

Performance Incentives with Multiple Instruments*

Erik Madsen[†]

Basil Williams[‡]

Andrzej Skrzypacz[§]

April 1, 2026

Abstract

We systematically study the tradeoff between monetary and non-monetary incentives as instruments to overcome moral hazard. Monetary incentives are provisioned through surplus-neutral payments but are restricted by limited liability. By contrast, non-monetary incentives impact total surplus by distorting the allotment of perks and workplace amenities, promotions, termination, project growth, or other promises about the future. An optimal incentive scheme punishes failure with reduced non-monetary allocations, while rewarding success with increased non-monetary allocations and possibly an accompanying monetary payment. As desired effort increases, optimal non-monetary incentives strengthen while monetary payments often respond in a non-monotone fashion.

Keywords: Moral hazard, performance pay, non-monetary incentives, intertemporal incentives, promotion and firing policies

JEL Classification: D82, D86, M51, M52

*We thank seminar participants and our colleagues and students for helpful comments and suggestions.

[†]Department of Economics, New York University (emadsen@nyu.edu)

[‡]Department of Finance, Imperial College Business School (basil.williams@imperial.ac.uk)

[§]Stanford Graduate School of Business, Stanford University (skrz@stanford.edu)

1 Introduction

This paper revisits the classic problem of designing performance incentives,¹ to understand how performance pay should be combined with general portfolios of non-monetary incentive instruments. Such instruments are common in principal-agent relationships, especially ongoing interactions within organizations: High-performing workers are often rewarded with promotions, preferential project assignments, or perquisites like the freedom to work remotely or flexibly schedule shift work; conversely, low-performing workers are threatened with termination or demotion. They are also prevalent in the gig economy, where freelance workers such as Uber drivers and Airbnb hosts are incentivized through adjustments to match quality and frequency, visibility to potential matches, and freedom to refuse matches.² We propose a general model of such incentive devices and study their use alongside money to overcome moral hazard.

In our model, a principal incentivizes hidden effort from an agent by tying a combination of monetary payments and non-monetary payouts to an observable performance metric (Section 2). To obtain generalizable insights, we avoid modeling the details of a dynamic principal-agent interaction. Instead, we augment a static performance pay problem (à la Holmström (1979)) with a reduced-form technology for delivering non-monetary payouts.

Monetary payments permit the perfect transfer of surplus between parties, but they are limited by a lower bound on allowed payments, as under limited liability. By contrast, non-monetary payouts incur a convex, non-monotone cost for the principal. In particular, modest payouts are *profitable* for the principal to deliver, meaning that they are optimally set at a positive baseline level absent a need to deliver incentives. Payouts beyond the baseline can be used to reward the agent for good performance, while payouts below the baseline can be used to punish him for bad performance. We show (Section 3) that this payoff specification arises from optimally applying a general portfolio of non-monetary incentive tools. Additionally, it can be interpreted as a promise of continuation utility in a dynamic relationship.

Our main result characterizes the cost-minimizing elicitation of an arbitrary target

¹For surveys of existing work on this topic, see Milgrom and Roberts (1992), Prendergast (1999), Gibbons and Roberts (2013), and Georgiadis (2024).

²We provide further detail on non-monetary incentives for freelance workers in Madsen, Williams, and Skrzypacz (2024).

effort level (Sections 4-5). The optimal contract combines rewards for success with punishments for failure, and it generically distorts non-monetary payouts away from the baseline at every performance level. It further exhibits a pecking order of incentive instruments: As total incentives grow, marginal incentives are at first exclusively non-monetary; eventually, however, non-monetary incentives are frozen and are topped up with monetary payments.

The optimal combination of rewards and punishments is determined by both the agent’s disutility of effort and the performance-monitoring structure (Section 6). As the agent becomes more costly to motivate, both rewards and punishments strengthen, eventually leading to the use of monetary payments. Meanwhile, under a binary performance metric, a shift from good-news toward bad-news monitoring is accompanied by a pivot away from rewards toward punishments. As a result, agents in “star” roles are optimally motivated by large rewards for success, including possibly monetary payments; while those in “guardian” roles are motivated mostly by severe punishments for failure, which are necessarily delivered through non-monetary incentives. These two forces interact to determine how incentives change with desired effort: Punishments grow monotonically with effort; in contrast, rewards (and therefore monetary payments) may respond to increased effort in a non-monotone fashion.

We demonstrate how the profitability of non-monetary incentives flows from the lower bound on monetary payments (Section 7). Because the agent cannot offer money in return for non-monetary payouts, the principal inefficiently reduces payouts in order to extract surplus. This inefficiency is alleviated to the extent the agent is “bought in” to the relationship via a binding participation constraint. As the constraint tightens, non-monetary payouts converge to the surplus-maximizing level. Once the agent is fully bought in, all incentives are provided through monetary payments.

Surprisingly, the principal’s surplus extraction motive may lead the agent to receive *negative* rents from performance incentives (Section 8). In contrast to a pure performance pay scheme, where rents are always positive due to limited liability, non-monetary punishments can drive the agent’s utility below what he would receive in an observable-effort benchmark. We provide sufficient conditions under which this force leads to negative rents under the optimal contract.

1.1 Related literature

Our interest in comparing the incentive properties of monetary and non-monetary instruments is shared by several existing studies. Chwe (1990) allows the principal to employ a combination of rewards and punishments to provide incentives, each of which incurs convex costs. Only one reward instrument is available, abstracting from a comparison between alternative reward channels.

Acemoglu and Wolitzky (2011) similarly allow the principal to both offer monetary rewards and threaten punishments. Punishments are costless for the principal to inflict, and their use is limited only by a binding participation constraint, abstracting from a comparison between the costs of rewards and punishments.

Wolitzky (2012) allows the principal to commit to both monetary payments and performance reports to a follow-on labor market. Performance reports are payoff-irrelevant to the principal, and so she optimally exhausts their incentive potential before resorting to monetary payments.

Che, Iossa, and Rey (2021) allow the principal to reward the agent with either money or preferential treatment in a follow-on procurement auction. Transfers are not restricted, but the agent possesses private information about her procurement costs, generating information rents which can be extracted only at the cost of allocative efficiency. We embed this efficiency/extraction tradeoff in a canonical reduced-form contracting setting, yielding general insights into the optimal combination of monetary and non-monetary incentives, as well as of rewards and punishments.

A literature studying status hierarchies (Auriol and Renault 2008; Besley and Ghatak 2008; Moldovanu, Sela, and Shi 2007; Dubey and Geanakoplos 2020) considers how conferral of status can be used alongside, or in place of, money as an incentive tool. Unlike the non-monetary instruments we study, the conferral of status is payoff-irrelevant to the principal. Its use is instead constrained by an externality: When one agent is awarded status, the value of status to other agents is reduced. We shut down this channel and focus instead on instruments for which the principal has allocative preferences.

Interpreting our non-monetary instrument as continuation value in a dynamic relationship connects our results to a large literature studying dynamic principal-agent problems. This literature has demonstrated that, in a variety of dynamic interactions, performance incentives are optimally delivered partly via intertemporal rewards and punishments. For instance, DeMarzo and Fishman (2007) and DeMarzo and Sannikov

(2006) consider deferred compensation; Zhu (2013) studies effort adjustment; Biais et al. (2010) allow for scaling of the agent’s activity; Ke, Li, and Powell (2018) model task reassignment; and Sannikov (2008, §4.3) contemplates investment in human capital.

In these environments, the principal’s profits from future interactions are typically hump-shaped in the agent’s promised utility. (Hörner (2013) provides a general explanation of this phenomenon. See also the references in Lee, Fuchs, and Dumav (2024), who embed a reduced-form continuation interaction in a delegated bandit exploration problem.) Our analysis systematically explores the implications of this continuation payoff structure for incentive provision in a static moral hazard problem. In particular, our comparative statics hold fixed the structure of future interactions and examine how changes in the initial interaction impact reliance on current versus future rewards.

2 Model

A principal contracts with an agent to perform a task subject to moral hazard. The agent’s hidden effort choice $e \in [0, 1]$ impacts the distribution of a performance signal y , which in our baseline analysis we take to be binary: $y \in \{y_H, y_L\}$. (We extend our analysis to a many-outcome model in Section 9.) As a normalization, we assume that the probability of high performance is affine in effort:

$$\Pr(y = y_H \mid e) = p(e) \equiv \underline{p} \cdot (1 - e) + \bar{p} \cdot e,$$

where $1 > \bar{p} > \underline{p} \geq 0$.³ The agent’s cost of exerting effort e is $h(e)$, where $h : [0, 1] \rightarrow \mathbb{R}_+$ is assumed to be C^1 , increasing, and strictly convex, with $h(0) = 0$. We do not explicitly model the value of effort to the principal. Instead, we focus on incentivizing a specified target effort level at minimal cost.

The principal can incentivize effort using two instruments: monetary payments, and a non-monetary tool we refer to as *credits*. Monetary payments perfectly transfer surplus from the principal to the agent. We impose a lower bound on payments, which we normalize to 0. Meanwhile, credits are a reduced-form stand-in for a portfolio of resource allocation decisions and promises about the future, which we microfound in

³To ensure uniqueness of the optimal contract, we formally exclude the case $\bar{p} = 1$. All of our results remain valid when $\bar{p} = 1$, with the caveat that the optimal contract is not unique when the principal induces maximal effort.

Section 3. As with payments, we impose a lower bound on credits, which we normalize to 0.

Unlike payments, credit allocations are not surplus-neutral: An allocation of Q credits generates utility Q (a harmless normalization) for the agent and profits $R(Q)$ to the principal. We assume that $R : \mathbb{R}_+ \rightarrow \mathbb{R}$ is C^1 , strictly concave, and satisfies $R(0) = 0$, $R'(0) > 0$, and $R'(\infty) < -1$. We will write $Q^0 \equiv (R')^{-1}(0) > 0$ to denote the unique maximizer of R , a quantity we refer to as the *baseline* credit allocation.

To summarize payoffs, if the agent exerts effort $e \in [0, 1]$ and receives T dollars and Q credits, his net utility is

$$U = T + Q - h(e)$$

while the principal's profits are

$$\Pi = R(Q) - T.$$

Recall that we do not explicitly model the value of effort to the principal, and so this profit specification captures only the value of the credit allocation and the cost of monetary payments.

The principal can commit to an incentive contract $\mathcal{C} = (Q_H, T_H, Q_L, T_L) \in \mathbb{R}_+^4$ which conditions her allocation of money and credits on observed performance. The principal's problem is to maximize expected profits conditional on eliciting a specified target effort level $e^* \in (0, 1]$:

$$\begin{aligned} & \max_{\mathcal{C}} p(e^*) \cdot (R(Q_H) - T_H) + (1 - p(e^*)) \cdot (R(Q_L) - T_L) \\ & \text{subject to} \\ & e^* \in \arg \max_{e \in [0, 1]} p(e) \cdot (Q_H + T_H) + (1 - p(e)) \cdot (Q_L + T_L) - h(e) \end{aligned} \tag{OC}$$

We will call a contract $\mathcal{C}^*$ *optimal* if it solves problem (OC).

3 Non-Monetary Instruments

Most elements of our model are standard in the principal-agent literature. Our sole innovation is the introduction of non-monetary credits as an additional instrument for

incentivizing effort. Our specification of credits includes two key features. First, the principal prefers to allocate some credits even absent incentive concerns. Second, she faces diminishing and eventually negative marginal returns from allocating credits.

In Section 3.1, we motivate these features of our model by connecting them to incentive tools available in real-world principal-agent relationships. In Section 3.2, we describe how these incentive tools can be interpreted as impacting the structure of future interactions in dynamic principal-agent relationships. In Section 3.3, we show how a portfolio of incentive tools can be reduced to a single-dimensional credit allocation as in our baseline model.

