

Dynamic contracts when the agent's quality is unknown

JULIEN PRAT

CNRS, CREST, and IAE, Barcelona GSE

BOYAN JOVANOVIĆ

Department of Economics, New York University

We solve a long-term contracting problem with symmetric uncertainty about the agent's quality and a hidden action of the agent. As information about quality accumulates, incentives become easier to provide because the agent has less room to manipulate the principal's beliefs. This result is opposite to that in the literature on "career concerns" in which incentives via short-term contracts become harder to provide as the agent's quality is revealed over time.

KEYWORDS. Principal–agent model, optimal contract, learning, private information, reputation, career.

JEL CLASSIFICATION. D82, D83, E24, J41.

1. INTRODUCTION

In an agency problem, the agent may have not just a hidden action, but also an unknown quality. Many relationships between firms and workers, between shareholders and CEOs, or between lenders and borrowers are of this kind. Yet, little is known about the optimal design of multiperiod contracts in such situations. For example, the question remains open as to whether quality uncertainty is a motivating factor.

When risk-neutral principals and agents deal in spot markets and quality is fixed over time, [Holmström \(1999\)](#) provides a clear answer: Quality uncertainty is good for incentives because it creates a reputational concern. [Gibbons and Murphy \(1992\)](#) confirm this result for one-period incentive contracts and risk-averse agents.

We find that the opposite is true under full commitment: Quality uncertainty harms incentives. Our conclusion differs from Holmström's because markets reward talent, whereas contracts are designed to extract effort; once committed to the relationship, it is never in the interest of the principal to discourage the agent by punishing him for having

Julien Prat: julien.prat@ensae.fr

Boyan Jovanovic: bj2@nyu.edu

Previous versions circulated under the title "Dynamic incentive contracts under parameter uncertainty." We thank Heski Bar-Isaac, Pieter Gautier, Robert Kohn, Albert Marcet, and Ennio Stacchetti as well as seminar participants at the IMT Institute in Lucca, the Tinbergen Institute, University of Cergy-Pontoise, University of Essex, Stockholm Institute for International Economic Studies (IIES), and the Minnesota Macro Workshop for comments and suggestions. We acknowledge the support of the Kauffman Foundation, of the NSF, of the Barcelona GSE and LABEX Investissements d'Avenir (ANR-11-IDEX-0003/Labex Ecodec/ANR-11-LABX-0047).

a low productivity. This creates an incentive for the agent to manipulate the principal's beliefs about quality.

An agent who has provided less effort than expected knows that output would have been higher had he taken the recommended action. Such private information drives a persistent wedge between the principal's and the agent's posteriors: A shirker remains more optimistic about quality than the principal. This motivates the manipulation an agent might undertake: By inducing the principal to underestimate his productivity, a shirker anticipates that he will benefit from overestimated inferences about his effort in future periods and thus higher rewards. Hence, of two agents with identical performance histories, the shirker will enjoy a higher expected future utility.

The benefits of manipulating the principal's belief downward is reminiscent of the "ratchet effect" discussed in [Laffont and Tirole \(1988\)](#).¹ To prevent such *belief manipulation*, contracts under quality uncertainty must link pay more tightly to performance, which lowers the welfare of the risk-averse agent. Since true quality is constant in our model, belief manipulation is more effective early on, as posteriors put more weight on new information, and the sensitivity of pay to performance declines over time.

We use a first-order approach to characterize the optimal contract. We focus on the necessary conditions for recommended effort to be incentive compatible and derive sufficient conditions under which the agent's problem is globally concave.

To circumvent the difficulties associated with the persistence of private information,² we impose a series of parametric restrictions. First, we cast our model in continuous time so as to use the optimization techniques originally introduced by [Schättler and Sung \(1993\)](#) and, in particular, to analytically verify whether incentive-compatible constraints are indeed sufficient. Second, we ensure that ability is drawn from the conjugate prior of the likelihood function by assuming that both noise and ability are normally distributed. It is common in Bayesian models to consider problems where priors and posteriors are conjugate distributions so as to avoid relying on numerical methods. Since posteriors are normal, the set of sufficient statistics boils down to the first two moments of beliefs. Third, we restrict our attention to linear technologies so that effort, ability, and noise affect output additively. This separability ensures that the value of private information does not directly depend on the posterior, and that we only need to keep track of cumulative output and effort to compute it.

¹In contrast to [Laffont and Tirole \(1988\)](#), our model features no adverse selection from the outset. Instead, we model a pure moral hazard problem where the principal and agent share the same prior. Asymmetric information can arise only off the equilibrium path through the persistent influence of past actions on posteriors.

²Establishing incentive compatibility when private information is fully persistent entails the following technical issue: As the duration of the relationship increases, the state space becomes unbounded because the entire history of actions matters for evaluating the agent's options off the equilibrium path. A recursive approach to the problem quickly becomes intractable since, as originally explained by [Fernandes and Phelan \(2000\)](#), it takes the beliefs of the agent and of the principal as separate states. The first-order approach bypasses this difficulty by focusing on the equilibrium path. Then the challenge consists in deriving sufficient conditions. To the best of our knowledge, the only proof in discrete time is by [Kapicka \(2006\)](#) and is rather specific to the reporting problem analyzed in his paper. One remedy is to numerically check the incentive compatibility of the contract, as in [Abraham and Pavoni \(2008\)](#).

Under these three premises, we can derive necessary and sufficient conditions for general time-separable preferences. Their analysis reveals that uncertainty about the agent's ability decreases the power of incentives. Having established the generality of our main finding, we investigate the model's implications when agents have constant absolute risk aversion. This specification neutralizes the wealth effect, making it possible to solve for the optimal contract in closed form.

Given the parametric approach, one may be inclined to think that our insights are driven by the choice of functional forms. To dispel this impression, we include in [Appendix C](#) a stylized two-period model that confirms that welfare rises when priors become more precise. We intentionally use parametric assumptions different from those in the main model to show that our prediction holds across a variety of environments. It is true, however, that dynamic contracts with more than two periods quickly become intractable when any of the restrictions discussed above is relaxed. Hence, the specific design of our setup should not be seen as a drawback detracting from the generality of its predictions, but instead as a contribution for its ability to produce an explicit solution to the otherwise intractable problem of contracting with learning. Its advantages are illustrated in the recent paper of [He et al. \(2012\)](#). They extend our framework by introducing hidden savings and effort costs that are convex instead of linear. They also restrict their attention to stationary learning, whereas we let posteriors' precision increase over time. Whether any of these assumptions should be interpreted as more judicious depends on the economic interpretation of the model. Given that our objective is to analyze the interactions between commitment and career concerns, we focus on nonstationary learning so as to capture the mechanism through which reputations are established. By contrast, [He et al. \(2012\)](#) are interested in characterizing the impact that learning has on effort in the canonical model of [Hölmstrom and Milgrom \(1987\)](#). Using optimality conditions similar to ours, [He et al. \(2012\)](#) show that, on average, recommended effort decreases over time.³ Another closely related paper is [DeMarzo and Sannikov \(2011\)](#). They study a problem that is similar in structure, but assume that agents are risk-neutral.

Our paper seems to be the first to study commitment in a repeated agency problem when the agent's quality is unknown and constant, and where the principal makes transfers to the risk-averse agent in each period. A few other papers have analyzed the interactions between quality and moral hazard, but under different assumptions about the structure of payments or the timing of actions. [Giat et al. \(2010\)](#) add initial private information to [Hölmstrom and Milgrom \(1987\)](#) so that there is a single transfer at the end of the contracting horizon. Conversely, [Hopenhayn and Jarque \(2007\)](#) analyze persistent unknown quality when the effort decision occurs solely in the first period. [Adrian and Westerfield \(2009\)](#) assume that principal and agent disagree about the resolution of uncertainty, knowing that the agent is dogmatic and, as such, never updates his prior. This eliminates the belief manipulation channel since the two parties agree to disagree.

³By contrast, in our model, effort is deterministic and equal to either zero or its first-best level. We reach different predictions because disutility is linear in effort, whereas [He et al. \(2012\)](#) consider losses that are convex in effort.

Our methodology borrows from [Williams \(2011\)](#), whose work extends the continuous time approach of [Sannikov \(2008\)](#) to a contracting environment with persistent private information. There are several differences between our paper and that of Williams: he assumes that the agent observes his productivity and that it evolves stochastically, whereas we keep productivity fixed but neither the agent nor the principal know its actual value. Second, Williams assumes that the initial ability is common knowledge. It, therefore, remains commonly known when the agent reports truthfully and the equilibrium features no learning. By contrast, we have a common learning process along the equilibrium path. Modeling it requires that we introduce contract duration as an additional state. Third, we use a proof strategy that does not rely on the stochastic maximum principle. We follow instead the approach proposed by [Cvitanic et al. \(2009\)](#), who use a variational argument to derive the first-order conditions.

Finally, our model is connected to the canonical work of [Hölmstrom and Milgrom \(1987\)](#). They also study a long-term contracting problem where information arrives continuously and agents have constant absolute risk aversion (CARA) preferences. Their approach differs in two important dimensions. First, they consider that all transfers are paid at the end of the contract. Their framework is, therefore, not amenable to comparison with the career concerns literature where agents are offered a sequence of spot payments. Moreover, ability is known in [Hölmstrom and Milgrom \(1987\)](#). We show that ability risk has a different impact on total surplus than does transitory risk. If optimal contracts treated ability risk in the same way as transitory output risk, the model's prediction would depend only on the time-path that *total* risk follows over the duration of the relationship. It turns out, however, that the two types of uncertainty cannot, in general, be so bundled because actions have a lasting effect solely when the agent can influence beliefs.⁴ As the history of play gradually reveals the agent's ability, the principal can lower the sensitivity of pay to performance over time and still maintain incentives.⁵

The paper proceeds as follows. [Section 2](#) lays out the model. The agent's necessary and sufficient conditions are derived in [Section 3](#). [Section 4](#) displays the contract under exponential utility that is optimal for the principal. It characterizes the set of parameters and initial beliefs under which the agent's first-order conditions represent a global optimum. [Section 5](#) discusses the properties of the optimal contract and equilibrium wage schedule. [Section 6](#) contrasts our full-commitment contract with the no-commitment model of [Holmström \(1999\)](#). [Section 7](#) sums up our findings, whereas the proofs of the main propositions and corollaries are in [Appendix A](#). We relegate the proof of a tangential claim to [Appendix B](#). Finally, we analyze in [Appendix C](#) a stylized two-period contract with learning.

⁴We establish this by means of a proposition in [Section 4](#), simulation in [Section 5](#), and, finally, in [Appendix C](#), where we use a stylized model to show that quality and output uncertainty are isomorphic only when there is a single transfer.

⁵[He et al. \(2012\)](#) thoroughly investigate how learning affects the contracting problem of [Hölmstrom and Milgrom \(1987\)](#). They find that, as opposed to [Hölmstrom and Milgrom \(1987\)](#), whose contract can be implemented by a constant equity share, the optimal contract under learning exhibits option-like features.

2. THE PROBLEM

Production process. Let $\{B_t\}_{t \geq 0}$ be a standard Brownian motion on a probability space $(\Omega, \mathcal{F}, P)$. The *cumulative output* Y_t of a match of duration t is observed by both parties and satisfies the stochastic integral equation

$$Y_t = \int_0^t (\eta + a_s) ds + \int_0^t \sigma dB_s. \quad (1)$$

The *time-invariant* productivity is denoted by η , whereas $a_t \in [0, 1]$ is the effort provided by the agent. The agent's action shifts average output without affecting its volatility.

Learning. No one knows η at the outset, and common priors are normal with mean m_0 and precision h_0 . Posteriors over η depend on Y_t and on cumulative effort $A_t \triangleq \int_0^t a_s ds$. Conditional on (Y_t, A_t, t) , they are also normal with mean

$$\hat{\eta}(Y_t - A_t, t) \triangleq E_t[\eta | Y_t, A_t] = \frac{h_0 m_0 + \sigma^{-2}(Y_t - A_t)}{h_t} \quad (2)$$

and with precision

$$h_t \triangleq h_0 + \sigma^{-2}t. \quad (3)$$

Focusing on normal priors over the mean of a normally distributed process lets one summarize all the statistically significant information with three variables: cumulative output Y , cumulative effort A , and elapsed time t . Especially useful for the characterization of optimal contracts is the fact that beliefs depend on the history of a through A alone. Hence it is sufficient to keep track of cumulative effort instead of the whole effort path.⁶

Beliefs. The principal does not observe the agent's effort and so has to assume that he takes his equilibrium action a_t^* . His beliefs are governed by (2) in which $A = A^*$ and by (3). By contrast, the agent's beliefs incorporate the actual level of effort a that only he knows. Thus his beliefs are governed by (2) in which A and not A^* enter. Let $\mathcal{F}_t^a \triangleq \sigma(Y_s, a_s; 0 \leq s \leq t)$ denote the filtration generated by (Y, a) and let $\mathbb{F}^a \triangleq \{\mathcal{F}_t^a\}_{t \geq 0}$ denote the P -augmentation of this natural filtration. Denote by Z_t the cumulative surprise of someone who believes that Y_t was accompanied by the effort sequence $\{a_s; 0 \leq s \leq t\}$. The filtering theorem of Fujisaki et al. (1972) implies that the *innovation process*

$$dZ_t \triangleq \frac{1}{\sigma} [dY_t - (\hat{\eta}(Y_t - A_t, t) + a_t) dt] \quad (4)$$

⁶This is why most of the literature on career concerns, Holmström's (1999) model included, focuses on the additive normal case. Dewatripont et al. (1999, p. 186) discuss in their remark the complications that arise when more general production functions are considered.

is a standard Brownian motion on the probability space $(\Omega, \mathcal{F}^a, P)$.⁷ Moreover, $\hat{\eta}$ is a P -martingale with decreasing variance:⁸

$$d\hat{\eta}(Y_t - A_t, t) = \frac{\sigma^{-1}}{h_t} dZ_t. \quad (5)$$

The agent is restricted to the class of control processes $\mathcal{A} \triangleq \{a : [0, T] \times \Omega \rightarrow [0, 1]\}$ that are $\mathbb{F}^a$ -predictable.⁹ Given that the principal does not observe actual effort, the information available to him is restricted to the filtration $\mathcal{F}_t^Y \triangleq \sigma(Y_s; 0 \leq s \leq t)$ generated by Y , whose augmentation we denote by $\mathbb{F}^Y \triangleq \{\mathcal{F}_t^Y\}_{t \geq 0}$. An effort path is an equilibrium path when recommended and equilibrium effort coincide, i.e., if $a_t = a_t^*$ for all (t, ω) .¹⁰

Contract. We assume that parties are able to commit to a long-term contract that lasts until date T and whose payments can depend on realized history in an arbitrary way. We follow the usual practice of adding *recommended effort* a^* to the contract definition. Accordingly, since a given output path is a random element of the space Ω , a contract is a mapping $(w, a^*) : [0, T] \times \Omega \rightarrow \mathbb{R} \times [0, 1]$ that associates, to any event $\omega \in \Omega$, a wage–effort pair that is $\mathbb{F}^Y$ -predictable as well as a terminal payment $W_T : \Omega \rightarrow \mathbb{R}$ that is $\mathcal{F}_T^Z$ -measurable. The mapping must be measurable based on information that the principal has and so can depend on past output but not on past effort. Otherwise contracts remain general since they can depend on the entire past and present $\{Y_s; 0 \leq s \leq t\}$ of the output process.¹¹

Preferences. The agent's preferences are time additive with discount rate $\rho > 0$. Flow utility is a concave and twice continuously differentiable function $u(w, a) \in \mathcal{C}^{2,2}(\mathbb{R} \times [0, 1])$, while the terminal utility is $U(W) \in \mathcal{C}^1(\mathbb{R})$. Thus the agent's preferences as of time 0 read

$$\mathcal{U}_0 \triangleq \int_0^T e^{-\rho t} u(w_t(\omega), a_t(\omega)) dt + e^{-\rho T} U(W_T(\omega)). \quad (6)$$

⁷As shown in Section 10.2 of Kallianpur (1980), the linearity of the filtering problem implies that the filtrations generated by the output and innovation processes coincide. More formally, for $\mathcal{F}_t^Z \triangleq \sigma(Z_s; 0 \leq s \leq t)$, we have $\mathcal{F}_t^a = \mathcal{F}_t^Z$.

⁸Equation (5) follows directly from Ito's lemma. Let $X_t \triangleq Y_t - A_t$ denote cumulative output net of cumulative effort so that

$$d\hat{\eta}(X_t, t) = \frac{\partial \hat{\eta}(X_t, t)}{\partial t} dt + \frac{\partial \hat{\eta}(X_t, t)}{\partial X_t} dX_t = -\frac{\sigma^{-2}}{h_t} \hat{\eta}(X_t, t) dt + \frac{\sigma^{-2}}{h_t} (\hat{\eta}(X_t, t) dt + \sigma dZ_t) = \frac{\sigma^{-1}}{h_t} dZ_t.$$

⁹A mapping is predictable when it is $\mathcal{P}$ -measurable, with $\mathcal{P}$ denoting the σ -algebra of predictable subsets of the product space $\mathbb{R}^+ \times \Omega$, i.e., the smallest σ -algebra on $\mathbb{R}^+ \times \Omega$ making measurable all left-continuous and adapted processes.

¹⁰Since a^* is $\mathbb{F}^Y$ -predictable, the two filtrations $\mathcal{F}_t^a$ and $\mathcal{F}_t^Y$ will coincide in equilibrium. This captures the fact that the principal and the agent share the same information sets when recommended effort is implemented. However, we have to consider out-of-equilibrium strategies so as to establish incentive compatibility. This is why we need to allow for the possibility that the filtrations $\mathcal{F}_t^a$ differ from $\mathcal{F}_t^Y$.

¹¹Given the diffusion property of the output process, one should think of $\Omega = C([0, T]; \mathbb{R})$ as the space of continuous functions $\omega : [0, T] \rightarrow \mathbb{R}$ and of the process defined in (4), $Z_t(\omega) = \omega(t)$, $0 \leq t \leq T$, as the coordinate mapping process with Wiener measure P on $(\Omega, \mathcal{F}_t^Y)$. Accordingly, a contract is a mapping $(w, a^*) : \mathbb{R}^+ \times C([0, T]; \mathbb{R}) \rightarrow \mathbb{R} \times [0, 1]$.

The principal is risk-neutral, also discounts at the rate ρ , and seeks to maximize the discounted flow of output net of wages and net of the terminal payment

$$\pi_0 \triangleq \int_0^T e^{-\rho t} dY_t - \int_0^T e^{-\rho t} w_t(\omega) dt - e^{-\rho T} W_T(\omega). \quad (7)$$

3. INCENTIVE-COMPATIBLE CONTRACTS

This section focuses on the agent's problem. We derive the necessary conditions for a given action to be optimal and then establish a restriction under which they are also sufficient. We impose a terminal date T on the contracting horizon. Until then, both principal and agent are *fully committed* to the relationship. The agent's continuation value at time t reads

$$v_t \triangleq \max_{a \in A} E \left[\int_t^T e^{-\rho(s-t)} u(w(\bar{Y}_s), a) ds + e^{-\rho(T-t)} U(W(\bar{Y}_T)) \middle| \mathcal{F}_t^a \right], \quad (8)$$

where $\bar{Y}_t \triangleq \{Y_s; 0 \leq s \leq t\}$ is the output history. The agent computes his continuation value by taking a conditional expectation under the filtration $\mathcal{F}_t^a$, which varies with the level of cumulative effort. The principal, on the other hand, does not observe actual actions. Thus he needs to keep track of continuation values for any potential level of cumulative effort. We shall simplify the problem by adopting a first-order approach: We focus on the continuation value along the equilibrium path and then establish conditions under which our solution is indeed globally optimal for the agent.

