\underline{x}^{\Pi_n}(c^*)$ is an increasing sequence upper bounded by $\underline{x}^\Pi(c^*)$. Since each Π_n is compact-valued, there exists a sequence $(v_n) \rightarrow v$ such that $(\underline{x}^{\Pi_n}(c^*), v_n) \in \Pi_n$ for all n . Thus, for any $m \in \mathbb{N}$, the compactness of $\Pi_m(c^*)$ and the monotonicity of (Π_n) imply that $(\lim \underline{x}^{\Pi_n}(c^*), v) \in \Pi_m(c^*)$. Hence, $\lim \underline{x}^{\Pi_n}(c^*) \geq \underline{x}^\Pi(c^*)$ and the result obtains. A similar argument establishes that $\underline{v}^{\Pi_n}(c^*) \nearrow \underline{v}^\Pi(c^*)$.

Step 2: $\cap_{n \in \mathbb{N}} B\Pi_n(c, E) = B\Pi(c, E)$ for any $c \in C$. By the monotonicity of (Π_n) , it suffices to show that any sequence σ_n of profiles E -admissible with respect to Π_n at c with value (x, v) has a subsequence converging to a profile E -admissible with respect to Π at c . To this end, notice that the continuation payoffs of any E -admissible profile lie in the compact set F and bonuses are upper bounded by $v^{\text{FB}}(u(r))\delta/(1 - \delta)$. Thus,

there exists a subsequence of σ_n converging to a profile σ . Since $\underline{x}^{\Pi_n}(c_y) \rightarrow \underline{x}^{\Pi}(c_y)$ and $\underline{v}^{\Pi_n}(c_y) \rightarrow \underline{v}^{\Pi}(c_y)$ for all y , σ is E -admissible with respect to Π at c .

Step 3: $\underline{x}(c, E, \Pi_n) \nearrow \underline{x}(c, E, \Pi)$ for any $c \in C$. By the monotonicity of (Π_n) , $\underline{x}(c, E, \Pi_n)$ is an increasing sequence upper bounded by $\underline{x}(c, E, \Pi)$. Consider any sequence $(x_n, v_n) \rightarrow (x, v)$ such that $(x_n, v_n) \in B\Pi_n(c, E)$ for all n . For any m , the monotonicity of (Π_n) implies that $(x, v) \in B\Pi_m(c, E)$, since $B\Pi_m(c, E)$ is compact. Thus, $(x, v) \in B\Pi(c, E)$ follows from Step 2, and $\lim \underline{x}(c, E, \Pi_n) \geq \underline{x}(c, E, \Pi)$, as required.

Step 4: $\underline{v}(c^*, c, \Pi_n; A) \nearrow \underline{v}(c^*, c, \Pi; A)$ for any $c^*, c \in C$. By the monotonicity of (Π_n) , $\underline{v}(c^*, c, \Pi_n; A)$ is an increasing sequence upper bounded by $\underline{v}(c^*, c, \Pi; A)$. Consider any sequence $(x_n, v_n) \rightarrow (x, v)$ such that $(x_n, v_n) \in B\Pi_n(c, E)$ and $x_n \geq \max\{u(r), \underline{x}(c^*, E, \Pi_n)\}$ for all n . It follows from Steps 2 and 3 that $(x, v) \in B\Pi(c, E)$ and $x \geq \max\{u(r), \underline{x}(c^*, E, \Pi)\}$. Thus, $\lim \underline{v}(c^*, c, \Pi_n; A) \geq \underline{v}(c^*, c, \Pi; A)$ in the event that $\underline{v}(c^*, c, \Pi_n; A)$ does not equal ∞ for any n ; otherwise, the inequality holds trivially.

Step 5: It follows from Steps 3 and 4 that $\underline{v}(c^*, c, \Pi_n) \nearrow \underline{v}(c^*, c, \Pi)$ for any $c^*, c \in C$. Thus, for any c^* ,

$$\lim_{n \rightarrow \infty} \sup_{c \in C} \underline{v}(c^*, c, \Pi_n) = \sup_{n \in \mathbb{N}} \sup_{c \in C} \underline{v}(c^*, c, \Pi_n) = \sup_{c \in C} \sup_{n \in \mathbb{N}} \underline{v}(c^*, c, \Pi_n) = \sup_{c \in C} \underline{v}(c^*, c, \Pi).$$

Thus, (16) implies that $\underline{v}^{B\Pi_n}(c^*) \rightarrow \underline{v}^{B\Pi}(c^*)$, as required.