3.1 Motivation

Our payoff specification for credits arises naturally whenever the principal delivers credits by adjusting a portfolio of non-monetary resource allocations and promises about the future. For instance, employees within an organization may care about amenities like free meals or the flexibility to work from home; the quality or difficulty of project assignments; the headcount they are assigned to manage; and the chances of future promotion or firing. Similarly, contractors providing services to an individual or firm may care about the volume or consistency of follow-up work; the quality of a review on a digital platform; and the number of recommendations to additional clients.

Each of these choices represents a tool for incentivizing effort. Some tools are always costly to utilize. For instance, free food and other in-office amenities have a direct financial cost, while firing a productive employee or terminating a productive contractor incurs replacement costs. Other tools may be initially profitable to allocate. For instance, a manager benefits from promising to promote her employee whenever he ends up being the best candidate for a future opening. Similarly, a client benefits from recommending her contractor to friends who ask for advice on whom to hire for a project.

All of these tools plausibly experience diminishing and eventually negative marginal returns. For instance, boosting an employee’s utility with increasingly lavish on-site meals may require ever-larger expenditures per unit of consumption value. Similarly, promising an employee a higher chance of promotion requires promoting her in circumstances when she is a poorer fit for a senior position on the margin.

In some applications, high performance may be a signal of suitability for non-monetary allocations. For instance, high-performing workers may be more qualified for promotion, and productive contractors may generate higher expected returns in a follow-on project. In that case, the marginal profitability of awarding additional credits would be rising in the agent’s performance: Formally, $R'_H > R'_L$. For expositional simplicity, we abstract from such dependence when deriving our main results. Nonetheless, all of our qualitative results continue to hold when non-monetary surplus is performance dependent, a point we revisit at several stages in our analysis.

3.2 Connection to dynamic relationships

Many of the tools described in Section 3.1 have a flavor of continuation promises about the structure of an ongoing relationship. For instance, working from home may impact the principal’s ability to monitor effort. Similarly, a client’s promise of follow-on work may impact the scale or cashflow diversion opportunities of a future interaction. (See also Marino and Zábajník (2008), who highlight that non-monetary amenities or perks may impact employee productivity.)

Whenever an incentive tool impacts the structure of future interactions, $R(Q)$ may be interpreted as a continuation profit to the principal generated from a promised continuation utility Q to the agent. Consistent with our specification, continuation profits are naturally concave and non-monotone in promised utility in a wide variety of dynamic principal-agent models. Importantly, this interpretation does not require an assumption that the principal can fully commit to the continuation contract governing future interactions. In an environment without dynamic commitment, the payoff function $R(Q)$ represents the maximal profits that can be supported by a self-enforcing continuation contract.

When Q is interpreted as continuation utility in a dynamic context, it need not be delivered solely through non-monetary channels. For instance, a component of continuation utility may be delivered through future payments. In such cases, our use of “non-monetary” to describe credits should be interpreted to mean “involving no immediate payments”.

Future payments may generate a continuation profit function which is not strictly concave in Q everywhere. However, whenever the agent discounts the future more heavily than the principal, $R'(Q) < -1$ where future payments are used to deliver

utility on the margin. Intuitively, money today is a cheaper way to deliver (discounted lifetime) utility than money tomorrow whenever the agent is relatively impatient. None of our results would be impacted by allowing R to be linear over a range of promises for which $R'(Q) < -1$, and we maintain an assumption of global strict concavity solely for expositional convenience.

3.3 Microfoundation

Suppose that the principal possesses a portfolio of incentive tools which can be used to adjust the agent's utility. Concretely, she possesses $A + B$ tools, with $A \geq 1$ and $B \geq 0$. Each tool can be applied with an intensity $q_i \in [0, \bar{q}_i]$, where $0 < \bar{q}_i \leq \infty$ and $\max_i \bar{q}_i = \infty$. Intensities are normalized so that q_i records the utility generated for the agent from application of tool i . The profit earned by the principal from applying tool i with intensity q_i is $R_i(q_i)$.

We assume that each R_i is C^1 and strictly concave, with $R'_i(0) > -1 > R'_i(\bar{q}_i)$. Further, $R'_i(0) > 0$ for each tool $1 \leq i \leq A$, while $R'_i(0) \leq 0$ for each tool $A + 1 \leq i \leq B$. In other words, the first A tools generate initially positive returns for the principal, while the next B tools are always costly. Since $A \geq 1$, at least one tool is initially profitable.

For each $Q \geq 0$, let

$$R(Q) \equiv \max_{\mathbf{q}} \sum_{i=1}^{A+B} R_i(q_i) \quad \text{s.t.} \quad \sum_{i=1}^{A+B} q_i = Q.$$

Standard arguments establish the following properties of this problem:

Remark. R is C^1 , strictly concave, and satisfies $R(0) = 0$, $R'(0) > 0$, and $R'(\infty) < -1$. Further, for each Q there exist unique maximizers $\mathbf{q}^*(Q)$ such that:

- Each q_i^* is continuous and nondecreasing in Q , and is further increasing whenever $q_i^*(Q) \in (0, \bar{q}_i)$,
- $q_i^*(Q) > 0$ for each $i \leq A$ and $Q > 0$.

This remark justifies our modeling of arbitrary non-monetary instruments using a one-dimensional specification. Credits can be interpreted as the aggregate usage of all instruments in tandem to delivery a desired non-monetary payout to the agent. In

applications, this payout is optimally delivered through a combination of instruments, with their intensity rising in tandem as the total payout increases.

In particular, tools which are initially costly ($i \geq A + 1$) to use are not generally utilized until total non-monetary payouts rise sufficiently high. Conversely, tools which are initially profitable ($i \leq A$) to invoke are always utilized. As we show in our main analysis, non-monetary instruments are used both to reward good performance as well as punish bad performance. Only initially-profitable tools are used to provide punishments, which involve intensities on the upward-sloping portion of each R_i . In contrast, both categories of tools may be used to provide rewards, which involve intensities on the downward-sloping portion of each R_i .

4 Utility Delivery

To characterize the optimal contract, we decompose the problem into two steps: an inner utility delivery problem, and an outer incentive provision problem. We first analyze the utility delivery problem, which identifies the optimal mix of incentive tools delivering a specified utility to the agent. In Section 5, we characterize the optimal utility promises which elicit the target effort level.

Recall that the agent's expected payoff under contract $\mathcal{C}$ and effort level e is

$$p(e) \cdot (Q_H + T_H) + (1 - p(e)) \cdot (Q_L + T_L) - h(e).$$

This payoff depends on $\mathcal{C}$ only through the gross utility (that is, ignoring effort costs) that the agent receives at each performance level:

$$\mathbf{W} = (W_H, W_L) \equiv (Q_H + T_H, Q_L + T_L).$$

The vector $\mathbf{W}$ therefore summarizes the agent's effort incentives, and all contracts delivering the same $\mathbf{W}$ induce the same effort.

The principal, on the other hand, is not indifferent between the various (Q, T) pairs which deliver a given gross utility W . Indeed, the principal's cost of delivering W is minimized at the (unique) solution to the auxiliary *utility delivery problem*

$$\min_{(Q, T) \geq 0} R(Q^0) - R(Q) + T \quad \text{s.t.} \quad W = Q + T \quad (\text{UD})$$

for every $W \geq 0$. (Recall that Q^0 is the credit allocation which maximizes the payoff function R .) By solving this problem, we can eliminate both money and credits from the profit-maximization problem (OC) and focus on the design of the utility promises W .

The following proposition characterizes the solution to the utility delivery problem along with key properties of its value $C(W)$, which we will refer to as the *utility cost function*. This function measures the reduction in the principal's profit from delivering W relative to the unconstrained optimum $R(Q^0)$. (All proofs are deferred to Appendix A.)

Proposition 1. *Problem (UD) has a unique solution $(Q^{**}(W), T^{**}(W))$ for each $W \geq 0$. There exists a $\bar{W} > Q^0$ such that:*

- Q^{**} and T^{**} are continuous and nondecreasing,
- Q^{**} is increasing and $T^{**} = 0$ on $[0, \bar{W}]$,
- Q^{**} is constant and T^{**} is increasing on $[\bar{W}, \infty)$.

Further, the value $C(W)$ of problem (UD) is:

- Convex and C^1 ,
- Uniquely minimized at $W = Q^0$ with $C(Q^0) = 0$,
- Strictly convex on $[0, \bar{W}]$,
- Linear with slope 1 on $[\bar{W}, \infty)$.

This proposition establishes a progression of instruments used to deliver incremental utility to the agent. For low utility levels, utility is optimally delivered by distorting non-monetary allocations. When $W < Q^0$, such distortions are trivially the cheapest way to deliver the specified utility due to limited liability. On the other hand, when $W > Q^0$ many combinations of money and credits could be used to deliver the utility increment $W - Q^0$. Since incremental non-monetary distortions incur initially zero but progressively higher marginal costs, incremental utility is provided at first with credits and eventually with money. We graphically depict this progression in Figure 1.

The threshold utility level at which incentive tools switch satisfies $\bar{W} = (R')^{-1}(-1)$, at which point marginal credit allocations cost the same as transferring money. One important question is under what conditions the principal promises a gross utility to the agent exceeding the level $\bar{W}$. In other words, when is money used as part of an optimal contract? We answer this question in Section 5 when we characterize the optimal contract.

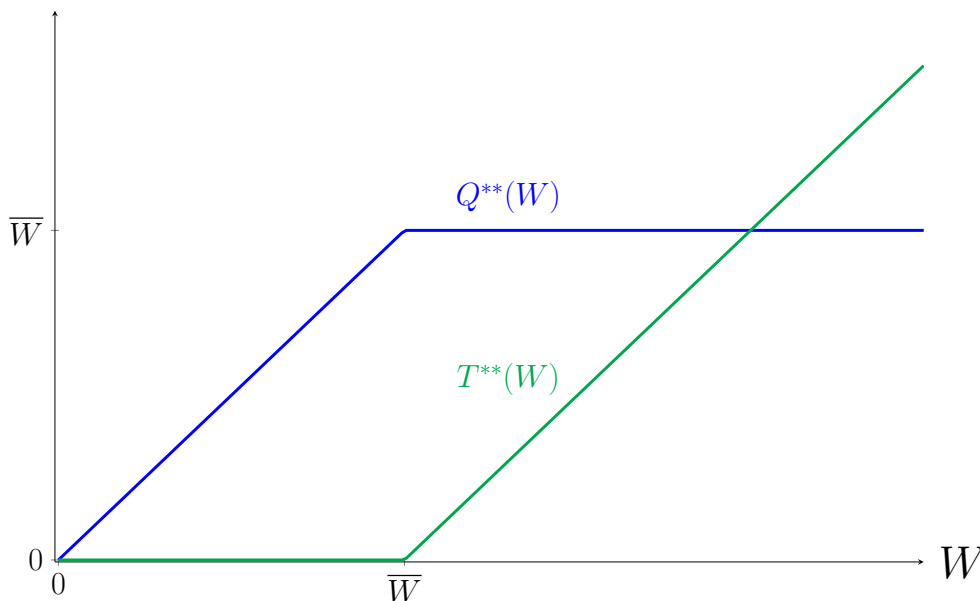


Figure 1: Optimal allocations as a function of promised utility

Another important feature of the utility delivery problem is that the utility cost function $C(W)$ is both convex and non-monotone. Unlike in a model involving only incentive pay, the principal's costs are not monotone increasing in the utility delivered to the agent. Instead, costs decrease until the point $W = Q^0$ at which the agent receives the (gross) utility he would enjoy under the baseline credit allocation. We will refer to utility promises $W < Q^0$ as *punishments* and $W > Q^0$ as *rewards*. Proposition 1 indicates that both punishments and rewards are costly to deliver and incur (weakly) increasing incremental costs. The main features of the utility cost function are depicted graphically in Figure 2.