3.1 Necessary conditions

The optimization problem (8) cannot be analyzed with standard methods because the objective function depends on the process w_t , which is *non-Markovian* since it depends on the whole output path $\bar{Y}_t$. We instead use a martingale approach. Faced with a contract (w, a^*) , the agent controls the distribution of wages through his choice of effort. Under this interpretation, the agent chooses the probability measure over realizations of w_t . This approach renders our optimization problem treatable with optimal control techniques because the Radon–Nikodym derivative associated with any effort path is a Markovian process.¹²

¹²More precisely, fix a probability measure Q such that cumulative output follows a martingale under Q and let Z^0 denote a standard Brownian motion under Q . By definition of Q , we have $dY_t = \sigma dZ_t^0$. A change in effort can be interpreted as a choice of probability measure. We explain in the [Appendix](#) that the probability measure Q^a associated to any arbitrary effort policy a is equivalent to Q with Radon–Nikodym derivative $dQ^a/dQ = \Lambda_{0,T}^a$ (see (32) in [Appendix A](#) for a formal definition of Λ^a). We can use Λ^a to relate the expectation operator $E^a[\cdot]$ under the probability measure Q^a to the expectation operator $E^0[\cdot]$ under Q as

$$V(a, t) = E_t^a \left[\int_t^T e^{-\rho(s-t)} u_s ds + e^{-\rho(T-t)} U_T \right] = E_t^0 \left[\int_t^T \Lambda_{s,T}^a e^{-\rho(s-t)} u_s ds + \Lambda_{T,T}^a e^{-\rho(T-t)} U_T \right].$$

Given that the construction of the measure Q^a ensures that the expectations $E_t[\cdot | \mathcal{F}_t^a]$ and $E_t^a[\cdot]$ coincide, the agent's problem consists in maximizing $V(a, t)$ subject to the laws of motion of Λ^a and Y . The key advantage of the weak formulation is that under our reference probability measure Q , the output process does not depend on a . Hence, we can treat it as fixed, which enables us to solve the optimization problem in spite of its non-Markovian structure.

The idea of using distributions as controls to solve principal–agent models goes back to [Mirrlees \(1974\)](#). The learning process complicates our problem, as past efforts affect not only current wages, but also future expectations. We show in the [Appendix](#) how this difficulty can be handled through an extension of the proof by [Cvitanic et al. \(2009\)](#), which leads to the necessary condition stated below.

PROPOSITION 1 (Necessary conditions). *There exists a unique decomposition for the agent's continuation value,*

$$\begin{aligned} dv_t &= [\rho v_t - u(w_t, a_t)] dt + \gamma_t \sigma dZ_t \\ v_T &= U(W_T), \end{aligned} \quad (9)$$

where γ is a square integrable predictable process. The necessary condition for a_t^* to be an optimal control reads

$$\left[\gamma_t + E_t \left[- \int_t^T e^{-\rho(s-t)} \gamma_s \frac{\sigma^{-2}}{h_s} ds \right] + u_a(w_t, a_t^*) \right] (a - a_t^*) \leq 0 \quad (10)$$

for all $a \in [0, 1]$.

An increase in current effort has two effects: it raises the promised value along the equilibrium path and increases cumulative effort. The first effect is proportional to the process γ , which measures the sensitivity of the agent's value to output surprises. The second effect is captured by the expectation term in (10). This term vanishes when η is known, since then $\sigma^{-2}/h_s = 0$ for all $s \geq t$. As a special case of our model, we then get the necessary condition in [Sannikov \(2008\)](#), which says that an optimal control must maximize $\gamma a + u(w, a)$.

Quality uncertainty leads to the addition of the expectation term on the left hand side of (10). To understand why, observe first that an increase in cumulative effort to-day lowers date- s posteriors over η by $\partial \hat{\eta}(Y_s - A_s, s) / \partial A_s = -\sigma^{-2}/h_s$. In other words, an upward deviation from recommended effort A_t^* creates a negative output surprise of $-\sigma^{-2}/h_s$ at all future dates $s > t$. The impact in utils is obtained by multiplying the output surprise $-\sigma^{-2}/h_s$ by the expected value of the sensitivity coefficient γ_s . Summing and discounting all these marginal effects yields the expected marginal returns of manipulating beliefs.

It is more convenient to rewrite condition (10) as

$$\left[\frac{\sigma^{-2}}{h_t} p_t + \gamma_t + u_a(w_t, a_t^*) \right] (a - a_t^*) \leq 0 \quad \text{for all } a \in [0, 1], \quad (11)$$

where

$$p_t \triangleq h_t E \left[- \int_t^T e^{-\rho(s-t)} \gamma_s \frac{1}{h_s} ds \middle| \mathcal{F}_t^a \right] \quad (12)$$

is a stochastic process capturing the value of private information. Since p is negative, it follows from (11) that for any recommended level of effort a_t^* and any given wage w_t , the

volatility, γ_t , of the agent's promised value must be higher with quality uncertainty than without it. The implication does not hinge on any particular specification for the utility function: An uncertain environment makes it harder to motivate the agent and so leads to greater exposure to risk.

The reformulated necessary condition (11) involves two stochastic variables, γ and p . This is a usual result for dynamic contracts with private information.¹³ First, we recover the now standard technique of using the promised value to encode past history. A related interpretation can be inferred for p by noticing that the incentive constraint (11) implies that

$$\gamma_t \geq -u_a(w_t, a_t^*) - \frac{\sigma^{-2}}{h_t} p_t \quad \text{whenever } a_t^* > 0. \quad (13)$$

Given that the agent is risk-averse, it is natural to conjecture that the principal will minimize the volatility parameter γ . Hence, as long as $a_s^* > 0$ for all $s \in [t, T]$, the necessary condition (13) will hold with equality almost everywhere on the equilibrium path. We show below that this is indeed true when the agent has exponential utility and precision exceeds a deterministic threshold. We, therefore, replace γ by its expression when (13) binds and, as shown in [Appendix B](#), obtain the solution

$$p_t = E \left[\int_t^T e^{-\rho(s-t)} u_a(w_s, a_s) ds \middle| \mathcal{F}_t^a \right]. \quad (14)$$

It follows that whenever $a_s^* > 0$ for all $s \in [t, T]$, the law of motion for p reads

$$\begin{aligned} dp_t &= [\rho p_t - u_a(w_t, a_t)] dt + \vartheta_t \sigma dZ_t \\ p_T &= 0. \end{aligned} \quad (15)$$

Along with γ , the coefficients ϑ is chosen by the principal so as to maximize expected returns.

Intuition behind p . The second state variable p is equal to the expected discounted marginal cost of future efforts. Multiplying it by the ratio σ^{-2}/h_t yields the marginal effect of cumulative effort on the continuation value. The intuition for this result can be laid out by considering *mimicking strategies*. Fix $\bar{Y}_t$ and decrease cumulative effort by $\delta > 0$. Then define a strategy enabling the agent to reproduce the payoffs of an agent with the reference level A_t^* of past effort. Let a_t^* denote the optimal effort at time t of the

¹³This feature was originally noticed by [Werning \(2001\)](#), who considered principal–agent problems with hidden savings. He proved that one has to introduce both continuation value and expected marginal utility from consumption. A general approach has been recently proposed by [Pavan et al. \(2010\)](#). They establish an envelope formula for the derivative of an agent's equilibrium payoff. When applied to adverse selection problems with Markovian types, the envelope formula leads to the definition of an additional recursive variable. To the best of our knowledge, [Williams \(2008\)](#) was the first to introduce two separate stochastic processes so as to solve dynamic incentive problems in continuous time. He also explains how one of them can be dispensed with when the utility function is exponential. We show in [Section 4](#) that a similar simplification holds in our setup.

reference policy with cumulative effort A_t^* . By providing $a_t^\delta = a_t^* - \delta\sigma^{-2}/h_t$,¹⁴ the agent with cumulative effort $A_t^* - \delta$ ensures that cumulative output will have the same drift as along the reference path; that is,

$$\begin{aligned}\hat{\eta}(Y_t - (A_t^* - \delta), t) + a_t^\delta &= \frac{h_0 m_0 + \sigma^{-2}(Y_t - (A_t^* - \delta))}{h_t} + a_t^* - \frac{\sigma^{-2}}{h_t} \delta \\ &= \hat{\eta}(Y_t - A_t^*, t) + a_t^*.\end{aligned}$$

Assume now that a similar strategy is employed afterward, so that $a_s^\delta = a_s^* - (\sigma^{-2}/h_t)\delta$ for all $s \geq t$. Cumulative effort will be $A_s^\delta = A_s^* - [1 + (\sigma^{-2}/h_t)(s - t)]\delta$, leading to the output drift

$$\begin{aligned}\hat{\eta}(Y_s - A_s^\delta, s) + a_s^\delta &= \frac{h_0 m_0 + \sigma^{-2}(A_s^* - [1 + (\sigma^{-2}/h_t)(s - t)]\delta)}{h_s} + a_s^* - \frac{\sigma^{-2}}{h_t} \delta \\ &= \hat{\eta}(Y_s - A_s^*, s) + a_s^* - \frac{\sigma^{-2}}{h_t h_s} \left[\underbrace{(h_t + \sigma^{-2}(s - t))}_{=h_s} - h_s \right] \\ &= \hat{\eta}(Y_s - A_s^*, s) + a_s^*.\end{aligned}$$

As desired, the mimicking strategy reproduces the distribution of Y_s for all $s \geq t$ and the product $-(\sigma^{-2}/h_t)p_t$ measures its expected discounted return in utils.¹⁵ It is positive because it took the agent with cumulative effort A_t^* more work to produce Y_t , implying that his productivity is likely to be lower. Returns decrease as t increases because the influence of output on beliefs is lower when η is known more precisely. This suggests that incentives become easier to provide over time, a result that we will discuss at length in Section 5.

3.2 Sufficient conditions

First-order conditions rely on the premise that the agent's objective is globally concave. Unfortunately, principal-agent problems do not always fulfill such a requirement. In our case, establishing concavity is complicated by the persistence of private information: Excluding one-shot deviation does not necessary rule out multiple deviations, because any departure from recommended effort drives a permanent wedge between the beliefs of the agent and those of the principal. Thus we have to verify the sufficiency of our necessary conditions. Only then can we be sure that the agent finds it indeed optimal to provide recommended effort when assigned the wage function satisfying the local incentive constraint (11).

¹⁴Such strategies are not feasible when the reference control is at the lower bound, i.e., when $a_t^* = 0$. One should, therefore, interpret our discussion of mimicking strategies as heuristic, the rigorous interpretation being that of the expectation term $E[-\int_t^T \gamma_s \frac{\sigma^{-2}}{h_s} ds | \mathcal{F}_t^a]$ laid out in the paragraphs above.

¹⁵The correction term σ^{-2}/h_t required to mimic the output distribution remains constant over time because of two countervailing mechanisms. One the one hand, as h_s increases, the impact of past deviations on posteriors decreases over time. On the other hand, the mimicking strategy involves repeated deviations so that the gap between A_s^* and A_s^δ widens over time. When the output distribution is normal, these two opposite forces offset each other.

How to establish incentive compatibility for discrete time contracts with persistent information remains an open question.¹⁶ By contrast, when the model is cast in continuous time, the sufficiency of the necessary conditions and thus the incentive compatibility of the effort path can be established from the concavity of the agent's Hamiltonian. This general mathematical result is summarized in Theorem 3.5.2 of [Yong and Zhou \(1999\)](#), and the next proposition is an application of that result to our model.¹⁷

PROPOSITION 2 (Sufficient conditions). *A control a^* is incentive compatible if (10) and*

$$-2u_{aa}(w_t, a_t^*) \geq e^{\rho t} \xi_t \sigma^2 h_t \quad (16)$$

are true for almost all t , where ξ is the predictable process defined uniquely by

$$E\left[-\int_0^T e^{-\rho s} \gamma_s \frac{\sigma^{-2}}{h_s} ds \middle| \mathcal{F}_t^{a^*}\right] - E\left[-\int_0^T e^{-\rho s} \gamma_s \frac{\sigma^{-2}}{h_s} ds \middle| \mathcal{F}_0^{a^*}\right] = \int_0^t \xi_s \sigma dZ_s \quad (17)$$

for all $t \in [0, T]$.

According to (14), the process ξ_t is the random fluctuation in the discounted sum of marginal utilities as evaluated from time 0. [Proposition 2](#) imposes stronger restrictions on ξ_t than required so that a control might violate them and nevertheless be incentive compatible. Moreover, (16) and (17) are stated in terms of (w_t, γ_t) , which are endogenous, implying that they have to be verified ex post for any given contract. In some cases, however, one can translate (16) and (17) into a requirement on the parameters of the model. Indeed, when the agent's utility function is exponential, as in (18), [Corollary 2](#) shows that the conditions of [Proposition 2](#) are fulfilled if (24) holds.

Some general results about local incentive constraints and their sufficiency have been established in this section. We have derived qualitative results on the interaction between quality uncertainty and incentive compatibility. It is difficult to make further progress without being more specific about the agent's preferences. This is why we hereafter restrict our attention to a particular class of utility functions.

4. OPTIMAL CONTRACT UNDER EXPONENTIAL UTILITY

We now explain how one can solve for the principal's problem and derive the optimal contract in closed form when the agent's utility is exponential. The main idea is to simplify the optimization program by eliminating two states: The first state is a component of the sufficient statistics for beliefs, $\hat{\eta}$; the second state is the value of private information, p . We now describe how each of these is dealt with.

¹⁶The difficulties arising in discrete time settings are thoroughly discussed by [Abraham and Pavoni \(2008\)](#). To circumvent them, they propose a numerical procedure verifying ex post the implementability of contracts with hidden effort and savings. See also [Kocherlakota \(2004\)](#) for a discussion of the problem and an analytical example.

¹⁷A similar approach has already been used in principal-agent settings by [Schättler and Sung \(1993\)](#), and more recently by [Williams \(2008\)](#). The sufficient condition is too stringent for some of the contracting problems considered in [Williams \(2008\)](#) because the agent's Hamiltonian is not concave. [Corollary 2](#) below shows that this is not the case in our model when posterior precision h_t exceeds a deterministic threshold.

Eliminating $\hat{\eta}$ from the list of states. For a given contract (w, a^*) , the principal expected utility reads

$$\begin{aligned}\Pi_t^T(w, a^*) &\triangleq E\left[\int_t^T e^{-\rho(s-t)}[\hat{\eta}(Y_s - A_s^*, s) + a_s^* - w_s]ds - e^{-\rho(T-t)}W_T \middle| \mathcal{F}_t^Y\right] \\ &= \frac{1 - e^{-\rho(T-t)}}{\rho} \hat{\eta}(Y_t - A_t^*, t) + E\left[\int_t^T e^{-\rho(s-t)}(a_s^* - w_s)ds - e^{-\rho(T-t)}W_T \middle| \mathcal{F}_t^Y\right].\end{aligned}$$

The second equality follows because the principal is risk-neutral and beliefs follow a martingale. Assuming that [Proposition 2](#) holds, so that the necessary condition is also sufficient, the principal's problem consists in solving for $\sup_{(w, a^*)} \Pi_t^T(w, a^*)$ subject to the two promise-keeping constraints (9) and (15), and to the incentive constraint (13). Given that the posterior mean $\hat{\eta}$ does not enter directly into any of the constraints, it can be dispensed with as a state, leaving only precision as a belief state. Furthermore, since h_t is deterministic, we may index precision by t . The fact that the expected value of η is immaterial to the principal's objective illustrates that incentives are designed to reward effort and not ability.

Eliminating p from the list of states. We now restrict our attention to exponential utility functions¹⁸ of the form

$$u(w, a) = -\exp(-\theta(w - \lambda a)) \quad \text{with } \lambda \in (0, 1), \theta > 0, \quad (18)$$

and $a \in [0, 1]$. Imposing $\lambda < 1$ ensures that $a = 1$ is the first-best action because the marginal utility of an additional unit of output exceeds the marginal cost of effort regardless of η .¹⁹ Our specification rules out agents with limited liability because utility is defined even for negative consumption, which occurs with positive probability in equilibrium.

When $u_a(w, a)$ is given by (18), the problem greatly simplifies because $u_a(w, a) = \theta \lambda u(w, a)$.²⁰ Reinserting this identity into (8) and (14) shows that whenever the incentive constraint binds for almost all $s \in [t, T]$, $p_t = \theta \lambda (v_t - e^{-\rho(T-t)} E_t[v_T])$. Given that the promised value is a one dimensional diffusion process, we can infer p_t from v_t . Furthermore, as the contracting horizon goes to infinity and the transversality condition $\lim_{T \rightarrow \infty} e^{-\rho T} v_T = 0$ holds, we have $p_t = \theta \lambda v_t$. This proportionality of v and p means that keeping track of one of the two states is sufficient.

¹⁸Even though the full characterization of the contract will hold only for utilities of the form (18), the optimality conditions derived in [Section 3](#) are true independently of this parametric restriction. One of its implications is that there is no wealth effect on leisure because $U_w(\cdot)/U_a(\cdot) = -\lambda^{-1}$ is equal to a constant that does not depend on w .

¹⁹Accordingly, one could interpret our model as resulting from a situation where the agent is able to divert cash flows $1 - a$ at the rate λ . As in [DeMarzo and Sannikov \(2011\)](#), setting λ below 1 ensures that cash diversion entails linear losses. Our problems differ because [DeMarzo and Sannikov \(2011\)](#) focus on risk-neutral agents, whereas we introduce risk aversion by taking a concave transformation of the agent's income net of his opportunity cost λa .

²⁰One can dispense with p as a state whenever the marginal disutility of effort is proportional to the flow utility. Thus the parametric restriction in (18) could be enlarged to encompass utility functions of the form $u(w, a) = -f(w) \exp(-\theta(g(w) - \lambda a))$ for some positive functions $f(\cdot)$ and $g(\cdot)$.

Termination utility. We now specify the termination utility. We assume that it lies on the Pareto frontier under zero effort in the future, i.e.,

$$U(W) = -\frac{\exp(-\theta\rho W)}{\rho}. \quad (19)$$

This expression captures a situation in which an infinitely lived agent retires at the termination date T of the contract and thereafter consumes the perpetual annuity derived from W while providing zero effort.²¹ In any case, we shall concentrate on problems where the contracting horizon goes to infinity, so that the specification of the terminal utility becomes immaterial to our results.²² Assuming this tractable form for U enables us to easily identify a sequence of finite horizon problems whose value functions converge to the infinite horizon solution described in [Proposition 4](#).

Incentives providing contracts. We focus on contracts whose recommended effort remains positive at all future dates. We call such contracts incentives-providing or, more concisely, incentive contracts. We will show in [Proposition 5](#) that they are indeed optimal when precision is high enough and preferences are described by (18) and (19). As explained above, we can omit the posterior mean $\hat{\eta}$ and recast the principal's optimization problem as

$$J_t^T \triangleq \max_{\{a, w, \gamma, \vartheta\}} E \left[\int_t^T e^{-\rho(s-t)} (a_s - w_s) ds - e^{-\rho(T-t)} W_T \middle| \mathcal{F}_t^Y \right],$$

subject to

$$\begin{aligned} dv_t &= [\rho v_t - u(w_t, a_t)] dt + \gamma_t \sigma dZ_t \quad \text{with } v_T = U(W_T) \\ dp_t &= [\rho p_t - u_a(w_t, a_t)] dt + \vartheta_t \sigma dZ_t \quad \text{with } p_T = 0 \\ \gamma_t &= -u_a(w_t, a_t) - \frac{\sigma^{-2}}{h_t} p_t, \end{aligned} \quad (20)$$

i.e., the two promise-keeping constraints with their associated terminal conditions and the necessary condition under which income volatility is minimized. Since the state

²¹Formal justifications for the Pareto-frontier assumption (6), in general, and for (19), in particular, are the following:

- (i) Parties commit to a long-term contract that last *forever*.
- (ii) Preferences of principal and agent are $\pi_0 = \int_0^\infty e^{-\rho t} dY_t - \int_0^\infty e^{-\rho t} w_t(\omega) dt$ and $\mathcal{U}_0 = \int_0^\infty e^{-\rho t} u(w_t(\omega), a_t(\omega)) dt$, respectively.
- (iii) A contract is constrained to be a fixed payment from T onward so that $w_t(\omega) = w_T(\omega)$ for $t \geq T$. If $W_T \triangleq w_T/\rho$, this leads to (7). The absence of incentives for the agent means that he exerts no effort beyond T , which leads to (6) and (19) with $U(W) = \int_0^\infty e^{-\rho t} u(\rho W, 0) dt = u(\rho W, 0)/\rho$, i.e., $U(W)$ is on the Pareto frontier. Notice that there does not exist a mapping between cumulative output Y_T and the termination value W because the latter is history dependent.

