

Dynamic Mechanisms without Money ^{*}

Yingni Guo[†], Johannes Hörner[‡]

March 2015

Abstract

We analyze the optimal design of dynamic mechanisms in the absence of transfers. The designer uses future allocation decisions to elicit private information. Values evolve according to a two-state Markov chain. We solve for the optimal allocation rule, which permits a simple implementation. Unlike with transfers, efficiency decreases over time, and both immiseration and its polar opposite are possible long-run outcomes. Considering the limiting environment in which time is continuous, we demonstrate that persistence hurts.

Keywords: Mechanism design. Principal-Agent. Token mechanisms.

JEL numbers: C73, D82.

1 Introduction

This paper is concerned with the dynamic allocation of resources when transfers are not allowed and information regarding their optimal use is private information to an individual. The informed agent is strategic rather than truthful.

We are searching for the social choice mechanism that would bring us closest to efficiency. Here, efficiency and implementability are understood to be Bayesian: both the individual and society understand the probabilistic nature of uncertainty and update based on it.

^{*}We thank Alex Frankel, Drew Fudenberg, Juuso Toikka and Alex Wolitzky for very useful conversations.

[†]Department of Economics, Northwestern University, 2001 Sheridan Rd., Evanston, IL 60201, USA, yingni.guo@northwestern.edu.

[‡]Yale University, 30 Hillhouse Ave., New Haven, CT 06520, USA, johannes.horner@yale.edu.

Both the societal decision not to allow money –for economic, physical, legal or ethical reasons– and the sequential nature of the problem are assumed. Temporal constraints apply to the allocation of goods, such as jobs, houses or attention, and it is often difficult to ascertain future demands.

Throughout, we assume that the good to be allocated is perishable.¹ Absent private information, the allocation problem is trivial: the good should be provided if and only if its value exceeds its cost.² However, in the presence of private information, and in the absence of transfers, linking future allocation decisions to current decisions is the only instrument available to society to elicit truthful information. Our goal is to understand this link.

Our main results are a characterization of the optimal mechanism and an intuitive indirect implementation for it. In essence, the agent should be granted an inside option, corresponding to a certain number of units of the good that he is entitled to receive “no questions asked.” This inside option is updated according to his choice: whenever the agent desires the unit, his inside option is reduced by one unit; whenever he forgoes it, his inside option is also revised, although not necessarily upward. Furthermore, we demonstrate that this results in very simple dynamics: an initial phase of random length in which the efficient choice is made during each round, followed by an irreversible shift to one of the two possible outcomes in the game with no communication, namely, the unit is either always supplied or never supplied again. These results contrast with those from static design with multiple units (*e.g.*, Jackson and Sonnenschein, 2007) and from dynamic mechanism design with transfers (*e.g.*, Battaglini, 2005).

Formally, our good can take one of two values during each round. Values are serially correlated over time. The binary assumption is certainly restrictive, but it is known that, even with transfers, the problem becomes intractable beyond binary types (see Battaglini and Lamba, 2014).³ We begin with the i.i.d. case, which suffices to illustrate many of the insights of our analysis, before proving the results in full generality. The cost of providing the good is fixed and known. Hence, it is optimal to assign the good during a given round if and only if the value is high. We cast our problem of solving for the efficient mechanism (given the values, cost and agent’s discount factor) as one faced by a disinterested principal with

¹Many allocation decisions involve goods or services that are perishable, such as how a nurse or a worker should divide time; which patients should receive scarce medical resources (blood or treatments); or which investments and activities should be approved by a firm.

²This is because the supply of the perishable good is taken as given. There is a considerable literature on the optimal ordering policy for perishable goods, beginning with Fries (1975).

³In Section 5.2, we consider the case of a continuum of types that are independently distributed over time.

commitment who determines when to supply the good as a function of the agent’s reports. There are no transfers, certification technology, or signals concerning the agent’s value, even *ex post*.

As mentioned above, we demonstrate that the optimal policy can be implemented through a “budget” mechanism in which the appropriate unit of account is the number of units that the agent is entitled to receive sequentially with “no questions asked.” While the updating process when the agent forgoes a unit is somewhat subtle, it is independent of the principal’s belief concerning the agent’s type. The only role of the prior belief is to specify the initial budget. This budget mechanism is not a token mechanism in the sense that the total (discounted) number of units the agent receives is not fixed. Depending on the sequence of reports, the agent might ultimately receive few or many units.⁴ Eventually, the agent is either granted the unit forever or never again. Hence, immiseration is not ineluctable.

In Section 5.1, we study the continuous time limit over which the flow value for the good changes according to a two-state Markov chain. This allows us to demonstrate that persistence hurts. As the Markov chain becomes more persistent, efficiency decreases, although the agent might actually benefit from this increased persistence.

Allocation problems in the absence of transfers are plentiful, and it is not our purpose to survey them here. We believe that our results can inform practices concerning how to implement algorithms to improve allocations. For example, consider nurses who must decide whether to take alerts that are triggered by the patients seriously. The opportunity cost of their time is significant. Patients, however, appreciate quality time with nurses irrespective of whether their condition necessitates it. This discrepancy produces a challenge with which every hospital must contend: ignore alarms and risk that a patient with a serious condition is not attended to or heed all alarms and overwhelm the nurses. “Alarm fatigue” is a serious problem that health care must confront (see, for instance, Sendelbach, 2012). We suggest the best approach for trading off the risks of neglecting a patient in need and attending to one who simply cries wolf.⁵

Related Literature. Our work is closely related to the bodies of literature on mechanism design with transfers and on “linking incentive constraints.” Sections 4.5 and 3.4 are devoted to these issues and explain why transfers (resp., the dynamic nature of the relationship)

⁴It is not a “bankruptcy” mechanism in the sense of Radner (1986) because the specific ordering of the reports matters.

⁵Clearly, our mechanism is much simpler than existing electronic nursing workload systems. However, none appears to seriously consider strategic agent behavior as a constraint.

matter, and hence, we will brief here.

The obvious benchmark work that considers transfers is Battaglini (2005),⁶ who considers our general model but allows transfers. Another important difference is his focus on revenue maximization, a meaningless objective in the absence of prices. The results of his work are diametrically opposed to ours. In Battaglini, efficiency necessarily improves over time (exact efficiency is eventually obtained with probability 1). Here, efficiency decreases over time, in the sense described above, with an asymptotic outcome that is at best the outcome of the static game. The agent’s utility can increase or decrease depending on the history that is realized: receiving the good forever is clearly the best possible outcome, while never receiving it again is the worst outcome. Krishna, Lopomo and Taylor (2013) provide an analysis of limited liability (though transfers are allowed) in a model closely related to that of Battaglini, suggesting that excluding the possibility of unlimited transfers affects both the optimal contract and dynamics. Note that there is an important exception to the quasi-linearity commonly assumed in the dynamic mechanism design literature, namely, Garrett and Pavan (2015).

“Linking incentive constraints” refers to the notion that as the number of identical decision problems increases, linking them allows the designer to improve on the isolated problem. See Fang and Norman (2006) and Jackson and Sonnenschein (2007) for papers specifically devoted to this idea (although they are much older, see Radner, 1981; Rubinstein and Yaari, 1983). Hortala-Vallve (2010) provides an interesting analysis of the unavoidable inefficiencies that must be incurred away from the limit, and Cohn (2010) demonstrates the suboptimality of the mechanisms that are commonly used, even regarding the rate of convergence. Our focus is on the exactly optimal mechanism for a fixed degree of patience, not on proving the asymptotic optimality of a specific mechanism (numerous mechanisms yield asymptotic optimality). This focus allows us to estimate the rate of convergence. Another important difference from most of these papers is that our problem is truly dynamic in the sense that the agent does not know future values but learns them as they come. Section 3.4 elaborates on this distinction.

The notion that virtual budgets could be used as intertemporal instruments to discipline agents with private information has appeared in several papers in economics. Möbius (2001) might be the first to suggest that tracking the difference in the number of favors granted (with two agents) and using it to decide whether to grant new favors is a simple but powerful way of sustaining cooperation in long-run relationships. See also Athey and Bagwell (2001),

⁶See also Zhang (2012) for an exhaustive analysis of Battaglini’s model as well as Fu and Krishna (2014).

Abdulkadiroğlu and Bagwell (2012) and Kalla (2010). While these token mechanisms are known to be suboptimal (as is clear from our characterization of the optimal mechanism), they have desirable properties nonetheless: properly calibrated, they yield an approximately efficient allocation as the discount factor approaches one. To our knowledge, Hauser and Hopenhayn (2008) come the closest to solving for the optimal mechanism (within the class of PPE). Their numerical analysis allows them to qualify the optimality of simple budget rules (according to which each favor is weighted equally, independent of the history), showing that this rule might be too simple (the efficiency cost can reach 30% of surplus). Remarkably, their analysis suggests that the optimal (Pareto-efficient) strategy shares many common features with the optimal policy that we derive in our one-player world: the incentive constraint always binds, and the efficient policy is followed unless it is inconsistent with promise keeping (meaning when promised utilities are too extreme). Our model can be regarded as a game with one-sided incomplete information in which the production cost of the principal is known to the second player. There are some differences, however. First, our principal has commitment and hence is not tempted to act opportunistically or bound by individual rationality. Second, this principal maximizes efficiency rather than his own payoff. Third, there is a technical difference: our limiting model in continuous time corresponds to the Markovian case in which flow values switch according to a Poisson process. In Hauser and Hopenhayn, the lump-sum value arrives according to a Poisson process, and the process is memoryless. Li, Matouschek and Powell (2015) solve for the perfect public equilibria in a model similar to our i.i.d. benchmark and allow for monitoring (public signals), demonstrating that better monitoring improves performance.

More generally, that allocation rights to other (or future) units can be used as a “currency” to elicit private information has long been recognized and dates to Hylland and Zeckhauser (1979), who first explained the extent to which this can be viewed as a pseudo-market. Casella (2005) develops a similar idea within the context of voting rights. Miralles (2012) solves a two-unit version of our problem with more general value distributions, but his analysis is not dynamic: both values are (privately) known at the outset. A dynamic two-period version of Miralles is analyzed by Abdulkadiroğlu and Loertscher (2007).

All versions considered in this paper would be trivial in the absence of imperfect observation of the values. If the values were perfectly observed, it would be optimal to assign the good if and only if the value is high. Due to private information, it is necessary to distort the allocation: after some histories, the good is provided independent of the report; after others, the good is never provided again. In this sense, the scarcity of goods provision is endogenously determined to elicit information. There is a large body of literature in opera-

tions research considering the case in which this scarcity is considered exogenous – there are only n opportunities to provide the good, and the problem is then when to exercise these opportunities. Important early contributions to this literature include Derman, Lieberman and Ross (1972) and Albright (1977). Their analyses suggest a natural mechanism that can be applied in our environment: the agent receives a certain number of “tokens” and uses them whenever he pleases.

Exactly optimal mechanisms have been computed in related environments. Frankel (2011) considers a variety of related settings. The most similar is his Chapter 2 analysis in which he also derives an optimal mechanism. While he allows for more than two types and actions, he restricts attention to the types that are serially independent over time (our starting point). More importantly, he assumes that the preferences of the agent are independent of the state, which allows for a drastic simplification of the problem. Gershkov and Moldovanu (2010) consider a dynamic allocation problem related to Derman, Lieberman and Ross in which agents possess private information regarding the value of obtaining the good. In their model, agents are myopic and the scarcity of the resource is exogenously assumed. In addition, transfers are allowed. They demonstrate that the optimal policy of Derman, Lieberman and Ross (which is very different from ours) can be implemented via appropriate transfers. Johnson (2013) considers a model that is more general than ours (he permits two agents and more than two types). Unfortunately, he does not provide a solution to his model.

A related literature considers the related problem of optimal stopping in the absence of transfers; see, in particular, Kováč, Krähmer and Tatur (2014). This difference reflects the nature of the good, namely, whether it is perishable or durable. When only one unit is desired and waiting is possible, this represents a stopping problem, as in their paper. With a perishable good, a decision must be made during every round. As a result, incentives (and the optimal contract) have hardly anything in common. In the stopping case, the agent might have an option value to forgo the current unit if the value is low and future prospects are good. This is not the case here – incentives to forgo the unit must be endogenously generated via promises. In the stopping case, there is only one history of outcomes that does not terminate the game. Here, policies differ not only in when the good is first provided but also thereafter.

Finally, while the motivations of the papers differ, the techniques for the i.i.d. benchmark that we use borrow numerous ideas from Thomas and Worrall (1990), as we explain in Section 3, and our intellectual debt is considerable.

Section 2 introduces the model. Section 3 solves the i.i.d. benchmark, introducing most of the ideas of the paper, while Section 4 solves the general model. Section 5 extends the

results to cases of continuous time or continuous types. Section 6 concludes.

2 The Model

Time is discrete and the horizon infinite, indexed by $n = 0, 1, \dots$. There are two parties, a disinterested principal and an agent. During each round, the principal can produce an indivisible unit of a good at a cost $c > 0$. The agent's value (or *type*) during round n , v_n is a random variable that takes value l or h . We assume that $0 < l < c < h$ such that supplying the good is efficient if and only if the value is high, but the agent's value is always positive.

The value follows a Markov chain as follows:

$$\mathbf{P}[v_{n+1} = h \mid v_n = h] = 1 - \rho_h, \quad \mathbf{P}[v_{n+1} = l \mid v_n = l] = 1 - \rho_l,$$

for all $n \geq 0$, where $\rho_l, \rho_h \in [0, 1]$. The (invariant) probability of h is $q := \rho_l / (\rho_h + \rho_l)$. For simplicity, we also assume that the initial value is drawn according to the invariant distribution, that is, $\mathbf{P}[v_0 = h] = q$. The (unconditional) expected value of the good is denoted $\mu := \mathbf{E}[v] = qh + (1 - q)l$. We make no assumptions regarding how μ compares to c .

Let $\kappa := 1 - \rho_h - \rho_l$ be a measure of the persistence of the Markov chain. Throughout, we assume that $\kappa \geq 0$ or, equivalently, $1 - \rho_h \geq \rho_l$; that is, the distribution over tomorrow's type conditional on today's type being h first-order stochastically dominates the distribution conditional on today's type being l .⁷ Two interesting special cases occur when $\kappa = 1$ and $\kappa = 0$. The former corresponds to perfect persistence, while the latter corresponds to independent values.

The agent's value is private information. Specifically, at the beginning of each round, the value is drawn and the agent is informed of it.

Players are impatient and share a common discount factor $\delta \in [0, 1)$.⁸ To exclude trivialities, we assume throughout that $\delta > l/\mu$ and $\delta > 1/2$.

Let $x_n \in \{0, 1\}$ refer to the supply decision during round n ; *e.g.*, $x_n = 1$ means that the good is supplied during round n .

⁷The role of this assumption, which is commonly adopted in the literature, and what occurs in its absence, when values are negatively serially correlated, is discussed at the end of Sections 4.3 and 4.4.

⁸The common discount factor is important. We view our principal as a social planner trading off the agent's utility with the social cost of providing the good as opposed to an actual player. As a social planner internalizing the agent's utility, it is difficult to understand why his discount rate would differ from the agent's.

Our focus is on identifying the (constrained) efficient mechanism defined below. Hence, we assume that the principal internalizes both the cost of supplying the good and the value of providing it to the agent. We solve for the principal's favorite mechanism.

Thus, given an infinite history $\{x_n, v_n\}_{n=0}^{\infty}$, the principal's realized *payoff* is defined as follows:

$$(1 - \delta) \sum_{n=0}^{\infty} \delta^n x_n (v_n - c),$$

where $\delta \in [0, 1)$ is a discount factor. The agent's realized *utility* is defined as follows:⁹

$$(1 - \delta) \sum_{n=0}^{\infty} \delta^n x_n v_n.$$

Throughout, payoff and utility refer to the expectation of these values given the relevant player's information. Note that the utility belongs to the interval $[0, \mu]$.

The (risk-neutral) agent seeks to maximize his utility. We now introduce or emphasize several important assumptions maintained throughout our analysis.

- There are no transfers. This is our point of departure from Battaglini (2005) and most of the previous research on dynamic mechanism design. Note also that our objective is efficiency, not revenue maximization. With transfers, there is a trivial mechanism that achieves efficiency: supply the good if and only if the agent pays a fixed price in the range (l, h) .
- There is no *ex post* signal regarding the realized value of the agent –not even payoffs are observed. Depending on the context, it might be realistic to assume that a (possibly noisy) signal of the value occurs at the end of each round, independent of the supply decision. In some other economic examples, it might make more sense to assume instead that this signal occurs only if the good is supplied (*e.g.*, a firm discovers the productivity of a worker who is hired). Conversely, statistical evidence might only occur from not supplying the good if supplying it averts a risk (a patient calling for care or police calling for backup). See Li, Matouschek and Powell (2014) for such an analysis (with “public shocks”) in a related context. Presumably, the optimal mechanism will differ according to the monitoring structure. Understanding what happens without *any* signal is the natural first step.
- We assume that the principal commits *ex ante* to a (possibly randomized) mechanism. This assumption brings our analysis closer to the literature on dynamic mechanism

⁹Throughout, the term payoff describes the principal's objective and utility describes the agent's.

design and distinguishes it from the literature on chip mechanisms (as well as Li, Matouschek and Powell, 2014), which assumes no commitment on either side and solves for the (perfect public) equilibria of the game.

- The good is perishable. Hence, previous choices affect neither feasible nor desirable future opportunities. If the good were perfectly durable and only one unit were demanded, the problem would be one of stopping, as in Kováč, Krähmer and Tatur (2014).

Due to commitment by the principal, we focus on policies in which the agent truthfully reports his type during every round and the principal commits to a (possibly random) supply decision as a function of this last report as well as of the entire history of reports without a loss of generality.

Formally, a direct mechanism or *policy* is a collection $(\mathbf{x}_n)_{n=0}^\infty$, with $\mathbf{x}_n : \{l, h\}^n \rightarrow [0, 1]$ (with $\{l, h\}^0 := \{\emptyset\}$),¹⁰ mapping a sequence of reports by the agent into a decision to supply the good during a given round. Our definition exploits the fact that, because preferences are time-separable, the policy may be considered independent of past realized supply decisions. A direct mechanism defines a decision problem for the agent who seeks to maximize his utility. A reporting strategy is a collection $(\mathbf{m}_n)_{n=0}^\infty$, where $\mathbf{m}_n : \{l, h\}^n \times \{l, h\} \rightarrow \Delta(\{l, h\})$ maps previous reports and the value during round n into a report for that round.¹¹ The policy is *incentive compatible* if truth-telling (that is, reporting the current value faithfully, independent of past reports) is an optimal reporting strategy.

Our first objective is to solve for the *optimal* (incentive-compatible) policy, that is, for the policy that maximizes the principal's payoff subject to incentive compatibility. The *value* is the resulting payoff. Second, we would like to find a simple indirect implementation of this policy. Finally, we wish to understand the payoff and utility dynamics under this policy.

3 The i.i.d. Benchmark

We begin our investigation with the simplest case in which values are i.i.d. over time; that is, $\kappa = 0$. This is a simple extension of Thomas and Worrall (1990), although the indivisibility caused by the absence of transfers leads to dynamics that differ markedly from theirs. See Section 4 for the analysis in the general case $\kappa \geq 0$.

¹⁰For simplicity, we use the same symbols l, h for the possible agent reports as for the values of the good.

¹¹Without a loss of generality, we assume that this strategy does not depend on past values, given past reports, as the decision problem from round n onward does not depend on these past values.

With independent values, it is well known that attention can be further restricted to policies that can be represented by a tuple of functions $U_l, U_h : [0, \mu] \rightarrow [0, \mu]$, $p_l, p_h : [0, \mu] \rightarrow [0, 1]$ mapping a utility U (interpreted as the continuation utility of the agent) onto a continuation utility $u_l = U_l(U)$, $u_h = U_h(U)$ beginning during the next round as well as the probabilities $p_h(U), p_l(U)$ of supplying the good during this round given the current report of the agent. These functions must be consistent in the sense that, given U , the probabilities of supplying the good and promised continuation utilities yield U as a given utility to the agent. This is “promise keeping.” We stress that U is the *ex ante* utility in a given round; that is, it is computed before the agent’s value is realized. The reader is referred to Spear and Srivastava (1987) and Thomas and Worrall (1990) for details.¹²

Because such a policy is Markovian with respect to the utility U , the principal’s payoff is also a function of U only. Hence, solving for the optimal policy and the (principal’s) value function $W : [0, \mu] \rightarrow \mathbf{R}$ amounts to a Markov decision problem. Given discounting, the optimality equation characterizes both the value and the (set of) optimal policies. For any fixed $U \in [0, \mu]$, the optimality equation states the following:

$$W(U) = \sup_{p_h, p_l, u_h, u_l} \{ (1 - \delta) (qp_h(h - c) + (1 - q)p_l(l - c)) + \delta (qW(u_h) + (1 - q)W(u_l)) \} \quad (\text{OBJ})$$

subject to incentive compatibility and promise keeping, namely,

$$(1 - \delta)p_h h + \delta u_h \geq (1 - \delta)p_l h + \delta u_l, \quad (IC_H)$$

$$(1 - \delta)p_l l + \delta u_l \geq (1 - \delta)p_h l + \delta u_h, \quad (IC_L)$$

$$U = (1 - \delta) (qp_h h + (1 - q)p_l l) + \delta (qu_h + (1 - q)u_l), \quad (PK)$$

$$(p_h, p_l, u_h, u_l) \in [0, 1] \times [0, 1] \times [0, \mu] \times [0, \mu].$$

The incentive compatibility and promise keeping conditions are denoted IC (IC_H, IC_L) and PK . This optimization program is denoted $\mathcal{P}$.

Our first objective is to calculate the value function W as well as the optimal policy. Obviously, the entire map might not be relevant once we account for the specific choice of the initial promise –some promised utilities might simply never arise for any sequence of reports. Hence, we are also interested in solving for the *initial promise* U^* , the maximizer of the value function W .

¹²Note that not every policy can be represented in this fashion, as the principal does not need to treat two histories leading to the same continuation utility identically. However, because they are equivalent from the agent’s viewpoint, the principal’s payoff must be maximized by some policy that does so.

3.1 Complete Information

Consider the benchmark case in which there is complete information: that is, consider $\mathcal{P}$ without the *IC* constraints. As the values are i.i.d., we can assume, without loss of generality, that p_l, p_h are constant over time. Given U , the principal chooses p_h and p_l to maximize

$$qp_h(h - c) + (1 - q)p_l(l - c),$$

subject to $U = qp_hh + (1 - q)p_ll$. It follows easily that

Lemma 1 *Under complete information, the optimal policy is*

$$\begin{cases} p_h = \frac{U}{qh}, & p_l = 0 & \text{if } U \in [0, qh], \\ p_h = 1, & p_l = \frac{U - qh}{(1 - q)l} & \text{if } U \in [qh, \mu]. \end{cases}$$

The value function, denoted $\bar{W}$, is equal to

$$\bar{W}(U) = \begin{cases} (1 - \frac{c}{h}) U & \text{if } U \in [0, qh], \\ (1 - \frac{c}{l}) U + cq (\frac{h}{l} - 1) & \text{if } U \in [qh, \mu]. \end{cases}$$

Hence, the initial promise (maximizing $\bar{W}$) is $U_0 := qh$.

That is, unless $U = qh$, the optimal policy (p_l, p_h) cannot be efficient. To deliver $U < qh$, the principal chooses to scale down the probability with which to supply the good when the value is high, maintaining $p_l = 0$. Similarly, for $U > qh$, the principal is forced to supply the good with positive probability even when the value is low to satisfy promise keeping.

While this policy is the only constant optimal one, there are many other (non-constant) optimal policies. We will encounter some in the sequel.

We call $\bar{W}$ the complete-information payoff function. It is piecewise linear (see Figure 1). Plainly, it is an upper bound to the value function under incomplete information.

3.2 The Optimal Mechanism

We now solve for the optimal policy under incomplete information in the i.i.d. case. We first provide an informal derivation of the solution. It follows from two observations (formally established below). First,

The efficient supply choice $(p_l, p_h) = (0, 1)$ is made “as long as possible.”

To understand this qualification, note that if $U = 0$ (resp., $U = \mu$), promise keeping allows no latitude in the choice of probabilities. The good cannot (or must) be supplied, independent of the report. More generally, if $U \in [0, (1 - \delta)qh)$, it is impossible to supply the good if the value is high while satisfying promise keeping. In this utility range, the observation must be interpreted as indicating that the supply choice is as efficient as possible given the restriction imposed by promise keeping. This implies that a high report leads to a continuation utility of 0, with the probability of the good being supplied adjusting accordingly. An analogous interpretation applies to $U \in (\mu - (1 - \delta)(1 - q)l, \mu]$.

These two peripheral intervals vanish as $\delta \rightarrow 1$ and are ignored for the remainder of this discussion. For every other promised utility, we claim that it is optimal to make the (“static”) efficient supply choice. Intuitively, there is never a better time to redeem part of the promised utility than when the value is high. During such rounds, the interests of the principal and agent are aligned. Conversely, there cannot be a worse opportunity to repay the agent what he is due than when the value is low because tomorrow’s value cannot be lower than today’s.

As trivial as this observation may sound, it already implies that the dynamics of the inefficiencies must be very different from those in Battaglini’s model with transfers. Here, inefficiencies are backloaded.

As the supply decision is efficient as long as possible, the high type agent has no incentive to pretend to be a low type. However,

Incentive compatibility of the low type agent always binds.

Specifically, without a loss of generality, assume that IC_L always binds and disregard IC_H . The reason that the constraint binds is standard: the agent is risk neutral, and the principal’s payoff must be a concave function of U (otherwise, he could offer the agent a lottery that the agent would be willing to accept and that would make the principal better off). Concavity implies that there is no gain in spreading continuation utilities u_l, u_h beyond what is required for IC_L to be satisfied.

