

Dynamic Mechanisms without Money

**Very* Preliminary. Comments welcome.*

Yingni Guo*, Johannes Hörner[†]

September 2014

Abstract

We analyze the optimal design of dynamic mechanisms in the absence of transfers. The designer uses future allocation decisions as a way of eliciting private information. Values evolve according to a two-state Markov chain. We solve for the optimal allocation rule, which admits 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 show 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 looking for the social choice mechanism that would get us closest to efficiency. Here, efficiency and implementability are understood to be Bayesian: both the individual and society understand the probabilistic nature of the uncertainty and update based on it.

*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.

Society’s decision not to allow for money –be it for physical, legal or ethical reasons– is taken for granted. So is the sequential nature of the problem: temporal constraints apply to the allocation of goods, whether jobs, houses or attention, and it is often difficult to ascertain future demands.

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

The allocation of perishable goods is relevant to many economic activities. Many allocation decisions involve services that are perishable by definition (how a nurse or a worker divides his time, for instance), and some of them involve maintenance costs (keeping a job vacancy open, etc.) that are high enough that treating them as non-durable might be a good approximation.

Formally, our good can take one of two values in each period. While this is certainly restrictive, it is known that, even with transfers, the problem is intractable beyond two types (see Battaglini, 2005).¹ We start with the i.i.d. case, then generalize to the case of a Markov chain. The cost of providing the good is fixed and known. It is optimal to assign the good if and only if the value is high.

We solve for the optimal policy and show a simple way of implementing it. Furthermore, in the Poisson limit, we can explicitly solve for the principal’s payoff. This allows us to show that, despite the absence of transfers, first-best is approached at a rate that is linear in the discount rate. Furthermore, persistence hurts: as the Markov chain becomes more persistent, efficiency decreases.

The optimal policy can be implemented by a “budget” mechanism. As we show, the appropriate unit of account is the number of units that the agent is entitled to get in a row, “no questions asked.” This naturally maps into a utility level for the agent, as a function of his current type. If he asks for the unit, the budget gets revised as follows: subtract from the low type’s utility level the unit’s value to a low type, and update the budget to what would give him tomorrow the equivalent of that utility today (if his type tomorrow is drawn according to the transition matrix of the low type). If he does not, the budget is updated the same way, except that no utility gets subtracted from the low type’s utility.

¹We believe that the analysis might remain tractable for renewal processes, an extension that would be interesting to consider.

Surprisingly, this updating is optimal independently of the principal’s belief. The only role of the prior belief is to pin down the initial budget. Furthermore, this policy is eventually absorbed into one of two possible long-run outcomes: either the agent is granted to get the unit forever, or never again. Immiseration does not necessarily result: given the optimal initial promise, both outcomes have positive probability

Relative to the literature on linking decisions in the absence of transfers, one important consequence of our characterization is that the interpretation of optimality as $\delta \rightarrow 1$ that is often associated with such linkages must be taken with a grain of salt: asymptotically as $t \rightarrow \infty$, the allocation is necessarily inefficient.

Relative to the literature on long-term contracting with Markovian consumers (Battaglini, 2005, in particular), we show that the dynamics of utility and efficiency depend on the absence or presence of transfers. In Battaglini, efficiency necessarily improves over time (in fact, exact efficiency obtains eventually). Here instead, efficiency decreases over time, in the sense described above, with an asymptotic outcome which is at best the outcome of the static game. As for the agent’s utility, it can go up or down, depending on the history that realizes: getting the good forever is clearly the best possible outcome from his point of view; never getting it again being the worst.

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 on how to implement algorithms to make better allocations. As an example, think about nurses that must decide whether to take seriously some alerts that are either triggered by sensors or by patients themselves. The opportunity cost of their time is significant. Patients, however, appreciate quality time with nurses whether or not their condition necessitates it. This gives rise to a challenge that every hospital must contend with: ignore alarms, and take the chance that a patient with a serious condition does not get attended to; pay heed to all of them, and end up with overwhelmed nurses. “Alarm fatigue” is a serious problem that health care must deal with (see, for instance, Sendelbach, 2012). We suggest the best way of trading off the two risks that come along with it: neglecting a patient in need of care, and one that simply cries wolf.

Related Literature: All the versions considered in this paper would be trivial in the absence of imperfect observation of the values. If values were perfectly observed, it would simply be optimal to assign the good if and only if the value is high. Because of private information, it is necessary to distort the allocation: after some histories, the good is provided independently of the report; after some others, it is never provided again. In this

sense, scarcity of good provision is endogenously determined, for the purpose of information elicitation. There is a large literature in operations research considering the case in which this scarcity is taken as exogenously given –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 are Derman, Lieberman and Ross (1972) and Albright (1977). Their analysis suggests a natural mechanism that can be applied in our environment: give the agent a certain number of “tokens,” and let him exercise them whenever he pleases.