²²The irrelevance in the limit is best illustrated by the observation that, as in [Williams \(2011\)](#), imposing standard transversality conditions allows one to discard the terminal utility and to nonetheless derive the same closed-form solution as in [Proposition 4](#).

variables v and p are Markovian, we can use an Hamilton–Jacobi–Bellman (HJB hereafter) equation to analyze the principal’s optimal control problem.²³ The necessary condition (20) allows us to express the coefficient γ as a function of t , p , w , and a . Thus, if we had to keep all three states (t, v, p) , the dynamic programming equation would read

$$\begin{aligned} \rho J_t^T = \max_{\{a, w, \vartheta\}} & \left\{ a - w + \frac{\partial J_t^T}{\partial t} + \frac{\partial J_t^T}{\partial v}(\rho v - u(w, a)) + \frac{\partial J_t^T}{\partial p}(\rho p - u_a(w, a)) \right. \\ & \left. + \frac{1}{2}\sigma^2 \left[\frac{\partial^2 J_t^T}{\partial v^2} \gamma(t, p, w, a)^2 + \frac{\partial^2 J_t^T}{\partial p^2} \vartheta^2 + 2 \frac{\partial^2 J_t^T}{\partial v \partial p} \gamma(t, p, w, a) \vartheta \right] \right\}. \end{aligned}$$

Instead of solving this equation, we use our parametric assumption (18) to reduce the dimensionality of the state space. As explained above, the fact that $u_a(\cdot)/u(\cdot) = \theta\lambda$ ensures that there exists a mapping between p and v as long as effort remains positive.²⁴ Thus we can drop p from the HJB equation to obtain

$$\rho J_t^T = \max_{\{a, w\}} \left\{ a - w + \frac{\partial J_t^T}{\partial t} + \frac{\partial J_t^T}{\partial v}(\rho v - u(w, a)) + \frac{\partial^2 J_t^T}{\partial v^2} \left(\frac{1}{2}\sigma^2 \right) \gamma(t, v, w, a)^2 \right\}. \quad (21)$$

The first-order conditions associated to (21) show that it cannot be profitable to recommend an interior action $a \in (0, 1)$. Due to the linear disutility of effort, corner solutions are optimal. Thus we can focus our attention on paths where a_t^* is equal to either 0 or 1, even though intermediate levels of effort remain feasible but not optimal.²⁵

CLAIM 1. *If the necessary condition (20) holds for almost all $s \geq t$, then recommended effort is set equal to its first-best level, i.e., $a_t^* = 1$.*

4.1 Contracts when η is known

To build intuition, we start by analyzing incentives providing contracts when the principal observes the ability of the agent. In this subsection only, we assume that η is known, but that volatility varies over time, i.e., that $\sigma = \sigma_t$. The time variation in σ_t will be useful when we discuss the dynamic of wages in Section 5.1 in which we distinguish the effect of learning from the effect of exogenous variations in uncertainty. Since there is no room for belief manipulation when η is known, the value of private information p is equal to 0 and the necessary condition (13) reads $\gamma_t = -u_a(w_t, a_t)$. This optimal control problem is closely related to the one analyzed by Sannikov (2008), who shows that when there is no persistent private information, the necessary condition is also sufficient.

²³We use a strong formulation for the principal’s problem even though we have used a weak formulation to solve for the agent’s problem. This change of solution method is usual for principal–agent models. Yet, as discussed in Cvitanic et al. (2009), it may lead to measurability issues if the optimal action directly depends on the Brownian motion. In our case, however, a^* turns out to be a function of time alone so that measurability of the optimal control will not be problematic.

²⁴In particular, when the horizon T is infinite, $p_t = \theta\lambda v_t$ and $J(t, v, p) = J(t, v, \theta\lambda v)$.

²⁵He et al. (2012) study a closely related problem with CARA utility. They introduce convex disutility so as to obtain effort levels that are interior with history dependent paths.

PROPOSITION 3 (Incentive contracts when η is known). *Assume that (i) quality is known, i.e., $h_t = \infty$ for all t , (ii) $u(w, a)$ and $U(W)$ are as specified in (18) and (19), and (iii) $a_t^* > 0$ for all t , so that the incentive constraint (13) binds for almost all $t \in [0, T]$. Then recommended effort is set at its first-best level $a_t^* = 1$ and the principal's value function reads*

$$\rho J_N^T(t, v) = F^T(t) + \frac{\ln(-\rho v)}{\theta}.$$

The function F^T is given by

$$F^T(t) = - \int_t^T e^{-\rho(s-t)} \left[\rho \left(1 - \lambda + \frac{\ln(K_s/\rho)}{\theta} \right) + \frac{1}{2} \theta (\lambda \sigma_s K_s)^2 \right] ds,$$

where $K_t \triangleq K(\sigma_t)$ is the positive root of the quadratic equation for K ,

$$(\theta \lambda \sigma_t)^2 K^2 + K - \rho = 0.$$

The value function derived in Proposition 3 holds even when the variance of output σ evolves over time in a deterministic fashion. The expression for J_N^T greatly simplifies when σ_t is constant over time: If $\sigma_t = \sigma$ for all $t \in [0, T]$, the coefficient K_t is also equal to the constant K and

$$\frac{F^T(t)}{1 - e^{-\rho(T-t)}} = F \triangleq 1 - \lambda + \frac{\ln(K/\rho)}{\theta} + \frac{\theta(\lambda \sigma K)^2}{2\rho}. \quad (22)$$

Let us compare the expression of J_N^T to its counterpart if effort were contractible. Observing actions allows the principal to elicit full effort while perfectly insuring the agent. The cost of delivering value v through a constant income stream is equal to $-\ln(-\rho v)/\theta$. The principal must add λ to the baseline remuneration so as to compensate the agent for his effort. Accordingly, first-best wages are $w^{\text{FB}}(v) = \lambda - \ln(-\rho v)/\theta$ and the value function is $\rho J_{\text{FB}}^T(t, v) = (1 - \lambda)(1 - e^{-\rho(T-t)}) + \ln(-\rho v)/\theta$. Comparing J_N^T under constant σ to J_{FB}^T , it is apparent that $1 - \lambda - F$ measures the per-period loss due to the action being hidden.²⁶

4.2 Contracts when η is unknown

4.2.1 Incentive contracts when η is unknown We turn our attention to cases where beliefs about η are imprecise. We first focus on incentive contracts so that, as shown in Claim 1, effort remains equal to its first-best level. We derive later a precision threshold such that it is indeed optimal to recommend this level of effort.

²⁶Since the value function cannot exceed its first-best level, it must hold true that

$$1 - \lambda - F = -\frac{\ln(K/\rho)}{\theta} - \frac{\theta(\lambda \sigma K)^2}{2\rho} > 0.$$

To see this, observe first that $\ln(K/\rho) < 0$. One still has to establish that its absolute value is higher than that of the second term on the right hand side (RHS). A Taylor approximation around 1 yields $\ln(K/\rho) < (K - \rho)/\rho$. Reinserting this inequality and using the definition of K , one finds that F is indeed smaller than $1 - \lambda$.

We solve in the [Appendix](#) for contracts with a finite horizon. [Proposition 4](#) focuses on the limit of the principal's value function as the retirement date diverges to infinity.²⁷

PROPOSITION 4 (Incentive contracts when η is unknown). *Assume that (i) $u(w, a)$ and $U(W)$ are as specified in (18) and (19), and (ii) $a_t^* > 0$ for all t , so that the incentive constraint (13) binds for almost all $s \geq t$. Then the recommended effort is set equal to the first-best level $a_t^* = 1$ and the sequence of the principal's value functions J_L^T converges pointwise to*

$$\rho J_L(t, v) = \lim_{T \rightarrow \infty} \rho J_L^T(t, v) = f(t) + \frac{\ln(-\rho v)}{\theta}.$$

The function $f(t)$ is given by

$$f(t) = - \int_t^\infty e^{-\rho(s-t)} \left[\rho \left(1 - \lambda + \frac{\ln(k_s/\rho)}{\theta} \right) - \frac{1}{2}(\sigma\lambda)^2 \theta \left(\left(\frac{\sigma^{-2}}{h_s} \right)^2 - k_s^2 \right) \right] ds,$$

where k_t is the positive root of the quadratic equation for k ,

$$(\theta\lambda\sigma)^2 k^2 + (1 + (\theta\lambda)^2 h_t^{-1})k - \rho = 0. \quad (23)$$

As discussed above for contracts when η is known, the term $1 - \lambda - f(t)$ measures the per-period loss due to both hidden effort and quality uncertainty. The following corollary shows that this loss decreases over time as quality uncertainty becomes less of a concern.

COROLLARY 1. *The function $f(t)$ is increasing over time and converges to the constant F defined in (22). The principal's expected profit as a function of the promised value v is therefore increasing in belief precision h_t .*

We still have to check whether the contract is incentive compatible, i.e., that the contract meets the conditions in [Proposition 2](#). Applying conditions (16) and (17) to the exponential-utility case (18) yields the following requirement:

COROLLARY 2. *First-best effort is incentive compatible (i.e., meets conditions (10) and (16)) when*

$$\rho\sigma^2 > \frac{1}{h_t} + 2 \left(\frac{\theta\lambda}{h_t} \right)^2. \quad (24)$$

Since precision h_t is increasing in t , the condition holds at all subsequent dates $s \geq t$.

²⁷[Proposition 4](#) does not assert that directly solving the infinite horizon problem yields $J_L(t, v)$. Instead, it shows that by increasing T sufficiently, one can make the difference between the finite horizon solutions $J_L^T(t, v)$ and $J_L(t, v)$ arbitrarily small. By restricting our attention to finite horizon problems, our approach circumvents the technical difficulties due to the fact that probability measures arising from different effort paths are mutually absolutely continuous for every finite T but not in the limit. In particular, the measures Q^a and Q defined in [Section 3.1](#) do not share the same sets of measure 0 events as T diverges to infinity. See Chapter 3 in [Jacod and Shiryaev \(1987\)](#) for an in-depth discussion of these issues and of the Novikov condition, which allows one to establish absolute continuity in the limit but which is unfortunately not satisfied in our setup.

The sufficient condition (24) is more likely to hold when both parties are impatient, output noise is high, the marginal cost of effort λ is low, the coefficient of absolute risk aversion θ is small, or precision is high. Indeed, (24) always holds in the limit case, $h_0 = \infty$, where quality is known because multiple deviations are no longer relevant.

We shall henceforth assume that our parameters satisfy (24). The condition is sufficient and not necessary, however, and our comparative statics results hold independently of it, showing that they are robust over a wider region of the parameter space.

4.2.2 Optimal contracts when η is unknown The derivation of the value function $J_L(t, v)$ was based on the premise that the incentive constraint always binds; see part (ii) of Proposition 4. But the principal has the option to perfectly insure the agent while recommending zero effort. As explained in the discussion of first-best contracts, implementing such a policy has a cost of $-\ln(-\rho v)/\theta$. By contrast, its return is 0 because the agent does not exert any effort. This suggests that the principal would rather insure the agent when $f(t)$ is negative and offer him an incentives-providing contract when $f(t)$ is positive. But this conclusion is misleading because it is based on a comparison between contracts that recommend full or zero effort at *every point* in the future. It may, instead, be optimal to insure the agent for a certain length of time and then to provide him with incentives to exert effort. In other words, the principal has the valuable option of delaying incentives provision.

Due to the absence of wealth effect, the value of the full-insurance option does not depend on the current belief about η but is instead deterministic. The marginal gains from delaying incentives are equal to $f'(t)$, while the costs due to discounting are given by $-\rho f(t)$. Hence, when $\psi(t) \triangleq \rho f(t) - f'(t) < 0$, the principal perfectly insures the agent. Conversely, when $\psi(t) \geq 0$, he offers the incentives-providing contract described in the previous subsection. Since $\psi(t)$ is increasing over time,²⁸ there is at most one precision level above which incentives provision is optimal, as illustrated in Figure 1 for the parameter values in Table 1. If quality is uncertain enough at the beginning of the relationship, the effort path starts at 0 and switches to 1 exactly at the time where $\psi(t)$ crosses the zero axis. Depending on the parameter constellation, it may also happen that $f(t)$, and consequently $\psi(t)$, remain negative at all t . In such cases, it is always optimal to perfectly insure the agent. Thus insurance is not always increasing as information about quality becomes more precise.

PROPOSITION 5 (Optimal contracts when η is unknown). *Let F be as defined in (22). (i) If $F > 0$, a^* is a step function, there exists a unique precision $\tilde{h}$ such that recommended effort $a_t^* = 0$ whenever $h_t < \tilde{h}$, and $a_t^* = 1$ otherwise. The principal's value function J^* is given by*

$$\rho J^*(t, v) = \begin{cases} e^{-\rho(\tau-t)} f(\tau) + \frac{\ln(-\rho v)}{\theta} & \text{when } h_t < h(\tau) = \tilde{h} \\ f(t) + \frac{\ln(-\rho v)}{\theta} & \text{when } h_t \geq \tilde{h}. \end{cases}$$

²⁸See the proof of Proposition 4.

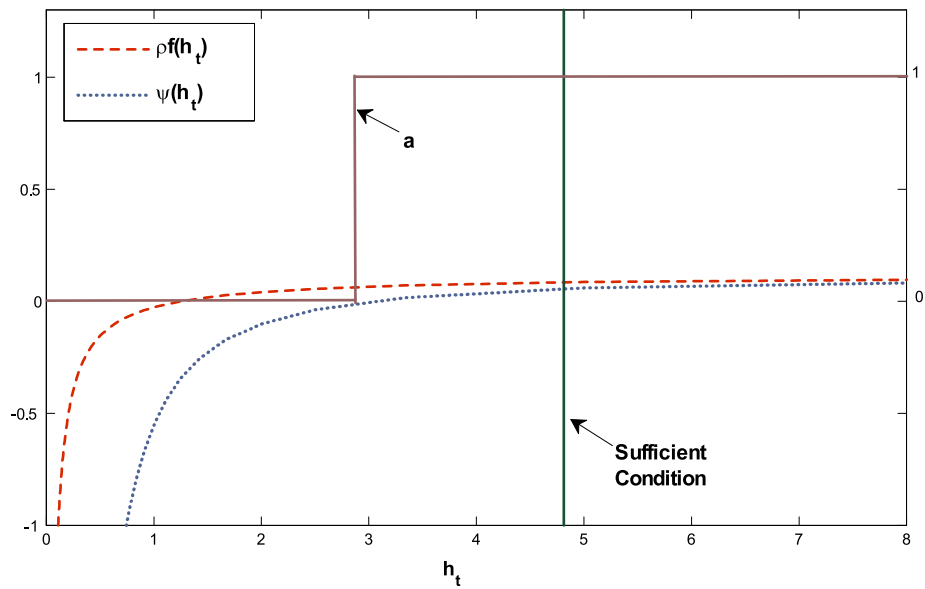


FIGURE 1. Recommended effort as a function of precision h_t . Parameters are as in Table 1.

ρ	σ^2	θ	λ	h_0
0.5	0.5	1	0.7	4.81

TABLE 1. Baseline parameters.

(ii) If $F \leq 0$, recommended effort $a_t^* = 0$ for all t and the principal's value function reads

$$\rho J^*(v) = \frac{\ln(-\rho v)}{\theta}.$$

This proposition completes our description of the optimal contract. Its properties and implications for wage dynamics are explored in the next section.

5. CHARACTERIZATION OF THE OPTIMAL CONTRACT

5.1 Wage dynamics

To isolate the impact that learning has on wages, it is useful to first analyze contracts when quality is observable and σ is constant. Optimal wages under incentives provision and known η are given by²⁹

$$w^*(v) = -\frac{\ln(-K(\sigma)v)}{\theta} + \lambda.$$

²⁹See the proof of Proposition 3.

Thus the evolution of wages mirrors that of the promised value. The law of motion of v follows by reinserting the optimal volatility $\gamma_t(v) = -\theta\lambda K v$ derived in the proof of [Proposition 3](#) into the stochastic differential equation (SDE) (9) to obtain

$$dv_t = v_t[(\rho - K(\sigma))dt - \theta\lambda K(\sigma)\sigma dZ_t]. \quad (25)$$

The sign of the deterministic trend $v_t(\rho - K(\sigma))$ indicates how utility is allocated over time. Given that $\rho > K(\sigma)$ and $v < 0$, the trend is negative.³⁰ Discounted expected utility drifts downward. To understand why it is in the interest of the principal to front load payments, consider the derivative of the incentive constraint (13) with respect to w : $\partial\gamma_t/\partial w = -u_{aw}(w_t^*, a_t^*) < 0$. For our choice of utility function, the marginal cost of effort is decreasing in income. This is why raising transfers today enables the principal to lower the volatility of the promised value and, consequently, to reduce transfers tomorrow.

The agent's immiserization is implied by the *inverse Euler equation* that can be established in the infinite-precision limit using Ito's lemma,

$$du_w(w_t^*, a_t^*)^{-1} = -\frac{\lambda\sigma}{v} dZ_t \quad \text{when } \sigma^{-2}/h_t = 0.$$

Under (18), $u_w(w_t, a_t)^{-1} = \exp(\theta[w - \lambda])/\theta$ is convex in w ; hence the immiserization. However, if we had solved the problem using preferences for which the inverse marginal utility of income is concave,³¹ the inverse Euler equation would imply that wages exhibit a positive trend. Immiserization is, therefore, specific to the exponential parametrization of the utility function.

5.1.1 Distinguishing the effects of learning and of decreasing output volatility As h_t rises, the precision of the agents' beliefs over dY_t also rises over time even though σ is constant. A similar reduction in output volatility could be imposed not through learning, but simply by assuming an exogenous decline in σ . What does learning add over and above a model in which η is known but $\sigma = \sigma_t$ declines over time deterministically? To answer this question, before analyzing wages under learning, we shall study this case first, and will then be able to distinguish the effect of transitory output risks from that of learning.

Wage dynamics when η is known and output volatility decreases. In contrast to contracts under constant σ , suppose now that fluctuations in wages are driven not only by changes in the promised value, but also by the evolution of the coefficient $K(\sigma_t)$ as

$$dw_t^* = \frac{1}{\theta} \left[-\frac{K'(\sigma_t)\dot{\sigma}_t}{K(\sigma_t)} dt - d\ln(-v_t) \right]. \quad (26)$$

³⁰The inequality $\rho - K > 0$ is always satisfied because K is the positive solution of $(K\sigma\theta\lambda)^2 + K - \rho = 0$.

³¹An example of such a utility function could be $U(w, a) = c(a)w^{1-\phi}/(1-\phi)$ with $\phi < 1$ and $c'(a) < 0$.

Reinserting (25) into (26) and applying Ito's lemma to the logarithmic transformation of v yields the “reduced form” for wage growth

$$dw_t^* = \frac{1}{\theta} \left[\underbrace{-\frac{1}{2}(K(\sigma_t)\theta\lambda\sigma_t)^2}_{\text{Immiserization}} - \underbrace{\frac{K'(\sigma_t)\dot{\sigma}_t}{K(\sigma_t)}}_{\text{Insurance}} \right] dt + \lambda K(\sigma_t)\sigma_t dZ_t.$$

Since $K'(\sigma_t) < 0$,³² a decrease in output variance adds a second negative component to the deterministic trend. As output becomes less volatile, the signal–noise ratio improves. The principal can extract full effort with less risk exposure and so trades lower wages in exchange for better insurance. Income stabilization is illustrated by the decrease of the volatility term $\lambda K(\sigma_t)\sigma_t$.