Case 2: $(x, v) \neq (u(r), 0)$.

Then for each n , there exists a profile c^n with value (x, v) admissible with respect to Π_n at c^* .

An inductive application of Lemma 10 shows that the optimal punishment contract c^{Π_n} for each Π_n can be chosen to satisfy

$$\bar{v}^{B\Pi_n}(c^{\Pi_n}) - \underline{v}^{B\Pi_n}(c^{\Pi_n}) \geq \eta(\Pi).$$

Thus, by Lemma 10, it is possible to choose optimal punishment contracts (c^{Π_n}) with the same first-period salaries, denoted by s_h^{Π} and s_l^{Π} . Moreover, $B\Pi_n(c^{\Pi_n}) \supseteq B\Pi_n(c_y^{\Pi_{n+1}})$ for any $y \in Y$, $n \in \mathbb{N}$. Thus, $c_h^{\Pi_{n+1}} = c_l^{\Pi_{n+1}} = c^{\Pi_n}$ without loss of generality. It follows that $c^{\Pi_n} \rightarrow c^{\Pi}$, where $c^{\Pi, t}(y^{t-1}, y) = s_y^{\Pi}$ for all $t \in \mathbb{N}$, $y^{t-1} \in Y^{t-1}$, and $y \in Y$.²⁰

For each n , the properties of the optimal punishment contract c^{Π_n} imply that $c_h^{n+1} = c_l^{n+1} = c^{\Pi_n}$ without loss of generality. Moreover, by Lemma 11, $s_y^n \leq \bar{s}$ for all $y \in Y$, $n \in \mathbb{N}$. Thus, (c^n) converges to a contract c with $c_h = c_l = c^{\Pi}$, taking a subsequence if necessary.

The admissibility of c^n implies that $(x, v) \in B\Pi_n(c^n, E)$ for all n . By the monotonicity of (Π_n) , $(x, v) \in B\Pi_n(c^m, E)$ for all $n \in \mathbb{N}$, $m \geq n$. Hence, the upper hemicontinuity of $B\Pi_n(\cdot, E)$ implies that $(x, v) \in B\Pi_n(c, E)$ for all n . It follows from Step 2 of Case 1 that $(x, v) \in B\Pi(c, E)$. Moreover, from Steps 3 and 5,

$$\begin{aligned} \underline{v}^{B\Pi}(c^*) &= \lim_{n \rightarrow \infty} \underline{v}^{B\Pi_n}(c^*) \\ \max\{u(r), \underline{x}(c^*, E, \Pi)\} &= \lim_{n \rightarrow \infty} \max\{u(r), \underline{x}(c^*, E, \Pi_n)\}. \end{aligned}$$

²⁰The convergence is in the product topology; see footnote 8.

Thus, a profile c with value (x, v) is admissible with respect to Π at c^* , as required. $\square$

COROLLARY 2. *We have $\mathcal{E} \in \mathcal{P}^+$.*

PROOF. In the notation of the proof of Proposition 8, $\Pi = \mathcal{E} \in \mathcal{P}$, so it suffices to show that c^Π is an optimal punishment contract for Π . The monotonicity of (Π_n) implies that $\underline{v}^{B\Pi_n}(c^{\Pi_n}) = 0$ for all $n \in \mathbb{N}$, $m \geq n$. Since $c^{\Pi_m} \rightarrow c^\Pi$, the upper hemicontinuity of $B\Pi_n$ implies that $\underline{v}^{B\Pi_n}(c^\Pi) = 0$ for all n . Hence, Step 5 from Case 1 yields $\underline{v}^{B\Pi}(c^\Pi) = 0$. A similar argument using Step 3 from Case 1 establishes that $\underline{x}(c^\Pi, E, \Pi) \leq u(r)$, as required. $\square$

To prove Theorem 2, recall from the proof of Proposition 8 that $\Pi_n \in \mathcal{P}^+$ for all n . Let $f^{\Pi_n} \in \mathcal{F}$ be the Pareto frontier of $\cup_c B\Pi_n(c, E) \cap F$ obtained from Lemma 9. Letting c^{Π_n} denote an optimal punishment contract for Π_n ,

$$\bigcup_{c \in C} B\Pi_{n+1}(c, E) \cap F = \bigcup_{c \in C} \{B\Pi_{n+1}(c, E) \cap F \mid c_h = c_l = c^{\Pi_n}\}.$$