In applications where non-monetary surplus is performance dependent, the utility cost function depends on realized performance. Formally, C_s must be indexed with $s \in \{H, L\}$. If $R'_H > R'_L$, so that marginal non-monetary surplus rises with performance, then $C'_H(W) \leq C'_L(W)$ for all W , with the inequality strict whenever $C'_H(W) < 1$. In

particular, $\bar{W}_H > \bar{W}_L$, and non-monetary instruments are used to deliver utility over a larger range following high performance rather than low. Additionally, $Q_H^0 > Q_L^0$, so that the baseline allocation is larger following high performance.

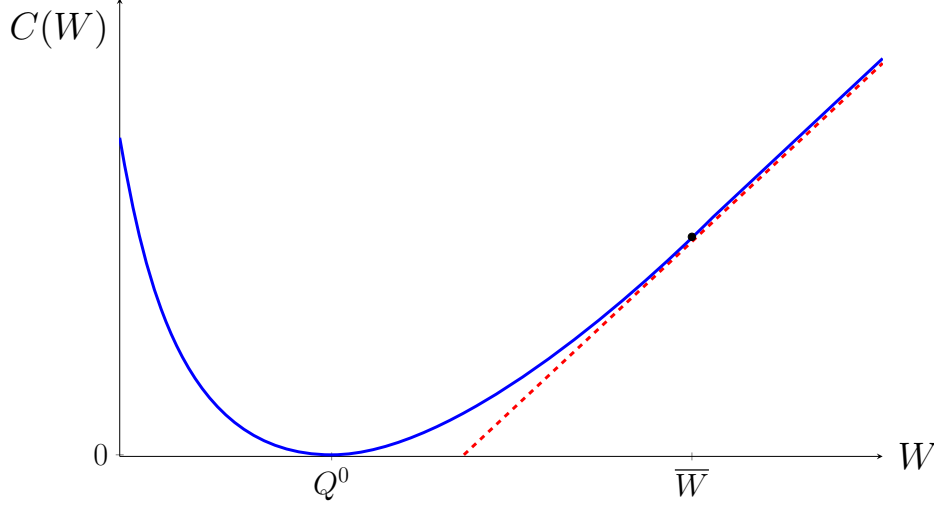


Figure 2: Shape of the utility cost function

5 Optimal Contract

We now characterize the principal's optimal contract. To do so, we pass to a reduced cost-minimization problem with respect to the vector of gross utilities $\mathbf{W}$ introduced in Section 4.

This reduction involves two steps. First, we use the utility cost function derived in Section 4 to replace payments and credit allocations with (gross) utility promises. Second, we observe that the agent's payoff is strictly concave in effort and pass from the incentive-compatibility constraint to its corresponding first-order condition:

$$W_H - W_L \geq \frac{h'(e^*)}{\Delta p} \quad (\text{IC})$$

where $\Delta p \equiv \bar{p} - \underline{p}$. Note that (IC) must hold as an equality to ensure incentive compatibility if $e^* < 1$. However, the complementary constraint is nonbinding, and we drop it from the problem to emphasize the direction in which the constraint binds.

Utilizing these reductions, the utility promises delivered under an optimal contract

can be characterized by solving the (convex) *cost-minimization problem*

$$\min_{\mathbf{W} \geq 0} p(e^*) \cdot C(W_H) + (1 - p(e^*)) \cdot C(W_L) \quad \text{s.t.} \quad (\text{IC}) \quad (\text{CM})$$

The following result formally verifies that a unique solution to problem (CM) exists and can be used to construct a contract solving the profit-maximization problem (OC). It additionally establishes an important ranking satisfied by the optimal utility promises.

Theorem 1. *Problem (CM) has a unique solution $\mathbf{W}^*$, which satisfies $W_H^* > Q^0 > W_L^*$.*

There exists a unique optimal contract $\mathcal{C}^ = (Q_s^*, T_s^*)_{s=L,H}$. It satisfies $Q_s^* = Q^{**}(W_s^*)$ and $T_s^* = T^{**}(W_s^*)$ for each $s = L, H$.*

Intuitively, incentives can be provisioned by promising rewards following high performance or by promising punishments following low performance. Both are costly to the principal, since C increases as W moves away from Q^0 in either direction. The unique optimal promises equalize the marginal costs of rewards and punishments subject to delivering the required magnitude of incentives. Importantly, since C is flat at Q^0 , small rewards or punishments incur no first-order cost. As a result, the optimal contract uses both types of incentive at the optimum.

It is instructive to compare this result with behavior under an optimal contract when only monetary incentives are available. In such a setting, even were the agent risk-averse over income, the utility cost function would be increasing everywhere and an optimal contract trivially involves $W_L^* = 0$. Since (IC) uniquely pins down W_H^* as a function of W_L^* , there are no remaining degrees of contractual freedom in that benchmark. By contrast, in problem (CM), the optimal utility promise W_L^* is not immediate. The substance of the optimization problem is to pin down the overall level of utility promises by balancing the incentive costs of punishments and rewards.

An immediate corollary of this result is that the optimal contract distorts non-monetary allocations no matter the agent's performance, but uses money only possibly following high performance:

Corollary. *The optimal contract satisfies $Q_H^* > Q^0 > Q_L^*$ and $T_L^* = 0$.*

Importantly, Theorem 1 is silent on whether the optimal contract uses monetary

incentives at all, i.e., whether $T_H^* > 0$. In the next section, we identify conditions under which money is used.

In applications where non-monetary surplus is performance-dependent, the objective function in problem (CM) must be modified to account for this dependence. In such environments, the optimal utility promises minimizes

$$p(e^*) \cdot C_H(W_H) + (1 - p(e^*)) \cdot C_L(W_L).$$

One new feature arising in this environment is that baseline performance incentives are free: Since $Q_H^0 > Q_L^0$, the effort level $\underline{e}$ characterized by

$$Q_H^0 - Q_L^0 = \frac{h'(\underline{e})}{\Delta p}$$

can be elicited without explicit performance incentives. For effort levels $e^* > \underline{e}$, the IC constraint continues to bind, and an optimal contract satisfies $W_H^* > Q_H^0 > Q_L^0 > W_L^*$. In other words, high performance is rewarded and low performance is punished, relative to the baseline allocation at each performance level.

6 Rewards versus Punishments

Theorem 1 demonstrated that an optimal contract employs a combination of rewards and punishments to incentivize the agent. We next examine how the optimal usage of each tool is influenced by features of the contractual environment.

In section 6.1, we show that as the agent becomes more difficult to incentivize, both rewards and punishments grow (weakly) stronger. A pecking order of instruments emerges, with marginal incentives shifting from credits to money as total incentives grow.

In section 6.2, we show that as the monitoring structure shifts from good news toward bad news, the balance of incentives shifts away from rewards and toward punishments.

Finally, in section 6.3, we show that as the target effort level increases, punishments grow stronger while rewards may exhibit non-monotonicity. In particular, the optimal monetary payment may locally decrease as target effort grows.

We formally establish these results assuming no performance dependence of non-

monetary surplus. However, all results continue to hold when non-monetary surplus is rising in performance, for all effort levels above the baseline effort $\underline{e}$ induced by the baseline credit allocation.

6.1 Disutility

Decompose the agent's effort function as $h(e) = \eta \cdot h_0(e)$, where $h_0(e)$ is a baseline effort function and $\eta > 0$ is a (known) scalar. As η grows, the agent becomes more difficult to incentivize. The following result characterizes how optimal incentives change with η , holding fixed all other model parameters (including target effort, an assumption we comment on below).

Proposition 2. *W_H^* is increasing in η while W_L^* is nonincreasing in η . Further, there exists an $\bar{\eta} > 0$ such that:*

- $W_H^* > \bar{W}$ iff $\eta > \bar{\eta}$,
- W_L^* is decreasing whenever it is positive on $[0, \bar{\eta}]$,
- W_L^* is independent of η on $[\bar{\eta}, \infty)$.

Mechanically, total incentives (measured by the gap $W_H^* - W_L^*$) must grow larger as η increases. Hence at least one of rewards or punishments must also strengthen. Proposition 2 demonstrates that the two optimally grow in tandem so long as total incentives are sufficiently weak—in particular, so long as money is not used at the optimum. In this regime, incentive costs grow on the margin for both types of incentive, and so they scale up together to maintain parity between their marginal costs. On the other hand, when total incentives are strong, the marginal cost of further non-monetary distortions reaches the cost of transferring money. In that regime, additional non-monetary incentives become unattractive, and all further incentives are provisioned with money. As a result, credit allocations become insensitive to further increases in required incentives.

Taken together, these comparative statics imply a *pecking order* of incentive tools: Non-monetary instruments are used to provide marginal incentives when overall incentives are low-powered, while money is used on the margin when incentives are high-powered. The following corollary summarizes this finding:

Corollary. *If $\eta < \bar{\eta}$, then:*

- $T_H^* = 0$,
- $Q_H^* - Q_L^*$ is increasing in η .

If $\eta > \bar{\eta}$, then:

- T_H^* is positive and increasing in η ,
- $Q_H^* - Q_L^*$ is independent of η .

Proposition 2 assumes that the principal maintains a fixed effort target as the agent becomes more difficult to incentivize. This comparative static is directly applicable to incentive problems in which effort is endogenous but discrete.⁴ In such environments, optimal effort is locally independent of η for generic model parameters, and our results characterize the local behavior of contractual incentives.

Even when effort is continuous, our analysis remains relevant for understanding how effort disutility shapes the optimal contract. In that case, Proposition 2 should be interpreted as identifying the *direct* effect of a change in η on the optimal contract. This effect can be combined with the corresponding *indirect* effect of a change in effort (which we characterize in Section 6.3) to understand how η impacts incentives when effort adjusts endogenously.

6.2 Monitoring

We now examine how contractual incentives respond to a change in the monitoring structure. Concretely, we examine the impact of an improvement in baseline performance, which we operationalize as an increase in $\underline{p}$ while the marginal impact of effort Δp is held fixed. We also hold fixed all other model parameters, including target effort. (The application of our findings to settings with endogenous effort involves considerations similar to those discussed in Section 6.1.)

An improvement in baseline performance effectively moves the monitoring structure from a good news environment, where failure is the norm and effort generates in-

⁴Our model can be modified to accommodate discrete effort by assuming that $h(e)$ is piecewise linear with convex kinks at each allowed effort level, and by replacing $h'(e^*)$ with $\lim_{e \uparrow e^*} h'(e)$ in (IC).

frequent successes, to a bad news one, where effort guards against infrequent failures.⁵ The following proposition characterizes how this change impacts optimal incentives.

Proposition 3. *Both W_H^* and W_L^* are nonincreasing in baseline performance, and both are decreasing whenever $W_L^* > 0$.*

In other words, optimal incentives shift away from rewards and toward punishments as baseline performance improves. This effect follows directly from the fact that as performance improves, rewards must be given more often while punishments can be inflicted less often. The principal therefore saves on incentive costs by shifting incentives toward the outcome which occurs less frequently—in other words, away from rewards and toward punishments.

A corollary of this result is that, as baseline performance improves, the principal shifts away from monetary incentives and toward non-monetary ones:

Corollary. *As baseline performance improves, T_H^* weakly decreases while $Q_H^* - Q_L^*$ weakly increases. Further, both changes are strict whenever $T_H^* > 0$.*

6.3 Effort

We now examine how optimal incentives depend on target effort e^* . In settings where effort is determined endogenously in a profit-maximization problem, an adjustment to target effort would reflect changes in the marginal profitability of effort. The comparative static is also relevant for understanding the indirect effect of changes in other environmental parameters, as discussed in Section 6.1.

Target effort appears in two places in the principal’s cost-minimization problem (CM): On the right-hand side of the incentive constraint (IC), and in the coefficient $p(e^*)$ appearing in the objective function. An increase in e^* tightens (IC) and simultaneously shifts probability away from punishments and toward rewards in the principal’s incentive costs. As a result, an increase in target effort is effectively a composite of an increase in effort disutility alongside an improvement in baseline performance.