Wage dynamics when η is unknown. When quality is gradually revealed, optimal wages under incentives provision are given by

$$w_t^*(v) = -\frac{\ln(-k_t v)}{\theta} + \lambda, \quad (27)$$

while the promised value satisfies the law of motion³³

$$dv_t = v_t \left[(\rho - k_t) dt - \theta \lambda \left(k_t + \frac{\sigma^{-2}}{h_t} \right) \sigma dZ_t \right]. \quad (28)$$

As in the case where η is known, utility is front loaded since the deterministic trend $v_t(\rho - k_t)$ is negative. Given that k_t increases over time, the principal resorts more intensively to front loading early in the relationship. Combining (27) with (28) and applying Ito's lemma leads to a stochastic differential equation that neatly sums up the three mechanisms that drive income dynamics:

$$dw_t^* = \frac{1}{\theta} \left(\underbrace{-\frac{1}{2}(k_t\theta\lambda\sigma)^2}_{\text{Immiserization}} - \underbrace{\frac{\dot{k}_t}{k_t}}_{\text{Insurance}} + \underbrace{\frac{1}{2}(\theta\lambda)^2 \left(\frac{\sigma^{-1}}{h_t} \right)^2}_{\text{Information Rent}} \right) dt + \lambda \left(k_t + \frac{\sigma^{-2}}{h_t} \right) \sigma dZ_t. \quad (29)$$

We recover the insurance and immiserization channels: (i) For a constant promised value, wages decrease over time due to better insurance and (ii) wages are driven downward by the agent's immiserization. Thus, of the three channels, only the last one is specific to learning. It measures the principal's effort to minimize the agent's information rent. Since the value of private information is equal to the expected discounted marginal cost of future efforts,³⁴ the principal can take advantage of the positive cross-derivative between income and effort. An increase in future transfers lowers the value of p and thus strengthens the agent's incentives in the current period. This is why the information rent channel is positive, which partially offsets the insurance and immiserization mechanisms.

³²The sign of the derivative follows from the quadratic equation $(K(\sigma)\sigma\theta\lambda)^2 + K(\sigma) - \rho = 0$.

³³The relations (27) and (28) are derived in the proof of Proposition 4.

³⁴See (14).

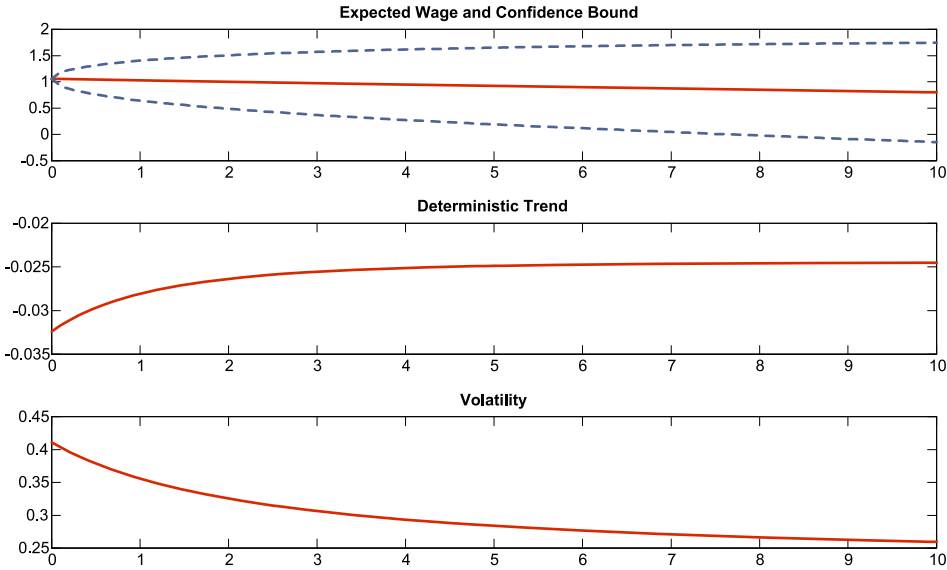


FIGURE 2. Wage dynamics as a function of time. Parameters are as in Table 1.

Williams (2011) proves qualitatively similar results in a reporting problem with persistent income shocks: Efficiency losses due to private information increase with the persistence of the endowment and, parallel to our result that the principal back loads payments more when h_t is lower, Williams also finds that persistence of shocks leads to a tendency to back load payments that is absent in reporting problems with independent and identically distributed (i.i.d.) shocks. He et al. (2012) analyze the effect of the information rent in a setup closely related to ours but with interior levels of optimal effort. This gives an additional margin of adjustment to the principal who decides to recommend decreasing efforts so as to minimize the information rent enjoyed by the agent. As in Williams’s model, persistence is stationary and, as a result, time is not a state, and contractual arrangements do not directly depend on the seniority of the worker as they do in our model.

It is also insightful to consider first-best contracts. If the principal could observe effort, he would totally insure the agent and wages would be constant whatever the actual ability. By contrast, when effort is hidden, output and rewards have to be positively correlated so as to provide incentives. Hence, although the contract does not directly condition on the posterior $\hat{\eta}$, an increase in cumulative output leads to higher wages. In other words, higher ability elicits greater transfers. This can be seen from (29), whose loading $\lambda(k_t + \sigma^{-2}/h_t)$ on dZ_t is positive, meaning that wages are increasing in the posterior $\hat{\eta}_t$ since $d\hat{\eta}_t = (\sigma^{-1}/h_t) dZ_t$.

The trend and volatility terms in (29) are both deterministic. We plot them in the second and third panels of Figure 2. The assumed parameter values are shown in Table 1. They are used as baseline values for all the simulations reported in the paper. The value $h_0 = 4.81$ is the smallest precision that satisfies the second-order condition (24) given the assumed values of the other parameters.

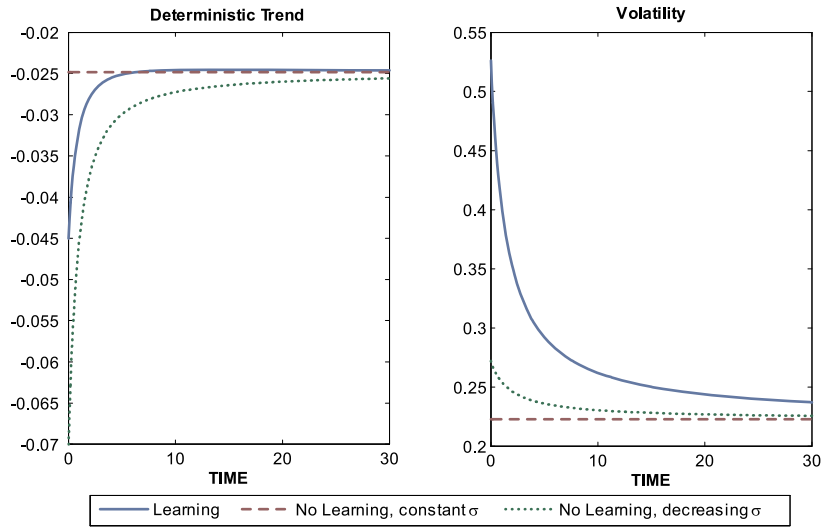


FIGURE 3. Trend and volatility of wages as a function of time. Parameters are as in Table 1.

The top panel of Figure 2 plots the mean wage and the 1-standard-deviation bands for the parameter values in Table 1. The stochastic term σdZ is the output surprise defined in (4), which means that the solution w_t to the SDE (29) is a normally distributed random variable. The distribution of wages at date t is the frequency distribution of wages among age- t workers with abilities randomly drawn from $\eta \sim N(0, h_0^{-1})$. By normality, the bands are equidistant from the mean and, hence, are symmetric. Furthermore k_t has a strictly positive limit K , implying that the volatility of the wage increments does not die off as $\lim_{t \rightarrow \infty} (k_t + \sigma^{-2} / h_t) = K > 0$. Since increments are independent, the cross-sectional variance of wages diverges to infinity. We sum up our findings in the corollary below, whereas Figure 2 illustrates them.

COROLLARY 3. *The volatility of the wage increments decreases to a positive limit, K , so that the cross-sectional variance of wages grows without bound. Provided that the sufficient condition (16) is satisfied, wages exhibit a negative trend.*

Comparison. We now compare contracts with learning to their counterparts when η is known. To make the comparison meaningful, we set $\sigma_t^2 = \sigma^2 + h_t^{-1}$ so that output variance is the same in both models. The simulations reported in Figure 3 highlight how the effect of transitory output risks differs from that of learning.

Consider first the simplest case where η is observable and σ is constant. Then both deterministic trend and volatility remain stable over time. As expected, their values are the limits of the model with decreasing σ_t as well as with learning. However, they converge at different rates: Wages are much more volatile and front loading is less significant under learning. The difference is due to the persistence of private information: To discourage the agent from manipulating his belief, the principal raises the agent's exposure to risk, hence the higher volatility. This adjustment is combined with an increase in the deterministic trend, which shrinks the differences between earnings today and

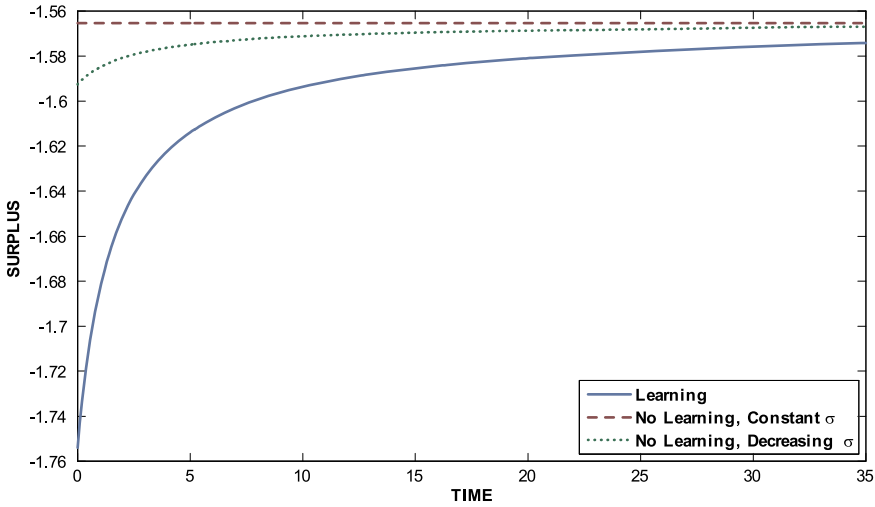


FIGURE 4. Surplus as a function of time. Parameters are as in Table 1.

tomorrow so as to lower the information rent. Figure 4, which compares the surpluses F , $F(t)$, and $f(t)$ of the three contracts, shows that welfare is substantially reduced by the increase in volatility needed to offset the belief manipulation channel.

Equations (23) and (29) summarize the implications for the wage data. They show that the vector of parameters $(\rho, \lambda, \theta, \sigma, h_0)$ is identified. Thus data sets that match employees and employers could provide one with estimates for the parameters vector, making it possible to identify the absolute and relative importance of the three channels described above.

5.2 Surplus

Instead of focusing on wage dynamics within a given match, we can use the model to compare the surplus associated with commitment across different environments. As stated in Corollary 2, the surplus is higher when priors are more accurate. The intuition for this result directly follows from Corollary 3: A tighter prior over η enables the principal to better stabilize income. As contracts get closer to the second best, the principal can deliver the promised value v at a lower expected cost.

Figure 5 plots the agent’s value as a function of the prior variance $1/h_0$ and of the marginal cost of effort parameter λ , holding the principal’s value constant at zero. The other parameters are as given in Table 1. We report on the horizontal plane a line that separates the regions where recommended effort is either 0 or 1.

We also include in Figure 5 a solid black line labeled “sufficient condition,” which identifies the maximal prior variance $1/h_0$ and λ above which incentive compatibility holds surely. In particular, (24) (which involves both λ and h) holds to the left of the line. For the parameter values used in the plot, (24) reads $1/4 > h_0^{-1} + 2(\lambda h_0^{-1})^2$, and so the

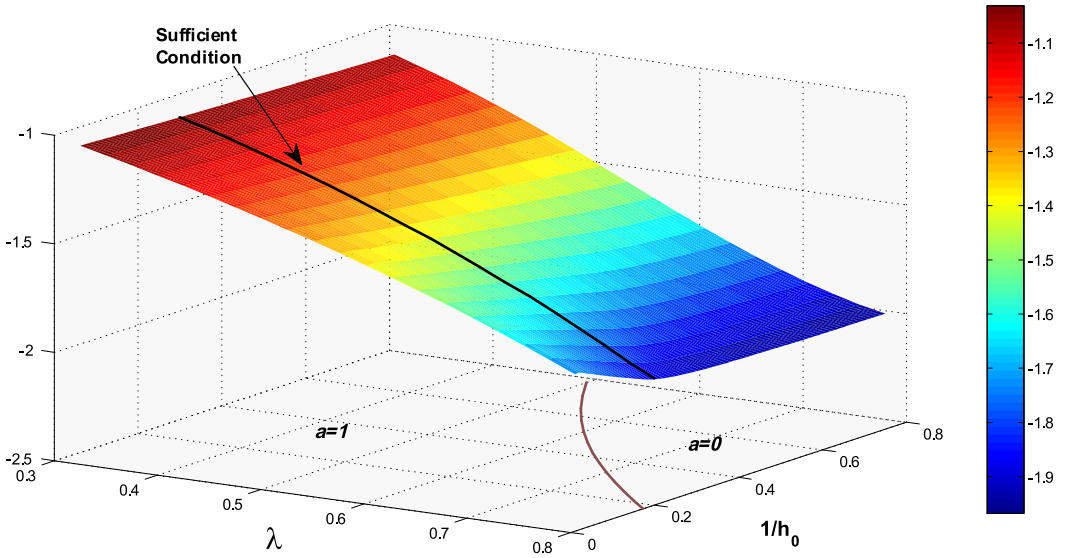


FIGURE 5. Agent's value as a function of $1/h_0$ and λ . Other parameters are as in Table 1.

maximal λ as a function of h_0 is given by

$$\lambda = \sqrt{\frac{1}{2}h_0\left(\frac{1}{4}h_0 - 1\right)}. \quad (30)$$

The RHS of this equation is positive only if $h_0 \geq 4$. In other words, (24) can be met only if $h_0^{-1} < 0.25$, and then more easily if λ is low enough. However, the RHS of (30) exceeds unity once $h_0^{-1} \leq 0.1830$. Then (24) holds for all $\lambda \in (0, 1)$.

As expected, the agent's value is decreasing in the prior variance $1/h_0$. Figure 5 also illustrates how an increase in λ lowers the surplus. This is what one should expect because the higher λ is, the more costly it is to provide effort. Hence an increase in λ intensifies the severity of the moral hazard problem, making it more costly for the principal to deliver a given utility.

6. COMMITMENT VERSUS SPOT MARKET

We now relate our model to the literature on reputations that typically adopts the interpretation that η is general ability. We focus on the canonical model of Holmström (1999; "H" hereafter), which assumes spot-market wages that may reflect the worker's history but cannot reflect current output.

In both Holmström's model and ours, the principal is risk-neutral. The agents' utility functions, however, differ because Holmström assumes that agents are risk-neutral. To make our analysis of commitment comparable to his analysis, we shall derive the spot-market equilibria of H for the case where the agent has period utility (18).

Holmström imposes zero expected profits for the principal after every history and at each date. In our model, the principal has full commitment and his profits will not

be zero at an arbitrary date. To compare our solution to H, it is natural to impose zero expected lifetime profits on the principal at the outset. Thus we shall assume that at date zero, the agent gets all the rents from the relationship.

We first show that the equilibrium behavior of spot-market wages and effort under risk aversion is essentially the same as in H: Reputational concerns are the only reason why the agent exerts any effort, and when information about η accumulates and as these concerns disappear, his effort converges to zero, just as in the risk-neutral case. Of itself this is not surprising. Rather, the result is useful because it enables us to isolate the role that full commitment plays in generating economic outcomes for the parties to the contract.

Employers cannot commit to paying wages that depend on performance, and competition among employers bids wages up to expected output. Denoting, as before, equilibrium actions by an asterisk, expected productivity reads

$$w_t^S = E[\eta + a_t | \mathcal{F}_t^Y] = \hat{\eta}(Y_t - A_t^*, t) + a_t^*, \quad (31)$$

where we have added an S superscript for spot wages. Effort is sustained by the market's imprecise knowledge of η and the agent's attempts to raise the market's expectation.

Definition of spot market equilibria. An equilibrium corresponds to a feasible strategy a^* that is $\mathbb{F}^Y$ -predictable and a wage process w^S such that (i) given a_t^* , the market sets a wage of the form given in equation (31); (ii) the continuation strategy a_s^* , for all $s \geq t$, maximizes the worker expected utility given the wage process in (31).

The solution concept is that of perfect Bayesian equilibrium since the strategy a^* is optimal given the law of motion of beliefs, and beliefs satisfy Bayes rule given the equilibrium action. We are restricting our attention to strategies that are adapted to the public signal Y . We do not prove that there always exists an equilibrium of this type or that it is unique.³⁵ But we can show that for any such equilibrium, the sequence of effort a_t^* eventually converges to zero.

PROPOSITION 6. *Assume that (i) $u(w, a)$ is as specified in (18) and (ii) the model admits at least one spot-market equilibrium. Then there exists a unique precision $\bar{h}$, whose value does not depend on the equilibrium strategy, and such that $a_t^* = 0$ whenever $h_t > \bar{h}$.*

Proposition 6 shows that although equilibrium strategies are not uniquely pinned down when precision is below $\bar{h}$, as soon as enough information is accumulated to reach this threshold, all equilibrium paths converge to zero and remain there thereafter.³⁶ Thus there is no feasible Pareto improvement over step profiles where effort is at its first-best level $a_t^* = 1$ whenever $h_t < \bar{h}$. Since the contracts described in the previous section

³⁵See Cisternas (2012) for a derivation of the conditions under which there exists a deterministic perfect public equilibrium when agents are risk-neutral.

³⁶Uniqueness is not ensured in our setup because the utility function (18) implies that the marginal cost of effort is decreasing in current income. As usual, this complementarity can be a source of equilibrium multiplicity.

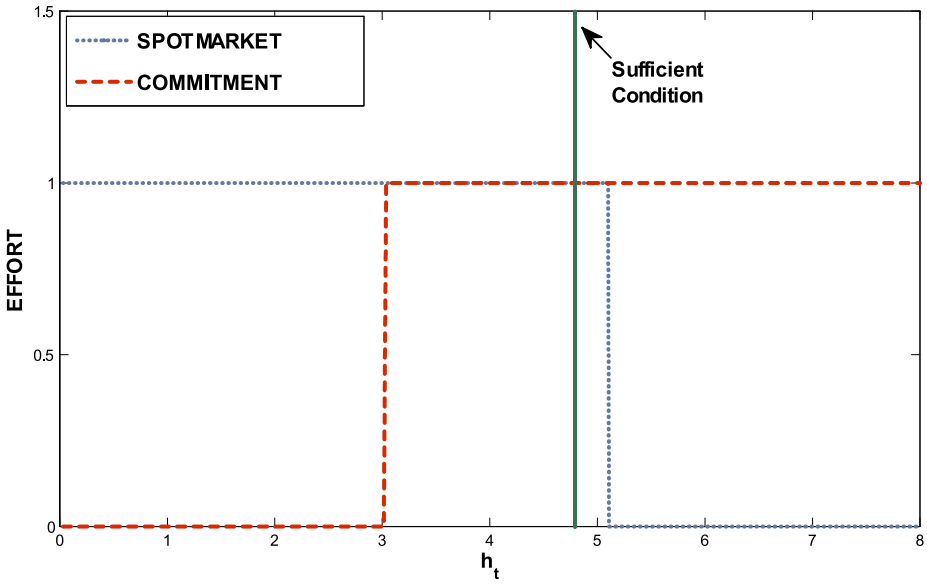


FIGURE 6. Effort as a function of parameter precision h_t . Parameters are as in Table 1.

lie on the Pareto frontier, it is natural to compare them to Pareto outcomes in the spot market. This ensures that we are not exaggerating the benefits of optimal contracting, but instead reporting its lower bound.

Figure 6 contains the effort path as a function of h_t when wages are set on the spot market along with its counterpart in the commitment scenario. It illustrates how uncertainty about general ability affects incentives in opposite directions. Spot markets elicit full effort when beliefs are imprecise because reputations have not yet been established. By contrast, under commitment, incentives are more costly to provide when precision is low. This is why the effort paths are almost mirror images of each other: It switches from 1 to 0 in the spot market and from 0 to 1 under commitment.³⁷ Their profiles are not smooth because the marginal cost of effort is *decreasing* in consumption. Hence, full effort can always be sustained through a less than proportional increase in wages whenever interior levels of effort are incentive compatible. Such a deviation is Pareto optimal and so dominates any equilibrium path with intermediate action.