Because we are left with two variables (u_l, u_h) and two constraints (IC_L and PK), we can immediately solve for the optimal policy. Algebra is not needed. Because the agent is always willing to state that his value is high, it must be the case that his utility can be computed *as if* he followed this reporting strategy, namely,

$$U = (1 - \delta)\mu + \delta u_h, \text{ or } u_h = \frac{U - (1 - \delta)\mu}{\delta}.$$

Because U is a weighted average of u_h and $\mu \geq U$, it follows that $u_h \leq U$. The promised utility necessarily decreases after a high report. To compute u_l , note that the reason that the high type agent is unwilling to pretend he has a low value is that he receives an incremental value $(1-\delta)(h-l)$ from obtaining the good relative to what would make him merely indifferent between the two reports. Hence, defining $\underline{U} := q(h-l)$, it holds that

$$U = (1-\delta)\underline{U} + \delta u_l, \text{ or } u_l = \frac{U - (1-\delta)\underline{U}}{\delta}.$$

Because U is a weighted average of $\underline{U}$ and u_l , it follows that $u_l \leq U$ if and only if $U \leq \underline{U}$. In that case, even a low report leads to a decrease in the continuation utility, albeit a smaller decrease than if the report had been high and the good provided.

The following theorem (proved in the appendix, as are all other results) summarizes this discussion with the necessary adjustments on the peripheral intervals.

Theorem 1 *The unique optimal policy is*

$$p_l = \max \left\{ 0, 1 - \frac{\mu - U}{(1-\delta)l} \right\}, \quad p_h = \min \left\{ 1, \frac{U}{(1-\delta)\mu} \right\}.$$

Given these values of (p_h, p_l) , continuation utilities are

$$u_h = \frac{U - (1-\delta)p_h\mu}{\delta}, \quad u_l = \frac{U - (1-\delta)(p_l l + (p_h - p_l)\underline{U})}{\delta}.$$

For reasons that will become clear shortly, this policy is not uniquely optimal for $U \leq \underline{U}$. We now turn to a discussion of the utility dynamics and of the shape of the value function, which are closely related. This discussion revolves around the following lemma.

Lemma 2 *The value function $W : [0, \mu] \rightarrow \mathbf{R}$ is continuous and concave on $[0, \mu]$, continuously differentiable on $(0, \mu)$, linear (and equal to $\bar{W}$) on $[0, \underline{U}]$, and strictly concave on $[\underline{U}, \mu]$. Furthermore,*

$$\lim_{U \downarrow 0} W'(U) = 1 - \frac{c}{h}, \quad \lim_{U \uparrow \mu} W'(U) = 1 - \frac{c}{l}.$$

Indeed, consider the following functional equation for W that we obtain from Theorem 1 (ignoring again the peripheral intervals for the sake of the discussion):

$$W(U) = (1-\delta)q(h-c) + \delta q W\left(\frac{U - (1-\delta)\mu}{\delta}\right) + \delta(1-q)W\left(\frac{U - (1-\delta)\underline{U}}{\delta}\right).$$

Hence, taking for granted the differentiability of W stated in the lemma,

$$W'(U) = qW'(U_h) + (1-q)W'(U_l).$$

In probabilistic terms, $W'(U_n) = \mathbf{E}[W'(U_{n+1})]$ given the information at round n . That is, W' is a bounded martingale and must therefore converge.¹³ This martingale was first uncovered by Thomas and Worrall (1990), and hence, we refer to it as the TW-martingale. Because W is strictly concave on $(\underline{U}, \mu)$, yet $u_h \neq u_l$ in this range, it follows that the process $\{U_n\}_{n=0}^\infty$ must eventually exit this interval. Hence, U_n must converge to either $U_\infty = 0$ or μ . However, note that, because $u_h < U$ and $u_l \leq U$ on the interval $(0, \underline{U}]$, this interval is a transient region for the process. Hence, if we began this process in the interval $[0, \underline{U}]$, the limit must be 0 and the TW-martingale implies that W' must be constant on this interval – hence the linearity of W .¹⁴

While $W'_n := W'(U_n)$ is a martingale, U_n is not. Because the optimal policy yields

$$U_n = (1 - \delta)qh + \delta \mathbf{E}[U_{n+1}],$$

utility drifts up or down (stochastically) according to whether $U = U_n$ is above or below qh . Intuitively, if $U > qh$, then the flow utility delivered is insufficient to honor the average promised utility. Hence, the expected continuation utility must be even larger than U .

This raises the question of the initial promise U^* : does it lie above or below qh , and where does the process converge given this initial value? The answer is provided by the TW-martingale. Indeed, U^* is characterized by $W'(U^*) = 0$ (uniquely so, given strict concavity on $[\underline{U}, \mu]$). Hence,

$$0 = W'(U^*) = \mathbf{P}[U_\infty = 0 \mid U_0 = U^*]W'(0) + \mathbf{P}[U_\infty = \mu \mid U_0 = U^*]W'(\mu),$$

where $W'(0)$ and $W'(\mu)$ are the one-sided derivatives given in the lemma. Hence,

$$\frac{\mathbf{P}[U_\infty = 0 \mid U_0 = U^*]}{\mathbf{P}[U_\infty = \mu \mid U_0 = U^*]} = \frac{(c - l)/l}{(h - c)/h}. \quad (1)$$

The initial promise is chosen to yield this ratio of absorption probabilities at 0 and μ . Remarkably, this ratio is independent of the discount factor (despite the discrete nature of the random walk, the step size of which depends on δ !). Hence, both long-run outcomes are possible irrespective of how patient the players are. However, depending on the parameters, U^* can be above or below qh , the first-best initial promise, as is easy to confirm through examples. In the appendix, we demonstrate that U^* is decreasing in the cost, which should

¹³It is bounded because W is concave, and hence, its derivative is bounded by its value at 0 and μ , given in the lemma.

¹⁴This yields multiple optimal policies on this range. As long as the spread is sufficiently large to satisfy IC_L , not so large as to violate IC_H , consistent with PK and contained in $[0, \underline{U}]$, it is an optimal choice.

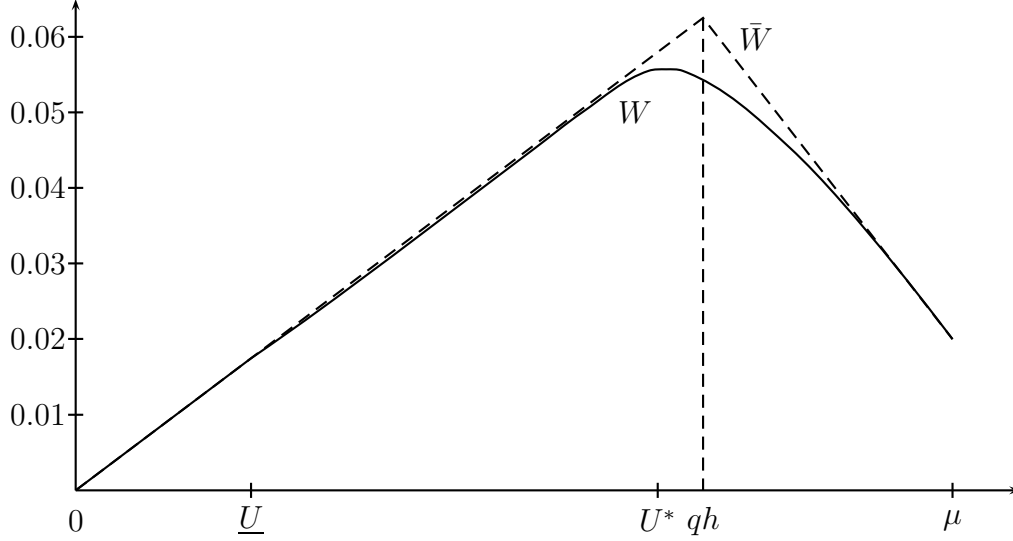


Figure 1: Value function for $(\delta, l, h, q, c) = (.95, .40, .60, .60, .50)$.

be clear, because the random walk $\{U_n\}$ only depends on c via the choice of initial promise U^* given by (1).

We record this discussion in the following lemma.

Lemma 3 *The process $\{U_n\}_{n=0}^\infty$ (with $U_0 = U^*$) converges to 0 or μ , a.s., with probabilities given by (1).*

3.3 Implementation

As mentioned above, the optimal policy is not a token mechanism because the number of units the agent receives is not fixed.¹⁵ However, the policy admits a very simple indirect implementation in terms of a budget that can be described as follows. Let $f := (1 - \delta)\underline{U}$, and $g := (1 - \delta)\mu - f = (1 - \delta)l$.

Provide the agent with an initial budget of U^* . At the beginning of every round, charge him a fixed fee equal to f . If the agent asks for the item, supply it and charge a variable fee g for it. Increase his budget by the interest rate $\frac{1}{\delta} - 1$ each round – provided that this is feasible.

¹⁵To be clear, this is not an artifact of discounting: the optimal policy in the finite-horizon undiscounted version of our model can be derived along the same lines (using the binding IC_L and PK constraints), and the number of units obtained by the agent is also history-dependent in that case.

This scheme might become infeasible for two reasons. First, his budget might no longer allow him to pay g for a requested unit. Then, award him whatever fraction his budget can purchase (at unit price g). Second, his budget might be so close to μ that it is no longer possible to pay him the interest rate on his budget. Then, return the excess to him, independent of his report, at a conversion rate that is also given by the price g .

For budgets below $\underline{U}$, the agent is “in the red,” and even if he does not buy a unit, his budget shrinks over time. If his budget is above $\underline{U}$, he is “in the black,” and forgoing a unit increases the budget. When doing so pushes the budget above $\mu - (1 - \delta)(1 - q)l$, the agent “breaks the bank” and reaches μ in case of another forgoing, which is an absorbing state.

This structure is somewhat reminiscent of results in research on optimal financial contracting (see, for instance, Biais, Mariotti, Plantin and Rochet, 2007), a literature that assumes transfers.¹⁶ In this literature, one obtains (for some parameters) an upper absorbing boundary (at which the agent receives the first-best outcome) and a lower absorbing boundary (at which the project is terminated). There are several important differences, however. Most importantly, the agent is not paid in the intermediate region: promises are the only source of incentives. In our environment, the agent receives the good if his value is high, achieving efficiency in this intermediate region.

3.4 A Comparison with Token Mechanisms as in Jackson and Sonnenschein (2007)

A discussion of the relationship of our results to those in environments with transfers is relegated to Section 4.5 because the environment considered in Section 4 is the counterpart to Battaglini (2005). However, because token mechanisms are typically introduced in i.i.d. environments, we make a few observations concerning the connection between our results and those of Jackson and Sonnenschein (2007) here to explain why our dynamic analysis is substantially different from the static one with many copies.

There are two conceptually distinct issues. First, are token mechanisms optimal? Second, is the problem static or dynamic? For the purpose of asymptotic analysis (when either the discount factor tends to 1 or the number of equally weighted copies $T < \infty$ tends to infinity), the distinctions are blurred: token mechanisms are optimal in the limit, whether the problem is static or dynamic. Because the emphasis in Jackson and Sonnenschein is on asymptotic

¹⁶There are other important differences in the set-up. They allow two instruments: downsizing the firm and payments. Additionally, this is a moral hazard-type problem because the agent can divert resources from a risky project, reducing the likelihood that it succeeds during a given period.

analysis, they focus on a static model and on a token mechanism and derive a rate of convergence for this mechanism (namely, the loss relative to the first-best outcome is of the order $\mathcal{O}(1/\sqrt{T})$), and they discuss the extension of their results to the dynamic case. We may then cast the comparison in terms of the agent's knowledge. In Jackson and Sonnenschein, the agent is a prophet (in the sense of stochastic processes, he knows the entire realization of the process from the beginning), whereas in our environment, the agent is a forecaster (the process of his reports must be predictable with respect to the realized values up to the current date).

Not only are token mechanisms asymptotically optimal regardless of whether the agent is a prophet or a forecaster, but also, the agent's information plays no role if we restrict attention to token mechanisms in a binary-type environment absent discounting. With binary values and a fixed number of units, it makes no difference whether one knows the realized sequence in advance. Forgoing low-value items as long as the budget does not allow all remaining units to be claimed is not costly, as subsequent units cannot be worth even less. Similarly, accepting high-value items cannot be a mistake.

However, the optimal mechanism in our environment is not a token mechanism. A report not only affects whether the agent obtains the current unit but also affects the total number he obtains.¹⁷ Furthermore, the optimal mechanism when the agent is a prophet is not a token one (even in the finite undiscounted horizon case). The optimal mechanism does not simply ask the agent to select a fixed number of copies that he would like but offers him a menu that trades off the risk in obtaining the units he claims are low or high and the expected number that he receives.¹⁸ The agent's private information pertains not only to whether a given unit has a high value but also to how many units are high. Token mechanisms do not elicit any information in this regard. Because the prophet has more information than the forecaster, the optimal mechanisms are distinct.

The question of how the two mechanisms compare (in terms of average efficiency loss) is ambiguous *a priori*. Because the prophet has more information regarding the number of high-value items, the mechanism must satisfy more incentive-compatibility constraints (which harms welfare) but might induce a better fit between the number of units he actually receives and the number he should receive. Indeed, it is not difficult to construct examples (say, for $T = 3$) in which the comparison could go either way according to the parameters. However, asymptotically, the comparison is clear, as the next lemma states.

¹⁷To be clear, token mechanisms are not optimal even without discounting.

¹⁸The characterization of the optimal mechanism in the case of a prophetic agent is somewhat peripheral to our analysis and is thus omitted.

Lemma 4 *It holds that*

$$|W(U^*) - q(h - c)| = \mathcal{O}(1 - \delta).$$

In the case of a prophetic agent, the average loss converges to zero at rate $\mathcal{O}(\sqrt{1 - \delta})$.

With a prophet, the rate is no better than with token mechanisms. Token mechanisms achieve rate $\mathcal{O}(\sqrt{1 - \delta})$ precisely because they do not attempt to elicit the number of high units. By the central limit theorem, this implies that a token mechanism is incorrect by an order of $\mathcal{O}(\sqrt{1 - \delta})$. The lemma indicates that the cost of incentive compatibility is sufficiently strong that the optimal mechanism performs little better, eliminating only a fraction of this inefficiency.¹⁹ The forecaster’s relative lack of information serves the principal. Because the former knows values only one round in advance, he gives the information away for free until absorption. His private information regarding the number of high units being of the order $(1 - \delta)$, the overall inefficiency is of the same order. Both rates are tight (see the proof of Lemma 4): indeed, were the agent to hold private information for the initial period only, there would already be an inefficiency of the order $1 - \delta$, and so welfare cannot converge faster than at that rate.

Hence, when interacting with a forecaster rather than a prophet, there is a real loss in using a token mechanism instead of the budget mechanism described above.

4 The General Markov Model

We now return to the general model in which types are persistent rather than independent.

As an initial exercise, consider the case of perfect persistence $\rho_h = \rho_l = 0$. If types never change, there is simply no possibility for the principal to use future allocations as an instrument to elicit truth-telling. We revert to the static problem for which the solution is to either always provide the good (if $\mu \geq c$) or never do so.

This case suggests that persistence plays an ambiguous role *a priori*. Because current types assign different probabilities of being (for example) high types tomorrow, one might hope that tying promised future utility to current reports might facilitate truth-telling. However, the case of perfectly persistent types suggests that correlation diminishes the scope for using future allocations as “transfers.” Utilities might still be separable over time, but the

¹⁹This result might be surprising given Cohn’s (2010) “improvement” upon Jackson and Sonnenschein. However, while Jackson and Sonnenschein cover our set-up, Cohn does not and features more instruments at the principal’s disposal. See also Eilat and Pauzner (2011) for an optimal mechanism in a related setting.

current type affects both flow and continuation utility. Definite comparative statics are obtained in the continuous time limit, see Section 5.1.

The techniques that served us well with independent values are no longer useful. We will not be able to rely on martingale techniques. Worse, *ex ante* utility is no longer a valid state variable. To understand why, note that with independent types, an agent of a given type can evaluate his continuation utility based only on current type, probabilities of trade as a function of his report, and promised utility tomorrow as a function of his report. However, if today's type is correlated with tomorrow's type, how can the agent evaluate his continuation utility without knowing how the principal intends to implement it? This is problematic because the agent can deviate, unbeknown to the principal, in which case the continuation utility computed by the principal, given his incorrect belief regarding the agent's type tomorrow, is not the same as the continuation utility under the agent's belief.

However, conditional on the agent's type tomorrow, his type today carries no information on future types by the Markovian assumption. Hence, tomorrow's promised *ex interim* utilities suffice for the agent to compute his utility today regardless of whether he deviates; that is, we must specify his promised utility tomorrow conditional on each possible report at that time. Of course, his type tomorrow is not observable. Hence, we must use the utility he receives from reporting his type tomorrow, conditional on truthful reporting. This creates no difficulty, as on path, the agent has an incentive to truthfully report his type tomorrow. Hence, he does so after having lied during the previous round (conditional on his current type and his previous report, his previous type does not affect the decision problem he faces). That is, the one-shot deviation principle holds here: when a player considers lying, there is no loss in assuming that he will report truthfully tomorrow. Hence, the promised utility pair that we use corresponds to his actual possible continuation utilities if he plays optimally in the continuation regardless of whether he reports truthfully today. We are obviously not the first to highlight the necessity of using as a state variable the vector of *ex interim* utilities, given a report today, as opposed to the *ex ante* utility when types are serially correlated. See Townsend (1982), Fernandes and Phelan (2000), Cole and Kocherlakota (2001), Doepke and Townsend (2006) and Zhang and Zenios (2008). Hence, to use dynamic programming, we must include as state variables the pair of utilities that must be delivered today as a function of the report. Nevertheless, this is insufficient to evaluate the payoff to the principal. Given such a pair of utilities, we must also specify his belief regarding the agent's type. Let ϕ denote the probability that he assigns the high type. This belief can take only three values depending on whether this is the initial round or whether the previous report was high or low. Nonetheless, we treat ϕ as an arbitrary element in the unit interval.

Another complication arises from the fact that the principal's belief depends on the history. For this belief, the last report is sufficient.

4.1 The Program

As discussed above, the principal's optimization program, cast as a dynamic programming problem, requires three state variables: the belief of the principal, $\phi = \mathbf{P}[v = h] \in [0, 1]$, and the pair of (*ex interim*) utilities that the principal delivers as a function of the current report, U_h, U_l . The highest utility μ_h (resp., μ_l) that can be given to a player whose type is high (or low) delivered by always supplying the good solves

$$\mu_h = (1 - \delta)h + \delta(1 - \rho_h)\mu_h + \delta\rho_h\mu_l, \quad \mu_l = (1 - \delta)l + \delta(1 - \rho_l)\mu_l + \delta\rho_l\mu_h;$$

that is,

$$\mu_h = h - \frac{\delta\rho_h(h - l)}{1 - \delta + \delta(\rho_h + \rho_l)}, \quad \mu_l = l + \frac{\delta\rho_l(h - l)}{1 - \delta + \delta(\rho_h + \rho_l)}.$$

We note that

$$\mu_h - \mu_l = \frac{1 - \delta}{1 - \delta + \delta(\rho_h + \rho_l)}(h - l).$$

The gap between the maximum utilities as a function of the type decreases in δ , vanishing as $\delta \rightarrow 1$.

A policy is now a pair $(p_h, p_l) : \mathbf{R}^2 \rightarrow [0, 1]^2$ mapping the current utility vector $U = (U_h, U_l)$ onto the probability with which the good is supplied as a function of the report and a pair $(U(h), U(l)) : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ mapping U onto the promised utilities $(U_h(h), U_l(h))$ if the report is h and $(U_h(l), U_l(l))$ if it is l . These definitions abuse notation, as the domain of $(U(h), U(l))$ should be those utility vectors that are feasible and incentive-compatible.

Define the function $W : [0, \mu_h] \times [0, \mu_l] \times [0, 1] \rightarrow \mathbf{R} \cup \{-\infty\}$ that solves the following program for all $U \in [0, \mu_h] \times [0, \mu_l]$, and $\phi \in [0, 1]$:

$$\begin{aligned} W(U, \phi) &= \sup \{ \phi((1 - \delta)p_h(h - c) + \delta W(U(h), 1 - \rho_h)) \\ &\quad + (1 - \phi)((1 - \delta)p_l(l - c) + \delta W(U(l), \rho_l)) \}, \end{aligned}$$

over $p_l, p_h \in [0, 1]$, and $U(h), U(l) \in [0, \mu_h] \times [0, \mu_l]$ subject to promise keeping and incentive compatibility, namely,

$$U_h = (1 - \delta)p_h h + \delta(1 - \rho_h)U_h(h) + \delta\rho_h U_l(h) \tag{2}$$

$$\geq (1 - \delta)p_l h + \delta(1 - \rho_h)U_h(l) + \delta\rho_h U_l(l), \tag{3}$$

and

$$U_l = (1 - \delta)p_l l + \delta(1 - \rho_l)U_l(l) + \delta\rho_l U_h(l) \quad (4)$$

$$\geq (1 - \delta)p_h l + \delta(1 - \rho_l)U_l(h) + \delta\rho_l U_h(h), \quad (5)$$

with the convention that $\sup W = -\infty$ whenever the feasible set is empty. Note that W is concave on its domain (by the linearity of the constraints in the promised utilities). An optimal policy is a map from (U, ϕ) into $(p_h, p_l, U(h), U(l))$ that achieves the supremum for some W .

4.2 Complete Information

Proceeding as with independent values, we briefly derive the solution under complete information, that is, dropping (3) and (5). Write $\bar{W}$ for the resulting value function. Ignoring promises, the efficient policy is to supply the good if and only if the type is h . Let v_h^* (resp., v_l^*) denote the utility that a high (low) type obtains under this policy. The pair (v_h^*, v_l^*) satisfies

$$v_h^* = (1 - \delta)h + \delta(1 - \rho_h)v_h^* + \delta\rho_h v_l^*, \quad v_l^* = \delta(1 - \rho_l)v_l^* + \delta\rho_l v_h^*,$$

which yields

$$v_h^* = \frac{h(1 - \delta(1 - \rho_l))}{1 - \delta(1 - \rho_h - \rho_l)}, \quad v_l^* = \frac{\delta h \rho_l}{1 - \delta(1 - \rho_h - \rho_l)}.$$

When a high type's promised utility U_h is in $[0, v_h^*]$, the principal supplies the good only if the type is high. Therefore, the payoff is $U_h(1 - c/h)$. When $U_h \in (v_h^*, \mu_h]$, the principal always supplies the good if the type is high. To fulfill the promised utility, the principal also produces the good when the agent's type is low. The payoff is $v_h^*(1 - c/h) + (U_h - v_h^*)(1 - c/l)$. We proceed analogously given U_l (notice that the problems of delivering U_h and U_l are uncoupled). In summary, $\bar{W}(U, \phi)$ is given by

$$\begin{cases} \phi \frac{U_h(h-c)}{h} + (1 - \phi) \frac{U_l(h-c)}{h} & \text{if } U \in [0, v_h^*] \times [0, v_l^*], \\ \phi \frac{U_h(h-c)}{h} + (1 - \phi) \left(\frac{v_l^*(h-c)}{h} + \frac{(U_l - v_l^*)(l-c)}{l} \right) & \text{if } U \in [0, v_h^*] \times [v_l^*, \mu_l], \\ \phi \left(\frac{v_h^*(h-c)}{h} + \frac{(U_h - v_h^*)(l-c)}{l} \right) + (1 - \phi) \frac{U_l(h-c)}{h} & \text{if } U \in [v_h^*, \mu_h] \times [0, v_l^*], \\ \phi \left(\frac{v_h^*(h-c)}{h} + \frac{(U_h - v_h^*)(l-c)}{l} \right) + (1 - \phi) \left(\frac{v_l^*(h-c)}{h} + \frac{(U_l - v_l^*)(l-c)}{l} \right) & \text{if } U \in [v_h^*, \mu_h] \times [v_l^*, \mu_l]. \end{cases}$$

For future purposes, note that the derivative of W (differentiable except at $U_h = v_h^*$ and $U_l = v_l^*$) is in the interval $[1 - c/l, 1 - c/h]$, as expected. The latter corresponds to the most efficient utility allocation, whereas the former corresponds to the most inefficient allocation. In fact, W is piecewise linear (a “tilted pyramid”) with a global maximum at $v^* = (v_h^*, v_l^*)$.

4.3 Feasible and Incentive-Feasible Payoffs

One difficulty in using *ex interim* utilities as state variables rather than *ex ante* utility is that the dimensionality of the problem increases with the cardinality of the type set. Another related difficulty is that it is not obvious which vectors of utilities are feasible given the incentive constraints. Promising to assign all future units to the agent in the event that his current report is high while assigning none if this report is low is simply not incentive compatible.

The set of *feasible* utility pairs (that is, the largest bounded set of vectors U such that (2) and (4) can be satisfied with continuation vectors in the set itself) is easy to describe. Because the two promise keeping equations are uncoupled, it is simply the set $[0, \mu_h] \times [0, \mu_l]$ itself (as was already implicit in Section 4.2).

What is challenging is to solve for the incentive-compatible, feasible (in short, incentive-feasible) utility pairs: these are *ex interim* utilities for which we can find probabilities and pairs of promised utilities tomorrow that make it optimal for the agent to report his value truthfully such that these promised utility pairs tomorrow satisfy the same property.