⁵Formally, this change decreases the informativeness of high performance relative to low performance. As we discuss further in Section 9, informativeness is measured by $I_s(e) = p'_s(e)/p_s(e)$, where $p_s(e)$ is the probability of performance level s given effort level e . Hence $|I_G(e^*)|/|I_B(e^*)| = (1 - p(e^*))/p(e^*)$ is decreasing in $\underline{p}$ when Δp is held fixed.

As we established in Proposition 2, an increase in disutility tends to strengthen both rewards and punishments; on the other hand, an improvement in baseline performance tends to shrink rewards and boost punishments. These effects work in the same direction for punishments, implying the following result:

Proposition 4. W_L^* is nonincreasing in e^* , and is decreasing whenever it is positive.

It follows easily that non-monetary incentives must also grow with effort. Indeed, if $T_H^* = 0$, then (IC) requires that $Q_H^* - Q_L^* = h'(e^*)/\Delta p$, which mechanically increases in e^* . And if $T_H^* > 0$, then $Q_H^* = (R')^{-1}(-1)$ is locally independent of e^* while $Q_L^* = W_L^*$ is nonincreasing. Hence:

Corollary. $Q_H^* - Q_L^*$ is nondecreasing in e^* , and it is increasing whenever $W_L^* > 0$ or $T_H^* = 0$.

However, the two effects exert opposing forces on the optimal reward W_H^* . Under auxiliary concavity assumptions, the baseline performance effect tends to dominate the disutility effect as effort increases. So long as effort is sufficiently effective and the non-negativity constraint on credits is non-binding for sufficiently high effort, the result is non-monotonicity of the optimal reward. (Once the optimal punishment involves allocating no credits, the baseline performance effect is muted and rewards must mechanically increase in effort in order to satisfy (IC).)

To formally establish this result, we augment the credit payoff function R with a free parameter which can be adjusted to shift R' horizontally. To that end, we suppose that

$$R(Q) = \bar{R}(Q - Q^0) - \bar{R}(-Q^0)$$

for some function $\bar{R} : \mathbb{R} \rightarrow \mathbb{R}$ which is C^1 , strictly concave, and satisfies $\bar{R}'(0) = 0$ and $\bar{R}'(-\infty) = \infty$ and $\bar{R}'(\infty) < -1$. The quantity $Q^0 > 0$ continues to represent the baseline allocation of credits, but we now take it to be an adjustable parameter rather than a derived property of R . Let

$$\bar{e} \equiv \sup\{e^* \in (0, 1] : W_L^* > 0\}$$

be the target effort level at which the non-negativity constraint on credits begins to bind. (It may equal 1 if $h'(1)$ is sufficiently small.)

Proposition 5. Suppose that $\bar{R}'$ and h' are concave and $h'(0) = 0$. Then:

- W_H^* is strictly quasiconcave in e^* on $(0, \bar{e}]$,
- W_H^* is non-monotone in e^* on $(0, \bar{e}]$ when $\bar{p}$ is sufficiently close to 1 and Q^0 is sufficiently large,
- If η is sufficiently large, then $\max_{e^* \in (0, \bar{e}]} W_H^* > \bar{W}$ when $\bar{p}$ is sufficiently close to 1 and Q^0 is sufficiently large.

Figure 3 graphically illustrates the comparative statics of Propositions 4 and 5 using a parametric example satisfying the concavity conditions of Proposition 5.

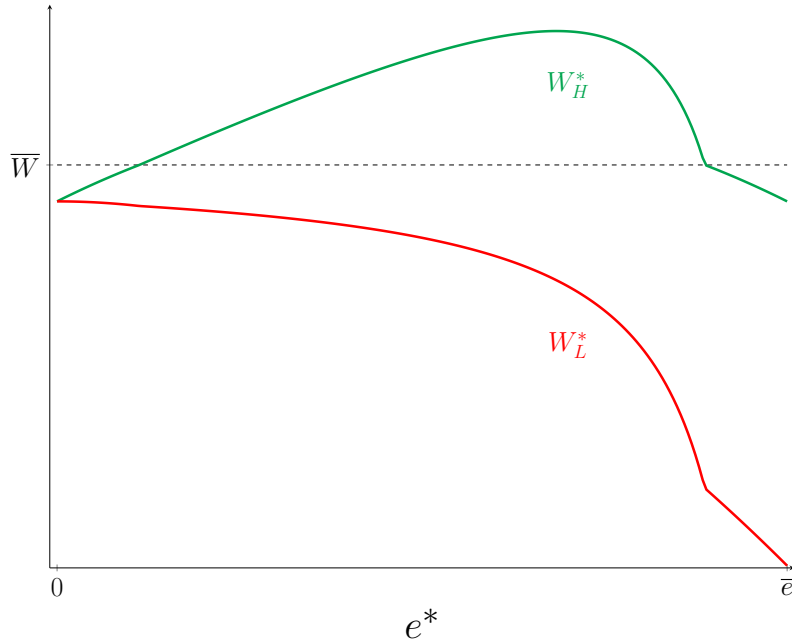


Figure 3: Optimal utility promises as a function of effort when $h(e) = e^2/2$ and $R(Q) = 10 \cdot Q \cdot (1 - Q/2)$ and $(\underline{p}, \bar{p}) = (0, 1)$.

The following corollary, which establishes analogous non-monotone behavior for the optimal monetary payment when the agent is sufficiently hard to motivate, is immediate:

Corollary. *Suppose that $\bar{R}'$ and h' are concave and $h'(0) = 0$. If η is sufficiently large, then T_H^* is quasiconcave and non-monotone in e^* on $(0, \bar{e}]$ when $\bar{p}$ is sufficiently close to 1 and Q^0 is sufficiently large.*

For the same parametric example, Figure 4a depicts the optimal allocations of money and credits, while Figure 4b depicts the optimal division of incentives between monetary and non-monetary instruments as a function of effort.

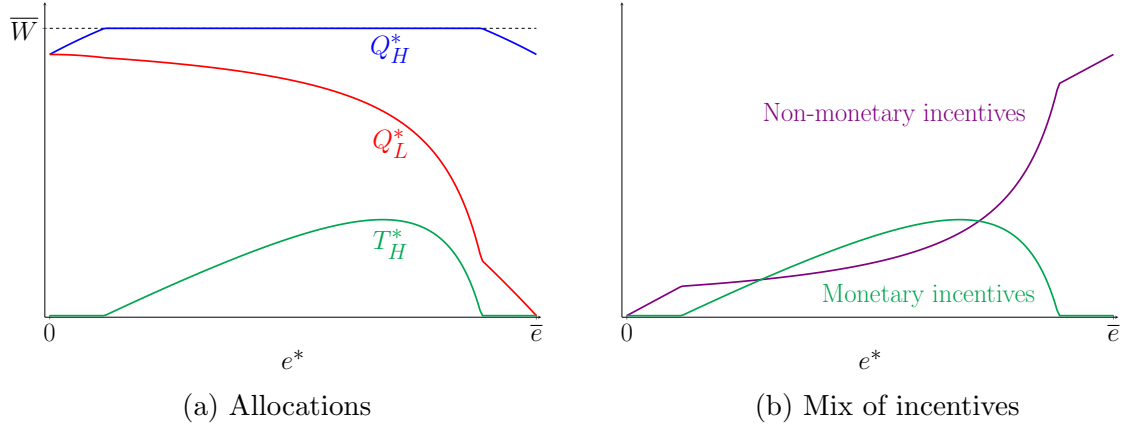


Figure 4: Optimal allocations and incentives as a function of effort when $h(e) = e^2/2$ and $R(Q) = 10 \cdot Q \cdot (1 - Q/2)$ and $(\underline{p}, \bar{p}) = (0, 1)$.

7 (In)efficiency of the Optimal Contract

A key distinction between monetary and non-monetary incentives is that monetary payments are a pure transfer between the principal and agent, while credit allocations impact total surplus. In this section we explore the efficiency properties of the credit allocation under an optimal contract. In Section 7.1, we establish that this allocation is generally inefficient and link the inefficiency to the lower bound on payments. In Section 7.2, we demonstrate that efficiency is restored as the agent's skin in the game rises and the principal's ability to extract surplus is diminished.

7.1 Non-monetary Surplus

When a quantity Q of credits is allocated to the agent, the agent gains utility Q while the principal gains profits $R(Q)$. Total surplus from this allocation is therefore

$$S(Q) = Q + R(Q).$$

This function is strictly concave and maximized when $R'(Q) = -1$. As noted in the discussion following Proposition 1, this condition also characterizes the utility level at which the principal begins delivering utility with money. In other words, S is maximized at $Q = \bar{W}$.

In light of this fact, the set of efficient incentive-compatible contracts satisfy

$$W_L \geq \bar{W}, \quad W_H = W_L + h'(e)/\Delta p.$$

The precise choice of W_L has no efficiency implications and can be adjusted to re-allocate surplus between the principal and agent. Under such contracts, credits are allocated efficiently, i.e., $Q_H = Q_L = \bar{W}$, and incentives are provisioned solely through monetary payments.

However, the principal does not choose to write such a contract. Indeed, Theorem 1 guarantees that $W_L^* < Q^0 < \bar{W}$, and so the optimal contract is invariably inefficient. This distortion arises because the lower bound on payments prevents the principal from efficiently extracting surplus. Instead, she misallocates credits as a second-best mechanism for surplus extraction.

7.2 Skin in the Game

To further illuminate the role of limited liability in distorting surplus, we study the impact of buying the agent into the relationship. Augment the moral hazard problem with a participation constraint requiring the principal to deliver a minimal reservation utility $U^0 \geq 0$. This requirement implies the additional constraint

$$p(e^*) \cdot W_H + (1 - p(e^*)) \cdot W_L - h(e^*) \geq U^0 \quad (\text{P})$$

which must be appended to problem (CM). The following result characterizes the solution to this augmented problem as a function of the agent's reservation utility.

Proposition 6. *There exists a unique solution $\mathbf{W}^*$ to problem (CM)+(P) for every $U^0 \geq 0$. There exists a unique optimal contract $\mathcal{C}^* = (Q_s^*, T_s^*)_{s=L,H}$ under (P), which satisfies $Q_s^* = Q_s^{**}(W_s^*)$ and $T_s^* = T_s^{**}(W_s^*)$ for each $s = L, H$.*

Further, there exist thresholds $\underline{U} > 0$ and $U^{eff} > \max\{\underline{U}, \bar{W}\}$ such that:

- (P) binds iff $U^0 > \underline{U}$,
- W_H^* and W_L^* are increasing in U^0 on $[\underline{U}, \infty)$,
- $W_L^* \geq \bar{W}$ iff $U^0 \geq U^{eff}$,
- Surplus is increasing in U^0 on $[\underline{U}, U^{eff}]$,

- *Surplus equals its efficient level when $U^0 \geq U^{eff}$.*

This result demonstrates how a binding participation constraint tends to relieve the distortion induced by the surplus extraction motive. As the agent is increasingly “bought in” to the relationship, the principal loses her ability to extract surplus. As a result, the benefit of distorting non-monetary allocations diminishes. Indeed, when the agent’s reservation utility exceeds U^{eff} , the credit allocation is efficient and U^0 affects only the split of surplus between the principal and agent.

The connection between efficiency and non-monetary distortions is reflected in the result that surplus becomes efficient exactly when all gross utilities exceed $\bar{W}$. At this point the agent receives credit allocation $\bar{W}$ unconditionally, and he is incentivized to exert effort solely by a performance-dependent monetary payment. Since the agent receives moral hazard rents from performance pay, the minimal reservation utility required for efficiency exceeds $\bar{W}$. For reservation utilities between this level and U^{eff} , the principal continues to distort non-monetary allocations in order to reduce the agent’s rents.

8 Moral hazard Rents

We now study the extent to which the agent earns rents from the presence of moral hazard. When only performance pay is used to provide incentives, these rents are invariably positive. By contrast, the use of punishments to provide incentives destroys surplus and may leave the agent worse off than under observable effort. In this section, we formalize this intuition and establish conditions under which the rents arising from moral hazard are negative.