Figure 6 does not accurately represent the distribution of lifetime gains that full commitment offers. That would be the distribution of the random variable $\mathcal{U}_0$ defined in (6), which we report in Figure 7.³⁸ While wages themselves are normally distributed, utilities are nonlinear and bounded above. This is why the distributions of $\mathcal{U}_0$ are skewed to the

³⁷Full effort in the spot market is incentive compatible for $h_t < \bar{h}$ since the function R_t defined in the proof of Proposition 6 is decreasing in h_t .

³⁸The distribution of lifetime utilities is obtained through Monte Carlo simulations. We simulate 10,000 sample paths and compute the resulting kernel densities. We verify the accuracy of the procedure by comparing the simulated and theoretical average utilities. The approximation error turns out to be around 10^{-3} in relative difference.

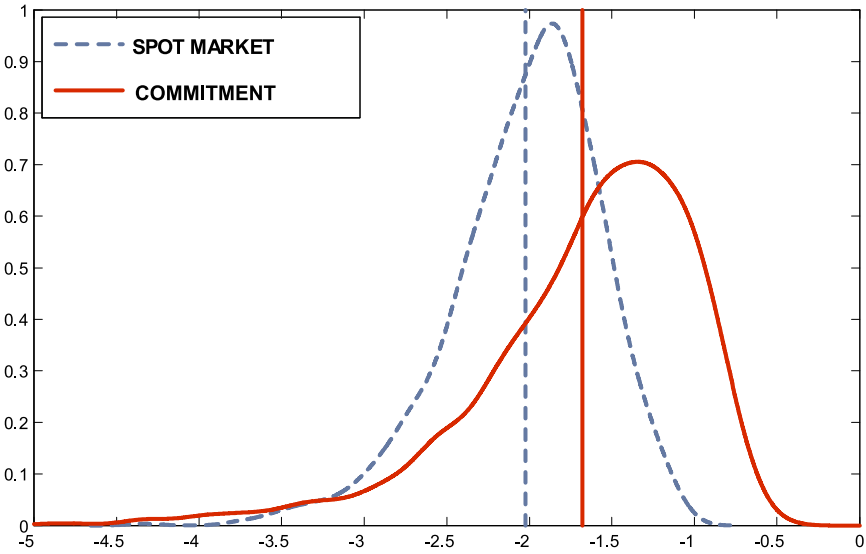


FIGURE 7. Distribution of lifetime utilities. Parameters are as in Table 1.

left with their means—as represented by the vertical lines—to the left of the modes. Figure 7 illustrates that commitment results in a noticeably higher expected lifetime utility $E_0[\mathcal{U}_0]$: Long-term contracts raise the agent's utility by 17.1%, a gain that is equivalent to a compensating variation of 26.4% in wages across first-best allocations.³⁹ Even though utilities derived from contracts exhibit more dispersion, they dominate from a stochastic point of view. In other words, not only the average worker, but most workers do benefit from contracting.

η as a match-specific ability. If, instead of denoting general ability, η were match-specific, then neither the optimal contract nor the Pareto frontier would change under full commitment. By contrast, spot markets would work poorly. The agent now has no reputational concern, implying that effort would remain constant at zero. The wage would equal $E_t[\eta]$ at all dates. The value of commitment is then even larger than in the case where ability is transferable.

Participation constraints. We have described two separate economies, each with its own wage-setting protocol. Our commitment solution is for a contract that would yield the principal zero expected profit at the outset, but after some histories his expected profit falls below zero. Similarly, the agent's continuation value may fall below the spot-market solution. An extension would add participation constraints as Rudanko (2009) and Lustig et al. (2011) have done for multi-agent environments without learning. We expect the risk of being fired to render belief manipulation less attractive.

³⁹The welfare gain is obtained by dividing the difference between the two expected utilities $E_0[\mathcal{U}_0]$ by the expected utility when wages are set on the spot market. To obtain the compensating variation, we first derive the wage such that $U(w, 1)/r = E_0[\mathcal{U}_0]$, which yields w^{Com} under commitment and w^{Spot} under spot market. The compensating variation follows by taking the difference between the two wages and dividing it by w^{Spot} .

In partial equilibrium settings without learning, there are more papers with limited commitment. Closely related to ours is the principal–agent model of [Sannikov \(2008\)](#), which, under some adjustments to the parametric form of the utility function, is encompassed in our framework as $h_t \rightarrow \infty$, i.e., when posteriors have converged to the true value of η . More precisely, Sannikov considers a utility function that is (i) defined over the positive real line, (ii) bounded from below, and (iii) separable in income and effort. By contrast, our utility function (18) is not bounded from below and, as a result, we do not have a low retirement point. Observe, however, that our characterization of the agent's necessary condition (10) does not depend on the parametric assumption (18) and so coincides with Sannikov's when precision becomes infinite.

Broader implications. Data show that early on in a career, a worker tends to engage in a job shopping phase, following which he settles down into a more permanent job. The model suggests an explanation for such behavior: as h grows, the gain to commitment rises, whereas incentives in the spot market decline.

Commitment should also raise incentives to accumulate human capital. [Becker \(1964\)](#) had argued that without commitment, a firm would not pay for general training. [Cisternas \(2012\)](#) modifies (1) and shows that in a spot market where human capital is general as in [Holmström \(1999\)](#), the worker's accumulation of skills is inefficiently low; our paper suggests that contracts with commitment could raise the level of accumulation, bringing it closer to the social optimum.⁴⁰

7. CONCLUSION

We have solved a contracting problem with quality uncertainty and explained why it worsens the trade-off between incentives and insurance. We developed an approach that works for any utility function when quality and noise are normally distributed. We found that the agent faces two opposite effects when considering a downward deviation from recommended effort. On the one hand, he will be punished by a lower promised value because of the decrease in observable output. On the other hand, he will benefit from higher expectations than the principal about the unknown productivity of the match. This second channel, which we label belief manipulation, is specific to problems under quality uncertainty. The extent to which it influences incentive provisions depends on the remaining length of the relationship. This is why it is not relevant in markets based on spot agreements.

Although the prospect of belief manipulation reduces the gains from commitment, our simulation shows that it does not eliminate them altogether. We found, in particular, that quality uncertainty makes it harder to reward effort under full commitment, in direct contrast to its tendency to stimulate effort in spot markets. Spot and full commitment settings are both highly stylized depictions of how markets operate in reality.

⁴⁰A positive deviation of effort in our model causes the principal to overestimate the agent's ability. The opposite is true in the case where effort raises accumulation of human capital: Here a positive deviation raises actual productivity and causes the principal to underestimate the agent's true level of human capital. This means that the value of private information, p , for human capital is not negative but positive, which encourages skill accumulation.

Thus a promising task would be to combine both environments in a model with limited commitment so as to evaluate how the two incentive channels interact.

APPENDIX A: PROOFS OF PROPOSITIONS AND COROLLARIES

PROOF OF PROPOSITION 1. Consider the Brownian motion Z^0 under some probability space with probability measure Q and $\mathbb{F}^{Z^0} \triangleq \{\mathcal{F}_t^{Z^0}\}_{0 \leq t \leq T}$ the suitably augmented filtration generated by Z^0 . Let

$$Y_t = \int_0^t \sigma dZ_s^0,$$

so that Y_t is also a Brownian motion under Q . Since expected output is linear in cumulative output,⁴¹ the exponential local martingale

$$\Lambda_{t,\tau}^a \triangleq \exp\left(\int_t^\tau \left(\frac{\hat{\eta}(Y_s - A_s, s) + a_s}{\sigma}\right) dZ_s^0 - \frac{1}{2} \int_t^\tau \left|\frac{\hat{\eta}(Y_s - A_s, s) + a_s}{\sigma}\right|^2 ds\right) \quad (32)$$

for $t \leq \tau \leq T$, is a martingale, i.e., $E_t[\Lambda_{t,T}^a] = 1$. Hence the Girsanov theorem holds and ensures that

$$Z_t^a \triangleq Z_t^0 - \int_0^t \left(\frac{\hat{\eta}(Y_s - A_s, s) + a_s}{\sigma}\right) ds$$

is a Brownian motion under the new probability measure $dQ^a/dQ \triangleq \Lambda_{0,T}^a$. Given that both measures are equivalent, the triple (Y, Z^a, Q^a) is a weak solution of the SDE

$$Y_t = \int_0^t (\hat{\eta}(Y_s - A_s, s) + a_s) ds + \int_0^t \sigma dZ_s^a.$$

Adopting a weak formulation allows us to view the choice of control a as determining the probability measure Q^a . So as to define the agent's optimization problem, let $R^a(t)$ denote the reward from time t onward so that

$$R^a(t) \triangleq e^{\rho t} \left[\int_t^T u(s, \bar{Y}_s, a_s) ds + U(T, \bar{Y}_T) \right],$$

where, with a slight abuse of notation, $u(s, \bar{Y}_s, a_s) \triangleq e^{-\rho s} u(w(\bar{Y}_s), a_s)$ and $U(T, \bar{Y}_T) \triangleq e^{-\rho T} U(\bar{Y}_T)$ are utilities at time t discounted from time 0. The agent's objective is to find an admissible control process that maximizes the expected reward $E^a[R^a(0)]$ over all admissible controls $a \in \mathcal{A}$. In other words, the agent solves the problem

$$v_t = \sup_{a \in \mathcal{A}} V^a(t) \triangleq \sup_{a \in \mathcal{A}} E_t^a[R^a(t)] \quad \text{for all } 0 \leq t \leq T.$$

⁴¹More formally, the martingale property holds true because

$$|\hat{\eta}(Y_t - A_t, t) + a_t| \leq C(1 + \|Z^0\|_t) \quad \text{for all } t \in [0, T],$$

with $C = \sigma^{-1}/h_0 + 1$ and $\|Z^0\|_t \triangleq \max_{0 \leq s \leq t} |Z^0(s)|$.

The objective function can be recast as

$$V^a(t) = E_t^a[R^a(t)] = E_t^0[\Lambda_{t,T}^a R^a(t)], \quad (33)$$

where the operators $E^a[\cdot]$ and $E^0[\cdot]$ are expectations under the probability measure Q^a and Q , respectively. One can see from (33) that varying a is indeed equivalent to changing the probability measure.

Our derivation of the necessary conditions builds on the variational argument in Cvitanić et al. (2009). Define the control perturbation

$$a^\varepsilon \triangleq a + \varepsilon \Delta a.$$

We assume that there exists an $\varepsilon_0 > 0$ for which any $\varepsilon \in [0, \varepsilon_0)$ satisfy $|a^\varepsilon|^4$, $|u^{a^\varepsilon}|^4$, $|u_{a^\varepsilon}^{a^\varepsilon}|^4$, $|\Lambda_{t,\tau}^{a^\varepsilon}|^4$, $(\mathcal{U}_{t,\tau}^{a^\varepsilon})^2$, and $(\partial_a \mathcal{U}_{t,\tau}^{a^\varepsilon})^2$ being uniformly integrable in $L^1(Q)$, where

$$\mathcal{U}_{t,\tau}^a \triangleq \int_t^\tau u(s, \bar{Y}_s, a_s) ds.$$

We introduce the following shorthand notations for “variations”:

$$\nabla \mathcal{U}_{t,\tau}^a \triangleq \int_t^\tau u_a(s, \bar{Y}_s, a_s) \Delta a_s ds \quad (34)$$

$$\nabla A_t \triangleq \int_0^t \Delta a_s ds \quad (35)$$

$$\begin{aligned} \nabla \Lambda_{t,\tau}^a &\triangleq \Lambda_{t,\tau}^a \left(\frac{1}{\sigma} \right) \left[\int_t^\tau \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^0 \right. \\ &\quad \left. - \int_t^\tau (\hat{\eta}_s + a_s) \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) ds \right] \\ &= \Lambda_{t,\tau}^a \left(\frac{1}{\sigma} \right) \int_t^\tau \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^a. \end{aligned} \quad (36)$$

Step 1. We first characterize the variations of the agent’s objective with respect to ε :

$$\begin{aligned} \frac{V^{a^\varepsilon}(t) - V^a(t)}{\varepsilon} &= E_t^0 \left[\left(\frac{\Lambda_{t,T}^{a^\varepsilon} - \Lambda_{t,T}^a}{\varepsilon} \right) R^{a^\varepsilon}(t) + \Lambda_{t,T}^a \left(\frac{R^{a^\varepsilon}(t) - R^a(t)}{\varepsilon} \right) \right] \\ &= E_t^0 \left[\nabla \Lambda_{t,T}^{a^\varepsilon} R^{a^\varepsilon}(t) + \Lambda_{t,T}^a \left(\frac{R^{a^\varepsilon}(t) - R^a(t)}{\varepsilon} \right) \right]. \end{aligned}$$

To obtain the limit of the first term as ε goes to zero, observe that

$$\nabla \Lambda_{t,T}^{a^\varepsilon} R^{a^\varepsilon}(t) - \nabla \Lambda_{t,T}^a R^a(t) = [\nabla \Lambda_{t,T}^{a^\varepsilon} - \nabla \Lambda_{t,T}^a] R^a(t) + \nabla \Lambda_{t,T}^{a^\varepsilon} [R^{a^\varepsilon}(t) - R^a(t)].$$

As shown in Cvitanić et al. (2009), for any $\varepsilon \in [0, \varepsilon_0)$, this expression is integrable uniformly with respect to ε and so

$$\lim_{\varepsilon \rightarrow 0} E_t^0[\nabla \Lambda_{t,T}^{a^\varepsilon} R^{a^\varepsilon}(t)] = E_t^0[\nabla \Lambda_{t,T}^a R^a(t)].$$

The limit of the second term reads

$$\lim_{\varepsilon \rightarrow 0} \frac{R^{a^\varepsilon}(t) - R^a(t)}{\varepsilon} = e^{\rho t} \nabla \mathcal{U}_{t,T}^a.$$

Due to the uniform integrability of $\Lambda_{t,T}^a(R^{a^\varepsilon}(t) - R^a(t))/\varepsilon$, the expectation is also well defined. Combining the two expressions above, we finally obtain

$$\lim_{\varepsilon \rightarrow 0} \frac{V^{a^\varepsilon}(t) - V^a(t)}{\varepsilon} = E_t^0[\nabla \Lambda_{t,T}^a R^a(t) + \Lambda_{t,T}^a e^{\rho t} \nabla \mathcal{U}_{t,T}^a] \triangleq \nabla V^a(t). \quad (37)$$

Step 2. We are now in a position to derive the necessary condition. Consider total earnings as of date 0:

$$I^a(t) \triangleq E_t^a \left[\int_0^T u(s, \bar{Y}_s, a_s) ds + U(T, \bar{Y}_T) \right] = \int_0^t u(s, \bar{Y}_s, a_s) ds + e^{-\rho t} V^a(t). \quad (38)$$

By definition, it is a Q^a -martingale. According to the extended martingale representation theorem⁴² of Fujisaki et al. (1972), all square integrable Q^a -martingales are stochastic integrals of $\{Z_t^a\}$ and there exists a unique process ζ in $L^2(Q^a)$ such that

$$I^a(T) = I^a(t) + \int_t^T \zeta_s \sigma dZ_s^a. \quad (39)$$

This decomposition allows us to solve for $\nabla V^a(t)$. Reinserting (34), (35), and (36) into (37) yields⁴³

$$\begin{aligned} \nabla V^a(t) &= E_t^0 \left[\Lambda_{t,T}^a R^a(t) \sigma^{-1} \int_t^T \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^a + \Lambda_{t,T}^a e^{\rho t} \left(\int_t^T u_a \Delta a_s ds \right) \right] \\ &= e^{\rho t} E_t^a \left[I^a(T) \sigma^{-1} \int_t^T \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^a + \int_t^T u_a \Delta a_s ds \right], \end{aligned}$$

where subscripts denote derivatives and arguments are omitted for brevity. Given the law of motion (39), applying Ito's rule to the first term yields

$$\begin{aligned} & d \left(I^a(\tau) \int_t^\tau \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^a \right) \\ &= \left[\zeta_\tau \sigma \left(-\left(\frac{\sigma^{-2}}{h_\tau} \right) \nabla A_\tau + \Delta a_\tau \right) \right] d\tau \\ &\quad + \left[\zeta_\tau \sigma \int_t^\tau \left(-\frac{\sigma^{-2}}{h_s} \nabla A_s + \Delta a_s \right) dZ_s^a + I^a(\tau) \left(-\left(\frac{\sigma^{-2}}{h_\tau} \right) \nabla A_\tau + \Delta a_\tau \right) \right] dZ_\tau^a. \end{aligned}$$

⁴²We cannot directly apply the standard martingale representation theorem because we are considering weak solutions, so that $\{Z_t^a\}$ does not necessarily generate $\{\mathcal{F}_t^Y\}$.

⁴³The additional expectation term vanishes because both $(h_\varepsilon/h_s)\nabla A_s$ and Δa_s are bounded, and so

$$\left(\int_0^t U(\tau, \bar{Y}_\tau, a_\tau) d\tau \right) E_t^a \left[\int_t^T \left(-\left(\frac{h_\varepsilon}{h_s} \right) \nabla A_s + \Delta a_s \right) dZ_s^a \right] = 0.$$

Hence $\nabla V^a(t)$ can be represented as

$$e^{-\rho t} \nabla V^a(t) = E_t^a \left[\int_t^T \Gamma_s^1 ds + \int_t^T \Gamma_s^2 dZ_s^a \right],$$

where

$$\begin{aligned} \Gamma_s^1 &\triangleq \zeta_s \left[-\frac{\sigma^{-2}}{h_s} \int_0^s \Delta a_\tau d\tau + \Delta a_s \right] + u_a(s, \bar{Y}_s, a_s) \Delta a_s \\ \Gamma_s^2 &\triangleq \zeta_s \left[\int_t^s \left(-\frac{\sigma^{-2}}{h_\tau} \int_0^\tau \Delta a_r dr + \Delta a_\tau \right) dZ_\tau^a \right] + I^a(s) \left(-\frac{\sigma^{-2}}{h_s} \int_0^s \Delta a_\tau d\tau + \Delta a_s \right). \end{aligned}$$

Given that Γ_s^2 is square integrable,⁴⁴ we have

$$E_t^a \left[\int_t^T \Gamma_s^2 dZ_s^a \right] = 0.$$

As for the deterministic term, collecting the effect of each perturbation Δa_s yields

$$e^{-\rho t} \nabla V^a(t) = E_t^a \left[\int_t^T \left(-\int_s^T \zeta_\tau \left(\frac{\sigma^{-2}}{h_\tau} \right) d\tau + \zeta_s + u_a(s, \bar{Y}_s, a_s) \right) \Delta a_s ds \right].$$

Finally, noticing that Δa_s was arbitrary leads to

$$\left(E_t^a \left[-\int_t^T \zeta_s \frac{\sigma^{-2}}{h_s} ds \right] + \zeta_t + u_a(t, \bar{Y}_t, a_t^*) \right) (a_t - a_t^*) \leq 0. \quad (40)$$

Step 3. We now rewrite our solution as a function of the promised value v_t . Differentiating (38) with respect to time yields

$$e^{-\rho t} dv_t - \rho e^{-\rho t} v_t + u(t, \bar{Y}_t, a_t) = dI^a(t) = \zeta_t \sigma dZ_t^a,$$

so that

$$dv_t = (\rho v_t - u(\bar{Y}_t, a_t)) dt + \gamma_t \sigma dZ_t^a,$$

with $\gamma_t \triangleq \zeta_t e^{\rho t}$. Collecting the exponential terms in (40) leads to (10). $\square$

PROOF OF PROPOSITION 2. The sufficient conditions are established by comparing the equilibrium path $\{a_t^*\}_{t=0}^T$ with an arbitrary effort path $\{a_t\}_{t=0}^T$. We define $\delta_t \triangleq a_t - a_t^*$ and $\Delta_t \triangleq \int_0^t \delta_s ds = A_t - A_t^*$ as the differences in current and cumulative effort between the arbitrary and recommended paths. We also attach an asterisk to denote the value of the $\mathbb{F}^Y$ -measurable stochastic processes along the equilibrium path. The Brownian motions generated by the two effort policies are related by

$$\begin{aligned} \sigma dZ_t^{a^*} &= \sigma dZ_t^a + [\hat{\eta}(Y_t - A_t, t) + a_t - \hat{\eta}(Y_t - A_t^*, t) - a_t^*] dt \\ &= \sigma dZ_t^a + \left[\delta_t - \frac{\sigma^{-2}}{h_t} \Delta_t \right] dt. \end{aligned}$$

⁴⁴Square integrability of Γ_s^2 can be established for any $\varepsilon \in [0, \varepsilon_0)$ following the same steps as in Lemma 7.3 of Cvitanić et al. (2009).