Definition 1 *The incentive-feasible set, $V \in \mathbf{R}^2$, is the set of ex interim utilities in round 0 that are obtained for some incentive-compatible policy.*

It is standard to show that V is the largest bounded set such that for each $U \in V$ there exists $p_h, p_l \in [0, 1]$ and two pairs $U(h), U(l) \in V$ solving (2)–(5).²⁰

Our first step toward solving for the optimal mechanism is to solve for V . To obtain some intuition regarding its structure, let us review some of its elements. Clearly, $0 \in V, \mu := (\mu_h, \mu_l) \in V$. It suffices to never or always supply the unit, independent of the reports, which is incentive compatible.²¹ More generally, for some integer $m \geq 0$, the principal can supply the unit for the first m rounds, independent of the reports, and never supply the unit after. We refer to such policies as pure *frontloaded* policies because they deliver a given number of units as quickly as possible. More generally, a (possibly mixed) frontloaded policy is one that randomizes over two pure frontloaded policies with consecutive integers $m, m + 1$.

²⁰Clearly, incentive-feasibility is closely related to self-generation (see Abreu, Pearce and Stacchetti, 1990), though it pertains to the different types of a single agent rather than to different players in the game. The distinction is not merely a matter of interpretation because a high type can become a low type and vice-versa, which represents a situation with no analogue in repeated games. Nonetheless, the proof of this characterization is identical.

²¹With some abuse, we write $\mu \in \mathbf{R}^2$, as it is the natural extension as an upper bound of the feasible set of $\mu \in \mathbf{R}$.

Similarly, we define a pure backloaded policy as one that does not supply the good for the first m rounds but does afterward, independent of the reports. (Mixed backloaded policies being defined in the obvious way.)

Suppose that we fix a backloaded and a frontloaded policy such that the high-value agent is indifferent between the two. Then, the low-value agent surely prefers the backloaded policy. The backloaded policy affords him “more time” to switch from his (initial) low value to a high value. Hence, given $U_h \in (0, \mu_h)$, the utility U_l obtained by the backloaded policy that awards U_h to the high type is higher than the utility U_l of the frontloaded policy that also yields U_h .

The utility pairs corresponding to backloading and frontloading are easily solved because they obey simple recursions. First, for $\nu \geq 0$, let

$$\bar{u}_h^\nu = \delta^\nu \mu_h - \delta^\nu (1 - q)(\mu_h - \mu_l)(1 - \kappa^\nu), \quad (6)$$

$$\bar{u}_l^\nu = \delta^\nu \mu_l + \delta^\nu q(\mu_h - \mu_l)(1 - \kappa^\nu), \quad (7)$$

and set $\bar{u}^\nu := (\bar{u}_h^\nu, \bar{u}_l^\nu)$. Second, for $\nu \geq 0$, let

$$\underline{u}_h^\nu = (1 - \delta^\nu)\mu_h + \delta^\nu (1 - q)(\mu_h - \mu_l)(1 - \kappa^\nu), \quad (8)$$

$$\underline{u}_l^\nu = (1 - \delta^\nu)\mu_l - \delta^\nu q(\mu_h - \mu_l)(1 - \kappa^\nu), \quad (9)$$

and set $\underline{u}^\nu := (\underline{u}_h^\nu, \underline{u}_l^\nu)$. The sequence $\bar{u}^\nu$ is decreasing (in both its arguments) as ν increases, with $\bar{u}^0 = \mu$ and with $\lim_{\nu \rightarrow \infty} \bar{u}^\nu = 0$. Similarly, $\underline{u}^\nu$ is increasing, with $\underline{u}^0 = 0$ and $\lim_{\nu \rightarrow \infty} \underline{u}^\nu = \mu$.

Backloading not only is better than frontloading for the low-value agent but also fixes the high-value agent’s utility. These policies yield the best and worst utilities. Formally,

Lemma 5 *It holds that*

$$V = \text{co}\{\bar{u}^\nu, \underline{u}^\nu : \nu \geq 0\}.$$

That is, V is a polygon with a countable infinity of vertices (and two accumulation points). See Figure 2 for an illustration. It is easily verified that

$$\lim_{\nu \rightarrow \infty} \frac{\bar{u}_l^{\nu+1} - \bar{u}_l^\nu}{\bar{u}_h^{\nu+1} - \bar{u}_h^\nu} = \lim_{\nu \rightarrow \infty} \frac{\underline{u}_l^{\nu+1} - \underline{u}_l^\nu}{\underline{u}_h^{\nu+1} - \underline{u}_h^\nu} = 1.$$

When the switching time ν is large, the change in utility from increasing this time has an impact on the agent’s utility that is essentially independent of his initial type. Hence, the slopes of the boundaries are less than 1 and approach 1 as $\nu \rightarrow \infty$. Because $(\mu_l - v_l^*)/(\mu_h -$

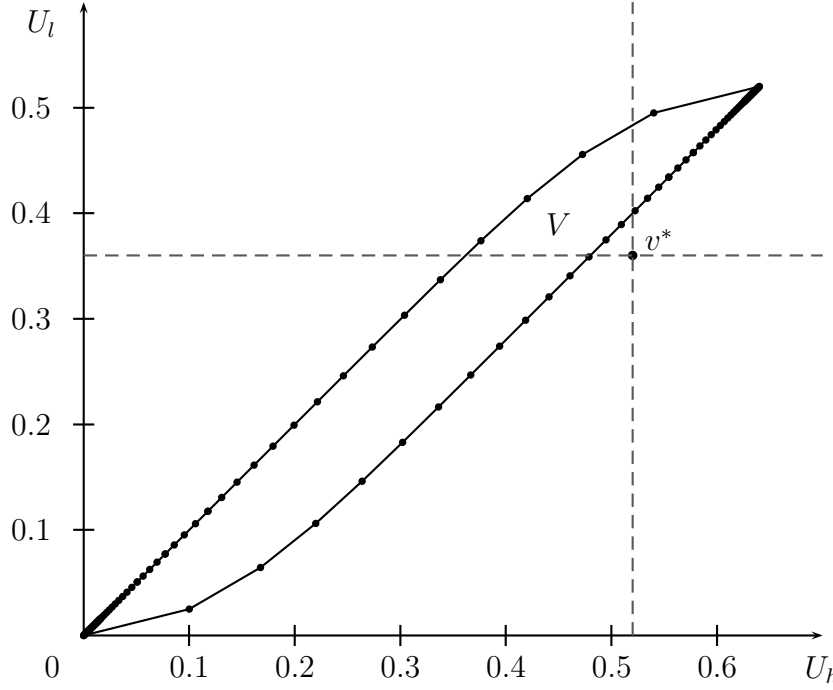


Figure 2: The set V for parameters $(\delta, \rho_h, \rho_l, l, h) = (9/10, 1/3, 1/4, 1/4, 1)$.

$v_h^*) > 1$, the vector v^* is outside V . For any ϕ , the complete-information value function $\bar{W}(U, \phi)$ increases as we vary U within V toward its lower boundary, either horizontally or vertically. This pattern should not be particularly surprising. Due to private information, the low-type agent derives information rents such that if the high-type agent's utility were first-best, the low-type agent's utility would be too high. It is instructive to study how the shape of V varies with persistence. When $\kappa = 0$ and values are independent over time, the lower type agent prefers to receive a larger fraction (or probability) of the good tomorrow rather than today (adjusting for discounting) but has no preference over later times. The case is analogous for the high type agent. As a result, all the vertices $\{\bar{u}^\nu\}_{\nu=1}^\infty$ (resp., $\{\underline{u}^\nu\}_{\nu=1}^\infty$) are aligned and V is a parallelogram with vertices $0, \mu, \bar{u}^1$ and $\underline{u}^1$. As κ increases, the imbalance between type utilities increases and the set V flattens. In the limit of perfect persistence, the low-type agent no longer feels differently about frontloading vs. backloading because no amount of time allows his type to change. See Figure 3.

The structure of V relies on the assumption $\kappa \geq 0$. If types were negatively correlated over time, then frontloading and backloading would not be policies spanning the boundary of V . This fact is easily seen by considering the case in which there is perfect negative serial correlation. Then, providing the unit if only if the round is odd (even) favors (hurts) the low-

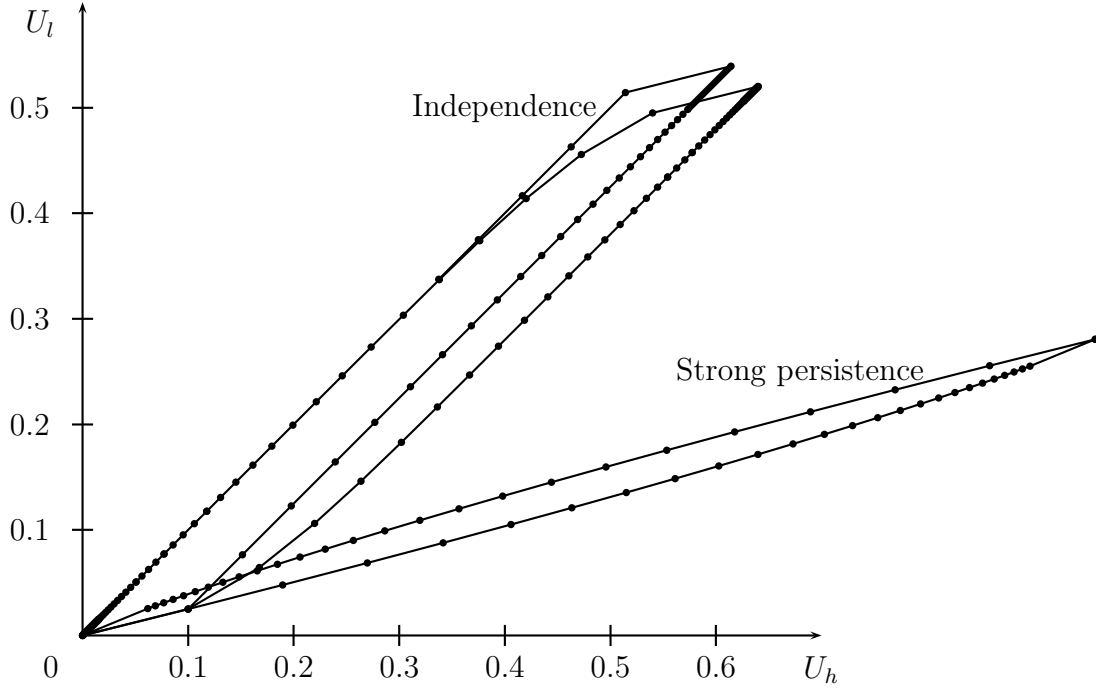


Figure 3: Impact of persistence, as measured by $\kappa \geq 0$.

type agent relative to the high-type agent. These two policies achieve the extreme points of V . According to whether higher or lower values of U_h are being considered, the other boundary points combine such alternation with frontloading or backloading. A negative correlation thus requires a separate (but analogous) treatment, motivating our focus on $\kappa \geq 0$.

Importantly, frontloading and backloading are not the only ways to achieve boundary payoffs. It is relatively clear that the lower locus corresponds to policies that (starting from this locus) assign as high a probability as possible to the good being supplied for every high report while promising continuation utilities that make IC_L always bind. Similarly, the upper boundary corresponds to those policies that assign as low a probability as possible to the good being supplied for low reports while promising continuation utilities that make IC_H always bind. Frontloading and backloading are representative examples of each class.

4.4 The Optimal Mechanism

Not every incentive-feasible utility vector is on path given the optimal policy. Irrespective of the sequence of reports, some vectors simply never arise. While it is necessary to solve for

the value function and optimal policy on the entire domain V , we first focus on the subset of V that is relevant given the optimal initial promise and resulting dynamics. We relegate discussion of the optimal policy for other utility vectors to the end of this section.

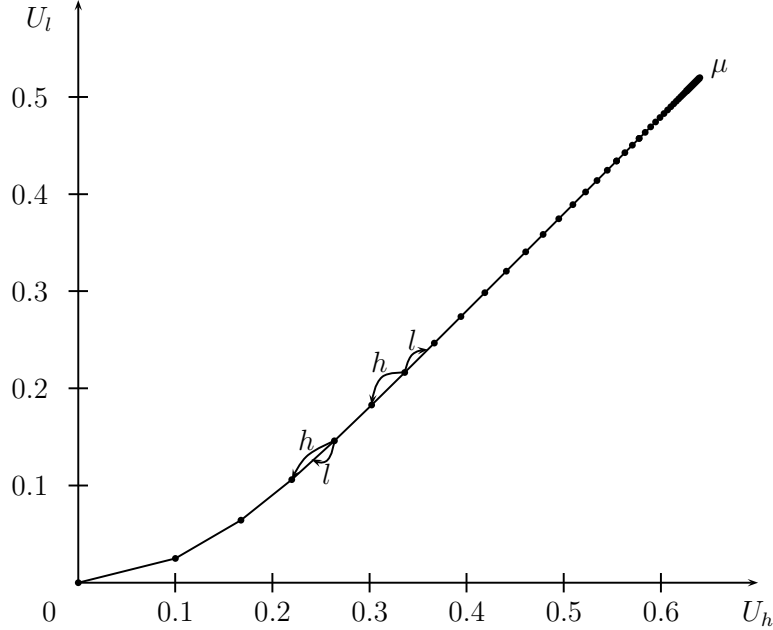


Figure 4: Dynamics of utility on the lower locus.

This subset is the lower locus—the polygonal chain spanned by pure frontloading. Furthermore, two observations from the i.i.d. case remain valid. First, the efficient choice is made as long as possible given feasibility, and second, the promises are chosen so the agent is indifferent between the two reports when his type is low.

To understand why such a policy yields utilities on the “frontloading” boundary (as mentioned at the end of Section 4.3), note that, because the low type is indifferent between the two reports, the agent is willing to always report high irrespective of his type. Because the principal then supplies the good, from the agent’s perspective, the pair of utilities can be computed as if frontloading were the policy being implemented.

From the principal’s perspective, however, it is important that this is not the actual policy. As in the i.i.d. case (a special case of the analysis), the payoff is higher under the efficient policy. Making the efficient choice, even if it involves delay, increases the principal’s payoff.

Hence, after a high report, as in the i.i.d. case, continuation utility declines.²² Specifically,

²²Because the lower boundary is upward sloping, the *ex interim* utilities of both types vary in the same

$U(h)$ is computed as under frontloading as the solution to the following system:

$$U_v = (1 - \delta)v + \delta \mathbf{E}_v[U(h)], \quad v = l, h,$$

where U is given. Here, $\mathbf{E}_v[U(h)]$ is the expectation of the utility vector $U(h)$ provided that the current type is v (e.g., for $v = h$, $\mathbf{E}_v[U(h)] = \rho_h U_l(h) + (1 - \rho_h)U_h(h)$).

The promised $U(l)$ does not admit such an explicit formula because it is specified by IC_L and the requirement that it lies on the lower boundary. In fact, $U(l)$ might be lower or higher than U (see Figure 4) depending on where U lies on the boundary. If it is high enough, $U(l)$ is higher; conversely, under certain conditions, $U(l)$ is lower than U when U is low enough.²³ The condition has a simple geometric interpretation: if the half-open line segment $(0, v^*]$ intersects the boundary,²⁴ then $U(l)$ is lower than U if and only if U lies below $\underline{U}$.²⁵ However, if there is no such intersection, then $U(l)$ is always higher than U . This intersection exists if and only if

$$\frac{h - l}{l} > \frac{1 - \delta}{\delta \rho_l}. \quad (10)$$

Hence, $U(l)$ is higher than U (for all U) if the low-type persistence is sufficiently high, which is intuitive. Utility declines even after a low report if U is so low that even the low-type agent expects to have sufficiently soon and often a high value such that the efficient policy would yield too high a utility. When the low-type persistence is high, this does not occur.²⁶ As in the i.i.d. case, the principal is able to achieve the complete-information payoff if and only if $U \leq \underline{U}$ (or $U = \mu$).

We summarize this discussion with the following theorem, a special case of the next.

Theorem 2 *The optimal policy consists of the constrained-efficient policy*

$$p_l = \max \left\{ 0, 1 - \frac{\mu_l - U_l}{(1 - \delta)l} \right\}, \quad p_h = \min \left\{ 1, \frac{U_h}{(1 - \delta)h} \right\}$$

in addition to a (specific) initially promised $U^0 > \underline{U}$ on the lower boundary of V and choices $(U(h), U(l))$ on this lower boundary such that IC_L always binds.

way. Accordingly, we use terms such as “higher” or “lower” utility and write $U < U'$ for the component-wise order.

²³As in the i.i.d. case, $U(l)$ is always higher than $U(h)$.

²⁴This line has the equation $U_l = \frac{\delta \rho_l}{1 - \delta(1 - \rho_l)} U_h$.

²⁵With some abuse, we write $\underline{U} \in \mathbf{R}^2$ because it is the natural extension of $\underline{U} \in \mathbf{R}$ as introduced in Section

3. Additionally, we set $\underline{U} = 0$ if the intersection does not exist.

²⁶This condition is satisfied in the i.i.d. case due to our assumption that $\delta > l/\mu$.

While the implementation in the i.i.d. case is described in terms of a “utility budget,” inspired by the use of (*ex ante*) utility as a state variable, the analysis of the Markov case strongly suggests the use of a more concrete metric – the number of units that the agent is entitled to claim in a row with “no questions asked.” The utility vectors on the boundary are parameterized by the number of rounds required to reach 0 under frontloading. Due to integer issues, we denote such a policy by a pair $(m, \lambda) \in (\mathbf{N}_0 \cup \{\infty\}) \times [0, 1]$, with the interpretation that the good is supplied with probability λ for m rounds and complementary probability $1 - \lambda$ for $m + 1$ rounds, and write $U_h(m, \lambda), U_l(m, \lambda)$. If $m = \infty$, the good is always supplied, yielding utility μ .

We may think of the optimal policy as follows. During a given round m , the agent is promised (m_n, λ_n) . If the agent asks for the unit (and this is feasible, that is, $m_n \geq 1$), the next promise (m_{n+1}, λ_{n+1}) is the solution to

$$\frac{U_l(m_n, \lambda_n) - (1 - \delta)l}{\delta} = \mathbf{E}_l [U_{v_{t+1}}(m_{n+1}, \lambda_{n+1})], \quad (11)$$

where $\mathbf{E}_l [U_{v_{t+1}}(m_{n+1}, \lambda_{n+1})] = (1 - \rho_l)U_l(m_{n+1}, \lambda_{n+1}) + \rho_l U_h(m_{n+1}, \lambda_{n+1})$ is the expected utility from tomorrow’s promise (m_{n+1}, λ_{n+1}) given that today’s type is low. If $m_n < 1$ and the agent claims to be high, he then receives the unit with probability $\tilde{q}$, which solves $U_l(m_n, \lambda_n) - \tilde{q}(1 - \delta)l = 0$.) However, claiming to be low simply leads to the revised promise

$$\frac{U_l(m_n, \lambda_n)}{\delta} = \mathbf{E}_l [U_{v_{n+1}}(m_{n+1}, \lambda_{n+1})], \quad (12)$$

provided that there exists a (finite) m_{n+1} and $\lambda_{n+1} \in [0, 1]$ that solve this equation.²⁷ The policy described by (11)–(12) reduces to that described in Section 3.3 (a special case of the Markovian case). The policy described in the i.i.d. case is obtained by taking the expectations of these dynamics with respect to today’s type.

It is perhaps surprising that the optimal policy can be derived but less surprising that comparative statics are difficult to obtain except by numerical simulations. By scaling both ρ_l and ρ_h by a common factor, $p \geq 0$, one varies the persistence of the value without affecting the invariant probability q , and hence, the value μ is also unaffected. Numerically, it appears that a decrease in persistence (an increase in p) leads to a higher payoff. When $p = 0$, types never change, and we are left with a static problem (for the parameters chosen here, it is best not to provide the good). When p increases, types change more rapidly, and the promised utility becomes a frictionless currency.

²⁷This is impossible if the promised (m_n, λ_n) is already too large (formally, if the corresponding payoff vector $(U_h(m, \lambda), U_l(m, \lambda)) \in V_h$). In that case, the good is provided with the probability that solves $\frac{U_l(m_n, \lambda_n) - \tilde{q}(1 - \delta)l}{\delta} = \mathbf{E}_l [U_{v_{n+1}}(\infty)]$.

As mentioned above, these comparative statics are merely suggested by simulations. As promised utility varies as a random walk with unequal step size on a grid that is itself a polygonal chain, there is little hope of establishing this result more formally here. To derive a result along these lines, see Section 5.1. Nonetheless, note that it is not persistence but positive correlation that is detrimental. It is tempting to believe that any type of persistence is bad because it endows the agent with private information that pertains not only to today's value but also to tomorrow's, and eliciting private information is usually costly in information economics. However, conditional on his knowledge regarding today's type, the agent's information regarding his future type is known (unlike in the case of a prophetic agent with i.i.d. values). Note that with perfectly negatively correlated types, the complete information payoff would be easy to achieve: offer the agent a choice between receiving the good in all odd or all even rounds. As $\delta > l/h$ (in fact, we assumed that $\delta > l/\mu$), the low-type agent would tell the truth. Just as in the case of a lower discount rate, a more negative correlation (or less positive correlation) makes future promises more effective incentives because preference misalignment is shorter-lived.

It is immediately apparent that given any initial choice of $U_0 \notin \underline{V} \cup \{\mu\}$, finitely many consecutive reports of l (or h) suffice for the promised utility to reach μ (or 0). As a result, both long-run outcomes have strictly positive probability under the optimal policy for any optimal initial choice. By the Borel-Cantelli lemma, this implies that absorption occurs almost surely. As in the i.i.d. model, the *ex ante* utility computed under the invariant distribution is a random process that drifts upward if and only if $qU_l + (1-q)U_h \geq qh$, where the right-hand side is the flow utility under the efficient policy. However, we are unable to derive the absorption probabilities beginning from U_0 because the Markov model admits no analogue to the TW-martingale.

4.5 A Comparison with Transfers as in Battaglini (2005)

As mentioned above, our model can be regarded as the no-transfer counterpart of Battaglini (2005).

Initially, the difference in results is striking. A main finding of Battaglini, “no distortion at the top,” has no counterpart in this model. With transfers, efficient provision occurs *forever* once the agent is revealed to be of the high type. Additionally, as noted above, with transfers, even along the history in which efficiency is not achieved in finite time, namely, an uninterrupted string of low reports, efficiency is asymptotically approached. As explained above, we necessarily obtain (with probability one) an inefficient outcome, which can be

implemented without further reports. Moreover, both outcomes (providing the good forever and never providing it again) can arise. In summary, in the presence of transfers, inefficiencies are frontloaded as much as possible, whereas here, they are backloaded to the greatest extent possible.

The difference can be understood as follows. First, and importantly, Battaglini’s results rely on revenue maximization being the objective function. With transfers, efficiency is trivial: simply charge c whenever the good must be supplied.

Once revenue maximization becomes the objective, the incentive constraints reverse with transfers: it is no longer the low type who would like to mimic the high type but the high type who would like to avoid paying his entire value for the good by claiming he is a low type. To avoid this, the high type must be given information rents, and his incentive constraint becomes binding. Ideally, the principal would like to charge for these rents *before* the agent has private information while the expected value of these rents to the agent remains common knowledge. When types are i.i.d., this poses no difficulty, and these rents can be expropriated one round ahead of time. With correlation, however, different types of agents value these rents differently, as their likelihood of being high in the future depends on their current type. However, when considering information rents sufficiently far in the future, the initial type exerts a minimal effect on the conditional expectation of the value of these rents. Hence, the value can “almost” be extracted. As a result, it is in the principal’s best interest to maximize the surplus and offer a nearly efficient contract at all dates that are sufficiently far away.

We observe that money plays two roles. First, because it is an instrument that allows promises to “clear” on the spot without allocative distortions, it prevents the occurrence of backloaded inefficiencies – a poor substitute for money in this regard. Even if payments could not be made “in advance,” this would suffice to restore efficiency if that were the objective. Another role of money, as highlighted by Battaglini, is that it allows value to be transferred from the agent to the principal before information becomes private. Hence, information rents no longer impede efficiency, at least with respect to the remote future. These future inefficiencies can be eliminated such that inefficiencies only arise over the short run.

A plausible intermediate case arises when money is available but the agent is protected by limited liability, meaning that payments can only be made in one direction: from the principal to the agent. The principal strives to maximize the social surplus net of any payments made.²⁸ In this case, we demonstrate in the appendix (see Lemma 11) that no transfers are made if (and only if) $c - l < l$. This condition can be interpreted as follows:

²⁸If payments do not matter for the principal, efficiency is easily achieved because he would pay c to the agent if and only if the report is low and nothing otherwise.

$c - l$ is the cost to the principal of incurring one inefficiency (supplying the good when the type is low), whereas l is the cost to the agent of forgoing a low-unit value. Hence, if it is costlier to buy off the agent than to supply the good when the value is low, the principal prefers to never use money as an instrument and to follow the optimal policy absent any money.

4.6 The General Solution

Theorem 2 follows from the analysis of the optimal policy on the entire domain V . Because only those values in V along the lower boundary are relevant, the reader might elect to skip this subsection, which completely solves the program in Section 4.1.

First, we further divide V into subsets and introduce two sequences of utility vectors for this purpose. Given $\underline{U}$, define the sequence $\{v^\nu\}_{\nu \geq 1}$ by

$$v_h^\nu = \delta^\nu ((1 - q)\underline{U}_l + q\underline{U}_h + (1 - q)\kappa^\nu(\underline{U}_h - \underline{U}_l)), v_l^\nu = \delta^\nu ((1 - q)\underline{U}_l + q\underline{U}_h - q\kappa^\nu(\underline{U}_h - \underline{U}_l)),$$

and define

$$\underline{V} = \text{co}\{\{0\} \cup \{v^\nu\}_{\nu \geq 0}\}. \quad (13)$$

See Figure 5. Note that $\underline{V}$ has a non-empty interior if and only if ρ_l is sufficiently large (see (10)). This set is the domain of utilities for which the complete information payoff can be achieved, as stated below.