Were effort observable, the principal could incentivize effort by delivering just enough gross utility to compensate the agent for his disutility of effort. If target effort is low, compliance can be enforced simply by threatening to withhold the baseline credit allocation. On the other hand, once target effort is sufficiently high, the agent must be promised additional rewards to make effort worthwhile.

The optimal gross utility which incentivizes observable effort is therefore

$$W^{obs} = \max\{Q^0, h(e^*)\}.$$

The corresponding net utility achieved by the agent under observable effort is

$$U^{obs} = \max\{Q^0 - h(e^*), 0\}.$$

Note that, under observable effort, net utility is initially decreasing in the disutility of effort. Essentially, effort serves as a channel by which the principal can extract the agent's surplus from non-monetary allocations. (As discussed in Section 7, the lower bound on payments prevents the principal from directly extracting this surplus through transfers.)

Let U^* be the agent's utility under an optimal contract and $\Delta U \equiv U^* - U^{obs}$ be the agent's *moral hazard rents*; i.e., the extra utility he receives from an incentive contract due to imperfectly-observed effort. As mentioned above, U^{obs} provides a lower bound on the agent's utility under moral hazard when only incentive pay is available to incentivize effort. However, the availability of non-monetary incentives means that ΔU is not guaranteed to be positive.

The following proposition identifies conditions under which the agent earns *negative* moral hazard rents under the optimal contract. The statement of the proposition utilizes the effort disutility parameterization $h(e) = \eta \cdot h_0(e)$ introduced in Section 6.1.

Proposition 7. *Suppose that R' is strictly concave near $Q = Q^0$. Then $\Delta U < 0$ for $\eta > 0$ sufficiently small.*

Intuitively, when η is small, the principal need not give any rents to the agent to elicit effort. By setting $W_H - W_L = h'(e^*)/\Delta p$ and $p(e^*)W_H + (1 - p(e^*))W_L = Q^0$, she can incentivize effort while delivering the same expected gross utility $W^{obs} = Q^0$ as under observable effort. However, she may wish to deliver the agent a higher or lower gross utility depending on the relative costs of rewards and punishments. When R' is locally concave, punishments are cheaper to deliver than rewards and the principal optimally reduces the agent's gross utility relative to the observable-effort benchmark. As a result, he earns negative moral hazard rents.

This logic relies on effort disutility being sufficiently small that $W_L > 0$ when the agent is delivered a gross utility of Q^0 in an incentive-compatible contract. Otherwise, punishments cannot be further increased and higher effort disutility must be overcome with larger rewards. In this regime, rents are increasing in η (as in a standard performance-pay model) and eventually positive. Indeed, when η is large enough

that $U^{obs} = 0$, then trivially $\Delta U = U^* > 0$. (Recall from Proposition 6 that the agent receives a positive utility under the optimal contract even absent a binding participation constraint.) Hence, moral hazard rents are always positive when the agent is sufficiently hard to incentivize.

9 Many-Outcome Model

We now extend our model to allow for more than two performance levels. We consider an environment with $S \geq 2$ possible realizations of the performance signal. Each performance level $s = 1, \dots, S$ arises with probability $p_s(e)$ when effort is $e \in [0, 1]$, where p_s is differentiable with respect to e and positive for every $e > 0$. In this framework, our baseline model corresponds to $S = 2$ and $p_2(e) = 1 - p_1(e) = p(e)$.

A key metric controlling optimal incentives is the (local) *informativeness* of each signal: $I_s(e) \equiv p'_s(e)/p_s(e)$. Fixing a target effort $e^* \in (0, 1]$, it is without loss to relabel performance levels so that $I_s(e^*)$ is nondecreasing in s . (In most applications, higher performance is a stronger signal of effort at all effort levels, in which case this ordering is independent of e^* .) To streamline the statement of results, we further assume that $I_s(e^*) \neq 0$ and $I_{s+1}(e^*) > I_s(e^*)$ for each s . In other words, all performance levels are informative and no two levels are equally informative.⁶

Letting $\sigma \equiv \min\{s : I_s(e^*) > 0\}$, we refer to all performance levels $s \geq \sigma$ as *positive signals* about effort and all other signals as *negative signals*. Since $\sum_{s=1}^S p_s(e) = 1$ for all e , it must be that $I_1(e^*) < 0$ while $I_S(e^*) > 0$. Hence there exists at least one positive and one negative signal. In our baseline model, high performance is a positive signal of effort while low performance is a negative signal.

It is well known that in many-outcome moral hazard models, the agent's utility is not guaranteed to be concave in effort under an arbitrary incentive contract. The validity of the first-order approach to incentive compatibility is therefore not guaranteed. To maintain tractability and facilitate comparison with our main results, we restrict attention to environments in which non-local incentive constraints are non-binding.

⁶Uninformative levels are given no reward or punishment in the optimal contract. If there are multiple equally-informative levels, then there exists an optimal contract which delivers the agent the same gross utility following all levels with the same informativeness. This contract is further uniquely optimal if $s = S$ is the unique most-informative positive signal, or else if monetary payments are not used in the optimal contract.

In this relaxed environment, the sole incentive-compatibility constraint is

$$\sum_{s=1}^S p_s(e^*) I_s(e^*) W_s \geq h'(e^*) \quad (\text{IC-S})$$

The principal chooses utility promises $\mathbf{W} = (W_1, \dots, W_S)$ to solve the cost-minimization problem

$$\min_{\mathbf{W} \geq 0} \sum_{s=1}^S p_s(e^*) C(W_s) \quad \text{s.t.} \quad (\text{IC-S}), \quad (\text{CM-S})$$

where C is the cost function characterized in Proposition 1. The following result generalizes Theorem 1 to the many-outcome environment.

Theorem 2. *Problem (CM-S) has a unique solution $\mathbf{W}^*$. It satisfies:*

- $W_{s+1}^* \geq W_s^*$ for each $s = 1, \dots, S-1$, with a strict inequality whenever $W_{s+1}^* > 0$,
- $\text{sign}(W_s^* - Q^0) = \text{sign}(I_s(e^*))$ for each $s = 1, \dots, S$,
- $\bar{W} > W_{S-1}^*$.

Whenever non-local incentive-compatibility constraints are non-binding, there exists a unique optimal contract $\mathcal{C}^ = (Q_s^*, T_s^*)_{s=1, \dots, S}$ satisfying $Q_s^* = Q^{**}(W_s^*)$ and $T_s^* = T^{**}(W_s^*)$ for each $s = 1, \dots, S$.*

This characterization naturally generalizes the key properties of the two-outcome contract. In particular, all positive signals are rewarded with utility levels beyond the level Q^0 associated with the baseline credit allocation, while all negative signals are punished with utility levels below Q^0 . As a corollary, non-monetary allocations are distorted at every performance level (since we have assumed that all signals are strictly informative). And since there exists at least one strictly informative positive and negative signal, both rewards and punishments are used in an optimal contract.

New to the many-outcome setting is a ranking *within* the classes of positive and negative signals. As the signal realization becomes more informative, the associated reward or punishment becomes correspondingly stronger. Further, monetary payments are made only possibly for the most informative positive signal, i.e., for performance level S . Intuitively, as incentives grow, the optimal reward associated

with performance level S reaches $\overline{W}$ before all other performance levels. Beyond this point, marginal incentives are optimally delivered by further rewarding that performance level, since marginal incentive costs are constant for that level but rising for all other levels. This bang-bang result echos a familiar finding in performance pay problems with a risk-neutral agent, where all rewards are optimally concentrated on the single most-informative positive signal.

The following proposition formalizes this reasoning through a comparative static involving the disutility of effort. It is a direct generalization of Proposition 2 to a many-outcome environment. As in that result, we adjust the agent's disutility by decomposing the effort cost function as $h(e) = \eta \cdot h_0(e)$, where η is a disutility parameter.

Proposition 8. *There exists an $\overline{\eta} > 0$ such that:*

- W_S^* is increasing in η and $W_S^* > \overline{W}$ iff $\eta > \overline{\eta}$,
- W_s^* is increasing in η on $[0, \overline{\eta}]$ and constant on $[\overline{\eta}, \infty)$ for each $s = \sigma, \dots, S - 1$
- W_s^* is nonincreasing in η , and decreasing whenever it is positive, on $[0, \overline{\eta}]$ and constant on $[\overline{\eta}, \infty)$ for each $s = 1, \dots, \sigma - 1$

This result implies a natural generalization of our pecking-order result to the many-signal case: Marginal incentives are provided initially with non-monetary instruments and eventually with money as total incentives grow stronger.

It is instructive to contrast these findings with the outcome when only performance pay is available. In that benchmark, the principal would make a monetary payment only following the most-informative outcome, concentrating all incentive power on that signal realization due to risk neutrality over money. By contrast, in our setting costly non-monetary distortions are used to provide incentives at *every* performance level, even when $\eta > \overline{\eta}$. Effectively, the availability of non-monetary incentives generates a form of risk aversion on the margin at low incentive levels, eventually giving way to risk neutrality on the margin at high incentive levels.

Finally, suppose that non-monetary surplus is performance dependent. All of the results just established continue to hold so long as higher performance is simultaneously a stronger signal of effort and of non-monetary surplus: that is, when $R'_{s+1} > R'_s$ using the informativeness ranking of outcomes. This assumption is natural in many

applications where output is a composite of effort and a latent type variable. In such environments, optimal utility promises rise with informativeness, and a monetary payment only possibly following the best performance realization.

The latter claim is not totally straightforward, because monetary payments are used sooner to deliver utility promises following poor performance. It can be established by writing the first-order conditions for optimality in problem (CM-S):

$$C'_s(W_s^*) = I_s(e^*) \cdot \lambda^*,$$

where λ^* is the Lagrange multiplier on (IC-S). The marginal costs of utility promises are therefore ranked according to informativeness. Critically, the threshold marginal cost for using money is 1 at all performance levels. Hence even when non-monetary surplus is performance dependent, money is paid out first (and only) following the highest performance. Performance-dependent non-monetary surplus amplifies dispersion between utility promises but does not overturn the qualitative features of the pecking order.

10 Concluding Remarks

In this paper we have proposed and analyzed a novel model of performance incentives, incorporating both standard monetary payments as well as a portfolio of non-monetary instruments. Our motivations were two-fold. First, given the ubiquity of such instruments as incentive tools in many real organizations, we wished to provide a conceptual framework for incorporating them into a standard moral hazard framework. And second, in light of the widely-recognized inefficiency of incentivizing through rewards such as career advancement, we wished to understand whether and to what extent their usage can be justified by contract-theoretic considerations.

Our analysis provides a simple, universal explanation for the use of non-monetary instruments as performance incentives: The allocative costs they incur are smaller than the costs of monetary incentives for providing baseline incentives. As a result, our model predicts that instruments such as career advancement should be used pervasively to provide incentives. Monetary incentives, on the other hand, should be used only as a supplemental source of incentives when effort is especially costly.

Our analysis also highlights that non-monetary incentives can be used both to re-

ward good performance as well as to *punish* bad performance. The optimal balance of rewards and punishments is an important factor in the incidence of monetary incentives. In particular, our analysis suggests that monetary rewards are most likely to be used in star roles, where workers strive for occasional successes for which they receive large rewards. By contrast, in guardian roles, where workers protect against rare disasters, incentives are optimally provisioned through large non-monetary punishments for failure, for instance demotion or firing.

We have focused on the design of performance incentives for a single agent, abstracting from any externalities across agents. In some applications, allocative rewards such as promotions may be in limited aggregate supply. In those environments, part of the cost of awarding a credit to one agent is the shadow cost of reducing the supply of credits available to other agents. In other words, the credit profit function R may be determined endogenously when designing incentive schemes across an entire organization. Our model provides an important building block for future work studying such interactions between resource allocation and incentive design problems within organizations.