By definition, the total reward from the optimal policy reads

$$\begin{aligned} I^{a^*}(T) &= \int_0^T u(t, \bar{Y}_t, a_t^*) dt + U(T, \bar{Y}_T) = V^{a^*}(0) + \int_0^T \xi_t^* \sigma dZ_t^{a^*} \\ &= V^{a^*}(0) + \int_0^T \xi_t^* \left[\delta_t - \frac{\sigma^{-2}}{h_t} \Delta_t \right] dt + \int_0^T \xi_t^* \sigma dZ_t^a. \end{aligned}$$

Hence, the total reward from the arbitrary policy is given by

$$\begin{aligned} I^a(T) &= \int_0^T [u(t, \bar{Y}_t, a_t) - u(t, \bar{Y}_t, a_t^*)] dt + I^{a^*}(T) \\ &= \int_0^T [u(t, \bar{Y}_t, a_t) - u(t, \bar{Y}_t, a_t^*)] dt + V^{a^*}(0) \\ &\quad + \int_0^T \xi_t^* \left[\delta_t - \frac{\sigma^{-2}}{h_t} \Delta_t \right] dt + \int_0^T \xi_t^* \sigma dZ_t^a. \end{aligned}$$

Let us focus on the third term on the right hand side,

$$\begin{aligned} - \int_0^T \xi_t^* \frac{\sigma^{-2}}{h_t} \Delta_t dt &= - \int_0^T \xi_t^* \frac{\sigma^{-2}}{h_t} \left(\int_0^t \delta_s ds \right) dt = \int_0^T \delta_t \left(- \int_t^T \xi_s^* \frac{\sigma^{-2}}{h_s} ds \right) dt \\ &= \int_0^T \delta_t \left(e^{-\rho t} \frac{\sigma^{-2}}{h_t} p_t^* + \int_t^T \xi_s^* \sigma dZ_s^{a^*} \right) dt, \end{aligned}$$

where the last equality follows from the definition of p and ξ .⁴⁵ Changing the Brownian motion and taking expectation yields

$$\begin{aligned} V^a(0) - V^{a^*}(0) &= E_0^a[I^a(T)] - V^{a^*}(0) \\ &= E_0^a \left[\int_0^T \left(u(t, \bar{Y}_t, a_t) - u(t, \bar{Y}_t, a_t^*) + \delta_t \left(\xi_t^* + e^{-\rho t} \frac{\sigma^{-2}}{h_t} p_t^* \right) \right) dt \right] \\ &\quad + E_0^a \left[\int_0^T \delta_t \left(\int_t^T \xi_s^* \left(\delta_s - \frac{\sigma^{-2}}{h_s} \Delta_s \right) ds \right) dt \right] \\ &= E_0^a \left[\int_0^T e^{-\rho t} \left(u(w_t, a_t) - u(w_t, a_t^*) + \delta_t \left(\gamma_t^* + \frac{\sigma^{-2}}{h_t} p_t^* \right) \right) dt \right] \\ &\quad + E_0^a \left[\int_0^T \xi_t^* \Delta_t \left(\delta_t - \frac{\sigma^{-2}}{h_t} \Delta_t \right) dt \right]. \end{aligned}$$

We know from the optimization property of a_t^* that the first expectation term is at most equal to zero. On the other hand, the sign of the second expectation term is ambiguous.

⁴⁵Observe that this additional step is linked to the introduction of private information. Then the volatility ζ of the continuation value will differ on and off the equilibrium path. To the contrary, in problems without private information, the volatility remains constant because it only depends on observable output and not on past actions. This is why sufficiency holds without any restriction in, e.g., Schättler and Sung (1993) or Sannikov (2008).

To bound it, we introduce the predictable process⁴⁶ $\chi_t^* \triangleq \gamma_t^* - e^{\rho t} \xi_t^* A_t^*$ and define the function⁴⁷

$$H(t, a, A; \chi^*, \xi^*, p^*) \triangleq u(w, a) + (\chi^* + e^{\rho t} \xi^* A) a - e^{\rho t} \xi^* \frac{\sigma^{-2}}{h_t} A^2 + \frac{\sigma^{-2}}{h_t} p^* a.$$

Taking a linear approximation of $H(\cdot)$ around A^* yields

$$\begin{aligned} H_t(a_t, A_t) - H_t(a_t^*, A_t^*) &= \frac{\partial H_t(a_t^*, A_t^*)}{\partial A} \Delta_t \\ &= u(w_t, a_t) - u(w_t, a_t^*) + \delta_t \underbrace{\left(\chi_t^* + e^{\rho t} \xi_t^* A_t^* + \frac{\sigma^{-2}}{h_t} p_t^* \right)}_{=\gamma_t^*} + e^{\rho t} \xi_t^* \Delta_t \left(\delta_t - \frac{\sigma^{-2}}{h_t} \Delta_t \right), \end{aligned}$$

so that

$$V^a(0) - V^{a^*}(0) = E_0^a \left[\int_0^T e^{-\rho t} \left(H_t(a_t, A_t) - H_t(a_t^*, A_t^*) - \frac{\partial H_t(a_t^*, A_t^*)}{\partial A} \Delta_t \right) dt \right]$$

is negative when $H(\cdot)$ is jointly concave. Given that the agent seeks to maximize expected returns, imposing concavity ensures that a^* dominates any alternative effort path. Concavity is established by considering the Hessian matrix of $H(\cdot)$,

$$\mathcal{H}(t, a, A) = \begin{pmatrix} u_{aa}(w_t, a_t) & e^{\rho t} \xi_t \\ e^{\rho t} \xi_t & -2e^{\rho t} \xi_t \frac{\sigma^{-2}}{h_t} \end{pmatrix},$$

which is negative semidefinite when $-2(\sigma^{-2}/h_t)u_{aa}(w_t, a_t) \geq e^{\rho t} \xi_t$, as stated in (16). $\square$

PROOF OF CLAIM 1. If $a^* = 0$, there is no need to provide any incentives and so satisfying the necessary condition (20) cannot be optimal. Thus recommended effort is necessarily positive under the premise that (20) holds for all $s \geq t$. Given that effort levels lie in a compact set, the following relationship holds for all positive a^* :

$$1 - \frac{\partial J^T(t, v)}{\partial v} u_a(w, a^*) + \sigma^2 \frac{\partial^2 J^T(t, v)}{\partial v^2} \gamma(t, v, w, a^*) \frac{\partial \gamma(t, v, w, a^*)}{\partial a} \geq 0.$$

By contrast, wages take value over the real line and so fulfill the optimality condition

$$-1 - \frac{\partial J^T(t, v)}{\partial v} u_w(w, a^*) + \sigma^2 \frac{\partial^2 J^T(t, v)}{\partial v^2} \gamma(t, v, w, a^*) \frac{\partial \gamma(t, v, w, a^*)}{\partial w} = 0. \quad (41)$$

Under our premise that the incentive constraint (13) holds with equality, we obtain $\partial \gamma / \partial w = -\lambda \partial \gamma / \partial a > -\partial \gamma / \partial a$, which implies, in turn, that when the optimality condition for wages binds, the one for effort is slack. It follows that optimal effort is constant and set equal to the upper bound $a^* = 1$. $\square$

⁴⁶ χ^* is predictable since both ξ^* and A^* are $\mathbb{F}^Y$ -predictable.

⁴⁷We use $H(\cdot)$ to denote the function because it is equivalent to the Hamiltonian of the optimal control problem that can be derived following Williams's (2008) method based on the stochastic maximum principle.

PROOF OF PROPOSITION 3. We seek a solution of the form

$$\rho j^T(t, v) = F_t^T + \frac{\ln(-\rho v)}{\theta}$$

and guess that the associated wage schedule is given by

$$w(v) = -\frac{\ln(-K_t v)}{\theta} + \lambda \quad \Rightarrow \quad u(w(v), 1) = K_t v,$$

where F_t^T and K_t are continuously differentiable functions that depend solely on time t . Given that we are focusing on cases where ability η is known, we can set the value of private information p equal to zero. This means that the incentive constraint boils down to

$$\gamma_t = -u_a(w_t, a_t) = -\theta \lambda K_t v_t.$$

The condition is similar to the one derived by Sannikov (2008). Since private information is not persistent, both parties share common expectations about changes in the promised value, even when the agent's promised value has deviated in the past. This implies that the necessary conditions are also sufficient.

The first-order condition (FOC) (41) for wages reads

$$-1 + \frac{\partial j^T(t, v)}{\partial v} \theta K_t v - \sigma_t^2 \frac{\partial^2 j^T(t, v)}{\partial v^2} \theta^3 (K_t \lambda v)^2 = \frac{1}{\rho} [-\rho + K_t + (K_t \sigma_t \theta \lambda)^2] = 0.$$

Thus it is satisfied when

$$K_t = \frac{\sqrt{1 + 4\rho(\sigma_t \theta \lambda)^2} - 1}{2(\sigma_t \theta \lambda)^2} > 0,$$

whereas the HJB equation

$$\begin{aligned} \rho j^T(t, v) &= 1 - w + \frac{\partial j^T(t, v)}{\partial t} + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w(v), 1)) + \frac{\partial^2 j^T(t, v)}{\partial v^2} \left(\frac{1}{2} \sigma_t^2\right) \gamma^2 \\ &= 1 + \frac{\ln(-K_t v)}{\theta} - \lambda + \frac{1}{\rho} \frac{dF^T(t)}{dt} + \frac{\rho - K_t}{\rho \theta} - \frac{\theta}{\rho} \left(\frac{1}{2} (\sigma_t \lambda K_t)^2\right), \end{aligned}$$

is satisfied when⁴⁸

$$\frac{dF^T(t)}{dt} - \rho F^T(t) = -\rho \left(1 - \lambda + \frac{\ln(K_t/\rho)}{\theta}\right) - \frac{1}{2} \theta (\sigma_t \lambda K_t)^2.$$

The function $F^T(t)$ is, therefore, given by

$$F^T(t) = -\int_t^T e^{-\rho(s-t)} \left[\rho \left(1 - \lambda + \frac{\ln(K_s/\rho)}{\theta}\right) + \frac{1}{2} \theta (\sigma_s \lambda K_s)^2 \right] ds.$$

⁴⁸The equality follows by replacing the identity $(K_t \sigma \theta \lambda)^2 = \rho - K_t$.

As desired, our guess $j^T \in \mathcal{C}^{1,2}([0, T[\times \mathbb{R}) \cap \mathcal{C}([0, T] \times \mathbb{R})$ is a solution of the dynamic programming equation

$$\rho j^T(t, v) = \frac{\partial j^T(t, v)}{\partial t} + \sup_{(w, a) \in \mathcal{A}} \left\{ a - w + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w(v), a)) + \frac{\partial^2 j^T(t, v)}{\partial v^2} \left(\frac{1}{2} \sigma_t^2 \right) \gamma(t, v, w, a)^2 \right\},$$

where $\mathcal{A} = \mathbb{R} \times (0, 1]$ denote the set of admissible controls (w, a) ,⁴⁹ with boundary condition

$$\rho j^T(T, v) = -\rho U^{-1}(v) = \frac{\ln(-\rho v)}{\theta}.$$

Furthermore, for each fixed $v \in \mathbb{R}$, the supremum in the expression

$$\sup_{(w, a) \in \mathcal{A}} \left\{ a - w + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w(v), a)) + \frac{\partial^2 j^T(t, v)}{\partial v^2} \left(\frac{1}{2} \sigma_t^2 \right) \gamma(t, v, w, a)^2 \right\}$$

is attained by $w_t^*(v) = -\ln(-K_t v)/\theta + \lambda$ and $a^*(v) = 1$.⁵⁰ Thus the verification theorem holds and the value function for the control problem $J_N^T(t, v) = j^T(t, v)$ while $w^*(v)$ and $a^*(v)$ are optimal Markovian controls. $\square$

PROOF OF PROPOSITION 4. We break down the derivation of the principal's value into six steps:

Step 1. Initial guess. We seek a solution to the HJB equation of the form

$$\rho j^T(t, v) = f_t^T + \frac{\ln(-\rho v)}{\theta},$$

and guess that the associated wage schedule and value of private information are given by

$$w(t, v) = -\frac{\ln(-k_t^T v)}{\theta} + \lambda \quad \Rightarrow \quad u(w(t, v), 1) = k_t^T v$$

$$p_t(v) = \theta \lambda \varphi_t^T v,$$

where f_t^T , k_t^T , and φ_t^T are continuously differentiable functions that depend solely on time t . We impose the boundary condition

$$\rho j^T(T, v) = \frac{\ln(-\rho v)}{\theta}.$$

Step 2. Incentive constraint. According to our guess, the incentive constraint reads

$$\gamma_t = -u_a(w(t, v), a_t) - \frac{\sigma^{-2}}{h_t} p_t = -\theta \lambda \left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right) v_t.$$

⁴⁹We exclude $a = 0$ from the set of admissible controls because we are focusing on incentive-providing contracts.

⁵⁰See the proof of Claim 1 for the optimality of full effort.

Differentiating it with respect to wages yields

$$\frac{\partial \gamma(t, v, w, a)}{\partial w} = -u_{aw}(w(t, v), a_t) = -\theta \gamma(t, v, w, a) - \theta^2 \lambda \frac{\sigma^{-2}}{h_t} \varphi_t^T v_t.$$

The FOC (41) for wages is, therefore, equivalent to

$$\begin{aligned} -1 + \frac{\partial j^T(\cdot)}{\partial v} \theta v k_t^T - \sigma^2 \frac{\partial^2 j^T(\cdot)}{\partial v^2} \theta^3 (\lambda v)^2 \left[\left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right)^2 - \left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right) \frac{\sigma^{-2}}{h_t} \varphi_t^T \right] \\ = \frac{1}{\rho} \left(-\rho + k_t^T + (\theta \lambda \sigma)^2 \left[k_t^T \left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right) \right] \right) = 0, \end{aligned}$$

implying the quadratic equation for k_t^T ,

$$(k_t^T \theta \lambda \sigma)^2 + k_t^T \left(1 + (\theta \lambda)^2 \frac{\varphi_t^T}{h_t} \right) - \rho = 0. \quad (42)$$

The relevant solution is given by the positive root because wages are not defined when k_t^T is negative.

Step 3. HJB equation. We now verify that the dynamic programming equation is indeed satisfied:

$$\begin{aligned} \rho j^T(t, v) &= 1 - w + \frac{\partial j^T(t, v)}{\partial t} + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w(t, v), 1)) + \frac{\partial^2 j^T(t, v)}{\partial v^2} \left(\frac{1}{2} \sigma^2 \right) \gamma^2 \\ &= 1 + \frac{\ln(-v)}{\theta} + \frac{\ln(k_t^T)}{\theta} - \lambda + \frac{1}{\rho} \frac{df^T(t)}{dt} + \frac{\rho - k_t^T}{\rho \theta} - \frac{(\sigma \lambda)^2 \theta}{2\rho} \left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right)^2. \end{aligned}$$

Replacing our guess for $j^T(t, v)$ on the left hand side, one sees that the HJB equation holds true for all promised value v as long as⁵¹

$$\frac{df^T(t)}{dt} - \rho f^T(t) = -\rho \left(1 - \lambda + \frac{\ln(k_t^T/\rho)}{\theta} \right) + \frac{1}{2} (\sigma \lambda)^2 \theta \left(\left(\frac{\sigma^{-2}}{h_t} \varphi_t^T \right)^2 - (k_t^T)^2 \right) \quad (43)$$

or

$$f^T(t) = - \int_t^T e^{-\rho(s-t)} \underbrace{\left[\rho \left(1 - \lambda + \frac{\ln(k_s^T/\rho)}{\theta} \right) - \frac{1}{2} (\sigma \lambda)^2 \theta \left(\left(\frac{\sigma^{-2}}{h_s} \varphi_s^T \right)^2 - (k_s^T)^2 \right) \right]}_{\triangleq \psi_s^T} ds.$$

Step 4. Verification of the guess for p_t . To derive the value of private information, we first have to characterize the law of motion of the promised value. Since

$$\gamma_t(v) = \Gamma_t^T v \triangleq -\theta \lambda \left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right) v,$$

⁵¹Equation (43) is obtained by reinserting

$$\rho - k_t^T = \frac{1}{2} (\sigma \theta \lambda)^2 \left[\left(k_t^T + \frac{\sigma^{-2}}{h_t} \varphi_t^T \right)^2 - \left(\left(\frac{\sigma^{-2}}{h_t} \varphi_t^T \right)^2 - (k_t^T)^2 \right) \right]$$

into the HJB equation.

reinserting its expression along with that of wages into the SDE for v_t yields

$$dv_t = [\rho v_t - u(w_t, a_t)] dt + \gamma_t \sigma dZ_t = v_t[(\rho - k_t^T) dt + \Gamma_t^T \sigma dZ_t].$$

Whenever Γ_t^T is bounded (a conjecture that will be justified below), we have $E_t[v_T] = v_t \exp(\int_t^T (\rho - k_s^T) ds)$ and so

$$p_t = \theta \lambda (v_t - e^{-\rho(T-t)} E_t[v_T]) = \theta \lambda \underbrace{(1 - e^{-\int_t^T k_s^T ds})}_{\triangleq \varphi_t^T} v_t.$$

Since k_t^T is always positive, we have $\varphi_t^T \in (0, 1)$ for all $t < T$. This and the fact that $k_t^T \in (0, K)$ implies that, as desired, Γ_t^T is bounded. Hence, we have verified our conjecture for the functional form of p_t . Notice that $p_T = \varphi_T^T = 0$, quality is revealed at time T , and there is no remaining informational rent derived from belief manipulation. In other words, the power of incentives increases as the horizon T approaches.

Step 5. Verification theorem. Our guess for the value function $j^T \in \mathcal{C}^{1,2}([0, T] \times \mathbb{R}) \cap \mathcal{C}([0, T] \times \mathbb{R})$ is a solution of the dynamic programming equation

$$\begin{aligned} \rho j^T(t, v) = & \frac{\partial j^T(t, v)}{\partial t} + \sup_{(w, a) \in \mathcal{A}} \left\{ a - w + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w, a)) \right. \\ & \left. + \frac{\partial^2 j^T(t, v)}{\partial v^2} \left(\frac{1}{2} \sigma^2 \right) \gamma(t, v, w, a)^2 \right\}, \end{aligned}$$

where $\mathcal{A} = \mathbb{R} \times (0, 1]$ denote the set of admissible controls (w, a) ,⁵² with boundary condition

$$\rho j^T(T, v) = \frac{\ln(-\rho v)}{\theta}.$$

Furthermore, for each fixed $(t, v) \in [0, T] \times \mathbb{R}$, the supremum in the expression

$$\sup_{(w, a) \in \mathcal{A}} \left\{ a - w + \frac{\partial j^T(t, v)}{\partial v} (\rho v - u(w, a)) + \frac{\partial^2 j^T(t, v)}{\partial v^2} \frac{\sigma^2}{2} \gamma(t, v, w, a)^2 \right\}$$

is attained by $w^*(t, v) = -\ln(-k_t^T v)/\theta + \lambda$ and $a^*(t, v) = 1$.⁵³ Thus the verification theorem holds and the value function for the control problem $J_L^T(t, v) = j^T(t, v)$ while $w^*(t, v)$ and $a^*(t, v)$ are optimal Markovian controls.

Step 6. Convergence as T goes to infinity. The solutions for k_t^T and φ_t^T are found by backward induction using the terminal condition $\varphi_T^T = 0$ and reinserting the expression for p_t into the quadratic equation (42), which defines k_t^T . Differentiating the explicit solution for k_t^T with respect to φ_t^T yields

$$\frac{dk_t^T}{d\varphi_t^T} = \frac{1}{2\sigma^2 h_t} \left[\frac{1 + (\theta\lambda)^2 (\varphi_t^T / h_t)}{\sqrt{(1 + (\theta\lambda)^2 (\varphi_t^T / h_t))^2 + 4\rho(\theta\lambda\sigma)^2}} - 1 \right] < 0.$$

⁵²Remember that we are focusing on incentive-providing contracts. This is why $a = 0$ is excluded from the set of admissible controls.