Lemma 6 *For all $U \in \underline{V} \cup \{\mu\}$ and all ϕ ,*

$$W(U, \phi) = \bar{W}(U, \phi).$$

Conversely, if $U \notin \underline{V} \cup \{\mu\}$, then $W(U, \phi) < \bar{W}(U, \phi)$ for all $\phi \in (0, 1)$.

Second, we define $\hat{u}^\nu := (\hat{u}_h^\nu, \hat{u}_l^\nu)$, $\nu \geq 0$ as follows:

$$\begin{aligned} \hat{u}_h^\nu &= \mu_h - (1 - \delta)h - \delta^{\nu+1} ((1 - q)l + qh + (1 - q)\kappa^{\nu+1}(\mu_h - \mu_l)), \\ \hat{u}_l^\nu &= \mu_l - (1 - \delta)l - \delta^{\nu+1} ((1 - q)l + qh - q\kappa^{\nu+1}(\mu_h - \mu_l)). \end{aligned}$$

We note that $\hat{u}^0 = 0$ and $\hat{u}^\nu$ is an increasing sequence (in both coordinates) contained in V , where $\lim_{\nu \rightarrow \infty} \hat{u}^\nu = \bar{u}^1$. The ordered sequence $\{\hat{u}^\nu\}_{\nu \geq 0}$ defines a polygonal chain P that divides $V \setminus \underline{V}$ into two further subsets, V_t and V_b , consisting of those points in $V \setminus \underline{V}$ that lie above or below P . It is readily verified that the points U on P are precisely those for which, assuming IC_H , the resulting $U(l)$ lies exactly on the lower boundary of V . We also let P_b ,

P_t be the (closure of the) polygonal chains defined by $\{\underline{u}^\nu\}_{\nu \geq 0}$ and $\{\overline{u}^\nu\}_{\nu \geq 0}$ that correspond to the lower and upper boundaries of V .

We now define a policy (which, as we will see below, is optimal) ignoring for the present the choice of the initial promise.

Definition 2 *For all $U \in V$, set*

$$p_l = \max \left\{ 0, 1 - \frac{\mu_l - U_l}{(1 - \delta)l} \right\}, \quad p_h = \min \left\{ 1, \frac{U_h}{(1 - \delta)h} \right\}, \quad (14)$$

and

$$U(h) \in P_b, \quad U(l) \in \begin{cases} P_b & \text{if } U \in V_b \\ P_t & \text{if } U \in P_t. \end{cases}$$

Furthermore, if $U \in V_t$, $U(l)$ is chosen such that IC_H binds.

For each continuation utility vector $U(h)$ or $U(l)$, this yields one constraint (either an incentive constraint or a constraint that the utility vector lies on one of the boundaries). In addition to the two promise keeping equations, this yields four constraints, which uniquely define the pair of points $(U(h), U(l))$. It is readily verified that the policy and choices of $U(l), U(h)$ also imply that IC_L binds for all $U \in P_b$. A surprising property of this policy is that it is independent of the principal's belief. That is, the principal's belief regarding the agent's value is irrelevant to the optimal policy given the promised utility. However, the initial choice of utility on the lower boundary depends on this belief, as does the pay-off. However, the monotonicity properties of the value function with respect to utilities are sufficiently strong and uniform that the constraints specify the policy.

Figure 5 illustrates the dynamics of the optimal policy. Given any promised utility vector in V , the vector $(p_h, p_l) = (1, 0)$ is played (unless U is too close to 0 or μ), and promised utilities depend on the report. A report of l shifts the utility to the right (toward higher utilities), whereas a report of h shifts utility to the left and toward the lower boundary. Below the polygonal chain, the l report also shifts us to the lower boundary (and IC_L binds), whereas above the chain, IC_H binds. In fact, note that if the utility vector is on the upper boundary, the continuation utility after l remains there.

For completeness, we also define the subsets over which promise keeping prevents the efficient choices $(p_h, p_l) = (1, 0)$ from being made. Let V_h be $\{(U_h, U_l) : (U_h, U_l) \in V, U_l \geq \overline{u}_l^1\}$ and V_l be $\{(U_h, U_l) : (U_h, U_l) \in V, U_h \leq \underline{u}_h^1\}$. It is easily verified that $(p_h, p_l) = (1, 0)$ is feasible at U given promise keeping if and only if $U \in V \setminus (V_h \cup V_l)$.

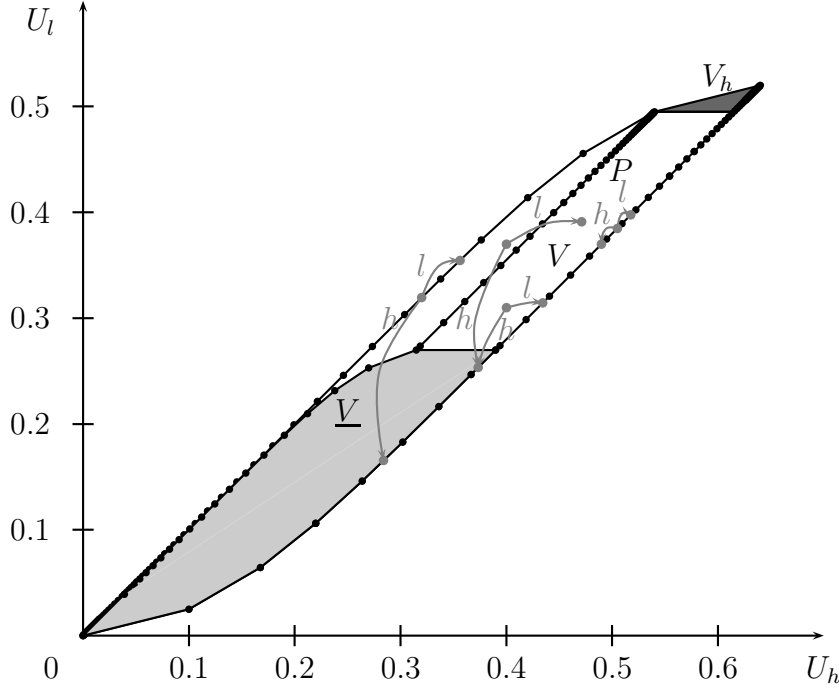


Figure 5: The set V and the optimal policy for $(\delta, \rho_h, \rho_l, l, h) = (9/10, 1/3, 1/4, 1/4, 1)$.

Theorem 3 Fix $U_0 \in V$; given U_0 , the policy stated above is optimal. The initial promise U^* is in $P_b \cap (V \setminus \underline{V})$, with U^* increasing in the principal's prior belief.

Furthermore, the value function $W(U_h, U_l, \phi)$ is weakly increasing in U_h along the rays $x = \phi U_h + (1 - \phi)U_l$ for any $\phi \in \{1 - \rho_h, \rho_l\}$.

Given that $U^* \in P_b$ and given the structure of the optimal policy, the promised utility vector never leaves P_b . It is also simple to verify that, as in the i.i.d. case (and by the same arguments), the (one-sided) derivative of W approaches the derivative of $\bar{W}$ as U approaches either μ or the set $\underline{V}$. As a result, the initial promise U^* is strictly interior.

5 Extensions

Two modeling choices deserve discussion. First, we have opted for a discrete time framework because it embeds the case of independent values – a natural starting point for which there is no counterpart in continuous time. However, this choice comes at a price. By varying the discount factor, we change both the patience of the players and the rate at which types change with independent values. This is not necessarily the case with Markovian types, but the analytical difficulties prevent us from deriving definitive comparative statics, a deficiency that we remedy below by resorting to continuous time.

Second, we have assumed that the agent's value is binary. As is well known (see Battaglini and Lamba, 2014, for instance), it is difficult to make progress with more types, even with transfers, unless strong assumptions are imposed. In the i.i.d. case, this is nonetheless possible. Below, we consider the case of a continuum of types, which allows us to evaluate the robustness of our various findings.

5.1 Continuous Time

To make further progress, we examine the limiting stochastic process of utility and payoff as transitions that are scaled according to the usual Poisson limit, when variable round length, $\Delta > 0$, is taken to 0, at the same time as the transition probabilities $\rho_h = \lambda_h \Delta$, $\rho_l = \lambda_l \Delta$. That is, we let $(v_t)_{t \geq 0}$ be a continuous time Markov chain (by definition, a right-continuous process) with values in $\{h, l\}$, initial probability q of h , and parameters $\lambda_h, \lambda_l > 0$. Let $T_0, T_1, T_2, \dots$ be the corresponding random times at which the value switches (setting $T_0 = 0$ if the initial state is l such that, by convention, $v_t = l$ on any interval $[T_{2k}, T_{2k+1})$). The initial type is h with probability $q = \rho_l / (\rho_h + \rho_l)$.

The optimal policy defines a tuple of continuous time processes that follow deterministic trajectories over any interval $[T_{2k}, T_{2k+1})$. First, the belief $(\mu_t)_{t \geq 0}$ of the principal takes values in the range $\{0, 1\}$. Namely, $\mu_t = 0$ over any interval $[T_{2k}, T_{2k+1})$, and $\mu_t = 1$ otherwise. Second, the utilities of the agent $(U_{l,t}, U_{h,t})_{t \geq 0}$ are functions of his type. Finally, the expected payoff of the principal, $(W_t)_{t \geq 0}$, is computed according to his belief μ_t .

The pair of processes $(U_{l,t}, U_{h,t})_{t \geq 0}$ takes values in V , obtained by considering the limit (as $\Delta \rightarrow 0$) of the formulas for $\{\underline{u}^\nu, \bar{u}^\nu\}_{\nu \in \mathbf{N}}$. In particular, one obtains that the lower bound is given in parametric form by

$$\begin{aligned}\underline{u}_h(\tau) &= (1 - e^{-r\tau})\mu_h + e^{-r\tau}(1 - e^{-(\lambda_h + \lambda_l)\tau})(1 - q)(\mu_h - \mu_l), \\ \underline{u}_l(\tau) &= (1 - e^{-r\tau})\mu_h - e^{-r\tau}(1 - e^{-(\lambda_h + \lambda_l)\tau})q(\mu_h - \mu_l),\end{aligned}$$

where $\tau \geq 0$ can be interpreted as the requisite time for promises to be fulfilled under the policy that consists of producing the good regardless of the reports until that time is elapsed. Here, as before,

$$\mu = \left(h - \frac{\lambda_h}{\lambda_h + \lambda_l + r}(h - l), l + \frac{\lambda_l}{\lambda_h + \lambda_l + r}(h - l) \right)$$

is the utility vector achieved by providing the good forever. The upper boundary is now

given by

$$\begin{aligned}\bar{u}_h(\tau) &= e^{-r\tau}\mu_h - e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau}(1 - q)(\mu_h - \mu_l)), \\ \bar{u}_l(\tau) &= e^{-r\tau}\mu_h - e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau}q(\mu_h - \mu_l)).\end{aligned}$$

Finally, the set $\underline{V}$ is either empty or defined by those utility vectors in V lying below the graph of the curve defined by

$$\begin{aligned}v_h(\tau) &= e^{-r\tau}((1 - q)\underline{U}_l + q\underline{U}_h) + e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau}(1 - q)(\underline{U}_h - \underline{U}_l)), \\ v_l(\tau) &= e^{-r\tau}((1 - q)\underline{U}_l + q\underline{U}_h) - e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau}q(\underline{U}_h - \underline{U}_l)),\end{aligned}$$

where $(\underline{U}_h, \underline{U}_l)$ are the coordinates of the largest intersection of the graph of $\bar{u} = (\bar{u}_h, \bar{u}_l)$ with the line $u_l = \frac{\lambda_l}{\lambda_l + r}u_h$. It is immediately verifiable that $\underline{V}$ has a non-empty interior iff (cf. (10))

$$\frac{h - l}{l} > \frac{r}{\lambda_l}.$$

Hence, the complete-information payoff cannot be achieved for any utility level (aside from 0 and μ) whenever the low state is too persistent. However, $\underline{V}$ is always non-empty when the agent is sufficiently patient.

Figure 6 illustrates this construction. Note that the boundary of V is smooth except at 0 and μ . It is also easy to verify that the limit of the chain defined by $\hat{u}^\nu$ lies on the lower boundary: V_b is asymptotically empty.

The great advantage of the Poisson system is that it allows us to explicitly solve for payoffs. We sketch the details of the derivation.

How does τ – the denomination of utility on the lower boundary – evolve over time? Along the lower boundary, it evolves continuously. On any interval over which h is continuously reported, it evolves deterministically, with increments

$$d\tau_h := -dt.$$

However, when l is reported, the evolution is more complicated. Algebra yields

$$d\tau_l := \frac{g(\tau)}{\mu - q(h - l)e^{-(\lambda_h+\lambda_l)\tau}}dt,$$

where

$$g(\tau) := q(h - l)e^{-(\lambda_h+\lambda_l)\tau} + le^{r\tau} - \mu,$$

and $\mu = qh + (1 - q)l$, as before.

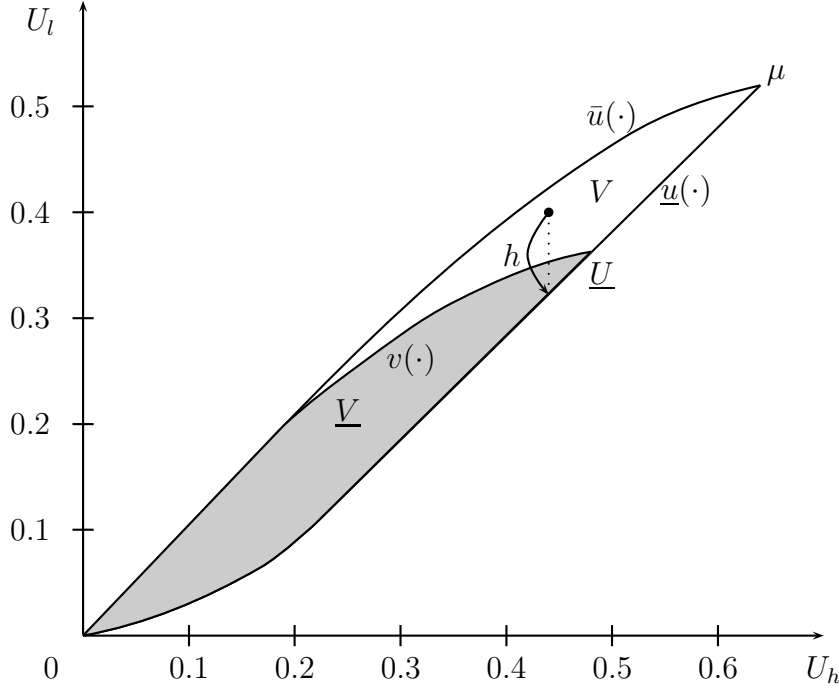


Figure 6: Incentive-feasible set for $(r, \lambda_h, \lambda_l, l, h) = (1, 10/4, 1/4, 1/4, 1)$.

The increment $d\tau_l$ is positive or negative depending upon whether τ maps onto a utility vector in $\underline{V}$. If $\underline{V}$ has a non-empty interior, we can identify the value of τ that is the intersection of the critical line and the boundary; call this value $\hat{\tau}$, which is simply the positive root (if any) of g . Otherwise, set $\hat{\tau} = 0$.

The evolution of utility is not continuous for utilities that are not on the lower boundary. A high report leads to a vertical shift in the utility of the low type down to the lower boundary. See Figure 6. This change is intuitive because by promise keeping, the utility of the high-type agent cannot jump because such an instantaneous report has only a minute impact on his flow utility. A low report, however, leads to a drift in the type's utility.

Our goal is to derive the principal's value function. Because his belief is degenerate, except at the initial instant, we write $W_h(\tau)$ (resp., $W_l(\tau)$) for the payoff when (he assigns probability one to the event that) the agent's valuation is currently high (or low). By definition of the policy that is followed, the value functions solve the following paired system of equations:

$$W_h(\tau) = rdt(h - c) + \lambda_h dt W_l(\tau) + (1 - rdt - \lambda_h dt) W_h(\tau + d\tau_h) + \mathcal{O}(dt^2),$$

and

$$W_l(\tau) = \lambda_l dt W_h(\tau) + (1 - rdt - \lambda_l dt) W_l(\tau + d\tau_l) + \mathcal{O}(dt^2).$$

Assume for now (as will be verified) that the functions W_h, W_l are twice differentiable. We then obtain the differential equations

$$(r + \lambda_h)W_h(\tau) = r(h - c) + \lambda_h W_l(\tau) - W_h'(\tau)$$

and

$$(r + \lambda_l)W_l(\tau) = \lambda_l W_h(\tau) + \frac{g(\tau)}{\mu - q(h - l)e^{-(\lambda_h + \lambda_l)\tau}} W_l'(\tau)$$

subject to the following boundary conditions.²⁹ First, at $\tau = \hat{\tau}$, the value must coincide with the value given by the first-best payoff $\bar{W}$ in that range. That is, $W_h(\hat{\tau}) = \bar{W}_h(\hat{\tau})$, and $W_l(\hat{\tau}) = \bar{W}_l(\hat{\tau})$. Second, as $\tau \rightarrow \infty$, it must hold that the payoff $\mu - c$ is approached. Hence,

$$\lim_{\tau \rightarrow \infty} W_h(\tau) = \mu_h - c, \quad \lim_{\tau \rightarrow \infty} W_l(\tau) = \mu_l - c.$$

Despite having variable coefficients, this system can be solved. See Section C.1 of the appendix for the solution upon which the following comparative statics are based.

Lemma 7 *The value $W(\tau) := qW_h(\tau) + (1 - q)W_l(\tau)$ decreases pointwise in persistence $1/p$, where $\lambda_h = p\bar{\lambda}_h$, $\lambda_l = p\bar{\lambda}_l$, for some fixed $\bar{\lambda}_h, \bar{\lambda}_l$ for all τ ,*

$$\lim_{p \rightarrow \infty} W(\tau) = \bar{W}(\tau), \quad \lim_{p \rightarrow 0} \max_{\tau} W(\tau) = \max\{\mu - c, 0\}.$$

The proof is in Appendix C.1. Hence, persistence hurts the principal's payoff, which is intuitive. With independent types, the agent's preferences are quasilinear in promised utility such that the only source of inefficiency derives from the bounds on this currency. When types are correlated, the promised utility is no longer independent of today's types in the agent's preferences, reducing the degree to which this mechanism can be used to efficiently provide incentives. With perfectly persistent types, there is no longer any leeway, and we are back to the inefficient static outcome. Figure 7 illustrates the value function for two levels of persistence and compares it to the complete-information payoff evaluated along the lower locus, $\bar{W}$ (the lower envelope of three curves).

What about the agent's utility? We note that the utility of both types is increasing in τ . Indeed, because a low type is always willing to claim that his value is high, we may compute his utility as the time over which he would obtain the good if he continuously claimed to be of the high type, which is precisely the definition of τ . However, persistence plays an ambiguous

²⁹To be clear, these are not HJB equations, as there is no need to verify the optimality of the policy that is being followed. This fact has already been established. The functions must satisfy these simple recursive equations.

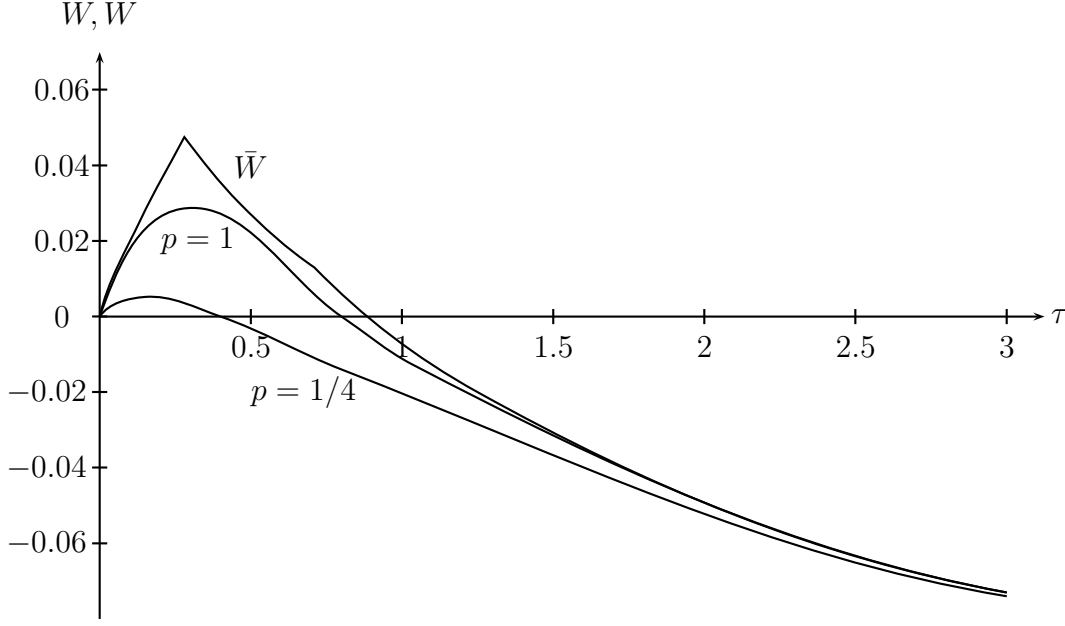


Figure 7: Value function and complete information payoffs as a function of τ (here, $(\lambda_l, \lambda_h, r, l, h, c) = (p/4, 10p/4, 1, 1/4, 1, 2/5)$ and $p = 1, 1/4$).

role in determining the agent's utility: perfect persistence is his favorite outcome if $\mu > c$. Hence, always providing the good is the best option in the static game. Conversely, perfect persistence is worse if $\mu < c$. Hence, persistence tends to improve the agent's situation when $\mu > c$.³⁰ As $r \rightarrow 0$, the principal's value converges to the complete information payoff $q(h-c)$. We conclude with a rate of convergence without further discussion given the comparison with Jackson and Sonnenschein (2007) described in Section 3.4.

Lemma 8 *It holds that*

$$|\max_{\tau} W(\tau) - q(h-c)| = \mathcal{O}(r).$$

5.2 Continuous Types

It is important to understand the role played by the assumption that there are only two types. To make progress, assume here that types are drawn i.i.d. from some distribution F with support $[\underline{v}, 1]$, $\underline{v} \in [0, 1)$ and density $f > 0$ on $[\underline{v}, 1]$. Below, we consider the specific case of a power distribution $F(v) = v^a$ with $a \geq 1$, but this example is not necessary yet. Let $\mu = \mathbf{E}[v]$ be the expected value of the type and hence the highest promised utility. As

³⁰However, this convergence is not necessarily monotone, which is easy to verify via examples.

before, we begin with the benchmark of complete information, with a lemma whose proof is straightforward and omitted.

Lemma 9 *The complete information payoff function $\bar{W}$ is strictly concave. The complete information policy is unique and of the threshold type, where the threshold v^* is continuously decreasing from 1 to 0 as U goes from 0 to $\bar{v}$. Furthermore, given the initial promise U , future utility remains constant at U .*

That is, given a promised utility $U \in [0, \mu]$, there exists a threshold v^* such that the good is provided if and only if the type is above v^* . Furthermore, utility does not evolve over time.

Returning to the case in which the agent privately observes his value, we prove the following theorem:³¹

Theorem 4 *The value function is strictly concave in U , continuously differentiable, and strictly below the complete information payoff (except for $U = 0, \mu$). Given $U \in (0, \mu)$, the optimal policy $p : [0, 1] \rightarrow [0, 1]$ is not a threshold policy.*

Once again, we note how the absence of money affects the structure of the allocation. One might have expected, given the linearity of the agent's utility and the principal's payoff, the solution to be “bang-bang” in p such that given some value of U , all types above a certain threshold receive the good supplied while those below receive it with probability zero. However, without transfers, incentive compatibility requires the continuation utility be distorted, and the payoff is not linear in the utility. Hence, consider a small interval of types around the indifferent candidate threshold type. From the principal's perspective, conditional on the agent being in this interval, the outcome is a lottery over $p = 0, 1$ and corresponding continuation payoffs. Replacing this lottery with its expected value would leave the agent virtually indifferent, but the substitution would certainly help the principal, as his continuation payoff is a strictly concave function of the continuation utility.

It is difficult to describe dynamics in the same level of detail as in the binary case. However, we recover the TW-martingale: W' is a bounded martingale, as U -a.e.,

$$W'(U) = \int_0^1 W'(\mathcal{U}(U, v)) dF(v),$$

where $\mathcal{U} : [0, \mu] \times [0, 1] \rightarrow [0, \bar{v}]$ is the optimal policy mapping current utility and reported type into continuation utility. Hence, because, except at $U = 0, \mu$, $\mathcal{U}(U, \cdot)$ is not constant (v -a.e.) and W is strictly concave, it must be that the limit is either 0 or μ and both must occur with positive probability. Hence,

³¹See the additional appendix.

Lemma 10 *Given any initial level U_0 , the utility process U_n converges to $\{0, \mu\}$, with both limits having strictly positive probability if $\underline{v} > 0$ (If $\underline{v} = 0$, 0 occurs a.s.).*

In Appendix C.2, we explain how the optimal policy may be found using control theory and prove the following proposition.

Proposition 1 *For a power distribution $F(v) = v^a$ with $a \geq 1$, there exists $U^{**} \in (0, \mu)$ such that*

1. *for any $U < U^{**}$, there exists v_1 such that $p(v) = 0$ for $v \in [0, v_1]$ and $p(v)$ is strictly increasing (and continuous) when $v \in (v_1, 1]$. The constraint $U(1) \geq 0$ binds, while the constraint $p(1) \leq 1$ does not.*
2. *for any $U \geq U^{**}$, there exists $0 \leq v_1 \leq v_2 \leq 1$ such that $p(v) = 0$ for $v \leq v_1$, $p(v)$ is strictly increasing (and continuous) when $v \in (v_1, v_2)$ and $p(v) = 1$ for $v \geq v_2$. The constraints $U(0) \leq \mu$ and $U(1) \geq 0$ do not bind.*

It is clear that indirect implementation is more difficult, as the agent is no longer making binary choices but is assigned the good with positive probability at certain values. Hence, at the very least, the variable fee of the two-part tariff that we describe must be extended to a nonlinear schedule in which the agent pays a price for each “share” of the good that he would like.