We have also abstracted from any uncertainty regarding the agent’s preferences for credits. In reality, workers in an organization may exhibit heterogeneous rates of substitution between money and non-monetary instruments such as promotions. In such environments, the organization may profitably screen workers for their rate of substitution by offering a menu of contracts with differing ratios of monetary and non-monetary incentives; and also potentially with differing overall incentives, depending on the nature of the organization’s effort goal. Future work analyzing such a model would shed light on the optimal use of parallel career tracks which involve similar work but differing prospects of career advancement.

A Proofs

A.1 Proof of Proposition 1

Using the constraint to eliminate T from the objective in problem (UD) yields the reduced objective

$$R(Q^0) - R(Q) + W - Q,$$

to be minimized over the domain $Q \in [0, W]$ given that $T = W - Q \geq 0$. This objective is continuous and strictly convex, and so has a unique minimum $Q^{**}(W)$ over the convex domain $[0, W]$. Letting $T^{**}(W) = W - Q^{**}(W)$, it follows that $(Q^{**}(W), T^{**}(W))$ is the unique solution to (UD).

Over the extended domain $\mathbb{R}_+$, the reduced objective is uniquely minimized at $Q = \bar{W}$, where $\bar{W} \equiv (R')^{-1}(-1)$. Hence if $W \geq \bar{W}$, then $Q^{**}(W) = \bar{W}$ and $T^{**}(W) = W - \bar{W}$. Otherwise, the reduced objective is uniquely minimized at $Q = W$ over the domain $[0, W]$, in which case $Q^{**}(W) = W$ and $T^{**}(W) = 0$. Hence Q^{**} and T^{**} satisfy the stated comparative statics in W .

The value C of the problem satisfies

$$C(W) = \begin{cases} R(Q^0) - R(W), & W \leq \bar{W} \\ R(Q^0) - R(\bar{W}) + W - \bar{W}, & W > \bar{W} \end{cases}$$

Hence for all $W \neq \bar{W}$ the utility cost function is differentiable and satisfies

$$C'(W) = \begin{cases} -R'(W), & W < \bar{W} \\ 1, & W > \bar{W} \end{cases}$$

The first expression is further valid for the left derivative at $W = \bar{W}$, while the second expression is valid for the right derivative at this point. Since $R'(\bar{W}) = -1$, the left- and right-hand derivative agree, and so C' exists and is continuous everywhere.

Since R' is decreasing, it follows that C' is nondecreasing everywhere and increasing on $[0, \bar{W}]$. Further, $C'(W) = 0$ when $W = Q^0 < \bar{W}$, so that Q^0 is the unique global minimum of C . At this point $Q^{**}(W) = Q^0$ and $T^{**}(W) = 0$, so that $C(Q^0) = 0$.

A.2 Proof of Theorem 1

Maximizing the principal's profits

$$p(e^*) \cdot (R(Q_H) - T_H) + (1 - p(e^*)) \cdot (R(Q_L) - T_L)$$

is equivalent to minimizing net costs

$$p(e^*) \cdot (R(Q^0) - R(Q_H) + T_H) + (1 - p(e^*)) \cdot (R(Q^0) - R(Q_L) + T_L).$$

Fix any IC contract $\mathcal{C} = (Q_s, T_s)_{s=L,H}$ and let $W_H = Q_H + T_H$ and $W_L = Q_L + T_L$. If $(Q_H, T_H) \neq (Q^{**}(W_H), T^{**}(W_H))$, the contract $\mathcal{C}' = (Q_L, T_L, Q^{**}(W_H), T^{**}(W_H))$ reduces the principal's net costs while delivering the same utility

$$p(e) \cdot W_H + (1 - p(e)) \cdot W_L - h(e)$$

to the agent for every effort level e . Hence $\mathcal{C}'$ is also IC and reduces the principal's net costs, so $\mathcal{C}$ could not have been optimal. A similar argument implies that $\mathcal{C}$ cannot be optimal if $(Q_L, T_L) \neq (Q^{**}(W_L), T^{**}(W_L))$. Hence it is without loss to consider contracts satisfying $\mathcal{C} = (Q^{**}(W_s), T^{**}(W_s))_{s=L,H}$ for some $\mathbf{W} = (W_L, W_H)$.

Given that

$$C(W) = R(Q^0) - R(Q^{**}(W)) + T(Q^{**}(W)),$$

an IC contract $\mathcal{C} = (Q^{**}(W_s), T^{**}(W_s))_{s=L,H}$ is cost minimizing iff $\mathbf{W}$ minimizes net costs

$$p(e^*) \cdot C(W_H) + (1 - p(e^*)) \cdot C(W_L)$$

subject to

$$e^* \in \arg \max_{e \in [0,1]} p(e) \cdot W_H + (1 - p(e)) \cdot W_L - h(e).$$

Since p is linear and h is C^1 and strictly convex and $e^* > 0$, the incentive-compatibility condition is satisfied iff e^* satisfies the FOC

$$W_H - W_L \geq \frac{h'(e)}{\Delta p},$$

with equality if $e^* < 1$.

Consider passing to the relaxed problem enforcing (IC), i.e., where the FOC is enforced as an inequality no matter the choice of $e^* \in (0, 1]$. Since the net cost function is convex, its minimum in this relaxed problem is the same as its unconstrained minimum in case (IC) is slack at the optimum. But its unconstrained minimum is $W_H = W_L = Q^0$, which violates (IC). Hence (IC) is active at the optimum of the relaxed problem, and so any solution to the relaxed problem also solves the unrelaxed problem, i.e., enforcing the true incentive-compatibility constraint. In other words, $\mathbf{W}$ is cost-minimizing iff it solves problem (CM).

Using the active constraint $W_H = W_L + h'(e^*)/\Delta p$ to eliminate W_H from the net

cost function yields the reduced objective

$$p(e^*) \cdot C(W_L + h'(e^*)/\Delta p) + (1 - p(e^*)) \cdot C(W_L).$$

This objective function is continuous and convex, and strictly so on $[0, \overline{W}]$. Further, it is increasing for $W_L \geq Q^0$. Hence it must have a unique minimum W_L^* on the domain $\mathbb{R}_+$, which lies in $[0, Q^0]$. Letting $W_H^* = W_L^* + h'(e^*)/\Delta p$, it must be that $\mathbf{W}^* = (W_H^*, W_L^*)$ is the unique solution to (CM).

Finally, we show that $W_L^* < Q^0 < Q_H^*$. The solution $\mathbf{W}^*$ must satisfy the first-order condition of the reduced objective

$$p(e^*)C'(W_H^*) + (1 - p(e^*))C'(W_L^*) \geq 0,$$

with equality if $W_L^* > 0$. Since also $W_H^* = W_L^* + h'(e^*)/\Delta p > W_L^*$, if $W_L^* \geq Q^0$, then the left-hand side of the FOC is positive, inconsistent with $W_L^* > 0$. So $W_L^* < Q^0$. And $W_L^* < Q^0$ implies $C'(W_L^*) < 0$, in which case the FOC requires $C'(W_H^*) > 0$, i.e., $W_H^* > Q^0$.

A.3 Proof of Proposition 2

Problem (CM) is a convex minimization problem involving only inequality constraints, and so there must exist a Lagrange multiplier $\lambda^*(\eta)$ on (IC) such that $\mathbf{W}^*(\eta)$ minimizes the Lagrangian

$$\mathcal{L}(\mathbf{W}; \lambda^*(\eta)) = p(e^*) \cdot C(W_H) + (1 - p(e^*)) \cdot C(W_L) - \lambda^*(\eta) \cdot (W_H - W_L - h'(e^*)/\Delta p).$$

Hence $\mathbf{W}^*(\eta)$ must satisfy the Lagrangian FOCs

$$C'(W_H^*(\eta)) = \frac{\lambda^*(\eta)}{p(e^*)}, \quad C'(W_L^*(\eta)) \geq -\frac{\lambda^*(\eta)}{1 - p(e^*)},$$

with the latter condition holding as an equality whenever $W_L^*(\eta) > 0$. We first establish that $\lambda^*(\eta)$ is bounded above by $p(e^*)$, uniquely defined, nondecreasing in η , increasing whenever $\lambda^*(\eta) < p(e^*)$, and continuous.

If $\lambda^*(\eta) > p(e^*)$, then the Lagrangian FOC for W_H cannot be satisfied by any choice of $W_H^*(\eta)$, a contradiction. So $\lambda^*(\eta) \leq p(e^*)$. If there are two distinct multipliers

$\lambda^1(\eta)$ and $\lambda^2(\eta)$ for which $\mathbf{W}^*(\eta)$ minimizes the Lagrangian, then the smaller of these multipliers must be less than $p(e^*)$. But then $W_H^*(\eta)$, which is unique, cannot satisfy the Lagrangian FOC for W_H for both multipliers, a contradiction. So $\lambda^*(\eta)$ is unique.

Next, suppose that $\lambda^*(\eta') < \lambda^*(\eta'')$ for some $\eta' > \eta''$. Then the Lagrangian FOCs imply that $W_H^*(\eta') < W_H^*(\eta'')$ while $W_L^*(\eta') \geq W_L^*(\eta'')$, meaning that $W_H^*(\eta') - W_L^*(\eta') < W_H^*(\eta'') - W_L^*(\eta'')$. But since (IC) is active at the optimum when $\eta = \eta''$, it follows that

$$W_H^*(\eta') - W_L^*(\eta') < \eta'' \cdot h'_0(e^*)/\Delta p < \eta' \cdot h'_0(e^*)/\Delta p,$$

i.e., $\mathbf{W}^*(\eta')$ violates (IC) when $\eta = \eta'$. This contradiction implies that λ^* must be nondecreasing in η . If $\lambda^*(\eta') = \lambda^*(\eta'') < p(e^*)$, then $W_H^*(\eta') = W_H^*(\eta'')$ and $W_L^*(\eta') = W_L^*(\eta'')$, again implying a violation of (IC) when $\eta = \eta'$. So λ^* must be increasing in η whenever it lies below $p(e^*)$.

Since λ^* is nondecreasing, it can have only jump-type discontinuities. Suppose that it is discontinuous at $\eta = \eta'$, say with $\lambda^*(\eta') > \lambda^*(\eta'-)$. Then since $\lambda^*(\eta'-) < p(e^*) < p(e^*)$, the Lagrangian FOCs imply that $W_H^*(\eta'-) < W_H^*(\eta')$ and $W_L^*(\eta'-) \leq W_L^*(\eta')$. But then the fact that (IC) is active at the optimum for all η implies

$$\eta' \cdot h'_0(e^*)/\Delta p = W_H^*(\eta') - W_L^*(\eta') > W_H^*(\eta'-) - W_L^*(\eta'-) = \eta' \cdot h'_0(e^*)/\Delta p,$$

a contradiction. A similar argument rules out a discontinuity of the form $\lambda^*(\eta'+) > \lambda^*(\eta')$. So λ^* must be continuous.

Now, let $\bar{\eta} \equiv \inf\{\eta : \lambda^*(\eta) = p^*(e)\}$. We first establish that $\bar{\eta} \in (0, \infty)$. Since (Q^0, Q^0) satisfies (IC) when $\eta = 0$, it follows that $\mathbf{W}^*(\eta) = (Q^0, Q^0)$. Hence $\lambda^*(0) = 0 < p^*(e)$ and $\bar{\eta} > 0$. If $\bar{\eta} = \infty$, then the Lagrangian FOC for W_H implies that $W_H^*(\eta) < \bar{W}$ for all η , which must eventually violate (IC) for sufficiently large η . So $\bar{\eta} < \infty$.

For all η , the Lagrangian FOC for $W_L(\eta)$ uniquely pins down W_L^* in terms of $\lambda^*(\eta)$. Given that λ^* is continuous and nondecreasing everywhere, it must be that W_L^* is continuous and nonincreasing everywhere. Additionally, on $[0, \bar{\eta}]$ the multiplier is increasing, so that W_L^* must be decreasing on this range whenever it is positive. Conversely, on $[\bar{\eta}, \infty)$ the multiplier is constant, so that W_L^* must be as well.