⁵³See the proof of Claim 1 and Step 3 above.

Given that $k_t^T > 0$ and $\varphi_t^T \in (0, 1)$ for all $t < T$, we have the bound

$$k_t^T > \frac{\sqrt{(1 + (\theta\lambda)^2/h_t)^2 + 4\rho(\theta\lambda\sigma)^2} - 1 - (\theta\lambda)^2/h_t}{2(\theta\lambda\sigma)^2} \triangleq k_t.$$

Furthermore, differentiating k_t with respect to time shows that $dk_t/dt > 0$. Hence, it follows that $\int_t^T k_s^T ds > \int_t^T k_s ds > k_t(T - t)$, which implies, in turn, that $\lim_{T \rightarrow \infty} \int_t^T k_s^T ds = \infty$ and so

$$\varphi_t \triangleq \lim_{T \rightarrow \infty} \varphi_t^T = \lim_{T \rightarrow \infty} (1 - e^{-\int_t^T k_s^T ds}) = 1.$$

Reinserting this limit into (42), we find that $\lim_{T \rightarrow \infty} k_t^T = k_t$. The pointwise convergence of k_t^T to k_t and φ_t^T to 1 ensures that ψ_t^T converges pointwise to

$$\psi_t \triangleq \lim_{T \rightarrow \infty} \psi_t^T = \rho \left(1 - \lambda + \frac{\ln(k_t/\rho)}{\theta} \right) - \frac{1}{2}(\sigma\lambda)^2\theta \left(\left(\frac{\sigma^{-2}}{h_t} \right)^2 - k_t^2 \right). \quad (44)$$

The convention that $\psi_t^T = 0$ for all $t > T$ preserves its convergence property and enables us to rewrite the definition of the sequence of functions f_t^T as $f_t^T = -\int_t^\infty e^{-\rho(s-t)} \psi_s^T ds$. Given that $|\psi_t^T|$ is bounded, $e^{-\rho(s-t)} \psi_s^T$ is dominated by some integrable function and we can apply the dominated convergence theorem to conclude that $\lim_{T \rightarrow \infty} f_t^T = f_t = -\int_t^\infty e^{-\rho(s-t)} \psi_s ds$. $\square$

PROOF OF COROLLARY 1. Differentiating the expression of $\psi(t)$ in (44) with respect to time yields⁵⁴

$$\psi'(t) = \left(\frac{\rho}{\theta} \right) \frac{\dot{k}_t}{k_t} - (\sigma\lambda)^2\theta \left(-\frac{\sigma^{-2}}{h_t^3} - \dot{k}_t k_t \right) > 0.$$

Observe that $\psi(t)$ has been defined so as to satisfy the differential equation $f'(t) = \rho f(t) - \psi(t)$. To reach a contradiction, assume that $\rho f(t) < \psi(t)$. Then $f'(t) < 0$ and so $\rho f(s) < \psi(t) < \psi(s)$ for all $s \geq t$. But this contradicts the boundary condition $\lim_{s \rightarrow \infty} \rho f(s) = \lim_{s \rightarrow \infty} \psi(s) > 0$. We can, therefore, conclude that $\rho f(t) > \psi(t)$, which implies, in turn, that $f'(t) > 0$. $\square$

PROOF OF COROLLARY 2. Letting the horizon T go to infinity allows us to replace the backward stochastic differential equations for the co-states by standard stochastic differential equations. To derive the law of motion of p_t , we introduce the auxiliary process

$$b_t \triangleq E \left[-\int_0^T e^{-\rho s} \gamma_s \left(\frac{h_\varepsilon}{h_s} \right) ds \middle| \mathcal{F}_t^a \right] = b_0 + \int_0^t \xi_s \sigma dZ_s \quad \text{for all } t \in [0, T],$$

where the second equality follows from (17). Then the definition of p_t in (14) implies that

$$p_t = e^{\rho t} \sigma^2 h_t \left[b_t + \int_0^t e^{-\rho s} \gamma_s \frac{\sigma^{-2}}{h_s} ds \right],$$

⁵⁴Remember that both $\dot{k}_t$ and k_t are negative.

and so as $T \rightarrow \infty$, p_t solves the SDE⁵⁵

$$dp_t = \left[\rho p_t + \frac{d(\sigma^2 h_t)}{dt} \frac{\sigma^{-2}}{h_t} p_t + \gamma_t \right] dt + e^{\rho t} \sigma^2 h_t db_t = \left[p_t \left(\rho + \frac{\sigma^{-2}}{h_t} \right) + \gamma_t \right] dt + \vartheta_t \sigma dZ_t,$$

with $\vartheta_t \triangleq e^{\rho t} \sigma^2 h_t \xi_t$.

Given that we are focusing on cases where $a_t^* = 1$, then $u_{aa}(w_t, a_t^*) = (\theta\lambda)^2 v k_t$ and $\vartheta_t^* = \theta\lambda\gamma_t^*(v) = -(\theta\lambda)^2 v(k_t + \sigma^{-2}/h_t)$, so that the sufficient conditions of [Proposition 2](#) are satisfied when

$$-2k_t v + v \left(k_t + \frac{\sigma^{-2}}{h_t} \right) = v \left(\frac{\sigma^{-2}}{h_t} - k_t \right) > 0 \quad \Leftrightarrow \quad k_t > \frac{\sigma^{-2}}{h_t}. \quad (45)$$

Differentiating the explicit solution of the quadratic equation for k_t yields

$$\frac{dk(t)}{dt} = \frac{1}{2} \left[\frac{1/(\theta\lambda\sigma)^2 + \sigma^{-2}/h_t}{\sqrt{(1/(\theta\lambda\sigma)^2 + \sigma^{-2}/h_t)^2 + 4\rho/(\theta\lambda\sigma)^2}} - 1 \right] \underbrace{\frac{d(\sigma^{-2}h_t^{-1})}{dt}}_{<0} > 0.$$

Since σ^{-2}/h_t is decreasing in t , condition (45) is satisfied for all $s \geq t$ provided that $k_t > \sigma^{-2}/h_t$, i.e.,

$$-\frac{1}{(\theta\lambda\sigma)^2} - 3 \left(\frac{\sigma^{-2}}{h_t} \right) + \sqrt{\left(\frac{1}{(\theta\lambda\sigma)^2} + \left(\frac{\sigma^{-2}}{h_t} \right) \right)^2 + \frac{4\rho}{(\theta\lambda\sigma)^2}} > 0,$$

which, after some straightforward simplifications, leads to requirement (24). $\square$

PROOF OF PROPOSITION 5. Consider an arbitrary strategy such that (13) does not hold over some time interval $[t, t + \varepsilon]$ with $\varepsilon > 0$.⁵⁶ This implies that recommended effort $a_\tau^* = 0$ for all $\tau \in [t, t + \varepsilon]$ and so there is no gain in letting wages fluctuate, as this reduces the agent's welfare without extracting any additional effort. It is, therefore, efficient for the principal to fully stabilize wages within the time frame $[t, t + \varepsilon]$. Let $w^\Delta(v)$ be defined as

$$w^\Delta(v) \triangleq -\frac{\ln(-\rho(v + \Delta))}{\theta} \quad \Rightarrow \quad u(w^\Delta(v), 0) = \rho(v + \Delta),$$

so that the promise-keeping constraint holds when

$$\int_t^{t+\varepsilon} e^{-\rho(t-s)} u(w^\Delta(v_t), 0) dt + e^{-\rho\varepsilon} v_{t+\varepsilon} = v,$$

⁵⁵The change with respect to time of σ^{-2}/h_t is given by

$$\frac{d(\sigma^{-2}/h_t)}{dt} = \frac{d(\sigma^{-2}(h_0 + t\sigma^{-2})^{-1})}{dt} = -\sigma^{-4}(h_0 + t\sigma^{-2})^{-2} = -\left(\frac{\sigma^{-2}}{h_t} \right)^2 < 0.$$

⁵⁶The proof easily extends to arbitrary strategies where the incentive constraint does not hold over a finite number of time intervals $[t_i, t_i + \varepsilon_i]$ with $\varepsilon_i > 0$, $t_{i+1} > t_i + \varepsilon_i$, and $0 < i \leq I < \infty$. One simply has to consider the last interval $[t_I, t_I + \varepsilon_I]$ and follow the logic of the proof to reach a contradiction.

that is, if

$$v_{t+\varepsilon} = e^{\rho\varepsilon}v_t - (e^{\rho\varepsilon} - 1)(v_t + \Delta) = v_t - (e^{\rho\varepsilon} - 1)\Delta.$$

The parameter Δ measures by how much wages differ from the certainty equivalent $w(v) = -\ln(-\rho v)/\theta$. We now show that it is optimal to set $\Delta = 0$ whenever the agent provides positive effort at every point in time following $t + \varepsilon$. Let $i(\Delta, \varepsilon; t, v)$ denote the expected profits of the principal for arbitrary values of Δ . It is equal to

$$\begin{aligned} i(\Delta, \varepsilon; t, v) &= \int_t^{t+\varepsilon} e^{-\rho(s-t)} w^\Delta(v) ds + e^{-\rho\varepsilon} J(t + \varepsilon, v - (e^{\rho\varepsilon} - 1)\Delta) \\ &= \frac{1}{\rho} \left[(1 - e^{-\rho\varepsilon}) \frac{\ln(-\rho(v + \Delta))}{\theta} + e^{-\rho\varepsilon} f(t + \varepsilon) + e^{-\rho\varepsilon} \frac{\ln(-\rho(v - (e^{\rho\varepsilon} - 1)\Delta))}{\theta} \right]. \end{aligned}$$

Totally differentiating this expression with respect to Δ shows that it is concave in Δ and that it reaches its maximum when $\Delta = 0$. We can, therefore, focus on $i(0, \varepsilon; t, v)$. Differentiating its expression with respect to the length ε of the interval where the worker is perfectly insured yields

$$\begin{aligned} \frac{\partial i(0, \varepsilon; t, v)}{\partial \varepsilon} &= \frac{\partial}{\partial \varepsilon} \left(\frac{1}{\rho} \left[e^{-\rho\varepsilon} f(t + \varepsilon) + \frac{\ln(-\rho v)}{\theta} \right] \right) \\ &= \frac{1}{\rho} [e^{-\rho\varepsilon} (f'(t + \varepsilon) - \rho f(t + \varepsilon))] = -\frac{e^{-\rho\varepsilon}}{\rho} \psi(t + \varepsilon), \end{aligned}$$

where $\psi(t)$ is defined in (44). Let us now consider the cases identified in the proposition.

- (i) $F > 0$: The boundary condition $\lim_{t \rightarrow \infty} f'(t) = 0$ implies that $\lim_{h \rightarrow \infty} \psi(h) = \rho F$. Furthermore, it follows from $\lim_{h \rightarrow 0} k(h) = 0$ that $\lim_{h \rightarrow 0} \psi(h) = -\infty$. Hence $\psi(h)$ must switch sign at least once when F is positive. However, there can only be one precision such that $\psi(h) = 0$ since we have shown in the proof of [Corollary 1](#) that $\psi'(h) > 0$. We can, therefore, conclude from the equation above that it is optimal to insure workers when $h_t < \tilde{h}$ as $\partial i(0, \varepsilon; t, v)/\partial \varepsilon > 0$. Conversely, when $h_t > \tilde{h}$, $\partial i(0, \varepsilon; t, v)/\partial \varepsilon < 0$, showing that it cannot be optimal to insure workers within any time interval of positive finite measure. We still have to consider cases where $\varepsilon \rightarrow \infty$ so that workers are perfectly insured after a given date t . But this is clearly suboptimal since $f(t) > \psi(t) > 0$ and so

$$\rho i(0, \infty; t, v) = \frac{\ln(-\rho v)}{\theta} < f(t) + \frac{\ln(-\rho v)}{\theta} = \rho J(t, v).$$

- (ii) $F \leq 0$: Then $\psi(h) < \lim_{h \rightarrow 0} \psi(h) = F \leq 0$, implying that $\partial i(0, \varepsilon; t, v)/\partial \varepsilon > 0$ for all t . In other words, provision of full insurance always maximizes profits. $\square$

PROOF OF COROLLARY 3. Reinserting the law of motion (28) for v into (27) and applying Itô's lemma yields

$$dw_t^* = -\left(\frac{1}{\theta}\right) \left[\left(\left(\frac{1}{k_t} \right) \frac{dk_t}{dt} - \frac{1}{2} (\theta \lambda \sigma)^2 \left(\left(\frac{\sigma^{-2}}{h_t} \right)^2 - k_t^2 \right) \right) dt - \theta \lambda \left(k_t + \frac{\sigma^{-2}}{h_t} \right) \sigma dZ_t \right].$$

The statement for the volatility term is established by reinserting $dk(\sigma^{-2}/h_t)/d(\sigma^{-2}/h_t)$ into

$$\frac{d(k(t) + \sigma^{-2}/h_t)}{dt} = \underbrace{\left[\frac{dk(\sigma^{-2}/h_t)}{d(\sigma^{-2}/h_t)} + 1 \right]}_{\in (-1/2, 0)} \underbrace{\frac{d(\sigma^{-2}/h_t)}{dt}}_{< 0} < 0.$$

The sign of the deterministic trend is established by remembering that the sufficient condition (16) holds if and only if $k_t > \sigma^{-2}/h_t$. Hence, $(\sigma^{-2}/h_t)^2 - k_t^2 < 0$, and so the trend is negative. $\square$

PROOF OF PROPOSITION 6. The proof proceeds in two steps.

Step 1. First, we shall establish that the necessary condition (10) for incentive compatibility is equivalent to

$$[R_t - \lambda \exp(-\theta(1 - \lambda)a_t^*)](a - a_t^*) \leq 0 \quad \text{for all } a \in [0, 1], \quad (46)$$

where

$$R_t \triangleq \int_t^\infty e^{-\rho(s-t)} \frac{\sigma^{-2}}{h_s} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}) - \theta(1 - \lambda)a_s^*\right) ds.$$

Notice that our focus on equilibrium paths whose recommended actions a^* depend on time alone allows us to treat R_t as a deterministic integral. By definition, the continuation value along the equilibrium path $\mathcal{U}(t, \hat{\eta}_t)$ is given by

$$\mathcal{U}(t, \hat{\eta}_t) = \int_t^{+\infty} e^{-\rho(s-t)} E_t[u(w^S(\hat{\eta}_s, a_s^*), a_s^*)] ds. \quad (47)$$

To solve for $\mathcal{U}(t, \hat{\eta}_t)$, we have to evaluate expected utilities. The market posterior $\hat{\eta}$ satisfies the law of motion $d\hat{\eta}_t = (\sigma^{-1}/h_t) dZ_t$. Hence $\hat{\eta}_s$ is normally distributed with mean $\hat{\eta}_t$ and variance $\text{Var}_t(\hat{\eta}_s) = E_t[(d\hat{\eta}_s)^2] = h_t^{-1} - h_s^{-1}$, which implies, in turn, that

$$\begin{aligned} E_t[u(w^S(\hat{\eta}_s, a_s^*), a_s^*)] &= -E_t[\exp(-\theta\hat{\eta}_s)] \exp(-\theta(1 - \lambda)a_s^*) \\ &= -\exp\left(-\theta\hat{\eta}_t + \frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}) - \theta(1 - \lambda)a_s^*\right). \end{aligned}$$

Reinserting this expression into (47) yields

$$\mathcal{U}(t, \hat{\eta}_t) = -\exp(-\theta\hat{\eta}_t) \int_t^\infty e^{-\rho(s-t)} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}) - \theta(1 - \lambda)a_s^*\right) ds. \quad (48)$$

Since $d\hat{\eta}_t = (\sigma^{-1}/h_t) dZ_t$, Ito's lemma implies that

$$d\mathcal{U}(t, \hat{\eta}_t) = \mathcal{U}(t, \hat{\eta}_t) \left[(\rho - u(w(\hat{\eta}_t, a_t^*), a_t^*)) dt - \frac{\theta\sigma^{-1}}{h_t} dZ_t \right].$$

Hence the volatility of the continuation value as defined in (9) is equal to $\gamma_t = -\mathcal{U}(t, \hat{\eta}_t)\theta\sigma^{-2}/h_t$, and the value of private information reads⁵⁷

$$\begin{aligned} \frac{\sigma^{-2}}{h_t} p_t &= -\sigma^{-2} E_t \left[\int_t^\infty e^{-\rho(s-t)} \frac{\gamma_s}{h_s} ds \right] = E_t \left[\int_t^\infty e^{-\rho(s-t)} \mathcal{U}(s, \hat{\eta}_s) \theta \left(\frac{\sigma^{-2}}{h_s} \right)^2 ds \right] \\ &= -\theta \exp(-\theta \hat{\eta}_t) \\ &\quad \times \int_t^\infty e^{-\rho(s-t)} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}) - \theta(1-\lambda)a_s^*\right) \left[\int_t^s \left(\frac{\sigma^{-2}}{h_\tau} \right)^2 d\tau \right] ds \\ &= -\theta \exp(-\theta \hat{\eta}_t) \\ &\quad \times \int_t^\infty e^{-\rho(s-t)} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}) - \theta(1-\lambda)a_s^*\right) \left(\frac{\sigma^{-2}}{h_t} - \frac{\sigma^{-2}}{h_s} \right) ds. \end{aligned}$$

Reinserting this expression into the necessary condition (10) for incentive compatibility yields

$$[\theta \exp(-\theta \hat{\eta}_t) R_t + u_a(w(\hat{\eta}_t, a_t^*), a_t^*)](a - a_t^*) \leq 0 \quad \text{for all } a \in [0, 1].$$

The expression of $u_a(w(\hat{\eta}_t, a_t^*), a_t^*)$ allows us to factor out $\exp(-\theta \hat{\eta}_t)$ and thus to obtain (46). The deterministic nature of effort follows because R_t is independent of the equilibrium belief $\hat{\eta}_t$.

Step 2. We now prove that there exists a precision $\bar{h}$ such that $a_t^* = 0$ if $h_t \geq \bar{h}$. Let R_t^0 be defined as

$$R_t^0 \triangleq \int_t^\infty e^{-\rho(s-t)} \frac{\sigma^{-2}}{h_s} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1})\right) ds,$$

so that $R_t = R_t^0$ if $a_s^* = 0$ for all $s > t$ or

$$\frac{\partial \mathcal{U}(t, \hat{\eta}_t)}{\partial a_t} = -\theta \exp(-\theta \hat{\eta}_t) [\lambda \exp(-\theta(1-\lambda)a_t^*) - R_t^0] \quad \text{if } a_s^* = 0 \text{ for all } s > t.$$

We wish to establish that R_t^0 is a decreasing function of time. Differentiating its expression with respect to t yields

$$\frac{dR_t^0}{dt} = R_t^0 \left[\rho - \frac{1}{2} \left(\frac{\theta \sigma^{-1}}{h_t} \right)^2 \right] - \frac{\sigma^{-2}}{h_t}. \quad (49)$$

When $\rho < \frac{1}{2}(\theta \sigma^{-1}/h_t)^2$, the derivative is obviously negative. To show that this is also true when $\rho > \frac{1}{2}(\theta \sigma^{-1}/h_t)^2$, we observe that

$$R_t^0 < \frac{\sigma^{-2}}{h_t} \int_t^\infty e^{-\rho(s-t)} \exp\left(\frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1})\right) ds = \frac{\sigma^{-2}}{h_t} \left[\rho - \frac{1}{2} \left(\frac{\theta \sigma^{-1}}{h_t} \right)^2 \right]^{-1}$$

whenever $\rho - \frac{1}{2}(\theta \sigma^{-1}/h_t)^2 > 0$. Reinserting this inequality into (49) shows that $dR_t^0/dt < 0$ with $\lim_{t \rightarrow \infty} R_t^0 = 0$. Hence there exists a unique precision $\bar{h}$ where

⁵⁷The third equality is obtained by inserting the expression of $\mathcal{U}(s, \hat{\eta}_s)$ given in (48), using the value of the expectation $E_t[\exp(-\theta \hat{\eta}_s)] = \exp(-\theta \hat{\eta}_t + \frac{1}{2}\theta^2(h_t^{-1} - h_s^{-1}))$, and interchanging the order of integration.