Markovian Types. Given the complexity of the problem, we see little hope for analytic results with additional types once independence is dropped. We note that deriving the incentive-feasible set is a difficult task. In fact, even with three types, an explicit characterization is lacking. It is intuitively clear that frontloading is the worst policy for the low type, given some promised utility to the high type, and backloading is the best, but what of maximizing a medium type’s utility given a pair of promised utilities to the low and high type? It appears that the convex hull of utilities from frontloading and backloading policies traces out the lowest utility that a medium type can obtain for any such pair, but the set of incentive-feasible payoffs has full dimension. The highest utility that he can receive occurs when one of his incentive constraints binds, but there are two possibilities, according to the incentive constraint. We obtain two hypersurfaces that do not appear to admit closed-form solutions. Additionally, the analysis of the i.i.d. case suggests that the optimal policy might follow a path of utility triples on such a boundary. One might hope that assuming that values follow a renewal process, as opposed to a general Markov process, might result in a lower-dimensional problem, but unfortunately, we fail to see how this is possible.

6 Concluding Comments

Here, we discuss a few obvious extensions.

Renegotiation-Proofness. The optimal policy, as described in Sections 3 and 4, is clearly not renegotiation-proof, unlike the case with transfers (see Battaglini, 2005). After a history of reports such that the promised utility would be zero, both agent and principal would be better off by renegeing and starting afresh. There are many ways to define renegotiation-proofness. Strong-renegotiation (Farrell and Maskin, 1989) would lead to a lower boundary on the utility vectors visited (except in the event that μ is sufficiently low that it makes the relationship altogether unprofitable such that $U^* = 0$.) However, the structure of the optimal policy can still be derived from the same observations. The low-type incentive-compatibility condition and promise keeping specify the continuation utilities, unless a boundary is reached regardless of whether it is the lower boundary (that must serve as a lower reflecting barrier) or the upper absorbing boundary μ .

Public Signals. While assuming no statistical evidence whatsoever allows us to clarify how the principal can exploit the repeated allocation decision to mitigate the inefficiency entailed by private information, there are many applications for which such evidence is available. This public signal depends on the current type and possibly on the action chosen by the principal. For instance, if we interpret the decision as filling a position (as in labor market applications), we might only receive feedback on the quality of the applicant if he is hired. If, provided that the good insures the agent against a risk with a cost that might be either high or low, the principal fails to provide a good, he might discover that the agent's claim was warranted.

Incomplete Information Regarding the Process. Thus far, we have assumed that the agent's type is drawn from a distribution that is common knowledge. This feature is obviously an extreme assumption. In practice, the agent might have superior information regarding the frequency with which high values arrive. If the agent knows the distribution from the beginning, the revelation principle applies, and it is a matter of revisiting the analysis from Section 3 with an incentive compatibility constraint at time 0.

Alternatively, the agent might not possess such information initially but be able to determine the underlying distribution from successive arrivals. This is the more challenging case in which the agent himself is learning about q (or, more generally, the transition matrix)

as time passes. In that case, the agent’s belief might be private (in the event that he has deviated in the past). Therefore, it is necessary to enlarge the set of reports. A mechanism is now a map from the principal’s belief μ (regarding the agent’s belief), a report by the agent of this belief, denoted ν , and his report on his current type (h or l) onto a decision of whether to allocate the good and the promised continuation utility. While we do not expect either token or budget mechanisms to be optimal in such environments, their simplicity and robustness suggest that they might provide valuable benchmarks.

References

- [1] Abdulkadiroğlu, A. and K. Bagwell (2012). “The Optimal Chips Mechanism in a Model of Favors,” working paper, Duke University.
- [2] Abdulkadiroğlu, A. and S. Loertscher (2007). “Dynamic House Allocations,” working paper, Duke University and University of Melbourne.
- [3] Abreu, D., D. Pearce, and E. Stacchetti (1990). “Toward a Theory of Discounted Repeated Games with Imperfect Monitoring,” *Econometrica*, **58**, 1041–1063.
- [4] Athey, S. and K. Bagwell (2001). “Optimal Collusion with Private Information,” *RAND Journal of Economics*, 428–465.
- [5] Battaglini, M. (2005). “Long-Term Contracting with Markovian Consumers,” *American Economic Review*, **95**, 637–658.
- [6] Battaglini, M. and R. Lamba (2014). “Optimal Dynamic Contracting: the First-Order Approach and Beyond,” working paper, Princeton University.
- [7] Benveniste, L.M. and J.A. Scheinkman (1979). “On the Differentiability of the Value Function in Dynamic Models of Economics,” *Econometrica*, **41**, 727–733.
- [8] Biais, B., T. Mariotti, G. Plantin and J.-C. Rochet (2007). “Dynamic Security Design: Convergence to Continuous Time and Asset Pricing Implications,” *Review of Economic Studies*, **74**, 345–390.
- [9] Casella, A. (2005). “Storable Votes,” *Games and Economic Behavior*, **51**, 391–419.
- [10] Cohn, Z. (2010). “A Note on Linked Bargaining,” *Journal of Mathematical Economics*, **46**, 238–247.

- [11] Cole, H.L. and N. Kocherlakota (2001). “Efficient Allocations with Hidden Income and Hidden Storage,” *Review of Economic Studies*, **68**, 523–542.
- [12] Derman, C. Lieberman, G.J. and S.M. Ross (1972). “A Sequential Stochastic Assignment Problem,” *Management Science*, **18**, 349–355.
- [13] Doepke, M. and R.M. Townsend (2006). “Dynamic Mechanism Design with Hidden Income and Hidden Actions,” *Journal of Economic Theory*, **126**, 235–285.
- [14] Fang, H. and P. Norman (2006). “To Bundle or Not To Bundle,” *RAND Journal of Economics*, **37**, 946–963.
- [15] Eilat, R. and A. Pautner (2011). “Optimal Bilateral Trade of Multiple Objects,” *Games and Economic Behavior*, **71**, 503–512.
- [16] Farrell, J. and E. Maskin (1989). “Renegotiation in Repeated Games,” *Games and Economic Behavior*, **1**, 327–360.
- [17] Fernandes, A. and C. Phelan (2000). “A Recursive Formulation for Repeated Agency with History Dependence,” *Journal of Economic Theory*, **91**, 223–247.
- [18] Frankel, A. (2011). “Contracting over Actions,” PhD Dissertation, Stanford University.
- [19] Fries, B.E. (1975). “Optimal ordering policy for a perishable commodity with fixed lifetime,” *Operations Research*, **23**, 46–61.
- [20] Fu, S. and R.V. Krishna (2014). “Dynamic Financial Contracting with Persistence,” working paper, Duke University.
- [21] Garrett, D. and A. Pavan (2015). “Dynamic Managerial Compensation: On the Optimality of Seniority-based Schemes,” *Journal of Economic Theory*, forthcoming.
- [22] Gershkov, A. and B. Moldovanu (2010). “Efficient Sequential Assignment with Incomplete Information,” *Games and Economic Behavior*, **68**, 144–154.
- [23] Hauser, C. and H. Hopenhayn (2008). “Trading Favors: Optimal Exchange and Forgiveness,” working paper, UCLA.
- [24] Hortala-Vallve, R. (2010). “Inefficiencies on Linking Decisions,” *Social Choice and Welfare*, **34**, 471–486.

- [25] Hylland, A., and R. Zeckhauser (1979). “The Efficient Allocation of Individuals to Positions,” *Journal of Political Economy*, **87**, 293–313.
- [26] Jackson, M.O. and H.F. Sonnenschein (2007). “Overcoming Incentive Constraints by Linking Decision,” *Econometrica*, **75**, 241–258.
- [27] Johnson, T. (2013). “Dynamic Mechanism Without Transfers,” working paper, Notre Dame.
- [28] Kováč, E., D. Krämer and T. Tatur (2014). “Optimal Stopping in a Principal-Agent Model With Hidden Information and No Monetary Transfers,” working paper, University of Bonn.
- [29] Krishna, R.V., G. Lopomo and C. Taylor (2013). “Stairway to Heaven or Highway to Hell: Liquidity, Sweat Equity, and the Uncertain Path to Ownership,” *RAND Journal of Economics*, **44**, 104–127.
- [30] Li, J., N. Matouschek and M. Powell (2015). “The Burden of Past Promises,” working paper, Northwestern University.
- [31] Miralles, A. (2012). “Cardinal Bayesian Allocation Mechanisms without Transfers,” *Journal of Economic Theory*, **147**, 179–206.
- [32] Möbius, M. (2001). “Trading Favors,” working paper, Harvard University.
- [33] Radner, R. (1981). “Monitoring Cooperative Agreements in a Repeated Principal-Agent Relationship,” *Econometrica*, **49**, 1127–1148.
- [34] Radner, R. (1986). “Repeated moral hazard with low discount rates,” *Uncertainty, information, and communication Essays in honor of Kenneth J. Arrow*, Vol. III, edited by W.P. Heller, R.M. Starr and D.A. Starrett, 25–64, Cambridge University Press.
- [35] Rubinstein, A. and M. Yaari (1983). “Repeated Insurance Contracts and Moral Hazard,” *Journal of Economic Theory*, **30**, 74–97.
- [36] Sendelbach, S. (2012). “Alarm Fatigue,” *Nursing Clinics of North America*, **47**, 375–382.
- [37] Spear, S.E. and S. Srivastava (1987). “On Repeated Moral Hazard with Discounting,” *Review of Economic Studies*, **54**, 599–617.

- [38] Thomas, J. and T. Worrall (1990). “Income Fluctuations and Asymmetric Information: An Example of the Repeated Principal-Agent Problem,” *Journal of Economic Theory*, **51**, 367–390.
- [39] Townsend, R.M. (1982). “Optimal Multiperiod Contracts and the Gain from Enduring Relationships under Private Information,” *Journal of Political Economy*, **90**, 1166–1186.
- [40] Zhang, H. and S. Zenios (2010). “A Dynamic Principal-Agent Model with Hidden Information: Sequential Optimality Through Truthful State Revelation,” *Operations Research*, **56**, 681–696.
- [41] Zhang, H. (2012). “Analysis of a Dynamic Adverse Selection Model with Asymptotic Efficiency,” *Mathematics of Operations Research*, **37**, 450–474.

A Missing Proof For Section 3

Proof of Theorem 1. Based on PK and the binding IC_L , we solve for u_h, u_l as a function of p_h, p_l and U :

$$u_h = \frac{U - (1 - \delta)p_h(qh + (1 - q)l)}{\delta}, \quad (15)$$

$$u_l = \frac{U - (1 - \delta)(p_h q(h - l) + p_l l)}{\delta}. \quad (16)$$

We want to show that an optimal policy is such that (i) either u_h as defined in (15) equals 0 or $p_h = 1$; and (ii) either u_l as defined in (16) equals $\bar{v}$ or $p_l = 0$. Write $W(U; p_h, p_l)$ for the maximum payoff from using p_h, p_l as probabilities of assigning the good, and using promised utilities as given by (15)–(16) (followed by the optimal policy from the period that follows). Substituting u_h and u_l into (OBJ), we get, from the fundamental theorem of calculus, for any fixed $p_h^1 < p_h^2$ such that the corresponding utilities u_h are interior,

$$W(U; p_h^2, p_l) - W(U; p_h^1, p_l) = \int_{p_h^1}^{p_h^2} \{(1 - \delta)q(h - c - (1 - q)(h - l)W'(u_l) - \bar{v}W'(u_h))\} dp_h.$$

This expression decreases (pointwise) in $W'(u_h)$ and $W'(u_l)$. Recall that $W'(u)$ is bounded from above by $1 - c/h$. Hence, plugging in the upper bound for W' , we obtain that $W(U; p_h^2, p_l) - W(U; p_h^1, p_l) \geq 0$. It follows that there is no loss (and possibly a gain) in increasing p_h , unless feasibility prevents this. An entirely analogous reasoning implies that $W(U; p_h, p_l)$ is nonincreasing in p_l .

It is immediate that $u_h \leq u_l$ and both u_h, u_l decreases in p_h, p_l . Therefore, either $u_h \geq 0$ binds or p_h equals 1. Similarly, either $u_l \leq \bar{v}$ binds or p_l equals 0. ■

Proof of Lemma 2. We start the proof with some notation and preliminary remarks. First, given any interval $I \subset [0, \mu]$, we write $I_h := \left[\frac{a-(1-\delta)\mu}{\delta}, \frac{b-(1-\delta)\mu}{\delta} \right] \cap [0, \mu]$ and $I_l := \left[\frac{a-(1-\delta)\underline{U}}{\delta}, \frac{b-(1-\delta)\underline{U}}{\delta} \right] \cap [0, \mu]$ where $I = [a, b]$; we also write $[a, b]_h$, etc. Furthermore we use the (ordered) sequence of subscripts to indicate the composition of such maps, *e.g.*, $I_{lh} = (I_l)_h$. Finally, given some interval I , we write $\ell(I)$ for its length.

Second, we note that, for any interval $I \subset [\underline{U}, \bar{U}]$, identically, for $U \in I$, it holds that

$$W(U) = (1 - \delta)(qh - c) + \delta q W\left(\frac{U - (1 - \delta)\mu}{\delta}\right) + \delta(1 - q) W\left(\frac{U - (1 - \delta)\underline{U}}{\delta}\right), \quad (17)$$

and hence, over this interval, it follows by differentiation that, a.e. on I ,

$$W'(U) = qW'(u_h) + (1 - q)W'(u_l).$$

Similarly, for any interval $I \subset [\bar{U}, \mu]$, identically, for $U \in I$,

$$W(U) = (1 - q)\left(U - c - (U - \mu)\frac{c}{l}\right) + (1 - \delta)q(\mu - c) + \delta q W\left(\frac{U - (1 - \delta)\mu}{\delta}\right), \quad (18)$$

and so a.e.,

$$W'(U) = -(1 - q)(c/l - 1) + qW'(u_h).$$

That is, the slope of W at a point (or an interval) is an average of the slopes at u_h, u_l , and this holds also on $[\bar{U}, \mu]$, with the convention that its slope at $u_l = \mu$ is given by $1 - c/l$. By weak concavity of W , if W is affine on I , then it must be affine on both I_h and I_l (with the convention that it is trivially affine at μ). We make the following observations.

1. For any $I \subseteq (\underline{U}, \mu)$ (of positive length) such that W is affine on I , $\ell(I_h \cap I) = \ell(I_l \cap I) = 0$. If not, then we note that, because the slope on I is the average of the other two, all three must have the same slope (since two intersect, and so have the same slope). But then the convex hull of the three has the same slope (by weak concavity). We thus obtain an interval $I' = \text{co}\{I_l, I_h\}$ of strictly greater length (note that $\ell(I_h) = \ell(I)/\delta$, and similarly $\ell(I_l) = \ell(I)/\delta$ unless I_l intersects μ). It must then be that I'_h or I'_l intersect I , and we can repeat this operation. This contradicts the fact the slope of W on $[0, \bar{U}]$ is $(1 - c/h)$, yet $W(\mu) = \mu - c$.
2. It follows that there is no interval $I \subseteq [\underline{U}, \mu]$ on which W has slope $(1 - c/h)$ (because then W would have this slope on $I' := \text{co}\{\{\underline{U}\} \cup I\}$, and yet I' would intersect I'_l .) Similarly, there cannot be an interval $I \subseteq [\underline{U}, \mu]$ on which W has slope $1 - c/l$.

3. It immediately follows from 2 that $W < \bar{W}$ on $(\underline{U}, \mu)$: if there is a $U \in (\underline{U}, \mu)$ such that $W(U) = \bar{W}(U)$, then by concavity again (and the fact that the two slopes involved are the two possible values of the slope of $\bar{W}$), W must either have slope $(1 - c/h)$ on $[0, U]$, or $1 - c/l$ on $[U, \mu]$, both being impossible.
4. Next, suppose that there exists an interval $I \subset [\underline{U}, \mu)$ of length $\varepsilon > 0$ such that W is affine on I . There might be many such intervals; consider the one with the smallest lower extremity. Furthermore, without loss, given this lower extremity, pick I so that it has maximum length, that W is affine on I , but on no proper superset of I . Let $I := [a, b]$. We claim that $I_h \in [0, \underline{U}]$. Suppose not. Note that I_h cannot overlap with I (by point 1). Hence, either I_h is contained in $[0, \underline{U}]$, or it is contained in $[\underline{U}, a]$, or $\underline{U} \in (a, b)_h$. This last possibility cannot occur, because W must be affine on $(a, b)_h$, yet the slope on $(a_h, \underline{U})$ is equal to $(1 - c/h)$, while by point 2 it must be strictly less on $(\underline{U}, b_h)$. It cannot be contained in $[\underline{U}, a]$, because $\ell(I_h) = \ell(I)/\delta > \ell(I)$, and this would contradict the hypothesis that I was the lowest interval in $[\underline{U}, \mu]$ of length ε over which W is affine.

We next observe that I_l cannot intersect I . Assume $b \leq \bar{U}$. Hence, we have that I_l is an interval over which W is affine, and such that $\ell(I_l) = \ell(I)/\delta$. Let $\varepsilon' := \ell(I)/\delta$. By the same reasoning as before, we can find $I' \subset [\underline{U}, \mu)$ of length $\varepsilon' > 0$ such that W is affine on I' , and such that $I'_h \subset [0, \bar{U}]$. Repeating the same argument as often as necessary, we conclude that there must be an interval $J \subset [\underline{U}, \mu)$ such that (i) W is affine on J , $J = [a', b']$, (ii) $b' \geq \bar{U}$, there exists no interval of equal or greater length in $[\underline{U}, \mu)$ over which W would be affine. By the same argument yet again, J_h must be contained in $[0, \underline{U}]$. Yet the assumption that $\delta > 1/2$ is equivalent to $\bar{U}_h > \underline{U}$, and so this is a contradiction. Hence, there exists no interval in $(\underline{U}, \mu)$ over which W is affine, and so W must be strictly concave.

This concludes the proof.

Differentiability follows from an argument that follows Benveniste and Scheinkman (1979), using some induction. We note that W is differentiable on $(0, \underline{U})$. Fix $U > \underline{U}$ such that $U_h \in (0, \underline{U})$. Consider the following perturbation of the optimal policy. Fix $\varepsilon = (p - \bar{p})^2$, for some $\bar{p} \in (0, 1)$ to be determined. With probability $\varepsilon > 0$, the report is ignored, the good is supplied with probability $p \in [0, 1]$ and the next value is U_l (Otherwise, the optimal policy is implemented). Because this event is independent of the report, the IC constraints are still satisfied. Note that, for $p = 0$, this yields a strictly lower utility than U to the agent, while it yields a strictly higher utility for $p = 1$. As it varies continuously, there is some critical value

–defined as $\bar{p}$ – that makes the agent indifferent between both policies. By varying p , we may thus generate all utilities within some interval $(U - \nu, U + \nu)$, for some $\nu > 0$, and the payoff $\tilde{W}$ that we obtain in this fashion is continuously differentiable in $U' \in (U - \nu, U + \nu)$. It follows that the concave function W is minimized by a continuously differentiable function $\tilde{W}$ –hence, it must be as well. ■

Proof of Lemma 4. We first consider the forecaster. We will rely on Lemma 8 from the continuous-time (Markovian) version of the game defined in Section 5.1. Specifically, consider a continuous-time model in which random shocks arrive according to a Poisson process at rate λ . Conditional on a shock, the agent’s value is h with probability q and l with the complementary probability. Both the shocks’ arrivals and the realized values are the agent’s private information. This is the same model as in Subsection 5.1 where $\lambda_h = \lambda(1-q)$, $\lambda_l = \lambda q$. The principal’s payoff W is the same as in Proposition 2. Let W^* denote the principal’s payoff if the shocks’ arrival times are *publicly* observed. Since the principal benefits from more information, his payoff weakly increases $W^* \geq W$. (The principal is guaranteed W by implementing the continuous-time limit of the policy specified in Theorem 2.) Given that both players are risk neutral, the model with random public arrivals is the same as the model in which shocks arrive at fixed intervals, $t = 1/\lambda, 2/\lambda, 3/\lambda, \dots$. This is effectively the discrete-time model with i.i.d. values in which the round length is $\Delta = 1/\lambda$ and the discount factor is $\delta = e^{-\frac{r}{\lambda}}$. Given that the loss is of the order $\mathcal{O}(r/\lambda)$ in the continuous-time private-shock model, the loss in the discrete-time i.i.d. model is of smaller order than $\mathcal{O}(1 - \delta)$.

Basing on the analysis above, we next show that the loss is of order $\mathcal{O}(1 - \delta)$. We consider an allocation problem in which the agent’s first-round type realization is private information whereas his type realization after the first round is public information. Let W^{**} denote the principal’s payoff in this problem, which is larger than the principal’s payoff in the benchmark model. Let U denote the promised utility before the first round and U_l, U_h the promised utilities after the agent reports l, h during round one. It is optimal to set $p_h = 1, p_l = 0$ during round one. From *PK* and binding *ICL*, we obtain

$$U_h = \frac{(\delta - 1)(qh - ql + l) + U}{\delta}, \quad U_l = \frac{(\delta - 1)q(h - l) + U}{\delta}.$$

The principal’s payoff given U is

$$(1 - \delta)q(h - c) + \delta (q\bar{W}(U_h) + (1 - q)\bar{W}(U_l)), \quad (19)$$

where $\bar{W}$ is the complete-information payoff function defined in Lemma 1. The principal’s payoff W^{**} is the maximum of (19) over U . It is easy to verify that the efficiency loss

$q(h - c) - W^{**}$ is proportional to $(1 - \delta)$. Therefore, the loss in the benchmark model has the order of $\mathcal{O}(1 - \delta)$.

We now consider the prophet. We divide the analysis in three stages. In the first two, we consider a fixed horizon $2N + 1$ and no discounting, as is usual. Let us start with the simplest case: a fixed number of copies $2N + 1$, and $q = 1/2$.³² Suppose that we relax the problem (so as to get a lower bound on the inefficiency). The number $m = 0, \dots, 2N + 1$, of high copies is drawn, and the information set $\{(m, 2N + 1 - m), (2N + 1 - m, m)\}$ is publicly revealed. That is, it is disclosed whether there are m high copies, of $N - m$ high copies (but nothing else).

The optimal mechanism consists of the collection of optimal mechanisms for each information set. We note that, because $q = 1/2$, both elements in the information set are equally likely. Hence, fixing $\{(m, 2N + 1 - m), (2N + 1 - m, m)\}$, with $m < N$, it must minimize the inefficiency

$$\min_{p_0, p_1, p_2} (1 - p_0)m(h - c) + (2N + 1 - 2m)\frac{(1 - p_1)(h - c) + p_1(c - l)}{2} + p_2m(c - l),$$

where p_0, p_1, p_2 are in $[0, 1]$. To understand this expression, we note that it is common knowledge that at least m units are high (hence, providing them with probability p_0 reduces the inefficiency $m(h - c)$ from these. It is also known that m are low, which if provided (with probability p_2) leads to inefficiency $m(c - l)$ and finally there are $2N + 1 - 2m$ units that are either high or low, and the choice p_1 in this respect implies one or the other inefficiency. This program is already simplified, as p_0, p_1, p_2 should be a function of the report (whether the state is $(m, 2N + 1 - m)$ or $(2N + 1 - m, m)$) subject to incentive-compatibility, but it is straightforward that both IC constraints bind and lead to the same choice of p_0, p_1, p_2 for both messages. In fact, it is also clear that $p_0 = 1$ and $p_2 = 0$, so for each information set, the optimal choice is given by the minimizer of

$$(2N + 1 - 2m)\frac{(1 - p_1)(h - c) + p_1(c - l)}{2} \geq (2N + 1 - 2m)\kappa,$$

where $\kappa = \min\{h - c, c - l\}$. Hence, the inefficiency is minimized by (adding up over all information sets)

$$\sum_{m=0}^N \binom{2N+1}{m} \left(\frac{1}{2}\right)^{2N+1} (2N + 1 - 2m)\kappa = \frac{\Gamma(N + \frac{3}{2})}{\sqrt{\pi}\Gamma(N + 1)}\kappa \rightarrow \frac{\sqrt{2N+1}}{\sqrt{2\pi}}\kappa.$$

³²We pick the number of copies as odd for simplicity. If not, let Nature reveal the event that all copies are high if this unlikely event occurs. This gives as lower bound for the inefficiency with $2N + 2$ copies the one we derive with $2N + 1$.