Since $\underline{x}(c^{\Pi_n}, E, \Pi_n) \leq u(r)$, $\underline{v}^{\Pi_{n+1}}(c^{\Pi_n}) = 0$, and $(u(r), f^{\Pi_{n+1}}(u(r))) \in B\Pi_n(c, E)$ for some c , it follows from admissibility that $\underline{x}^{\Pi_{n+1}}(c^{\Pi_n}) = u(r)$. Since the worst payoffs in $\Pi_{n+1}(c^{\Pi_n})$ for both players equal their respective outside options, it follows from E -admissibility that $f^{\Pi_{n+1}} = Tf^{\Pi_n}$. Since (Π_n) is monotone, (f^{Π_n}) is monotone in the sense that $f^{\Pi_{n+1}} \leq f^{\Pi_n}$ on $\text{dom } f^{\Pi_{n+1}}$ and $\text{dom } f^{\Pi_{n+1}} \subseteq \text{dom } f^{\Pi_n}$. Thus, f^{Π_n} converges pointwise to some $f^\mathcal{E}$. Using Step 3 of the proof of Proposition 8 gives

$$\bigcup_{c \in C} B\mathcal{E}(c, E) \cap F = \bigcup_{c \in C} \bigcap_{n \in \mathbb{N}} B\Pi_n(c, E) \cap F = \bigcap_{n \in \mathbb{N}} \bigcup_{c \in C} B\Pi_n(c, E) \cap F.$$

Hence, $f^\mathcal{E}$ gives the Pareto frontier of $\cup_c B\mathcal{E}(c, E) \cap F$. Moreover, since $\mathcal{E} \in \mathcal{P}$, Lemma 9 implies that $f^\mathcal{E} \in \mathcal{F}$.

It follows from admissibility and Proposition 7 that (x, v) is an efficient equilibrium payoff of subgame $c^* \in C$ if and only if $x \in \text{dom } f^\mathcal{E}$, $v = f^\mathcal{E}(x)$, $x \geq \max\{u(r), \underline{x}(c^*, E, \Pi)\}$, and $v \geq \underline{v}(c^*)$. Theorem 2 then follows from $\underline{x}(\mathbf{0}, E, \mathcal{E}) = -\infty$, implied by the existence of an equilibrium of subgame $(\mathbf{0}, E)$ where no bonuses are paid.

REFERENCES

- Abreu, Dilip (1988), “On the theory of infinitely repeated games with discounting.” *Econometrica*, 383–396. [1501]
- Abreu, Dilip, David Pearce, and Ennio Stacchetti (1990), “Toward a theory of discounted repeated games with imperfect monitoring.” *Econometrica*, 58, 1041–1063. [1476]
- Baker, George, Robert Gibbons, and Kevin J. Murphy (1994), “Subjective performance measures in optimal incentive contracts.” *The Quarterly Journal of Economics*, 109, 1125–1156. [1493]