Meanwhile, continuity of W_L^* plus the active constraint (IC) implies that W_H^* is continuous in η . on $[0, \bar{\eta})$ the Lagrangian FOC for W_L uniquely pin downs $W_H^* < \bar{W}$ in terms of $\lambda^*(\eta)$. Given that λ^* is increasing on this interval, W_H^* must be as well.

Meanwhile, on $[\bar{\eta}, \infty)$ the Lagrangian FOC implies that $W_H^* \geq \bar{W}$, and constancy of W_L^* on this interval plus the active constraint (IC) implies that W_H^* is increasing. In particular, $W_H^*(\eta) > \bar{W}$ for all $\eta > \bar{\eta}$. Finally, continuity of λ^* implies that $\lambda^*(\bar{\eta}) = p(e^*)$, so that $W_H^*(\bar{\eta}-) = \bar{W}$. Continuity of W_H^* therefore implies $W_H^*(\bar{\eta}) = \bar{W}$.

A.4 Proof of Proposition 3

Recall from the proof of Theorem 1 that W_L^* is the unique minimizer of

$$p(e^*) \cdot C(W_L + h'(e^*)/\Delta p) + (1 - p(e^*)) \cdot C(W_L).$$

Since this objective is convex, $W_L^* < Q^0$ is uniquely characterized by the first-order condition

$$p(e^*) \cdot C'(W_L^* + h'(e^*)/\Delta p) + (1 - p(e^*)) \cdot C'(W_L^*) \geq 0,$$

with equality whenever $W_L^* > 0$. Since $p(e^*)$ increases in baseline performance while $1 - p(e^*)$ decreases, the left-hand side of this condition is increasing in baseline performance. Hence W_L^* is nonincreasing in baseline performance, and is decreasing whenever positive. The active constraint (IC) implies that W_H^* obeys the same comparative static.

A.5 Proof of Proposition 4

Fix target effort levels e' and $e'' > e'$, disutility η' , and baseline performance $\underline{p}'$. Note that the objective and constraint in problem (CM) under parameters $(e^*, \eta, \underline{p}) = (e'', \eta', \underline{p}')$ are the same as under parameters $(e^*, \eta, \underline{p}) = (e', \eta'', \underline{p}'')$, where $\eta'' \equiv h'(e'')/h'(e') \cdot \eta' > \eta'$ and $\underline{p}'' \equiv \underline{p}' + \Delta p \cdot (e'' - e') > \underline{p}'$. (Throughout this proof, Δp is held fixed.) Hence the optimal contract is also the same across these two environments.

Proposition 2 implies that W_L^* does not increase when passing from parameters $(e', \eta', \underline{p}')$ to parameters $(e', \eta'', \underline{p}')$. Further, Proposition 3 implies that W_L^* does not increase when passing from parameters $(e', \eta'', \underline{p}')$ to $(e', \eta'', \underline{p}'')$, and decreases if W_L^* is initially positive. Hence W_L^* does not increase, and decreases if initially positive, when passing from parameters $(e', \eta', \underline{p}')$ to $(e', \eta'', \underline{p}'')$. Equivalently, this result holds when passing to parameters $(e'', \eta', \underline{p}')$.

A.6 Proof of Proposition 5

Recall from the proof of Theorem 1 that $W_L^*(e^*)$ is the unique minimizer of

$$p(e^*) \cdot C(W_L + h'(e^*)/\Delta p) + (1 - p(e^*)) \cdot C(W_L).$$

For $e^* < \bar{e}$ it must be that $W_L^*(e^*)$ satisfies the first-order condition

$$\frac{p(e^*)}{1 - p(e^*)} \cdot C'(W_L^*(e^*) + h'(e^*)/\Delta p) + C'(W_L^*(e^*)) = 0.$$

Since this first-order condition is continuous in e^* , the same result must also hold at $e^* = \bar{e}$. Equivalently, for all $e^* \in (0, \bar{e}]$ the utility promise $W_H^*(e^*)$ must satisfy

$$\frac{p(e^*)}{1 - p(e^*)} \cdot C'(W_H^*(e^*)) + C'(W_H^*(e^*) - h'(e^*)/\Delta p) = 0.$$

Define

$$\Phi(W_H; e^*) \equiv \frac{p(e^*)}{1 - p(e^*)} \cdot C'(W_H) - R'(W_H - h'(e^*)/\Delta p).$$

This function is increasing in W_H and vanishes at $W_H^*(e^*)$ given that $W_L^* = W_H^* - h'(e^*)/\Delta p < Q^0$ and therefore $C'(W_L^*) = -R'(W_L^*)$. Hence $W_H^*(e^*)$ is uniquely characterized by the condition

$$\Phi(W_H^*(e^*); e^*) = 0$$

for $e^* \in (0, \bar{e}]$.

If R' is concave, then $-R'$ is convex. If additionally h' is concave, then the fact that $-R'$ is an increasing function implies that the composite function $-R'(W_H - h'(e^*)/\Delta p)$ is convex in e^* . Meanwhile, $p(e^*)/(1 - p(e^*))$ is strictly convex in e^* . Hence $\Phi(W_H; e^*)$ is strictly convex in e^* for each $W_H > Q^0$.

Fix target effort levels e' and $e'' > e'$ in $(0, \bar{e}]$, and let $W' \equiv W_H^*(e')$ and $W'' \equiv W_H^*(e'')$. Choose $\lambda \in (0, 1)$ and let $e''' \equiv \lambda e' + (1 - \lambda)e''$ and $W''' \equiv W_H^*(e''')$. If $W' \leq W''$, then $\Phi(W'; e'') \leq 0 = \Phi(W'; e')$, in which case strict convexity of Φ in e^* ensures that

$$\Phi(W'; e''') < \lambda \Phi(W'; e') + (1 - \lambda) \Phi(W'; e'') \leq 0.$$

Hence $W''' > W'$. On the other hand, if $W' > W''$, then $\Phi(W''; e') < 0 = \Phi(W''; e'')$.

In that case, strict convexity of Φ in e^* implies that

$$\Phi(W''; e''') < \lambda \Phi(W''; e') + (1 - \lambda) \Phi(W''; e'') < 0.$$

Hence $W''' > W''$. In both cases, $W''' > \min\{W', W''\}$, establishing strict quasiconcavity.

We next derive sufficient conditions for non-monotonicity. Since h is convex, $h' \leq h'(1)$. If further h' is concave, then $h'(1) < \infty$. Then since $W_H^* > Q^0$, it must be that $W_L^* > Q^0 - h'(1)/\Delta p$. If $Q^0 > h'(1)/\Delta p$, then $W_L^* > 0$ for all e^* and $\bar{e} = 1$. Let $p^* \in (\underline{p}, 1)$ be any lower bound on $\bar{p}$. Then if $Q^0 > h'(1)/(p^* - \underline{p})$, we have $W_L^* > 0$ for all e^* and any $\bar{p} \geq \bar{p}^*$. Going forward, we impose this lower bound on Q^0 and restrict attention to $\bar{p} \geq \bar{p}^*$. Under this lower bound, $W_H^*(e^*)$ satisfies $\Phi(W_H^*(e^*); e^*) = 0$ for every $e^* \in (0, 1]$.

The assumption $h'(0) = 0$ ensures that vanishing incentives are necessary for small effort levels and therefore $\lim_{e^* \rightarrow 0} W_H^*(e^*) = Q^0$. Then since $W_H^*(e^*) > Q^0$ for all $e^* > 0$, it must be that W_H^* is initially increasing. To obtain non-monotonicity of W_H^* in e^* over this range, it is therefore sufficient to establish that $C'(W_H^*(1)) < C'(W_H^*(1/2))$. Since $W_L^*(1) \in [0, Q^0]$ and $p(1) = \bar{p}$, it follows that

$$\lim_{\bar{p} \rightarrow 1} C'(W_H^*(1)) = \lim_{\bar{p} \rightarrow 1} \frac{1 - \bar{p}}{\bar{p}} \cdot R'(W_L^*(1)) = 0.$$

Meanwhile, Φ is continuous in $(W_H, \bar{p})$, meaning that $\lim_{\bar{p} \rightarrow 1} W_H^*(1/2)$ exists and is the unique solution to the equation $\lim_{\bar{p} \rightarrow 1} \Phi(W_H; 1/2) = 0$, which may be written

$$\frac{1 + \underline{p}}{1 - \underline{p}} \cdot C'(W_H) = R'(W_H - h'(1/2)/(1 - \underline{p})).$$

Since this solution must be greater than Q^0 , it follows that $\lim_{\bar{p} \rightarrow 1} C'(W_H^*(1/2)) > 0$. In other words, $C'(W_H^*(1)) < C'(W_H^*(1/2))$ for $\bar{p}$ sufficiently close to 1, as desired.

Finally, we establish conditions under which $\max_{e^* \in (0, \bar{e}]} W_H^*(e^*) > \bar{W}$. Supposing that $\bar{W} > h'(e^*)/\Delta p$, the condition $W_H^*(e^*) > \bar{W}$ is equivalent to

$$\frac{p(e^*)}{1 - p(e^*)} < R'(\bar{W} - \eta \cdot h'_0(e^*)/\Delta p).$$

Letting $\hat{W} \equiv (\bar{R}')^{-1}(-1) = \bar{W} - Q^0$, this condition may be written in terms of $\bar{R}$ as

$$\frac{p(e^*)}{1 - p(e^*)} < \bar{R}'(\hat{W} - \eta \cdot h'_0(e^*)/\Delta p).$$

Fix any $e^0 \in (0, 1)$. Since $\bar{R}'(-\infty) = \infty$, it follows that the condition just stated is satisfied for $e^* = e^0$ whenever η is sufficiently large, independent of the value of $\bar{p} > \underline{p}$. Now, the bound $\bar{W} > h'(e^0)/\Delta p$ may be written in terms of Q^0 as $Q^0 > \eta \cdot h'_0(e^0)/\Delta p - \hat{W}$. Hence, for any choice of η and $p^* > \underline{p}$, the bound $\bar{W} > h'(e^0)/\Delta p$ may be satisfied for all $\bar{p} \geq p^*$ by setting $Q^0 > \eta \cdot h'_0(e^0)/(p^* - \underline{p}) - \hat{W}$. In other words, when η is sufficiently large, $W_H^*(e^0) > \bar{W}$ for sufficiently large Q^0 and $\bar{p}$ sufficiently close to 1.

A.7 Proof of Proposition 6

The fact that a contract is feasible and profit-maximizing if and only if it satisfies $(Q_s^*, T_s^*)_{s=L,H} = (Q_s^{**}(W_s^*), T_s^{**}(W_s^*))_{s=L,H}$ for $\mathbf{W}^*$ solving (CM)+(P) follows along similar lines to the proof of Theorem 1.

Let $\mathbf{W}^0$ be the solution to problem (CM). Let

$$\underline{U} \equiv p(e^*) \cdot W_H^0 + (1 - p(e^*)) \cdot W_L^0 - h(e^*)$$

be the utility the agent enjoys from these utility promises. Since $W_H^0 = W_L^0 + h'(e^*)/\Delta p$, we have

$$\underline{U} = W_L^0 + \left(\frac{p}{\Delta p} + e^* \right) \cdot h'(e^*) - h(e^*).$$

Strict convexity of h implies that

$$e^* \cdot h'(e^*) > \int_0^{e^*} h'(e) de = h(e^*),$$

so that $\underline{U} > 0$.