We now move on to the case where $q \neq 1/2$. Without loss of generality, assume $q > 1/2$. Consider the following public disclosure rule. Given the realized draw of high and lows, for any high copy, Nature publicly reveals it with probability $\lambda = 2 - 1/q$. Low copies are not revealed. Hence, if a copy is not revealed, the principal's posterior belief that it is high is

$$\frac{q(1-\lambda)}{q(1-\lambda) + (1-q)} = \frac{1}{2}.$$

Second, Nature reveals among the undisclosed balls (say, N' of those) whether the number of highs is m or $N' - m$, namely it discloses the information set $\{(m, N' - m), (N' - m, m)\}$, as before. Then the agent makes a report, etc. Conditional on all publicly revealed information, and given that both states are equally likely, the principal's optimal rule is again to pick a probability p_1 that minimizes

$$(N' - 2m) \frac{(1-p_1)(h-c) + p_1(c-l)}{2} \geq (N' - 2m)\kappa.$$

Hence, the total inefficiency is

$$\sum_{m=0}^{2N+1} \binom{2N+1}{m} q^m (1-q)^{2N+1-m} \left(\sum_{k=0}^m \binom{m}{k} \lambda^k (1-\lambda)^{m-k} |2N+1-k-2(m-k)| \right) \kappa,$$

since with k balls revealed, $N' = 2N+1-k$, and the uncertainty concerns whether there are (indeed) $m-k$ high values or low values. Alternatively, because the number of undisclosed copies is a compound Bernoulli, it is a Bernoulli random variable as well with parameter $q\lambda$, and so we seek to compute

$$\frac{1}{\sqrt{2N+1}} \sum_{m=0}^{2N+1} \binom{2N+1}{m} (q\lambda)^m (1-q\lambda)^{N+1-m} \frac{\Gamma(N-m+\frac{3}{2})}{\sqrt{\pi}\Gamma(N-m+1)} \kappa.$$

We note that

$$\begin{aligned} & \lim_{N \rightarrow \infty} \frac{1}{\sqrt{2N+1}} \sum_{m=0}^{2N+1} \binom{2N+1}{m} (q\lambda)^m (1-q\lambda)^{N+1-m} \frac{\Gamma(N-m+\frac{3}{2})}{\sqrt{\pi}\Gamma(N-m+1)} \\ &= \lim_{N \rightarrow \infty} \sum_{m=0}^{2N+1} \binom{2N+1}{m} (q\lambda)^m (1-q\lambda)^{N+1-m} \frac{\sqrt{2N-1-m}}{2\sqrt{N\pi}} \\ &= \sup_{\alpha > 0} \lim_{N \rightarrow \infty} \sum_{m=0}^{2N+1} \binom{2N+1}{m} (q\lambda)^m (1-q\lambda)^{N+1-m} \frac{\sqrt{2N-1-(2N+1)q\lambda(1+\alpha)}}{2\sqrt{N\pi}} \\ &= \sup_{\alpha > 0} \frac{\sqrt{1-(1+\alpha)q\lambda}}{\sqrt{2\pi}} = \frac{\sqrt{1-q\lambda}}{\sqrt{2\pi}} = \frac{\sqrt{1-q}}{\sqrt{\pi}}, \end{aligned}$$

hence the inefficiency converges to

$$\sqrt{2N+1} \frac{\sqrt{1-q}}{\sqrt{\pi}} \kappa.$$

Third, we consider the case of discounting. Note that, because the principal can always treat items separately, facing a problem with k i.i.d. copies, whose value l, h is scaled by a factor $1/k$ (along with the cost) is worth at least as much as one copy with a weight 1. Hence, if say, $\delta^m = 2\delta^k$, then modifying the discounted problem by replacing the unit with weight δ^m by two i.i.d. units with weight δ^k each makes the principal better off. Hence, we fix some small $\alpha > 0$, and consider N such that $\delta^N = \alpha$, i.e., $N = \ln \alpha / \ln \delta$. The principal's payoff is also increased if the values of all units after the N -th one are revealed for free. Hence, assume as much. Replacing each copy $k = 1, \dots, N$ by $\lfloor \delta^k / \delta^N \rfloor$ i.i.d. copies each with weight δ^N gives us as lower bound to the loss to the principal

$$\sup_{\alpha} \frac{\delta^N}{\sqrt{\sum_{k=1}^N \lfloor \delta^k / \delta^N \rfloor}},$$

and the right-hand side tends to a limit in excess of $\frac{1}{2\sqrt{1-\delta}}$ (use $\alpha = 1/2$ for instance). ■

B Missing Proof For Section 4

Proof of Lemma 5. Let W denote the set $\text{co}\{\bar{u}^\nu, \underline{u}^\nu : \nu \geq 0\}$. The point $\bar{u}^0$ is supported by $(p_h, p_l) = (1, 1), U(h) = U(l) = (\mu_h, \mu_l)$. For $\nu \geq 1$, $\bar{u}^\nu$ is supported by $(p_h, p_l) = (0, 0), U(h) = U(l) = \bar{u}^{\nu-1}$. The point $\underline{u}^0$ is supported by $(p_h, p_l) = (0, 0), U(h) = U(l) = (0, 0)$. For $\nu \geq 1$, $\underline{u}^\nu$ is supported by $(p_h, p_l) = (1, 1), U(h) = U(l) = \underline{u}^{\nu-1}$. Therefore, we have $W \subset \mathcal{B}(W)$. This implies that $\mathcal{B}(W) \subset V$.

We define four sequences as follows. First, for $\nu \geq 0$, let

$$\begin{aligned} \bar{w}_h^\nu &= \delta^\nu (1 - \kappa^\nu) (1 - q) \mu_l, \\ \bar{w}_l^\nu &= \delta^\nu (1 - q + \kappa^\nu q) \mu_l, \end{aligned}$$

and set $\bar{w}^\nu = (\bar{w}_h^\nu, \bar{w}_l^\nu)$. Second, for $\nu \geq 0$, let

$$\begin{aligned} \underline{w}_h^\nu &= \mu_h - \delta^\nu (1 - \kappa^\nu) (1 - q) \mu_l, \\ \underline{w}_l^\nu &= \mu_l - \delta^\nu (1 - q + \kappa^\nu q) \mu_l, \end{aligned}$$

and set $\underline{w}^\nu = (\underline{w}_h^\nu, \underline{w}_l^\nu)$. For any $\nu \geq 1$, $\bar{w}^\nu$ is supported by $(p_h, p_l) = (0, 0), U(h) = U(l) = \bar{w}^{\nu-1}$, and $\underline{w}^\nu$ is supported by $(p_h, p_l) = (1, 1), U(h) = U(l) = \underline{w}^{\nu-1}$. The sequence $\bar{w}^\nu$ starts

at $\bar{w}^0 = (0, \mu_l)$ with $\lim_{\nu \rightarrow \infty} \bar{w}^\nu = 0$. Similarly, $\underline{w}^\nu$ starts at $\underline{w}^0 = (\mu_h, 0)$ and $\lim_{\nu \rightarrow \infty} \underline{w}^\nu = \mu$. We define a set sequence as follows:

$$W^\nu = \text{co} \left(\{\bar{u}^k, \underline{u}^k : 0 \leq k \leq \nu\} \cup \{\bar{w}^\nu, \underline{w}^\nu\} \right).$$

It is obvious that $V \subset \mathcal{B}(W^0) \subset W^0$. To prove that $V = W$, it suffices to show that $W^\nu = \mathcal{B}(W^{\nu-1})$ and $\lim_{\nu \rightarrow \infty} W^\nu = W$.

For any $\nu \geq 1$, we define the supremum score in direction (λ_1, λ_2) given $W^{\nu-1}$ as $K((\lambda_1, \lambda_2), W^{\nu-1}) = \sup_{p_h, p_l, U(h), U(l)} (\lambda_1 U_h + \lambda_2 U_l)$, subject to (2)–(5), $p_h, p_l \in [0, 1]$, and $U(h), U(l) \in W^{\nu-1}$. The set $\mathcal{B}(W^{\nu-1})$ is given by

$$\bigcap_{(\lambda_1, \lambda_2)} \{(U_h, U_l) : \lambda_1 U_h + \lambda_2 U_l \leq K((\lambda_1, \lambda_2), W^{\nu-1})\}.$$

Without loss of generality, we focus on directions $(1, -\lambda)$ and $(-1, \lambda)$ for all $\lambda \geq 0$. We define three sequences of slopes as follows:

$$\begin{aligned} \lambda_1^\nu &= \frac{(1-q)(\delta\kappa-1)\kappa^\nu(\mu_h-\mu_l)-(1-\delta)(q\mu_h+(1-q)\mu_l)}{q(1-\delta\kappa)\kappa^\nu(\mu_h-\mu_l)-(1-\delta)(q\mu_h+(1-q)\mu_l)} \\ \lambda_2^\nu &= \frac{1-(1-q)(1-\kappa^\nu)}{q(1-\kappa^\nu)} \\ \lambda_3^\nu &= \frac{(1-q)(1-\kappa^\nu)}{q\kappa^\nu+(1-q)}. \end{aligned}$$

It is easy to verify that

$$\lambda_1^\nu = \frac{\bar{u}_h^\nu - \bar{u}_h^{\nu+1}}{\bar{u}_l^\nu - \bar{u}_l^{\nu+1}} = \frac{\underline{u}_h^\nu - \underline{u}_h^{\nu+1}}{\underline{u}_l^\nu - \underline{u}_l^{\nu+1}}, \quad \lambda_2^\nu = \frac{\bar{u}_h^\nu - \bar{w}_h^\nu}{\bar{u}_l^\nu - \bar{w}_l^\nu} = \frac{\underline{u}_h^\nu - \underline{w}_h^\nu}{\underline{u}_l^\nu - \underline{w}_l^\nu}, \quad \lambda_3^\nu = \frac{\bar{w}_h^\nu - 0}{\bar{w}_l^\nu - 0} = \frac{\underline{w}_h^\nu - \mu_h}{\underline{w}_l^\nu - \mu_l}.$$

When $(\lambda_1, \lambda_2) = (-1, \lambda)$, the supremum score as we vary λ is

$$K((-1, \lambda), W^{\nu-1}) = \begin{cases} (-1, \lambda) \cdot (0, 0) & \text{if } \lambda \in [0, \lambda_3^\nu] \\ (-1, \lambda) \cdot \bar{w}^\nu & \text{if } \lambda \in [\lambda_3^\nu, \lambda_2^\nu] \\ (-1, \lambda) \cdot \bar{u}^\nu & \text{if } \lambda \in [\lambda_2^\nu, \lambda_1^{\nu-1}] \\ (-1, \lambda) \cdot \bar{u}^{\nu-1} & \text{if } \lambda \in [\lambda_1^{\nu-1}, \lambda_1^{\nu-2}] \\ \dots & \\ (-1, \lambda) \cdot \bar{u}^0 & \text{if } \lambda \in [\lambda_1^0, \infty) \end{cases}$$

Similarly, when $(\lambda_1, \lambda_2) = (1, -\lambda)$, we have

$$K((1, -\lambda), W^{\nu-1}) = \begin{cases} (1, -\lambda) \cdot (\mu_h, \mu_l) & \text{if } \lambda \in [0, \lambda_3^\nu] \\ (1, -\lambda) \cdot \underline{w}^\nu & \text{if } \lambda \in [\lambda_3^\nu, \lambda_2^\nu] \\ (1, -\lambda) \cdot \underline{u}^\nu & \text{if } \lambda \in [\lambda_2^\nu, \lambda_1^{\nu-1}] \\ (1, -\lambda) \cdot \underline{u}^{\nu-1} & \text{if } \lambda \in [\lambda_1^{\nu-1}, \lambda_1^{\nu-2}] \\ \dots & \\ (1, -\lambda) \cdot \underline{u}^0 & \text{if } \lambda \in [\lambda_1^0, \infty). \end{cases}$$

Therefore, we have $W^\nu = \mathcal{B}(W^{\nu-1})$. Note that this method only works when parameters are such that $\lambda_3^\nu \leq \lambda_2^\nu \leq \lambda_1^{\nu-1}$ for all $\nu \geq 1$. If $\rho_l/(1 - \rho_h) \geq l/h$, the proof stated above applies. Otherwise, the following proof applies.

We define four sequences as follows. First, for $0 \leq m \leq \nu$, let

$$\begin{aligned} \overline{w}_h(m, \nu) &= \delta^{\nu-m} (q\mu_h (1 - \delta^m) + (1 - q)\mu_l) - (1 - q)(\delta\kappa)^{\nu-m} (\mu_h ((\delta\kappa)^m - 1) + \mu_l), \\ \overline{w}_l(m, \nu) &= \delta^{\nu-m} (q\mu_h (1 - \delta^m) + (1 - q)\mu_l) + q(\delta\kappa)^{\nu-m} (\mu_h ((\delta\kappa)^m - 1) + \mu_l), \end{aligned}$$

and set $\overline{w}(m, \nu) = (\overline{w}_h(m, \nu), \overline{w}_l(m, \nu))$. Second, for $0 \leq m \leq \nu$, let

$$\begin{aligned} \underline{w}_h(m, \nu) &= \frac{(1 - q)\delta^\nu \kappa^\nu (\mu_h (\delta^m \kappa^m - 1) + \mu_l) + \kappa^m (\mu_h \delta^m - \delta^\nu (q\mu_h (1 - \delta^m) + (1 - q)\mu_l))}{\delta^m \kappa^m}, \\ \underline{w}_l(m, \nu) &= \frac{-q\delta^\nu \kappa^\nu (\mu_h (\delta^m \kappa^m - 1) + \mu_l) + \kappa^m (\mu_l \delta^m - \delta^\nu (q\mu_h (1 - \delta^m) + (1 - q)\mu_l))}{\delta^m \kappa^m}, \end{aligned}$$

and set $\underline{w}(m, \nu) = (\underline{w}_h(m, \nu), \underline{w}_l(m, \nu))$. Fixing ν , the sequence $\overline{w}(m, \nu)$ is increasing (in both its arguments) as m increases, with $\lim_{\nu \rightarrow \infty} \overline{w}(\nu - m, \nu) = \overline{u}^m$. Similarly, fixing ν , $\underline{w}(m, \nu)$ is decreasing as m increases, $\lim_{\nu \rightarrow \infty} \underline{w}(\nu - m, \nu) = \underline{u}^m$.

Let $\overline{W}(\nu) = \{\overline{w}(m, \nu) : 0 \leq m \leq \nu\}$ and $\underline{W}(\nu) = \{\underline{w}(m, \nu) : 0 \leq m \leq \nu\}$. We define a set sequence as follows:

$$W(\nu) = \text{co}(\{(0, 0), (\mu_h, \mu_l)\} \cup \overline{W}(\nu) \cup \underline{W}(\nu)).$$

Since $W(0)$ equals $[0, \mu_h] \times [0, \mu_l]$, it is obvious that $V \subset \mathcal{B}(W(0)) \subset W(0)$. To prove that $V = W := \text{co}\{\overline{u}^\nu, \underline{u}^\nu : \nu \geq 0\}$, it suffices to show that $W(\nu) = \mathcal{B}(W(\nu - 1))$ and $\lim_{\nu \rightarrow \infty} W(\nu) = W$. The rest of the proof is similar to the first part and hence omitted. ■

Proof of Lemma 6. It will be useful in this proof and those that follows to define the operator $\mathcal{B}_{ij}$, $i, j = 0, 1$. Given an arbitrary $A \subset [0, \mu_h] \times [0, \mu_l]$, let

$$\mathcal{B}_{ij}(A) := \{(U_h, U_l) \in [0, \mu_h] \times [0, \mu_l] : U(h) \in A, U(l) \in A \text{ solving (2)–(5) for } (p_h, p_l) = (i, j)\},$$

and similarly $\mathcal{B}_i(A), \mathcal{B}_j(A)$ when only p_h or p_l is constrained.

The first step is to compute V_0 , the largest set such that $V_0 \subset \mathcal{B}_0(V_0)$. Plainly, this is a proper subset of V , because any promise $U_l \in (\delta\rho_l\mu_h + \delta(1 - \rho_l)\mu_l, \mu_l]$ requires that p_l be strictly positive.

Note that the sequence $\{v^\nu\}$ solves the system of equations, for all $\nu \geq 1$:

$$\begin{cases} v_h^{\nu+1} = \delta(1 - \rho_h)v_h^\nu + \delta\rho_h v_l^\nu \\ v_l^{\nu+1} = \delta(1 - \rho_l)v_l^\nu + \delta\rho_l v_h^\nu, \end{cases}$$

and $v_l^1 = v_l^0$ (From $v_l^1 = v_l^0$ and the second equation for $\nu = 0$, we obtain that v^0 lies on the line $U_l = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h$.) In words, the utility vector $v^{\nu+1}$ obtains by setting $p_h = p_l = 0$, choosing as a continuation payoff vector $U(l) = v^\nu$, and assuming that IC_H binds (so that the high type's utility can be derived from the report l). To prove that these vectors are incentive feasible using such a scheme, it remains to exhibit $U(h)$ and show that it satisfies IC_L . In addition, we must argue that $U(h) \in \mu$. We prove by construction. Pick any v^ν such that $\nu \geq 1$. Once we fix a $p_h \in [0, 1]$, PK_H requires that $U(h)$ must lie on the line $\delta(1 - \rho_h)U_h(h) + \delta\rho_h U_l(h) = v_h^\nu - \delta p_h h$. There exists a unique p_h , denoted p_h^ν , such that v^ν lies on the same line as $U(h)$ does, that is

$$\delta(1 - \rho_h)U_h(h) + \delta\rho_h U_l(h) = v_h^\nu - \delta p_h^\nu h = \delta(1 - \rho_h)v_h^\nu + \delta\rho_h v_l^\nu.$$

It is easy to verify that

$$p_h^\nu = \delta^\nu (1 - (1 - q)(1 - \kappa^\nu)) \frac{v_h^0}{v_h^*}.$$

Given that $v_h^0 \leq v_h^*$, we have $p_h^\nu \in [0, 1]$. Substituting p_h^ν into PK_H and IC_L , we want to show that there exists $U(h) \in \mu$ such that both PK_H and IC_L are satisfied. It is easy to verify that the intersection of PK_H and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ is below the intersection of the binding IC_L and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$. Therefore, the intersection of PK_H and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ satisfies both PK_H and IC_L . In addition, the constructed PK_H goes through the boundary point v^ν , so the intersection of PK_H and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ is inside μ .

Finally, we must show that the point v^0 can itself be obtained with continuation payoffs in μ . That one is obtained by setting $(p_h, p_l) = (1, 0)$, set IC_L as a binding constraint, and $U(l) = v^0$ (again one can check as above that $U(h)$ is in μ and that IC_H holds). This suffices to show that $\mu \subseteq V_0$, because this establishes that the extreme points of μ can be sustained with continuation payoffs in the set, and all other utility vectors in μ can be written as a convex combination of these extreme points.

The proof that $V_0 \subset \mu$ follows the same lines as determining the boundaries of V in the proof of Lemma 5: one considers a sequence of (less and less) relaxed programs, setting $\hat{W}^0 = V$ and defining recursively the supremum score in direction (λ_1, λ_2) given $\hat{W}^{\nu-1}$ as $K((\lambda_1, \lambda_2), W^{\nu-1}) = \sup_{p_h, p_l, U(h), U(l)} \lambda_1 U_h + \lambda_2 U_l$, subject to (2)–(5), $p_h, p_l \in [0, 1]$, and $U(h), U(l) \in \hat{W}^{\nu-1}$. The set $\mathcal{B}(\hat{W}^{\nu-1})$ is given by

$$\bigcap_{(\lambda_1, \lambda_2)} \left\{ (U_h, U_l) \in V : \lambda_1 U_h + \lambda_2 U_l \leq K((\lambda_1, \lambda_2), W^{\nu-1}) \right\},$$

and the set $\hat{W}^\nu = \mathcal{B}(\hat{W}^{\nu-1})$ obtains by considering an appropriate choice of λ_1, λ_2 . More precisely, we always set $\lambda_2 = 1$, and for $\nu = 0$, pick $\lambda_1 = 0$. This gives $\hat{W}^1 = V \cap \{U : U_l \leq v_l^0, U_l \geq \frac{v_l^1 - v_l^2}{v_h^1 - v_h^2}(U_h - v_h^2)\}$. We then pick (for every $\nu \geq 1$) as direction λ the vector $(\lambda_1 1, 1) \cdot (1, (v_l^\nu - v_l^{\nu+1})/(v_h^\nu - v_h^{\nu+1}))$, and as result obtain that

$$\mu \subseteq \hat{W}^{\nu+1} = \hat{W}^\nu \cap \left\{ U : U_l \geq \frac{v_l^{\nu+1} - v_l^{\nu+2}}{v_h^{\nu+1} - v_h^{\nu+2}}(U_h - v_h^{\nu+2}) \right\}.$$

It follows that $\mu \subseteq \text{co}\{(0, 0)\} \cup \{v^\nu\}_{\nu \geq 0}$.

Next, we argue that this achieves the complete-information payoff. First, note that $\mu \subseteq V \cap \{U : U_l \leq v_l^*\}$. In this region, it is clear that any policy that never gives the unit to the low type while delivering the promised utility to the high type must be optimal. This is a feature of the policy that we have described to obtain the boundary of V (and plainly it extends to utilities U below this boundary).

Finally, one must show that above it the complete-information payoff cannot be achieved. It follows from the definition of μ as the largest fixed point of $\mathcal{B}_0$ that starting from any utility vector $U \in V \setminus \mu$, $U \neq \mu$, there is a positive probability that the unit is given (after some history that has positive probability) to the low type. This implies that the complete-information payoff cannot be achieved in case $U \leq v^*$. For $U \geq v^*$, achieving the complete-information payoff requires that $p_h = 1$ for all histories, but it is not hard to check that the smallest fixed point of $\mathcal{B}_1$ is not contained in $V \cap \{U : U \geq v^*\}$, from which it follows that suboptimal continuation payoffs are collected with positive probability. ■

Proof of Theorem 2 and 3. We start the proof by defining the function $W : V \times \{\rho_l, 1 - \rho_h\} \rightarrow \mathbf{R} \cup \{-\infty\}$, that solves the following program, for all $(U_h, U_l) \in V$, and $\mu \in \{\rho_l, 1 - \rho_h\}$,

$$\begin{aligned} W(U_h, U_l, \mu) &= \sup \left\{ \mu ((1 - \delta)p_h(h - c) + \delta W(U_h(h), U_l(h), 1 - \rho_h)) \right. \\ &\quad \left. + (1 - \mu) ((1 - \delta)p_l(l - c) + \delta W(U_h(l), U_l(l), \rho_l)) \right\}, \end{aligned}$$

over $(p_l, p_h) \in [0, 1]^2$, and $U(h), U(l) \in V$ subject to PK_H, PK_L, IC_L . Note that IC_H is dropped so this is a relaxed problem. We characterize the optimal policy and value function for this relaxed problem and relate the results to the original optimization problem. Note that for both problems the optimal policy for a given (U_h, U_l) is independent of μ as μ appears in the objective function additively and does not appear in constraints. Also note that the first best is achieved when $U \in \underline{V}$. So, we focus on the subset $V \setminus \underline{V}$.

1. We want to show that for any U , it is optimal to set p_h, p_l as in (14) and to choose $U(h)$ and $U(l)$ that lie on P_b . It is feasible to choose such a $U(h)$ as the intersection of IC_L and PK_H lies above P_b . It is also feasible to choose such a $U(l)$ as IC_H is dropped. To show that it is optimal to choose $U(h), U(l) \in P_b$, we need to show that $W(U_h, U_l, 1 - \rho_h)$ (resp., $W(U_h, U_l, \rho_l)$) is weakly increasing in U_h along the rays $x = (1 - \rho_h)U_h + \rho_h U_l$ (resp., $y = \rho_l U_h + (1 - \rho_l)U_l$). Let $\tilde{W}$ denote the value function from implementing the policy above.
2. Let $(U_{h1}(x), U_{l1}(x))$ be the intersection of P_b and the line $x = (1 - \rho_h)U_h + \rho_h U_l$. We define function $w_h(x) := \tilde{W}(U_{h1}(x), U_{l1}(x), 1 - \rho_h)$ on the domain $[0, (1 - \rho_h)\mu_h + \rho_h \mu_l]$. Similarly, let $(U_{h2}(y), U_{l2}(y))$ be the intersection of P_b and the line $y = \rho_l U_h + (1 - \rho_l)U_l$. We define $w_l(y) := \tilde{W}(U_{h2}(y), U_{l2}(y), \rho_l)$ on the domain $[0, \rho_l \mu_h + (1 - \rho_l)\mu_l]$. For any U , let $X(U) = (1 - \rho_h)U_h + \rho_h U_l$ and $Y(U) = \rho_l U_h + (1 - \rho_l)U_l$. We want to show that (i) $w_h(x)$ (resp., $w_l(y)$) is concave in x (resp., y); (ii) w'_h, w'_l is bounded from below by $1 - c/l$ (derivatives have to be understood as either right- or left-derivatives, depending on the inequality); and (iii) for any U on P_b

$$w'_h(X(U)) \geq w'_l(Y(U)). \quad (20)$$

Note that we have $w'_h(X(U)) = w'_l(Y(U)) = 1 - c/h$ when $U \in \mu$. For any fixed $U \in P_b \setminus (\mu \cup V_h)$, a high report leads to $U(h)$ such that $(1 - \rho_h)U_h(h) + \rho_h U_l(h) = (U_h - (1 - \delta)h)/\delta$ and $U(h)$ is lower than U . Also, a low report leads to $U(l)$ such that $\rho_l U_h(l) + (1 - \rho_l)U_l(l) = U_l/\delta$ and $U(l)$ is higher than U if $U \in P_b \setminus (\mu \cup V_h)$. Given the definition of w_h, w_l , we have

$$\begin{aligned} w'_h(x) &= (1 - \rho_h)U'_{h1}(x)w'_h\left(\frac{U_{h1}(x) - (1 - \delta)h}{\delta}\right) + \rho_h U'_{l1}(x)w'_l\left(\frac{U_{l1}(x)}{\delta}\right) \\ w'_l(y) &= \rho_l U'_{h2}(y)w'_h\left(\frac{U_{h2}(y) - (1 - \delta)h}{\delta}\right) + (1 - \rho_l)U'_{l2}(y)w'_l\left(\frac{U_{l2}(y)}{\delta}\right). \end{aligned}$$

If x, y are given by $X(U), Y(U)$, it follows that $(U_{h1}(x), U_{l1}(y)) = (U_{h2}(y), U_{l2}(y))$ and hence

$$\begin{aligned} w'_h \left(\frac{U_{h1}(x) - (1 - \delta)h}{\delta} \right) &= w'_h \left(\frac{U_{h2}(y) - (1 - \delta)h}{\delta} \right) \\ w'_l \left(\frac{U_{l1}(x)}{\delta} \right) &= w'_l \left(\frac{U_{l2}(y)}{\delta} \right). \end{aligned}$$

Next, we want to show that for any $U \in P_b$ and $x = X(U), y = Y(U)$

$$\begin{aligned} (1 - \rho_h)U'_{h1}(x) + \rho_h U'_{l1}(x) &= \rho_l U'_{h2}(y) + (1 - \rho_l)U'_{l2}(y) = 1 \\ (1 - \rho_h)U'_{h1}(x) - \rho_l U'_{h2}(y) &\geq 0. \end{aligned}$$

This can be shown by assuming that U is on the line segment $U_h = aU_l + b$. For any $a > 0$, the equalities/inequality above hold. The concavity of w_h, w_l can be shown by taking the second derivative

$$\begin{aligned} w''_h(x) &= (1 - \rho_h)U'_{h1}(x)w''_h \left(\frac{U_{h1}(x) - (1 - \delta)h}{\delta} \right) \frac{U'_{h1}(x)}{\delta} + \rho_h U'_{l1}(x)w''_l \left(\frac{U_{l1}(x)}{\delta} \right) \frac{U'_{l1}(x)}{\delta} \\ w''_l(y) &= \rho_l U'_{h2}(y)w''_h \left(\frac{U_{h2}(y) - (1 - \delta)h}{\delta} \right) \frac{U'_{h2}(y)}{\delta} + (1 - \rho_l)U'_{l2}(y)w''_l \left(\frac{U_{l2}(y)}{\delta} \right) \frac{U'_{l2}(y)}{\delta}. \end{aligned}$$

Here, we use the fact that $U_{h1}(x), U_{l1}(x)$ (resp., $U_{h2}(y), U_{l2}(y)$) are piece-wise linear in x (resp., y). For any fixed $U \in P_b \cap V_h$ and $x = X(U), y = Y(U)$, we have

$$\begin{aligned} w'_h(x) &= (1 - \rho_h)U'_{h1}(x)w'_h \left(\frac{U_{h1}(x) - (1 - \delta)h}{\delta} \right) + \rho_h U'_{l1}(x) \frac{l - c}{l} \\ w'_l(y) &= \rho_l U'_{h2}(y)w'_h \left(\frac{U_{h2}(y) - (1 - \delta)h}{\delta} \right) + (1 - \rho_l)U'_{l2}(y) \frac{l - c}{l}. \end{aligned}$$

Inequality (20) and the concavity of w_h, w_l can be shown similarly. To sum up, if w_h, w_l satisfy properties (i), (ii) and (iii), they also do after one iteration.