$R_t^0 \leq \lambda \exp(-\theta(1 - \lambda))$ if $h_t \geq \bar{h}$. But then the fact that $R_t \leq R_t^0$ for all possible equilibrium paths and (46) imply, in turn, that $a_t^* = 0$ is the only incentive compatible level of effort for all t such that $h_t \geq \bar{h}$. In other words, no positive level effort can be sustained in equilibrium whenever $h_t \geq \bar{h}$. $\square$

APPENDIX B: ADDITIONAL RESULT

DERIVATION OF (14). We first change variables and define $\tilde{p}_t \triangleq (\sigma^{-2}/h_t)p_t$. Then $\tilde{p}_t = E[-\int_t^T e^{-\rho(s-t)} \gamma_s(\sigma^{-2}/h_s) ds]$, so that differentiating with respect to time leads to

$$\frac{d\tilde{p}_t}{dt} = \rho\tilde{p}_t + \frac{\sigma^{-2}}{h_t} \gamma_t = \rho\tilde{p}_t - \frac{\sigma^{-2}}{h_t} (u_a(w_t, a_t) + \tilde{p}_t),$$

where the second equality follows after substitution of $\gamma_t = -u_a(w_t, a) - \tilde{p}_t$. Integrating this expression, we obtain

$$\tilde{p}_t = E_a \left[\int_t^T e^{[-\rho(s-t) + \int_t^s (\sigma^{-2}/h_\tau) d\tau]} \frac{\sigma^{-2}}{h_s} u_a(w_s, a_s) ds \right].$$

To simplify the integral in the exponent, we observe that

$$\frac{\sigma^{-2}}{h_\tau} = \frac{\sigma^{-2}}{h_0 + \tau\sigma^{-2}} = \frac{d \ln h_t}{d\tau} \implies \exp\left(\int_t^s \frac{\sigma^{-2}}{h_\tau} d\tau\right) = \exp(\ln h_s - \ln h_t) = \frac{h_s}{h_t}.$$

Therefore,

$$\tilde{p}_t = E_a \left[\int_t^T e^{-\rho(s-t)} \left(\frac{h_s}{h_t}\right) \left(\frac{\sigma^{-2}}{h_s}\right) u_a(w_s, a_s) ds \right] = \frac{\sigma^{-2}}{h_t} E_a \left[\int_t^T e^{-\rho(s-t)} u_a(w_s, a_s) ds \right],$$

which, given the definition of $\tilde{p}_t$, is equivalent to (14). Observe, however, that when $a_t = 0$ for some t , then (12) is not representable as (14).

APPENDIX C: TWO-PERIOD EXAMPLE

This section focuses on two-period contracts. Shortening the horizon enables us to clarify how parameter uncertainty affects incentive provision. Given that the two-period model is mostly illustrative, we set it up in the most parsimonious fashion and restrict our attention to contracts that elicit full or maximal effort in both periods.⁵⁸ We use the model to establish the following three points. First, quality uncertainty modifies the nature of the contracting problem solely when the relationship is repeated over time.

⁵⁸Many of the assumptions below can easily be generalized without modifying the model's main message. For example, our findings are robust to the introduction of a general distribution for output such that $y = 1$ with probability $g(a, \theta)$. Similarly, the cost function for effort does not have to be linear, but could, instead, be convex. This may lead to optimal effort being interior so that one would have to solve the model using compensating variations. This significantly complicates the analysis, so we restrict our attention to linear cost functions and full effort along the equilibrium path.

Second, the strategy space is increasing in the number of periods because the principal has to discourage multiple deviations. Third, the more uncertain is the environment, the more costly it is to motivate the agent; see [Claim 2](#) below. Finally, the two-period model features a multiplicative form between η and a , showing that our finding that quality uncertainty harms incentives under commitment does not depend on the additive form $\eta + a$ used in the body of the paper

Setup. Both parties have the ability to commit. The principal is risk-neutral. The agent cannot borrow or lend and so consumes his wage in each period. He is risk-averse with utility function $u(w) - Ca$, where w denotes wages and $a \in [0, 1]$ is effort. Output y is observed by both parties at the end of each period. Wages are paid at the beginning of the next period and can be contingent on realizations. Output is either high or low with the probabilities

$$y = \begin{cases} 1 & \text{with probability } a\eta \\ 0 & \text{with probability } 1 - a\eta. \end{cases} \quad (50)$$

The parameter $\eta \in [0, 1]$ is the productivity of the relationship, which is unknown to both parties. This is the counterpart of η in the body of the paper.

Notation on beliefs. Let $\mu_1(\eta)$ denote the common prior over η . Output is observed by both parties. Let the one-step-ahead Bayes map conditional on $\{a, y\}$ be $\mu_2(\eta) = b^a(y, \mu_1)$, so that

$$b^a(y, \mu)(\eta) = \frac{(a\eta)^y(1 - a\eta)^{1-y}\mu(\eta)}{\int_{\Theta} (x\eta)^y(1 - x\eta)^{1-y} d\mu(x)} = \begin{cases} (\frac{1-a\eta}{1-a\eta_\mu})\mu(\eta) & \text{when } y = 0 \\ (\frac{\eta}{\eta_\mu})\mu(\eta) & \text{when } y = 1, \end{cases} \quad (51)$$

where $\eta_\mu \triangleq \int_{\Theta} \eta d\mu(\eta)$ is the expected value of η . Observe that b is independent of a when $y = 1$ because a raises the log likelihood of observing $y = 1$ equally for any value of η . This feature is specific to the log-linear form of the probability distribution $(a\eta)^y(1 - a\eta)^{1-y}$. Finally we let $\pi(y|a, \mu)$ denote the probability of y given a and μ , i.e., $\pi(1|a, \mu) \triangleq \int_{\Theta} a\eta d\mu(\eta) = a\eta_\mu$ and $\pi(0|a, \mu) = 1 - \pi(1|a, \mu)$.

One-period contract. We first show that quality uncertainty does not fundamentally modify the problem when the relationship is not repeated. In one-period contracts, the agent's incentive constraint is satisfied when

$$\beta \sum_{y \in \{0,1\}} \pi(y|1, \mu_1) u(w(y)) - C \geq \beta \sum_{y \in \{0,1\}} \pi(y|a, \mu_1) u(w(y)) - Ca, \quad \forall a \in [0, 1],$$

where $\beta \in (0, 1)$ is the agent's discount rate.⁵⁹ It is easily verified that, due to the linearity of the cost function, the incentive constraint holds everywhere if it holds comparing $a = 1$ to $a = 0$.⁶⁰ The participation constraint, for an exogenously given reservation

⁵⁹Remember that wages are paid at the beginning of the following period.

⁶⁰More precisely, $a = 1$ is incentive compatible whenever $\eta_{\mu_1}[U(w(1)) - U(w(0))] \geq C/\beta$.

value $\underline{V}$, reads

$$\beta \sum_{y \in \{0,1\}} \pi(y|1, \mu) u(w(y)) \geq \underline{V} + C.$$

Solving the Lagrangian with λ_{IC} attached to the incentive constraint and λ_{PC} attached to the participation constraint yields

$$\frac{1}{u'(w(y))} = \lambda_{PC} + \lambda_{IC} \left[1 - \frac{\pi(y|0, \mu_1)}{\pi(y|1, \mu_1)} \right],$$

which is equivalent to the standard first-order condition in problems without parameter uncertainty. One would simply replace (50) by

$$y = \begin{cases} 1 & \text{with probability } a\eta_\mu \\ 0 & \text{with probability } 1 - a\eta_\mu. \end{cases}$$

In other words, parameter uncertainty can be bundled with effort uncertainty when designing optimal contracts in single period settings.

Two-period contract and multiple deviations. We use a recursive formulation. The agent's value function in the second and the first periods are

$$\begin{aligned} V^2(\mu_2) &\triangleq \max_{a_2 \in [0,1]} \left\{ -Ca_2 + \beta \sum_{y_2 \in \{0,1\}} \pi(y_2|a_2, \mu_2) u(w^2(y_1, y_2)) \right\} \\ V^1(\mu_1) &\triangleq \max_{a_1 \in [0,1]} \left\{ -Ca_1 + \beta \sum_{y_1 \in \{0,1\}} \pi(y_1|a_1, \mu_1) [u(w^1(y_1)) + V^2(b^{a_1}(y_1, \mu_1))] \right\}. \end{aligned}$$

The second-period value function depends on past effort through the Bayesian operator $b^{a_1}(y_1, \mu_1)$: Conditional on the same output y_1 , agents with diverging actions entertain different beliefs about η . As a result, they may find it optimal to provide different levels of effort in the second period. To deal with such multiple deviations, [Fernandes and Phelan \(2000\)](#) propose attaching an additional threat-keeping constraint to each off-the-equilibrium path. In our problem, however, effort takes value over a continuum and so their methodology cannot be applied. We show, instead, that multiple deviations are not a concern because the relationship ends after the second period. To see this, notice that $a_2 = 1$ is incentive compatible when

$$\eta_{b^{a_1}(y_1, \mu_1)} [u(w^2(1, y_1)) - u(w^2(0, y_1))] \geq C/\beta. \quad (52)$$

It is easily verified from (51) that $\eta_{b^{a_1}(y_1, \mu_1)} \geq \eta_{b^1(y_0, \mu_1)}$ for all $a_1 \in [0, 1]$, $y_1 \in \{0, 1\}$, and initial prior μ_1 . It follows that the incentive constraint (52) holds for any $a_1 \in [0, 1]$ as long as it holds when $a_1 = 1$. It is, therefore, sufficient to exclude one-shot deviations so as to establish the incentive compatibility of the equilibrium path with full effort in both periods.

In other words, multiple deviations are not a concern because shirkers are more optimistic about their probability of success: They have provided less effort while observing the same output realization y_1 . Given that the relationship lasts until the end of the

second period, they face similar costs but higher returns and so are also inclined to provide full effort. This result is specific to our setup and, in particular, to the two-period horizon. If further periods were added to the contract, workers might find it attractive to deviate several times *before* reaching termination.

Incentive constraint in the first period. When the contract implements full effort, i.e., $a_1 = 1$, the value function is equivalent to

$$V^1(\mu_1) = -C + \beta[\gamma(\mu_1) + u(w^1(0)) + V^2(b^1(0, \mu_1))],$$

where

$$\gamma(\mu) \triangleq \eta_\mu[u(w^1(1)) + V^2(b^1(1, \mu)) - u(w^1(0)) - V^2(b^1(0, \mu))].$$

Thus $\gamma(\mu_0)$ measures the volatility in expected utilities as of date 0. The following claim shows that the minimum degree of risk exposure that implements full effort is higher when the prior is more dispersed.

CLAIM 2. *The incentive constraint for full effort in the first period $a_1 = 1$ is satisfied when*

$$\gamma(\mu_1) \geq \frac{C}{\beta} \left[1 + \beta^{-1} \left(\frac{\text{Var}_{\mu_1}[\eta]}{\eta_{\mu_1} - E_{\mu_1}[\eta^2]} \right) \right]. \quad (53)$$

PROOF. The set of incentive constraints reads

$$\begin{aligned} & \eta_{\mu_1}[u(w^1(1)) + V^2(b^1(1, \mu_1))] + (1 - \eta_{\mu_1})[u(w^1(0)) + V^2(b^1(0, \mu_1))] - \frac{C}{\beta} \\ & \geq a\eta_{\mu_1}[u(w^1(1)) + V^2(b^a(1, \mu_1))] + (1 - a\eta_{\mu_1})[u(w^1(0)) + V^2(b^a(0, \mu_1))] - \frac{Ca}{\beta} \end{aligned}$$

for all $a \in [0, 1]$. It can be simplified by noticing that (i) $b^a(1, \mu_1)$ is independent of a and (ii) any given IC is equivalent to

$$(1 - a)\gamma(\mu_1) - (1 - a\eta_{\mu_1})(V^2(b^a(0, \mu_1)) - V^2(b^1(0, \mu_1))) \geq C(1 - a)/\beta. \quad (54)$$

We have shown before that all agents find it optimal to provide full effort in the second period so that

$$V^2(b^a(0, \mu_1)) = -C + \beta^{-1}[u(w^2(0, 0)) + \eta_{b^a(0, \mu_1)}(u(w^2(1, 0)) - u(w^2(0, 0)))].$$

Using the second-period incentive constraint (IC)

$$\eta_{b^1(0, \mu_1)}[u(w^2(1, 0)) - u(w^2(0, 0))] \geq C/\beta$$

and the expression of $\eta_{b^a(0, \mu_1)}$, we obtain

$$\begin{aligned} \beta[V^2(b^a(0, \mu_1)) - V^2(b^1(0, \mu_1))] &= (\eta_{b^a(0, \mu_1)} - \eta_{b^1(0, \mu_1)})[u(w^2(1, 0)) - u(w^2(0, 0))] \\ &= (1 - a) \left(\frac{\text{Var}_{\mu_1}[\eta]}{1 - a\eta_{\mu_1}} \right) \frac{C}{\beta(\eta_{\mu_1} - E_{\mu_1}[\eta^2])}. \end{aligned}$$

Replacing this solution into (54) and simplifying yields

$$\gamma(\mu_1) \geq \beta^{-1} \left[C + \left(\frac{\text{Var}_{\mu_1}[\eta]}{\eta_{\mu_1} - E_{\mu_1}[\eta^2]} \right) \frac{C}{\beta} \right].$$

Hence, if the IC binds at $a = 1$, it binds at any $a \in (0, 1)$, which establishes condition (53). $\square$

If η is known, the IC reads $\gamma(\mu_1) \geq C/\beta$, which is lower than the expression in (53). Since $\gamma(\mu_1)$ measures the volatility in expected utility, Claim 2 illustrates that parameter uncertainty raises risk exposure. To understand why, compare expected utilities at the beginning of the second period. Under our parametric assumptions, we have seen that following a positive output realization, $y_1 = 1$, posteriors do not depend on effort. On the other hand, following a low realization, $y_1 = 0$, shirkers have higher expectations about η . As shown in the proof of Claim 2, they enjoy a higher expected utility,⁶¹

$$V^2(b^a(0, \mu_1)) - V^2(b^1(0, \mu_1)) = (1 - a) \left(\frac{\text{Var}_{\mu_1}[\eta]}{\eta_{\mu_1} - E_{\mu_1}[\eta^2]} \right) \frac{C}{\beta^2(1 - a\eta_{\mu_1})} > 0.$$

This implies that agents find it more attractive to deviate in the first period because they will face higher continuation values. When designing the contract, the principal takes into account this additional channel and discourages shirking by raising income volatility.

This result can also be interpreted from the agent's standpoint. A shirker's expected income in the second period is higher than that of a complying agent because he is more optimistic about η and so can at least mimic the equilibrium payoffs while putting in less effort. This motivates the manipulation he might undertake: If the principal thinks η is low due to low output in the first period, he attributes high output in the second period to hard work and rewards it more generously. We refer to this mechanism as *belief manipulation*. This channel is, by definition, shut down in the last period of the contract. More insurance can then be provided, making it cheaper to deliver utils and, consequently, optimal to back load payments.

REFERENCES

- Abraham, Arpad and Nicola Pavoni (2008), "Efficient allocations with moral hazard and hidden borrowing and lending: A recursive formulation." *Review of Economic Dynamics*, 11, 781–803. [866, 875]
- Adrian, Tobias and Mark Westerfield (2009), "Disagreement and learning in a dynamic contracting model." *Review of Financial Studies*, 22, 3873–3906. [867]
- Becker, Gary (1964), *Human Capital: A Theoretical and Empirical Analysis, With Special Reference to Education*. University of Chicago Press, Chicago. [892]

⁶¹Observe that the difference collapses to zero when η is known.

- Cisternas, Gonzalo (2012), “Shock persistence, endogenous skills and career concerns.” Working paper, Princeton University. [889, 892]
- Cvitanić, Jakša, Xuhu Wan, and Jianfeng Zhang (2009), “Optimal compensation with hidden action and lump-sum payment in a continuous-time model.” *Applied Mathematics & Optimization*, 59, 99–146. [868, 872, 878, 894, 896]
- DeMarzo, Peter M. and Yuliy Sannikov (2011), “Learning in dynamic incentive contracts.” Unpublished manuscript, Princeton University. [867, 876]
- Dewatripont, Mathias, Ian Jewitt, and Jean Tirole (1999), “The economics of career concerns, part I: Comparing information structures.” *Review of Economic Studies*, 66, 183–198. [869]
- Fernandes, Ana and Christopher Phelan (2000), “A recursive formulation for repeated agency with history dependence.” *Journal of Economic Theory*, 91, 223–247. [866, 910]
- Fujisaki, Masatoshi, Gopinath Kallianpur, and Hiroshi Kunita (1972), “Stochastic differential equation for the non linear filtering problem.” *Osaka Journal of Mathematics*, 9, 19–40. [869, 895]
- Giat, Yahel, Steve Hackman, and Ajay Subramanian (2010), “Investment under uncertainty, heterogeneous beliefs, and agency conflicts.” *Review of Financial Studies*, 23, 1360–1404. [867]
- Gibbons, Robert and Kevin J. Murphy (1992), “Optimal incentive contracts in the presence of career concerns: Theory and evidence.” *Journal of Political Economy*, 100, 468–505. [865]
- He, Zhiguo, Bin Wei, and Jianfeng Yu (2012), “Optimal long-term contracting with learning.” Unpublished manuscript, University of Chicago. [867, 868, 878, 885]
- Holmström, Bengt (1999), “Managerial incentive problems: A dynamic perspective.” *Review of Economic Studies*, 66, 169–182. [865, 868, 869, 888, 892]
- Hölmstrom, Bengt and Paul R. Milgrom (1987), “Aggregation and linearity in the provision of intertemporal incentives.” *Econometrica*, 55, 303–328. [867, 868]
- Hopenhayn, Hugo and Arantxa Jarque (2007), “Moral hazard and persistence.” Working Paper 07-07, FRB Richmond. [867]
- Jacod, Jean and Albert Shiryaev (1987), *Limit Theorems for Stochastic Processes*. Springer-Verlag, Berlin. [880]
- Kallianpur, Gopinath (1980), *Stochastic Filtering Theory*. Springer-Verlag, New York. [870]
- Kapicka, Marek (2006), “Efficient allocations in dynamic private information economies with persistent shocks: A first-order approach.” Unpublished manuscript, University of California, Santa Barbara. [866]
- Kocherlakota, Narayana (2004), “Figuring out the impact of hidden savings on optimal unemployment insurance.” *Review of Economic Dynamics*, 7, 541–554. [875]

- Laffont, Jean-Jacques and Jean Tirole (1988), “The dynamics of incentive contracts.” *Econometrica*, 56, 1153–1175. [866]
- Lustig, Hanno, Chad Syverson, and Stijn Van Nieuwerburgh (2011), “Technological change and the growing inequality in managerial compensation.” *Journal of Financial Economics*, 99, 601–627. [891]
- Mirrlees, James (1974), “Notes on welfare economics, information and uncertainty.” In *Essays on Economic Behavior Under Uncertainty* (Michael Balch, Daniel McFadden, and Shih-Yen Wu, eds.), North-Holland, Amsterdam. [872]
- Pavan, Alessandro, Ilya Segal, and Juuso Toikka (2010), “Dynamic mechanism design: Revenue equivalence, profit maximization, and information disclosure.” Unpublished manuscript, Stanford University. [873]
- Rudanko, Leena (2009), “Labor market dynamics under long-term wage contracting.” *Journal of Monetary Economics*, 56, 170–183. [891]
- Sannikov, Yuliy (2008), “A continuous-time version of the principal–agent problem.” *Review of Economic Studies*, 75, 957–984. [868, 872, 878, 892, 897, 899]
- Schättler, Heinz and Jaeyoung Sung (1993), “The first-order approach to the continuous-time principal–agent problem with exponential utility.” *Journal of Economic Theory*, 61, 331–371. [866, 875, 897]
- Werning, Iván (2001), “Moral hazard with unobserved endowments: A recursive approach.” Unpublished manuscript, University of Chicago. [873]
- Williams, Noah (2008), “On dynamic principal–agent models in continuous time.” Unpublished manuscript, UW Madison. [873, 875, 898]
- Williams, Noah (2011), “Persistent private information.” *Econometrica*, 79, 1233–1275. [868, 877, 885]
- Yong, Jiongmin and Xun Yu Zhou (1999), *Stochastic Controls*. Springer-Verlag, New York. [875]