If $\underline{U} \geq U^0$, then $\mathbf{W}^0$ must be the unique solution to problem (CM)+(P). Otherwise, since (CM)+(P) is a convex optimization problem, (P) must be active at any solution to the problem. If there exists a solution for which (IC) is slack, then convexity of the problem implies that this solution must also solve the relaxed problem

without (IC). But the unique solution to that problem is $W_H = W_L = U^0 + h(e^*)$, which violates IC. So (IC) must also be active at any solution to the problem. The system of active constraints (IC)+(P) admits a unique solution

$$W_L^* = U^0 - \left(\frac{p}{\Delta p} + e^* \right) \cdot h'(e^*) + h(e^*), \quad W_H^* = U^0 + \left(\frac{1-p}{\Delta p} - e^* \right) \cdot h'(e^*) + h(e^*),$$

which must trivially be the unique solution to the problem (CM)+(P).

These solutions are increasing in U^0 by inspection, and W_L^* crosses $\bar{W}$ at

$$U^{eff} \equiv \bar{W} + \left(\frac{p}{\Delta p} + e^* \right) \cdot h'(e^*) - h(e^*).$$

Since $e^* \cdot h'(e^*) > h(e^*)$, this threshold satisfies $U^{eff} > \bar{W}$. Additionally, since $W_L^0 < Q^0 < \bar{W}$ and $W_L^* = W_H^*$ when $U^0 = \underline{U}$, it must be that $U^{eff} > \underline{U}$.

Under the optimal contract, the credit allocation at each performance level is $Q_s^* = \min\{W_s^*, \bar{W}\}$. Hence expected non-monetary surplus under the optimal contract is

$$S^* = p(e^*) \cdot S(\min\{W_H^*, \bar{W}\}) + (1 - p(e^*)) \cdot S(\min\{W_L^*, \bar{W}\}).$$

S is increasing on $[0, \bar{W}]$, while W_H^* and W_L^* are nondecreasing in U^0 and W_L^* is increasing and bounded above by $\bar{W}$ on $[\underline{U}, U^{eff}]$. Hence S^* is increasing in U^0 on $[\underline{U}, U^{eff}]$ and equal to $S(\bar{W})$ afterward. Since S is maximized at $\bar{W}$, non-monetary surplus is efficient for $U^0 \geq U^{eff}$.

A.8 Proof of Proposition 7

For η sufficiently small, the observable-effort utility level is $U^{obs} = Q^0 - h(e^*)$. Hence

$$\Delta U = p(e^*) \cdot W_H^* + (1 - p(e^*)) \cdot W_L^* - Q^0.$$

The proof of Proposition 2 established that $\mathbf{W}^*$ is continuous in η . Then since $W_H^* = W_L^* = Q^0$ when $\eta = 0$, it follows that $W_H^* < \bar{W}$ and $W_L^* > 0$ for η sufficiently small. In that regime, the characterization of $\mathbf{W}^*$ obtained in the proof of Theorem 1 implies that the optimal utility promises $\mathbf{W}^*$ satisfy

$$p(e^*) \cdot R'(W_H^*) + (1 - p(e^*)) \cdot R'(W_L^*) = 0.$$

Since $\lim_{\eta \rightarrow 0} W_H^* = \lim_{\eta \rightarrow 0} W_L^* = Q^0$, continuity of $\mathbf{W}^*$ in η combined with strict concavity of R' near $Q = Q^0$ implies that

$$R'(p(e^*) \cdot W_H^* + (1 - p(e^*)) \cdot W_L^*) > 0$$

when η is sufficiently small. As a result,

$$p(e^*) \cdot W_H^* + (1 - p(e^*)) \cdot W_L^* < Q^0.$$

In other words, $\Delta U < 0$.

A.9 Proof of Theorem 2

The fact that a contract is feasible and profit-maximizing if and only if it satisfies $(Q_s^*, T_s^*)_s = (Q_s^{**}(W_s^*), T_s^{**}(W_s^*))_s$ for $\mathbf{W}^*$ solving (CM)+(P) follows along similar lines to the proof of Theorem 1.

Since problem (CM-S) is convex and involves only inequality constraints, for every solution $\mathbf{W}^*$ there exists a multiplier λ^* on (IC-S) such that $\mathbf{W}^*$ minimizes the Lagrangian

$$\mathcal{L}(W; \lambda^*) = \sum_{s=1}^S p_s(e^*) \left(C(W_s) - \lambda^* I_s(e^*) W_s \right) + \lambda^* \cdot h'(e^*).$$

We establish existence and uniqueness of a solution to problem (CM-S) by showing that there exists a unique multiplier λ for which the minimizer of $\mathcal{L}(W; \lambda)$ satisfies complementary slackness.

If $\lambda > 1/I_S(e^*)$, then unboundedly negative values of $\mathcal{L}$ can be achieved by taking W_S large, and there exists no minimizer of the Lagrangian. So we search for solutions on the range $[0, 1/I_S(e^*)]$. For each $\lambda < 1/I_S(e^*)$, the Lagrangian is concave in each W_s and has a unique minimizer $W_s^{**}(\lambda)$ characterized by the first-order conditions

$$C'(W_s^{**}(\lambda)) \geq I_s(e^*) \cdot \lambda, \quad s = 1, \dots, S,$$

with equality if $W_s^{**}(\lambda) > 0$. Each W_s^{**} for $s \geq \sigma$ is therefore continuous and increasing in λ , while each W_s^{**} for $s < \sigma$ is continuous and nonincreasing, and decreasing

whenever positive. Thus

$$\zeta(\lambda) \equiv \sum_{s=1}^S p_s(e^*) I_s(e^*) W_s^{**}(\lambda)$$

is continuous and increasing in λ . Additionally, $W^{**}(0) = Q^0$ for all s , so that $\zeta(0) = Q^0 \cdot \sum_{s=1}^S p_s(e^*) I_s(e^*) = 0$.

If $\lambda = 1/I_S(e^*)$, then the Lagrangian has a unique minimizer $W_s^{**}(\lambda)$ for each $s \leq S-1$ defined by the same first-order condition as in the $\lambda < 1/I_S(e^*)$ case. Meanwhile the set of optimal W_S is the set $[\overline{W}, \infty)$.

If $\zeta(1/I_S(e^*)) > h'(e^*)$, then the unique λ for which a minimizer of the Lagrangian satisfies complementary slackness is the unique solution $\lambda^* \in (0, 1/I_S(e^*))$ to the equation $\zeta(\lambda) = h'(e^*)$. Hence there exists a unique solution to problem (CM-S). If $\zeta(1/I_S(e^*)) \leq h'(e^*)$, then no choice of $\lambda < 1/I_S(e^*)$ produces a minimizer satisfying complementary slackness. On the other hand, if $\lambda = 1/I_S(e^*)$, then there exists a unique minimizer of the Lagrangian satisfying complementary slackness; namely, the unique $W_S^* \geq \overline{W}$ satisfying

$$p_S(e^*) I_S(e^*) (W_S^* - \overline{W}) = h'(e) - \zeta(1/I_S(e^*)).$$

Hence in this case too there exists a unique solution to problem (CM-S), with corresponding multiplier $\lambda^* = 1/I_S(e^*)$.

The Lagrangian FOCs characterizing $\mathbf{W}^*$ imply that, no matter the optimal multiplier λ^* , utility promises are ordered by informativeness, and the ordering is strict whenever promises are non-negative. Further, the sign of $C'(W_s^*)$ equals the sign of $I_s(e^*)$, implying the sign dependence of $W_s^* - Q^0$ in the theorem statement. Finally, $\lambda^* \leq 1/I_S(e^*)$ implies that $W_{S-1}^* < \overline{W}$.

A.10 Proof of Proposition 8

This result can be established along lines very similar to the proof of Proposition 2. The optimal multiplier λ^* used to characterized optimal utility promises in the proof of Theorem 2 can be shown to be continuous and nondecreasing in η , increasing whenever it lies below $1/I_S(e^*)$, and to eventually reach $1/I_S(e^*)$. The critical threshold $\overline{\eta}$ is the value of η at which λ^* reaches $1/I_S(e^*)$.

References

- Acemoglu, Daron and Alexander Wolitzky (2011). “The Economics of Labor Coercion”. *Econometrica* 79 (2), pp. 555–600.
- Auriol, Emmanuelle and Régis Renault (2008). “Status and incentives”. *The RAND Journal of Economics* 39 (1), pp. 305–326.
- Besley, Timothy and Maitreesh Ghatak (2008). “Status Incentives”. *The American Economic Review* 82 (2). Proceedings of the One Hundred Twentieth Annual Meeting of the American Economic Association (May, 2008), pp. 206–211.
- Biais, Bruno, Thomas Mariotti, Jean-Charles Rochet, and Stéphane Villeneuve (2010). “Large Risks, Limited Liability, and Dynamic Moral Hazard”. *Econometrica* 78 (1), pp. 73–118.
- Che, Yeon-Koo, Elisabetta Iossa, and Patrick Rey (2021). “Prizes versus Contracts as Incentives for Innovation”. *The Review of Economic Studies* 88 (5), pp. 2149–2178.
- Chwe, Michael Suk-Young (1990). “Why Were Workers Whipped? Pain in a Principal-Agent Model”. *The Economic Journal* 100 (403), pp. 1109–1121.
- DeMarzo, Peter M. and Michael J. Fishman (Sept. 2007). “Optimal Long-Term Financial Contracting”. *The Review of Financial Studies* 20 (6), pp. 2079–2128.
- DeMarzo, Peter M. and Yuliy Sannikov (2006). “Optimal Security Design and Dynamic Capital Structure in a Continuous-Time Agency Model”. *The Journal of Finance* 61 (6), pp. 2681–2724.
- Dubey, Pradeep and John Geanakoplos (2020). “Money and Status in a Meritocracy”. Unpublished.
- Georgiadis, George (2024). “Contracting with moral hazard”. In: *Elgar Encyclopedia on the Economics of Competition and Regulation*. Ed. by Michael Noel, pp. 24–37.
- Gibbons, Robert and John Roberts (2013). “2. Economic Theories of Incentives in Organizations”. In: *The Handbook of Organizational Economics*. Ed. by Robert Gibbons and John Roberts. Princeton: Princeton University Press, pp. 56–99.
- Holmström, Bengt (1979). “Moral Hazard and Observability”. *The Bell Journal of Economics* 10 (1), pp. 74–91.
- Hörner, Johannes (2013). “Comments on ‘Contracts’”. In: *Advances in Economics and Econometrics: Tenth World Congress*. Ed. by Daron Acemoglu, Manuel Arellano,

- and Eddie Dekel. Econometric Society Monographs. Cambridge University Press, pp. 172–176.
- Ke, Rongzhu, Jin Li, and Michael Powell (2018). “Managing Careers in Organizations”. *Journal of Labor Economics* 36 (1), pp. 197–252.
- Lee, Jangwoo, William Fuchs, and Martin Dumav (2024). “Exploitation Payoffs and Incentives for Exploration”. Unpublished.
- Madsen, Erik, Basil Williams, and Andrzej Skrzypacz (2024). “Reward Schemes for Autonomous Workers”. Unpublished.
- Marino, Anthony M. and Ján Zábajník (2008). “Work-related perks, agency problems, and optimal incentive contracts”. *The RAND Journal of Economics* 39 (2), pp. 565–585.
- Milgrom, Paul and John Roberts (1992). *Economics, Organization and Management*. Pearson.
- Moldovanu, Benny, Aner Sela, and Xianwen Shi (2007). “Contests for Status”. *Journal of Political Economy* 115 (2), pp. 338–363.
- Prendergast, Canice (1999). “The Provision of Incentives in Firms”. *Journal of Economic Literature* 37 (1), pp. 7–63.
- Sannikov, Yuliy (2008). “A Continuous- Time Version of the Principal: Agent Problem”. *The Review of Economic Studies* 75 (3), pp. 957–984.
- Wolitzky, Alexander (2012). “Career Concerns and Performance Reporting in Optimal Incentive Contracts”. *The B.E. Journal of Theoretical Economics* 12 (1).
- Zhu, John Y. (2013). “Optimal Contracts with Shirking”. *The Review of Economic Studies* 80.2 (283), pp. 812–839.