3. Let $\mathcal{W}$ be the set of $W(U_h, U_l, 1 - \rho_h)$ and $W(U_h, U_l, \rho_l)$ such that

- (a) $W(U_h, U_l, 1 - \rho_h)$ (resp., $W(U_h, U_l, \rho_l)$) is weakly increasing in U_h along the rays $x = (1 - \rho_h)U_h + \rho_h U_l$ (resp., $y = \rho_l U_h + (1 - \rho_l)U_l$);
- (b) $W(U_h, U_l, 1 - \rho_h)$ and $W(U_h, U_l, \rho_l)$ coincide with $\tilde{W}$ on P_b .
- (c) $W(U_h, U_l, 1 - \rho_h)$ and $W(U_h, U_l, \rho_l)$ coincide with $\tilde{W}$ on μ ;

If we pick $W_0(U_h, U_l, \mu) \in \mathcal{W}$ as the continuation value function, the conjectured policy is optimal. Note that it is optimal to choose p_h, p_l according to (14) because w'_h, w'_l are in the interval $[1 - c/l, 1 - c/h]$. We want to show that the new value function W_1 is also in $\mathcal{W}$. Property (b) and (c) are trivially satisfied. We need to prove property (a) for $\mu \in \{1 - \rho_h, \rho_l\}$. That is,

$$W_1(U_h + \varepsilon, U_l, \mu) - W_1(U_h, U_l, \mu) \geq W_1(U_h, U_l + \frac{1 - \rho_h}{\rho_h} \varepsilon, \mu) - W_1(U_h, U_l, \mu). \quad (21)$$

We start with the case in which $\mu = 1 - \rho_h$. The left-hand side equals

$$\delta(1 - \rho_h) \left(W_0(\tilde{U}_h(h), \tilde{U}_l(h), 1 - \rho_h) - W_0(U_h(h), U_l(h), 1 - \rho_h) \right), \quad (22)$$

where $\tilde{U}(h)$ and $U(h)$ are on P_b and

$$\begin{aligned} (1 - \delta)h + \delta \left((1 - \rho_h)\tilde{U}_h(h) + \rho_h\tilde{U}_l(h) \right) &= U_h + \varepsilon, \\ (1 - \delta)h + \delta \left((1 - \rho_h)U_h(h) + \rho_hU_l(h) \right) &= U_h. \end{aligned}$$

For any fixed $U \in V \setminus (\mu \cup V_h)$, the right-hand side equals

$$\delta\rho_h \left(W_0(\tilde{U}_h(l), \tilde{U}_l(l), \rho_l) - W_0(U_h(l), U_l(l), \rho_l) \right), \quad (23)$$

where $\tilde{U}(l)$ and $U(l)$ are on P_b and

$$\begin{aligned} \delta \left(\rho_l\tilde{U}_h(l) + (1 - \rho_l)\tilde{U}_l(l) \right) &= U_l + \frac{1 - \rho_h}{\rho_h} \varepsilon, \\ \delta \left(\rho_lU_h(l) + (1 - \rho_l)U_l(l) \right) &= U_l. \end{aligned}$$

We need to show that (28) is greater than (29). Note that $U(h), \tilde{U}(h), U(l), \tilde{U}(l)$ are on P_b , so only the properties of w_h, w_l are needed. Inequality (21) is equivalent to

$$w'_h \left(\frac{U_h - (1 - \delta)h}{\delta} \right) \geq w'_l \left(\frac{U_l}{\delta} \right), \quad \forall (U_h, U_l) \in V \setminus (\mu \cup V_h \cup V_l). \quad (24)$$

The case in which $\mu = \rho_l$ leads to the same inequality as above. Given that w_h, w_l are concave, w'_h, w'_l are decreasing. Therefore, we only need to show that inequality (24) holds when (U_h, U_l) are on P_b . This is true given that (i) w_h, w_l are concave; (ii) inequality (20) holds; (iii) $(U_h - (1 - \delta)h)/\delta$ corresponds to a lower point on P_b than U_l/δ does. When $U \in V_h$, the right-hand side of (21) is given by $(1 - \rho_h)\varepsilon(1 - c/l)$. Inequality (21) is equivalent to $w'_h((U_h - (1 - \delta)h)/\delta) \geq 1 - c/l$, which is obviously true. Similar analysis applies to the case in which $U \in V_l$.

This shows that the optimal policy for the relaxed problem is indeed the conjectured policy and $\tilde{W}$ is the value function. The maximum is achieved on P_b and the continuation utility never leaves P_b . Given that this optimal mechanism does not violate IC_H , it is the optimal mechanism of our original problem.

We are back to the original optimization problem. The first observation is that we can decompose the optimization problem into two sub-problems: (i) choose $p_h, U(h)$ to maximize $(1 - \delta)p_h(h - c) + \delta W(U_h(h), U_l(h), 1 - \rho_h)$ subject to PK_H and IC_L ; (ii) choose $p_l, U(l)$ to maximize $(1 - \delta)p_l(l - c) + \delta W(U_h(l), U_l(l), \rho_l)$ subject to PK_L and IC_H . We want to show that the conjecture policy with respect to $p_h, U(h)$ is the optimal solution to the first sub-problem. This can be shown by taken the value function $\tilde{W}$ as the continuation value function. We know that the conjecture policy is optimal given $\tilde{W}$ because (i) it is always optimal to choose $U(h)$ that lies on P_b due to property (a); (ii) it is optimal to set p_h to be 1 because w'_h lies in $[1 - c/l, 1 - c/h]$. The conjecture policy solves the first sub-problem because (i) $\tilde{W}$ is weakly higher than the true value function point-wise; (ii) $\tilde{W}$ coincides with the true value function on P_b . The analysis above also implies that IC_H binds for $U \in V_t$. Next, we show that the conjecture policy is the solution to the second sub-problem.

For a fixed $U \in V_t$, PK_L and IC_H determines $U_h(l), U_l(l)$ as a function of p_l . Let γ_h, γ_l denote the derivative of $U_h(l), U_l(l)$ with respect to p_l

$$\gamma_h = \frac{(1 - \delta)(l\rho_h - h(1 - \rho_l))}{\delta(1 - \rho_h - \rho_l)}, \quad \gamma_l = \frac{(1 - \delta)(h\rho_l - l(1 - \rho_h))}{\delta(1 - \rho_h - \rho_l)}.$$

It is easy to verify that $\gamma_h < 0$ and $\gamma_h + \gamma_l < 0$. We want to show that it is optimal to set p_l to be zero. That is, among all feasible $p_l, U_h(l), U_l(l)$ satisfying PK_L and IC_H , the principal's payoff from the low type, $(1 - \delta)p_l(l - c) + \delta W(U_h(l), U_l(l), \rho_l)$, is the highest when $p_l = 0$. It is sufficient to show that within the feasible set

$$\gamma_h \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_h} + \gamma_l \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_l} \leq \frac{(1 - \delta)(c - l)}{\delta}, \quad (25)$$

where the left-hand side is the directional derivative of $W(U_h, U_l, \rho_l)$ along the vector (γ_h, γ_l) . We first show that (25) holds for all $U \in V_b$. For any fixed $U \in V_b$, we have

$$W(U_h, U_l, \rho_l) = \rho_l \left((1 - \delta)(h - c) + \delta w_h \left(\frac{U_h - (1 - \delta)h}{\delta} \right) \right) + (1 - \rho_l) \delta w_l \left(\frac{U_l}{\delta} \right).$$

It is easy to verify that $\partial W / \partial U_h = \rho_l w'_h$ and $\partial W / \partial U_l = (1 - \rho_l) w'_l$. Using the fact that $w'_h \geq w'_l$ and $w'_h, w'_l \in [1 - c/l, 1 - c/h]$, we prove that (25) follows. Using similar arguments, we can show that (25) holds for all $U \in V_h$.

Note that $W(U_h, U_l, \rho_l)$ is concave on V . Therefore, its directional derivative along the vector (γ_h, γ_l) is monotone. For any fixed (U_h, U_l) on P_b , we have

$$\begin{aligned} & \lim_{\varepsilon \rightarrow 0} \frac{\gamma_h \frac{\partial W(U_h + \gamma_h \varepsilon, U_l + \gamma_l \varepsilon, \rho_l)}{\partial U_h} + \gamma_l \frac{\partial W(U_h + \gamma_h \varepsilon, U_l + \gamma_l \varepsilon, \rho_l)}{\partial U_l} - \left(\gamma_h \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_h} + \gamma_l \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_l} \right)}{\varepsilon} \\ &= \gamma_h^2 \frac{\rho_l}{\delta} w_h'' \left(\frac{U_h - (1 - \delta)h}{\delta} \right) + \gamma_l^2 \frac{1 - \rho_l}{\delta} w_l'' \left(\frac{U_l}{\delta} \right) \leq 0. \end{aligned}$$

The last inequality follows as w_h, w_l are concave. Given that (γ_h, γ_l) points towards the interior of V , (25) holds within V .

For any $x \in [0, (1 - \rho_h)\mu_h + \rho_h\mu_l]$, let $z(x)$ be $\rho_l U_{h1}(x) + (1 - \rho_l)U_{l1}(x)$. The function $z(x)$ is piecewise linear with z' being positive and increasing in x . Let μ_0 denote the prior belief of the high type. We want to show that the maximum of $\mu_0 W(U_h, U_l, 1 - \rho_h) + (1 - \mu_0)W(U_h, U_l, \rho_l)$ is achieved on P_b for any prior μ_0 . Suppose not. Suppose $(\tilde{U}_h, \tilde{U}_l) \in V \setminus P_b$ achieves the maximum. Let U^0 (resp., U^1) denote the intersection of P_b and $(1 - \rho_h)U_h + \rho_h U_l = (1 - \rho_h)\tilde{U}_h + \rho_h \tilde{U}_l$ (resp., $\rho_l U_h + (1 - \rho_l)U_l = \rho_l \tilde{U}_h + (1 - \rho_l)\tilde{U}_l$). It is easily verified that $U^0 < U^1$. Given that $(\tilde{U}_h, \tilde{U}_l)$ achieves the maximum, it must be true that

$$\begin{aligned} W(U_h^1, U_l^1, 1 - \rho_h) - W(U_h^0, U_l^0, 1 - \rho_h) &< 0 \\ W(U_h^1, U_l^1, \rho_l) - W(U_h^0, U_l^0, \rho_l) &> 0. \end{aligned}$$

We show that this is impossible by arguing that for any $U^0, U^1 \in P_b$ and $U^0 < U^1$, $W(U_h^1, U_l^1, 1 - \rho_h) - W(U_h^0, U_l^0, 1 - \rho_h) < 0$ implies that $W(U_h^1, U_l^1, \rho_l) - W(U_h^0, U_l^0, \rho_l) < 0$. It is without loss to assume that U^0, U^1 are on the same line segment $U_h = aU_l + b$. It follows that

$$\begin{aligned} W(U_h^1, U_l^1, 1 - \rho_h) - W(U_h^0, U_l^0, 1 - \rho_h) &= \int_{s^0}^{s^1} w_h'(s) ds \\ W(U_h^1, U_l^1, \rho_l) - W(U_h^0, U_l^0, \rho_l) &= z'(s) \int_{s^0}^{s^1} w_l'(z(s)) ds, \end{aligned}$$

where $s^0 = (1 - \rho_h)U_h^0 + \rho_h U_l^0$ and $s^1 = (1 - \rho_h)U_h^1 + \rho_h U_l^1$. Given that $w_h'(s) \geq w_l'(z(s))$ and $z'(s) > 0$, $\int_{s^0}^{s^1} w_h'(s) ds < 0$ implies that $z'(s) \int_{s^0}^{s^1} w_l'(z(s)) ds < 0$.

The optimal U_0 is chosen such that $X(U_0)$ maximizes $\mu_0 w_h(x) + (1 - \mu_0)w_l(z(x))$ which is concave in x . Therefore, at $x = X(U_0)$ we have

$$\mu_0 w_h'(X(U_0)) + (1 - \mu_0)w_l'(z(X(U_0)))z'(X(U_0)) = 0.$$

According to (20), we know that $w_h'(X(U_0)) \geq 0 \geq w_l'(z(X(U_0)))$. Therefore, the derivative above is weakly positive for any $\mu'_0 > \mu_0$ and hence U_0 increases in μ_0 . ■

C Missing Proof for Section 5

C.1 Continuous Time

We directly work with the expected payoff $W(\tau) = qW_h(\tau) + (1 - q)W_l(\tau)$. Let τ_0 denote the positive root of

$$w_0(\tau) := \mu e^{-r\tau} - (1 - q)l.$$

As is easy to see, this root always exists and is strictly above $\hat{\tau}$, with $w_0(\tau) > 0$ iff $\tau < \hat{\tau}$. Finally, let

$$f(\tau) := r - (\lambda_h + \lambda_l) \frac{w_0(\tau)}{g(\tau)} e^{r\tau}.$$

It is then straightforward to verify (though not quite as easy to obtain) that³³

Proposition 2 *The value function of the principal is given by*

$$W(\tau) = \begin{cases} \bar{W}_1(\tau) & \text{if } \tau \in [0, \hat{\tau}), \\ \bar{W}_1(\tau) - w_0(\tau) \frac{h-l}{hl} c r \mu \frac{\int_{\hat{\tau}}^{\tau} \frac{e^{-\int_{\tau_0}^t f(s) ds}}{w_0^2(t)} dt}{\int_{\hat{\tau}}^{\infty} \frac{\lambda_h + \lambda_l}{g(t)} e^{2rt - \int_{\tau_0}^t f(s) ds} dt} & \text{if } \tau \in [\hat{\tau}, \tau_0), \\ \bar{W}_1(\tau) + w_0(\tau) \frac{h-l}{hl} c \left(1 + r \mu \frac{\int_{\tau}^{\infty} \frac{e^{-\int_{\tau_0}^t f(s) ds}}{w_0^2(t)} dt}{\int_{\hat{\tau}}^{\infty} \frac{\lambda_h + \lambda_l}{g(t)} e^{2rt - \int_{\tau_0}^t f(s) ds} dt} \right) & \text{if } \tau \geq \tau_0, \end{cases}$$

where

$$\bar{W}_1(\tau) := (1 - e^{-r\tau})(1 - c/h)\mu.$$

It is straightforward to derive the closed-form expressions for complete-information payoff, which we omit here.

Proof of Lemma 7. The proof has three steps. We recall that $W(\tau) = qW_h(\tau) + (1 - q)W_l(\tau)$. Using the system of differential equations, we get

$$\begin{aligned} & (e^{r\tau}l + q(h-l)e^{-(\lambda_h + \lambda_l)\tau} - \mu) ((r + \lambda_h)W'(\tau) + W''(\tau)) \\ &= (h-l)q\lambda_h e^{-(\lambda_h + \lambda_l)\tau} W'(\tau) + \mu(r(\lambda_h + \lambda_l)W(\tau) + \lambda_l W'(\tau) - r\lambda_l(h-c)). \end{aligned}$$

It is easily verified that the function W given in Proposition 2 solves this differential equation, and hence is the solution to our problem. Let $w := W - \bar{W}_1$. By definition, w solves a homogeneous second-order differential equation, namely,

$$k(\tau)(w''(\tau) + rw'(\tau)) = r\mu w(\tau) + e^{r\tau}w_0(\tau)w'(\tau), \quad (26)$$

³³As $\tau \rightarrow t_0$, the integrals entering in the definition of W diverge, although not W itself, given that $\lim_{\tau \rightarrow t_0} w_0(\tau) \rightarrow 0$. As a result, $\lim_{\tau \rightarrow t_0} W(\tau)$ is well-defined, and strictly below $W_1(t_0)$.

with boundary conditions $w(\hat{\tau}) = 0$ and $\lim_{\tau \rightarrow \infty} w(\tau) = -(1 - l/h)(1 - q)c$. Here,

$$k(\tau) := \frac{q(h-l)e^{-(\lambda_h + \lambda_l)\tau} + le^{r\tau} - \mu}{\lambda_h + \lambda_l}.$$

By definition of $\hat{\tau}$, $k(\tau) > 0$ for $\tau > \hat{\tau}$. First, we show that k increases with persistence $1/p$, where $\lambda_h = p\bar{\lambda}_h$, $\lambda_l = p\bar{\lambda}_l$, for some $\bar{\lambda}_h, \bar{\lambda}_l$ fixed independently of $p > 0$. Second, we show that $r\mu w(\tau) + e^{r\tau}w_0(\tau)w'(\tau) < 0$, and so $w''(\tau) + rw'(\tau) < 0$ (see (26)). Finally we use these two facts to show that the payoff function is pointwise increasing in p . We give the arguments for the case $\hat{\tau} = 0$, the other case being analogous.

1. Differentiating k with respect to p (and without loss setting $p = 1$) gives

$$\frac{dk(\tau)}{dp} = \frac{\mu}{\bar{\lambda}_h + \bar{\lambda}_l} - \frac{e^{-(\bar{\lambda}_h + \bar{\lambda}_l)\tau}(h-l)\bar{\lambda}_l(1 + (\bar{\lambda}_l + \bar{\lambda}_h)\underline{\tau})}{(\bar{\lambda}_h + \bar{\lambda}_l)^2} - \frac{l}{\bar{\lambda}_h + \bar{\lambda}_l}e^{r\tau}.$$

Evaluated at $\tau = \hat{\tau}$, this is equal to 0. We majorize this expression by ignoring the term linear in τ (underlined in the expression above). This majorization is still equal to 0 at 0. Taking second derivatives with respect to τ of the majorization shows that it is concave. Finally, its first derivative with respect to τ at 0 is equal to

$$h\frac{\bar{\lambda}_l}{\bar{\lambda}_h + \bar{\lambda}_l} - l\frac{r + \bar{\lambda}_l}{\bar{\lambda}_h + \bar{\lambda}_l} \leq 0,$$

because $r \leq \frac{h-l}{l}\bar{\lambda}_l$ whenever $\hat{\tau} = 0$. This establishes that k is decreasing in p .

2. For this step, we use the explicit formulas for W (or equivalently, w) given in Proposition 2. Computing $r\mu w(\tau) + e^{r\tau}w_0(\tau)w'(\tau)$ over the two intervals $(\hat{\tau}, \tau_0)$ and (τ_0, ∞) yields on both intervals, after simplification,

$$-\frac{\frac{h-l}{hl}c}{\int_{\hat{\tau}}^{\infty} \frac{\bar{\lambda}_h + \bar{\lambda}_l}{rv g(t)} e^{2rt - \int_{\tau_0}^t f(s)ds} dt} e^{r\tau} e^{-\int_{\hat{\tau}}^{\tau} f(s)ds} < 0.$$

[The fraction can be checked to be negative. Alternatively, note that $W \leq \bar{W}_1$ on $\tau < \tau_0$ is equivalent to this fraction being negative, yet $\bar{W}_1 \geq \bar{W}$ ($\bar{W}_1$ is the first branch of the complete-information payoff), and because W solves our problem it has to be less than $\bar{W}_1$.]

3. Consider two levels of persistence, $p, \tilde{p}$, with $\tilde{p} > p$. Write $\tilde{w}, w$ for the corresponding solutions to the differential equation (26), and similarly $\tilde{W}, W$. Note that $\tilde{W} \geq W$ is equivalent to $\tilde{w} \geq w$, because $\bar{W}_1$ and w_0 do not depend on p . Suppose that there

exists τ such that $\tilde{w}(\tau) < w(\tau)$ yet $\tilde{w}'(\tau) = w'(\tau)$. We then have that the right-hand sides of (26) can be ranked for both persistence levels, at τ . Hence, so must be the left-hand sides. Because $k(\tau)$ is lower for $\tilde{p}$ than for p (by our first step), because $k(\tau)$ is positive and because the terms $w''(\tau) + rw'(\tau)$, $\tilde{w}''(\tau) + r\tilde{w}'(\tau)$ are negative, and finally because $\tilde{w}'(\tau) = w'(\tau)$, it follows that $\tilde{w}''(\tau) \leq w''(\tau)$. Hence, the trajectories of w and $\tilde{w}$ cannot get closer: for any $\tau' > \tau$, $w(\tau) - \tilde{w}(\tau) \leq w(\tau') - \tilde{w}(\tau')$. This is impossible, because both w and $\tilde{w}$ must converge to the same value, $-(1-l/h)(1-q)c$, as $\tau \rightarrow \infty$. Hence, we cannot have $\tilde{w}(\tau) < w(\tau)$ yet $\tilde{w}'(\tau) = w'(\tau)$. Note however that this means that $\tilde{w}(\tau) < w(\tau)$ is impossible, because if this were the case, then by the same argument, since their values as $\tau \rightarrow \infty$ are the same, it is necessary (by the intermediate value theorem) that for some τ such that $\tilde{w}(\tau) < w(\tau)$ the slopes are the same.

■

Proof of Lemma 8. The proof is divided into two steps. First we show that the difference in payoffs between $W(\tau)$ and the complete-information payoff computed at the same level of utility $\underline{u}(\tau)$ converges to 0 at a rate linear in r , for all τ . Second, we show that the distance between the closest point on the graph of $\underline{u}(\cdot)$ and the complete-information payoff maximizing pair of utilities converges to 0 at a rate linear in r . Given that the complete-information payoff is piecewise affine in utilities, the result follows from the triangle inequality.

1. We first note that the complete-information payoff along the graph of $\underline{u}(\cdot)$ is at most equal to $\max\{\bar{W}_1(\tau), \bar{W}_2(\tau)\}$, where $\bar{W}_1$ is defined in Proposition 2 and

$$\bar{W}_2(\tau) = (1 - e^{-r\tau})(1 - c/l)\mu + q(h/l - 1)c.$$

These are simply two of the four affine maps whose lower envelope defines $\bar{W}$, see Section 3.1 (those for the domains $[0, v_h^*] \times [0, v_l^*]$ and $[0, \mu_h] \times [v_l^*, \mu_l]$). The formulas obtain by plugging in $\underline{u}_h, \underline{u}_l$ for U_h, U_l , and simplifying. Fix $z = r\tau$ (note that as $r \rightarrow 0$, $\hat{\tau} \rightarrow \infty$, so that changing variables is necessary to compare limiting values as $r \rightarrow 0$), and fix z such that $le^z > \mu$ (that is, such that $g(z/r) > 0$ and hence $z \geq r\hat{\tau}$ for small enough r). Algebra gives

$$\lim_{r \rightarrow 0} f(z/r) = \frac{(e^z - 1)\lambda_h l - \lambda_l h}{le^z - \mu},$$

and similarly

$$\lim_{r \rightarrow 0} w_0(z/r) = (qh - (e^z - 1)(1 - q)l)e^{-z},$$

as well as

$$\lim_{r \rightarrow 0} g(z/r) = le^z - \mu.$$

Hence, fixing z and letting $r \rightarrow 0$ (so that $\tau \rightarrow \infty$), it follows that $\frac{w_0(\tau) \int_{\hat{\tau}}^{\tau} \frac{e^{-\int_{\tau_0}^t f(s) ds}}{w_0^2(t)} dt}{\int_{\hat{\tau}}^{\infty} \frac{\lambda_h + \lambda_l}{g(t)} e^{2rt - \int_{\tau_0}^t f(s) ds} dt}$ converge to a well-defined limit. (Note that the value of τ_0 is irrelevant to this quantity, and we might as well use $r\tau_0 = \ln(\mu/((1-q)l))$, a quantity independent of r). Denote this limit κ . Hence, for $z < r\tau_0$, because

$$\lim_{r \rightarrow 0} \frac{\bar{W}_1(z/r) - W(z/r)}{r} = \frac{h-l}{hl} c\kappa,$$

it follows that $W(z/r) = \bar{W}_1(z/r) + \mathcal{O}(r)$. On $z > r\tau_0$, it is immediate to check from the formula of Proposition 2 that

$$W(\tau) = \bar{W}_2(\tau) + w_0(\tau) \frac{h-l}{hl} c r \mu \frac{\int_{\hat{\tau}}^{\tau} \frac{e^{-\int_{\tau_0}^t f(s) ds}}{w_0^2(t)} dt}{\int_{\hat{\tau}}^{\infty} \frac{\lambda_h + \lambda_l}{g(t)} e^{2rt - \int_{\tau_0}^t f(s) ds} dt}.$$

[By definition of τ_0 , $w_0(\tau)$ is now negative.] By the same steps it follows that $W(z/r) = \bar{W}_2(z/r) + \mathcal{O}(r)$ on $z > r\tau_0$. Because $W = \bar{W}_1$ for $\tau < \hat{\tau}$, this concludes the first step.