- Battigalli, Pierpaolo and Giovanni Maggi (2008), “Costly contracting in a long-term relationship.” *The RAND Journal of Economics*, 39, 352–377. [1493]
- Bernheim, B. Douglas, and Debraj Ray (1989), “Collective dynamic consistency in repeated games.” *Games and Economic Behavior*, 1, 295–326. [1487]
- Bernheim, B. Douglas, and Michael D. Whinston (1998), “Incomplete contracts and strategic ambiguity.” *American Economic Review*, 88, 902–932. [1483, 1493]
- Brennan, James R. and Joel Watson (2013), “The renegotiation-proofness principle and costly renegotiation.” *Games*, 4, 347–366. [1493]
- Bull, Clive (1987), “The existence of self-enforcing implicit contracts.” *The Quarterly Journal of Economics*, 102, 147–160. [1492]
- Che, Yeon-Koo and Seung-Weon Yoo (2001), “Optimal incentives for teams.” *American Economic Review*, 91, 525–541. [1493]
- Farrell, Joseph and Eric Maskin (1989), “Renegotiation in repeated games.” *Games and Economic Behavior*, 1, 327–360. [1487]
- Guriev, Sergei and Dmitriy Kvasov (2005), “Contracting on time.” *American Economic Review*, 95, 1369–1385. [1493]
- Hart, Oliver and John Moore (1988), “Incomplete contracts and renegotiation.” *Econometrica*, 45, 755–785. [1493]
- Hermalin, Benjamin E. and Michael L. Katz (1991), “Moral hazard and verifiability: The effects of renegotiation in agency.” *Econometrica*, 59, 1735–1753. [1493]
- Hermalin, Benjamin E., Larry Li, and Tony Naughton (2013), “The relational underpinnings of formal contracting and the welfare consequences of legal system improvement.” *Economics Letters*, 119, 72–76. [1493]
- Huberman, Gur and Charles Kahn (1988a), “Limited contract enforcement and strategic renegotiation.” *American Economic Review*, 78, 471–484. [1493]
- Huberman, Gur and Charles Kahn (1988b), “Strategic renegotiation.” *Economics Letters*, 28, 117–121. [1493]
- Iossa, Elisabetta and Giancarlo Spagnolo (2011), “Contracts as threats: On a rationale for rewarding a while hoping for b.” Discussion Paper DP8195, Centre for Economic Policy Research. [1493]
- Itoh, Hideshi and Hodaka Morita (2015), “Formal contracts, relational contracts, and the threat-point effect.” *American Economic Journal: Microeconomics*, 7, 318–346. [1493]
- Kvaløy, Ola and Trond E. Olsen (2009), “Endogenous verifiability and relational contracting.” *American Economic Review*, 99, 2193–2208. [1493]
- Levin, Jonathan (2003), “Relational incentive contracts.” *American Economic Review*, 93, 835–857. [1473, 1487, 1492, 1493]

- MacLeod, W. Bentley and James M. Malcomson (1989), “Implicit contracts, incentive compatibility, and involuntary unemployment.” *Econometrica*, 57, 447–480. [1493]
- MacLeod, W. Bentley (2003), “Optimal contracting with subjective evaluation.” *American Economic Review*, 93, 216–240. [1493]
- Mailath, George J. and Larry Samuelson (2006), *Repeated Games and Reputations – Long-Run Relationships*. Oxford University Press. [1477]
- Maskin, Eric and John Moore (1999), “Implementation and renegotiation.” *Review of Economic Studies*, 66, 39–56. [1493]
- Miller, David A. and Joel Watson (2013), “A theory of disagreement in repeated games with bargaining.” *Econometrica*, 81, 2303–2350. [1473, 1488, 1489]
- Pearce, David G. (1987), “Renegotiation-proof equilibria: Collective rationality and intertemporal cooperation.” Discussion Paper 855, Cowles Foundation. [1487]
- Pearce, David G. and Ennio Stacchetti (1998), “The interaction of implicit and explicit contracts in repeated agency.” *Games And Economic Behavior*, 23, 75–96. [1472, 1473, 1486, 1491, 1493]
- Rey, Patrick and Bernard Salanie (1990), “Long-term, short-term and renegotiation: On the value of commitment in contracting.” *Econometrica*, 58, 597–619. [1493]
- Safronov, Mikhail and Bruno Strulovici (2018), “Contestable norms.” Working Paper <http://faculty.wcas.northwestern.edu/~bhs675/SSN.pdf>. [1473, 1487]
- Schmidt, Klaus M. and Monika Schnitzer (1995), “The interaction of explicit and implicit contracts.” *Economics letters*, 48, 193–199. [1493]
- Segal, Ilya and Michael D. Whinston (2002), “The Mirrlees approach to mechanism design with renegotiation (with applications to hold-up and risk sharing).” *Econometrica*, 70, 1–45. [1493]
- Thomas, Jonathan and Tim Worrall (1988), “Self-enforcing wage contracts.” *Review of Economic Studies*, 55, 541–554. [1492, 1493]
- Thomas, Jonathan P. and Tim Worrall (2018), “Dynamic relational contracts under complete information.” *Journal of Economic Theory*, 175, 624–651. [1493]
- Watson, Joel (2013), “Contract and game theory: Basic concepts for settings with finite horizons.” *Games*, 4, 457–496. [1473, 1488]
- Watson, Joel, David A. Miller, and Trond E. Olsen (2020), “Relational contracting, negotiation, and external enforcement.” *American Economic Review*, 110, 2153–2197. [1471, 1473, 1488, 1492]

Co-editor Thomas Mariotti handled this manuscript.

Manuscript received 14 July, 2018; final version accepted 27 January, 2021; available online 3 February, 2021.