2. For the second step, note that the utility pair maximizing complete-information payoff is given by $v^* = \left(\frac{r+\lambda_l}{r+\lambda_l+\lambda_h} h, \frac{\lambda_l}{r+\lambda_l+\lambda_h} h \right)$. (Take limits from the discrete game.) We evaluate $\underline{u}(\tau) - v^*$ at a particular choice of τ , namely

$$\tau^* = \frac{1}{r} \ln \frac{\mu}{(1-q)l}.$$

It is immediate to check that

$$\frac{\underline{u}(\tau^*) - v_l^*}{qr} = -\frac{\underline{u}_h(\tau^*) - v_h^*}{(1-q)r} = \frac{l + (h-l) \left(\frac{(1-q)l}{\mu} \right)^{\frac{r+\lambda_l+\lambda_h}{r}}}{r + \lambda_l + \lambda_h} \rightarrow \frac{l}{\lambda_l + \lambda_h},$$

and so $\|\underline{u}(\tau^*) - v^*\| = \mathcal{O}(r)$. It is also easily verified that this gives an upper bound on the order of the distance between the polygonal chain and the point v^* . This concludes the second step.

■

C.2 Continuous Types

Proof of Theorem 4. By the principle of optimality, letting $S := W - U$,

$$S(U) = \delta \int S(\mathcal{U}(U, v)) dF - (1 - \delta)c\mathbf{E}_F[q(v)],$$

over $q : [\underline{v}, 1] \rightarrow [0, 1]$ and $\mathcal{U} : [0, \mu] \times [\underline{v}, 1] \rightarrow [0, \mu]$, subject to

$$U = \int ((1 - \delta)q(v)v + \delta\mathcal{U}(U, v)) dF,$$

and, for all $v, v' \in [\underline{v}, 1]$,

$$(1 - \delta)q(v)v + \delta\mathcal{U}(U, v) \geq (1 - \delta)q(v')v + \delta\mathcal{U}(U, v').$$

Note that the dependence of q on U is omitted. By the usual arguments, it follows that q is nondecreasing and differentiable a.e., with $(1 - \delta)q'(v)v + \delta \frac{\partial \mathcal{U}(U, v)}{\partial v} = 0$, and so

$$\mathcal{U}(U, v) = \mathcal{U}(U, \underline{v}) - \frac{1 - \delta}{\delta} \left(vq(v) - \underline{v}q(\underline{v}) - \int_{\underline{v}}^v q(s) ds \right).$$

This formula is also correct if q is discontinuous. Promise keeping becomes

$$U = \delta\mathcal{U}(U, \underline{v}) + (1 - \delta) \left(\underline{v}q(\underline{v}) + \int_{\underline{v}}^1 (1 - F(v))q(v) dv \right).$$

So, the objective $S(U)$ equals

$$\sup \left\{ \delta \int S \left(\frac{U}{\delta} - \frac{1 - \delta}{\delta} \left(vq(v) - \int_{\underline{v}}^v q(s) ds - \int_{\underline{v}}^1 (1 - F(v))q(v) ds \right) \right) dF - (1 - \delta)c\mathbf{E}_F[q(v)] \right\},$$

over $q : [\underline{v}, 1] \rightarrow [0, 1]$, nondecreasing, and the feasibility restriction

$$\forall v \in [\underline{v}, 1] : U - (1 - \delta) \left(vq(v) - \int_{\underline{v}}^v q(s) ds - \int_{\underline{v}}^1 (1 - F(v))q(v) ds \right) \in [0, \delta\mu].$$

We note that, by the envelope theorem,

$$S'(U) = \int S'(\mathcal{U}(U, v)) dF.$$

We restrict attention to the case in which $q(\underline{v}) = 0$, $q(1) = 1$, slight adjustments might be necessary otherwise.

Again, let us suppose contrary to the assumption that S' is constant over some interval I . Pick two points in this interval, $U_1 < U_2$. Given $U = \lambda U_1 + (1 - \lambda)U_2$, $\lambda \in (0, 1)$, consider the

policy $q_\lambda, \mathcal{U}_\lambda$ that uses $q_\lambda = \lambda q_1 + (1 - \lambda)q_2$, and similarly $\mathcal{U}_\lambda(\cdot, v) = \lambda \mathcal{U}_1(\cdot, v) + (1 - \lambda)\mathcal{U}_2(\cdot, v)$. To be clear, this is the strategy that consists, for every report v , in giving the agent the good with probability $q_\lambda(v) = \lambda q_1(v) + (1 - \lambda)q_2(v)$, and transiting to the utility the averages of the utility after v under the policy starting at U_1 and U_2 (more generally, the weighted average given the sequence of reports). We note that, given risk neutrality of the agent, this policy induces the agent to report truthfully (since he does both at U_1 and at U_2), and gives him utility U , by construction.

We claim that, given U , and for any given v , $S'(\mathcal{U}(U_1, v)) = S'(\mathcal{U}(U_2, v))$. If not, then there exists v such that $S'(\mathcal{U}(U_1, v)) \neq S'(\mathcal{U}(U_2, v))$ and some $U' = \lambda \mathcal{U}(U_1, v) + (1 - \lambda)\mathcal{U}(U_2, v)$ in between these two values such that $S(U') > \lambda S(\mathcal{U}(U_1, v)) + (1 - \lambda)S(\mathcal{U}(U_2, v))$. Then, consider using the policy that uses $q_\lambda, \mathcal{U}_\lambda$ for one step and then reverts to the optimal policy. Because it does at least as well as the average of the two policies for all values of v , and does strictly better for v , it is a strict improvement, a contradiction.

Hence, we may assume that $S'(\mathcal{U}(U_1, v)) = S'(\mathcal{U}(U_2, v))$. We next claim that this implies that, without loss, $q_1(\cdot) = q_2(\cdot)$. Indeed, we can divide $[\underline{v}, 1]$ into those (maximum length) intervals over which $S'(\mathcal{U}(U_1, v)) = S'(\mathcal{U}(U_2, v))$ and those over which $S'(\mathcal{U}(U_i, v)) > S'(\mathcal{U}(U_{-i}, v))$, for some $i = 1, 2$. On any interval of values of v over which $\mathcal{U}(U_1, v) = \mathcal{U}(U_2, v)$, it follows from the formula above, namely,

$$\mathcal{U}(U, v) = \mathcal{U}(U, v') - \frac{1 - \delta}{\delta} \int_{v'}^v s dq(s),$$

that the variation is the same for q_1 and q_2 (Since the function $\mathcal{U}(U, v)$ is the same). Over intervals of values of v over which S' is independent of i , S must be affine over the ranges $[\min_i \{\mathcal{U}(U_i, v)\}, \max_i \{\mathcal{U}(U_i, v)\}]$, so that, because S is affine, it follows from the Bellman equation that U does not matter for the optimal choice of q either.

Hence, $q_1(\cdot) = q_2(\cdot)$. It follows that, if for some v , $\mathcal{U}(U_1, v) = \mathcal{U}(U_2, v)$, it must also be that $\mathcal{U}(U_1, \cdot) = \mathcal{U}(U_2, \cdot)$. This is impossible, because then $U_1 = U_2$, by promise-keeping. Hence, there is no v such that $\mathcal{U}(U_1, v) = \mathcal{U}(U_2, v)$, and S is affine on the entire range of $\mathcal{U}(U_1, \cdot), \mathcal{U}(U_2, \cdot)$. In fact, the values of $\mathcal{U}(U_1, \cdot)$ must be translates of those of $\mathcal{U}(U_2, \cdot)$. Without loss, we might take $[U_1, U_2]$ to be the largest interval over which S is affine. Given that $q(\underline{v}) = 0 < q(1) = 1$, neither $\mathcal{U}(U_1, \cdot)$ nor $\mathcal{U}(U_2, \cdot)$ is degenerate (that is, constant). Therefore, the only possibility is that both the range of $\mathcal{U}(U_1, \cdot)$ and that of $\mathcal{U}(U_2, \cdot)$ are in $[U_1, U_2]$. This is impossible given promise keeping and that $q_1(\cdot) = q_2(\cdot)$. ■

For clarity of exposition, we assume that the agent's value v is drawn from $[\underline{v}, \bar{v}]$ (instead of $[\underline{v}, 1]$) according to F with $\underline{v} \in [0, \bar{v}]$. Let $x_1(v) = p(v)$ and $x_2(v) = \mathcal{U}(U, v)$. The optimal

policy x_1, x_2 is the solution to the control problem

$$\max_u \int_{\underline{v}}^{\bar{v}} (1 - \delta)x_1(v)(v - c) + \delta W(x_2(v))dF,$$

subject to the law of motion $x'_1 = u$ and $x'_2 = -(1 - \delta)vu/\delta$. The control is u and the law of motion captures the incentive compatibility constraints. We define a third state variable x_3 to capture the promise-keeping constraint

$$x_3(v) = (1 - \delta)vx_1(v) + \delta x_2(v) + (1 - \delta) \int_v^{\bar{v}} x_1(s)(1 - F(s))ds.$$

The law of motion of x_3 is $x'_3(v) = (1 - \delta)x_1(v)(F(v) - 1)$.³⁴ The constraints are

$$\begin{aligned} u &\geq 0 \\ x_1(\underline{v}) &\geq 0, \quad x_1(\bar{v}) \leq 1 \\ x_2(\underline{v}) &\leq \bar{v}, \quad x_2(\bar{v}) \geq 0 \\ x_3(\underline{v}) &= U, \quad x_3(\bar{v}) - (1 - \delta)\bar{v}x_1(\bar{v}) - \delta x_2(\bar{v}) = 0. \end{aligned}$$

Let $\gamma_1, \gamma_2, \gamma_3$ be the costate variables and μ_0 the multiplier for $u \geq 0$. For the rest of this sub-section the dependence on v is omitted when no confusion arises. The Lagrange is

$$\mathcal{L} = ((1 - \delta)x_1(v - c) + \delta W(x_2))f + \gamma_1 u - \gamma_2 \frac{1 - \delta}{\delta} vu + \gamma_3 (1 - \delta)x_1(F - 1) + \mu_0 u.$$

The first-order conditions are

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial u} &= \gamma_1 - \gamma_2 \frac{1 - \delta}{\delta} v + \mu_0 = 0 \\ \dot{\gamma}_1 &= -\frac{\partial \mathcal{L}}{\partial x_1} = (1 - \delta)(\gamma_3(1 - F) - f(v - c)) \\ \dot{\gamma}_2 &= -\frac{\partial \mathcal{L}}{\partial x_2} = -\delta f W'(x_2) \\ \dot{\gamma}_3 &= -\frac{\partial \mathcal{L}}{\partial x_3} = 0. \end{aligned}$$

³⁴Note that the promise-keeping constraint can be rewritten as

$$U = (1 - \delta)\underline{v}x_1(\underline{v}) + \delta x_2(\underline{v}) + (1 - \delta) \int_{\underline{v}}^{\bar{v}} x_1(s)(1 - F(s))ds.$$

The transversality conditions are

$$\begin{aligned}
\gamma_1(\underline{v}) &\leq 0, \quad \gamma_1(\bar{v}) + (1 - \delta)\bar{v}\gamma_3(\bar{v}) \leq 0, \\
\gamma_1(\underline{v})x_1(\underline{v}) &= 0, \quad (\gamma_1(\bar{v}) + (1 - \delta)\bar{v}\gamma_3(\bar{v}))(1 - x_1(\bar{v})) = 0, \\
\gamma_2(\underline{v}) &\geq 0, \quad \gamma_2(\bar{v}) + \delta\gamma_3(\bar{v}) \geq 0, \\
\gamma_2(\underline{v})(\bar{v} - x_2(\underline{v})) &= 0, \quad (\gamma_2(\bar{v}) + \delta\gamma_3(\bar{v}))x_2(\bar{v}) = 0, \\
\gamma_3(\underline{v}) \text{ and } \gamma_3(\bar{v}) &\text{ free.}
\end{aligned}$$

The first observation is that $\gamma_3(v)$ is constant, denoted γ_3 . Moreover, given γ_3 , $\dot{\gamma}_1$ involves no endogenous variables. Therefore, for a fixed $\gamma_1(\underline{v})$, the trajectory of γ_1 is fixed. Whenever $u > 0$, we have $\mu_0 = 0$. The first-order condition $\frac{\partial \mathcal{L}}{\partial u} = 0$ implies that

$$\gamma_2 = \frac{\delta\gamma_1}{(1 - \delta)v} \text{ and } \dot{\gamma}_2 = \frac{\delta(\gamma_1 - v\dot{\gamma}_1)}{(\delta - 1)v^2}.$$

Given that $\dot{\gamma}_2 = -\delta fW'(x_2)$, we could determine the state x_2

$$x_2 = (W')^{-1} \left(\frac{v\dot{\gamma}_1 - \gamma_1}{(\delta - 1)fv^2} \right). \quad (27)$$

The control u is given by $-\dot{x}_2\delta/((1 - \delta)v)$. As the promised utility varies, we conjecture that the solution can be one of the three cases.

Case one occurs when U is intermediate: There exists $\underline{v} \leq v_1 \leq v_2 \leq \bar{v}$ such that $x_1 = 0$ for $v \leq v_1$, x_1 is strictly increasing when $v \in (v_1, v_2)$ and $x_1 = 1$ for $v \geq v_2$. Given that $u > 0$ iff $v \in (v_1, v_2)$, we have

$$x_2 = \begin{cases} (W')^{-1} \left(\frac{v\dot{\gamma}_1 - \gamma_1}{(\delta - 1)fv^2} \right) \Big|_{v=v_1} & \text{if } v < v_1 \\ (W')^{-1} \left(\frac{v\dot{\gamma}_1 - \gamma_1}{(\delta - 1)fv^2} \right) & \text{if } v_1 \leq v \leq v_2 \\ (W')^{-1} \left(\frac{v\dot{\gamma}_1 - \gamma_1}{(\delta - 1)fv^2} \right) \Big|_{v=v_2} & \text{if } v > v_2, \end{cases}$$

and

$$x_1 = \begin{cases} 0 & \text{if } v < v_1 \\ -\frac{\delta}{1 - \delta} \int_{v_1}^v \frac{\dot{x}_2}{s} ds & \text{if } v_1 \leq v \leq v_2 \\ 1 & \text{if } v > v_2. \end{cases}$$

The continuity of x_1 at v_2 requires that

$$-\frac{\delta}{1 - \delta} \int_{v_1}^{v_2} \frac{\dot{x}_2}{s} ds = 1. \quad (28)$$

The trajectory of γ_2 is given by

$$\gamma_2 = \begin{cases} \frac{\delta\gamma_1}{(1-\delta)v_1} + \delta(F(v_1) - F(v)) \frac{v_1\dot{\gamma}_1(v_1) - \gamma_1(v_1)}{(\delta-1)f(v_1)v_1^2} & \text{if } v < v_1 \\ \frac{\delta\gamma_1}{(1-\delta)v} & \text{if } v_1 \leq v \leq v_2 \\ \frac{\delta\gamma_1}{(1-\delta)v_2} - \delta(F(v) - F(v_2)) \frac{v_2\dot{\gamma}_1(v_2) - \gamma_1(v_2)}{(\delta-1)f(v_2)v_2^2} & \text{if } v > v_2. \end{cases}$$

If $(W')^{-1} \left(\frac{v_1\dot{\gamma}_1(v_1) - \gamma_1(v_1)}{(\delta-1)f(v_1)v_1^2} \right) < \bar{v}$ and $(W')^{-1} \left(\frac{v_2\dot{\gamma}_1(v_2) - \gamma_1(v_2)}{(\delta-1)f(v_2)v_2^2} \right) > 0$, the transversality condition requires that

$$\frac{\delta\gamma_1(v_1)}{(1-\delta)v_1} + \delta F(v_1) \frac{v_1\dot{\gamma}_1(v_1) - \gamma_1(v_1)}{(\delta-1)f(v_1)v_1^2} = 0 \quad (29)$$

$$\frac{\delta\gamma_1(v_2)}{(1-\delta)v_2} - \delta(1 - F(v_2)) \frac{v_2\dot{\gamma}_1(v_2) - \gamma_1(v_2)}{(\delta-1)f(v_2)v_2^2} = -\delta\gamma_3. \quad (30)$$

We have four unknowns $v_1, v_2, \gamma_3, \gamma_1(\bar{v})$ and four equations, (28)–(30) and the promise-keeping constraint. Alternatively, for a fixed v_1 , (28)–(30) determine the three other unknowns $v_2, \gamma_3, \gamma_1(\bar{v})$. We need to verify that all inequality constraints are satisfied.

Case two occurs when U is close to 0: There exists v_1 such that $x_1 = 0$ for $v \leq v_1$ and x_1 is strictly increasing when $v \in (v_1, \bar{v}]$. The $x_1(\bar{v}) \leq 1$ constraint does not bind. This implies that $\gamma_1(\bar{v}) + (1-\delta)\bar{v}\gamma_3 = 0$. When $v > v_1$, the state x_2 is pinned down by (27).

From the condition that $\gamma_1(\bar{v}) + (1-\delta)\bar{v}\gamma_3(\bar{v}) = 0$, we have that $W'(x_2(\bar{v})) = 1 - c/\bar{v}$. Given strict concavity of W and $W'(0) = 1 - c/\bar{v}$, we have $x_2(\bar{v}) = 0$. The constraint $x_2(\bar{v}) \geq 0$ binds, so (30) is replaced with

$$\frac{\delta\gamma_1(\bar{v})}{(1-\delta)\bar{v}} + \delta\gamma_3 \leq 0,$$

which is always satisfied given that $\gamma_1(\bar{v}) \leq 0$. From (29), we can solve γ_3 in terms of v_1 . Lastly, the promise-keeping constraint pins down the value of v_1 . Note that the constraint $x_1(\bar{v}) \leq 1$ does not bind. This requires that

$$-\frac{\delta}{1-\delta} \int_{v_1}^{\bar{v}} \frac{\dot{x}_2}{s} ds \leq 1. \quad (31)$$

There exists a v_1^* such that this inequality is satisfied if and only if $v_1 \geq v_1^*$. When $v_1 < v_1^*$, we move to case one. We would like to prove that the left-hand side increases as v_1 decreases. Note that γ_3 measures the marginal benefit of U , so it equals $W'(U)$.

Case three occurs when $\underline{v} > 0$ and U is close to μ : There exists v_2 such that $x_1 = 1$ for $v \geq v_2$ and x_2 is strictly increasing when $v \in [\underline{v}, v_2)$. The $x_1(\underline{v}) \geq 0$ constraint does not bind.

This implies that $\gamma_1(\underline{v}) = 0$. When $v < v_2$, the state x_2 is pinned down by (27). From the condition that $\gamma_1(\underline{v}) = 0$, we have that $W'(x_2(\underline{v})) = 1 - c/\underline{v}$. Given strict concavity of W and $W'(\bar{v}) = 1 - c/\bar{v}$, we have $x_2(\underline{v}) = \bar{v}$. The constraint $x_2(\underline{v}) \leq 1$ binds, so (29) is replaced with

$$\frac{\delta\gamma_1(\underline{v})}{(1-\delta)\underline{v}} \leq 0,$$

which is always satisfied given that $\gamma_1(\underline{v}) \leq 0$. From (30), we can solve γ_3 in terms of v_2 . Lastly, the promise-keeping constraint pins down the value of v_2 . Note that the constraint $x_1(\underline{v}) \geq 0$ does not bind. This requires that

$$-\frac{\delta}{1-\delta} \int_{\underline{v}}^{v_2} \frac{\dot{x}_2}{s} ds \leq 1. \quad (32)$$

There exists a v_2^* such that this inequality is satisfied if and only if $v_2 \leq v_2^*$. When $v_2 > v_2^*$, we move to case one.

Proof of Proposition 1. To illustrate, we assume that v is uniform on $[0, 1]$. The proof for $F(v) = v^a$ with $a > 1$ is similar. We start with case two. From condition (29), we solve for $\gamma_3 = 1 + c(v_1 - 2)$. Substituting γ_3 into $\gamma_1(v)$, we have

$$\gamma_1(v) = \frac{1}{2}(1-\delta)(1-v)(v(c(v_1-2)+2)-cv_1).$$

The transversality condition $\gamma_1(0) \leq 0$ is satisfied. The first-order condition $\frac{\partial \mathcal{L}}{\partial u} = 0$ is also satisfied for $v \leq v_1$. Let G denote the function $((W')^{-1})'$. We have

$$\begin{aligned} -\frac{\delta}{1-\delta} \int_{v_1}^1 \frac{\dot{x}_2}{s} ds &= -\frac{\delta}{(1-\delta)} \int_{v_1}^1 G\left(1-c+\frac{c}{2}\left(v_1-\frac{v_1}{s^2}\right)\right) \frac{cv_1}{s^3} \frac{1}{s} ds \\ &= -\frac{\delta}{(1-\delta)} \int_{v_1-1/v_1}^0 G\left(1-c+\frac{c}{2}x\right) \frac{c}{2} \sqrt{1-\frac{x}{v_1}} dx. \end{aligned}$$

The last equality is obtained by the change of variables. As v_1 decreases, $v_1 - 1/v_1$ decreases and $\sqrt{1 - x/v_1}$ increases. Therefore, the left-hand side of (31) indeed increases as v_1 decreases.

We continue with case one. From (29) and (30), we can solve for γ_3 and $\gamma_1(v)$

$$\begin{aligned} \gamma_3 &= 1 + c \left(\frac{v_1(2v_2-1)}{v_2^2} - 2 \right), \\ \gamma_1(v) &= \frac{1}{2}(\delta-1) \left(v \left((v-2) \left(c \left(\frac{v_1(2v_2-1)}{v_2^2} - 2 \right) + 1 \right) - 2c + v \right) + cv_1 \right). \end{aligned}$$

It is easily verified that $\gamma_1(0) \leq 0$, $\gamma_1(1) \leq 0$, and the first-order condition $\frac{\partial \mathcal{L}}{\partial u} = 0$ is satisfied. Equation (28) can be rewritten as

$$-\frac{\delta}{1-\delta} \int_{v_1}^{v_2} \frac{\dot{x}_2}{s} ds = -\frac{\delta}{(1-\delta)} \int_{v_1}^{v_2} G \left(1 - c + \frac{c}{2} \left(\frac{v_1(2v_2-1)}{v_2^2} - \frac{v_1}{s^2} \right) \right) \frac{cv_1}{s^3} \frac{1}{s} ds = 1.$$

For any $v_1 \leq v_1^*$, there exists $v_2 \in (v_1, 1)$ such that (28) is satisfied. ■

Transfers with Limited Liability. Here, we consider the case in which transfers are allowed but the agent is protected by limited liability. Therefore, only the principal can pay the agent. The principal maximizes his payoff net of payments. The following lemma shows that transfers occur on the equilibrium path when the ratio c/l is higher than 2.

Lemma 11 *The principal makes transfers on path if and only if $c - l > l$.*

Proof. We first show that the principal makes transfers if $c - l > l$. Suppose not. The optimal mechanism is the same as the one characterized in Theorem 1. When U is sufficiently close to μ , we want to show that it is “cheaper” to provide incentives using transfers. Given the optimal allocation (p_h, u_h) and (p_l, u_l) , if we reduce u_l by ε and make a transfer of $\delta\varepsilon/(1-\delta)$ to the low type, the *IC/PK* constraints are satisfied. When u_l is sufficiently close to μ , the principal’s payoff increment is close to $\delta(c/l - 1)\varepsilon - \delta\varepsilon = \delta(c/l - 2)$, which is strictly positive if $c - l > l$. This contradicts the fact that the allocation (p_h, u_h) and (p_l, u_l) is optimal. Therefore, the principal makes transfers if $c - l > l$.

If $c - l \leq l$, we first show that the principal never makes transfers if $u_l, u_h < \mu$. With abuse of notation, let t_m denote the current-period transfer after m report. Suppose $u_m < \mu$ and $t_m > 0$. We can increase u_m ($m = l$ or h) by ε and reduce t_m by $\delta\varepsilon/(1-\delta)$. This adjustment has no impact on *IC/PK* constraints and strictly increases the principal’s payoff given that $W'(U) > 1 - c/l$ when $U < \mu$.³⁵ Suppose $u_l = \mu$ and $t_l > 0$. We can always replace p_l, t_l with $p_l + \varepsilon, t_l - \varepsilon l$. This adjustment has no impact on *IC/PK* and (weakly) increases the principal’s payoff. If $u_l = \mu, p_l = 1$, we know that the promised utility to the agent is at least μ . The optimal scheme is to provide the unit forever. ■

³⁵It is easy to show that the principal’s complete-information payoff, if $U \in [0, \mu]$ and $c - l \leq l$, is the same as $\bar{W}$ in Lemma 1.