The idea that tokens could be used as intertemporal “budgets” to discipline agents with private information has appeared in several papers in economics before. Möbius (2001) might well be the first who suggests that keeping track of the difference in the number of favors granted (with two agents) and granting favors or not as a function of this difference might be a simple but powerful way of sustaining cooperation in long-run relationships. See also Abdulkadiroğlu and Bagwell (2012) and Kalla (2010). While these mechanisms are known to be suboptimal (as is clear from our characterization of the optimal one), they have desirable properties nonetheless: properly calibrated, they yield an approximately efficient allocation as the discount factor tends to one. To our knowledge, Hauser and Hopenhayn (2008) is the paper that comes 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, independently 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 bind, and the efficient policy is followed unless it is inconsistent with promise-keeping (so, when promised utilities are too extreme). Our model can be viewed as a game with one-sided incomplete information, in which the production cost of the principal is the known value to the second player. There are some differences, however: first, our principal has commitment, so he is not tempted to act opportunistically, and is not 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, so that the process is memoryless.

A second related strand of literature might be referred to as “linking incentive constraints.” The idea that, as the number of identical copies of a decision problem increases, tying them together might allow the designer to improve on the isolated problem appears in

a number of papers, under various degrees of generality. See Fang and Norman (2006), and Jackson and Sonnenschein (2007). But many mechanisms work to establish this asymptotic result, and one of our goals is precisely to discriminate among them for fixed discounting (so that no single copy is negligible). Hortala-Vallve (2010) provides an interesting analysis of the unavoidable inefficiencies that must be incurred away from the limit, and Cohn (2010) shows the suboptimality of the mechanisms that are commonly used, even in terms of the rate of convergence. The best rate of convergence is derived in Lemma 11.

A third related branch of literature could be referred to as the literature on “immiseration.” Thomas and Worrall (1990) is one of the early papers studying the problem of insurance under incomplete information, showing how the utility dynamics inexorably take the utility of the agent to minus infinity. No such immiseration occurs here. In both cases, the spread in continuation payoffs requires payoffs to converge (if ever) to one of the boundaries, but the assumptions that Thomas and Worrall make on the utility function rule out any other boundary than immiseration.

That allocation rights to other (or future) units can be used as a “currency” for eliciting private information is long known. It goes back at least to Hylland and Zeckhauser (1979), who are the first to explain to what extent 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).

Exactly optimal mechanisms have been computed in related environments. Frankel (2011) considers a variety of related settings. Closest is his analysis in his Chapter 2, where he also derives an optimal mechanism. While he allows for more than two types and actions, he restricts attention to the case of 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 have private information regarding the value of obtaining the good. In their model, agents are myopic, and the scarcity in the resource is exogenously assumed. In addition, transfers are allowed. They show 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 strictly more general than ours (he allows two agents, and more than two types). Unfortunately, he does not provide a solution to his model.

A related literature considers the problem of optimal stopping in the absence of transfers,

see in particular Kováč, Krähmer and Tatur (2014). The major difference between these two problems is that we are not dealing with a stopping problem: a decision must be taken in every period. 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 foregoing the current unit, in case the value is low and the future prospects are good. Not here – his incentives to forego the unit must be endogenously generated via the promises. In the stopping case, there is only one history of outcomes that does not terminate the game. Here instead, policies differ not only in when they first provide the good, but what happens afterwards.

Finally, the relevant benchmark with transfers, as already mentioned, is provided by Battaglini (2005). He shows that the solution is non-stationary (with infinite memory) but nonetheless admits a simple description (using a state variable). Supply converges to efficiency (very much unlike what happens without transfers, since asymptotic allocation is necessarily inefficient), although the convergence is history-dependent. A detailed comparison of our results with his is relegated to Section 3.6. Krishna, Lopomo and Taylor (2013) provide an analysis with limited liability (though transfers are allowed) in a model closely related to Battaglini, suggesting that, indeed, ruling out unlimited transfers matters for both the optimal contract and the dynamics.

2 The Baseline Model

We start our investigation with the simplest case, in which there is one agent, with only two possible values (or types) in a given period, and values are i.i.d. over time. Section 3 relaxes the independence assumption.

2.1 Set-up

Time is discrete and the horizon infinite. There are two parties, a principal and an agent. In each period, the principal can produce a unit of good at a cost of $c > 0$. The agent's value (or type) is either l or h such that $0 < l < c < h$. The probability of the high (h) type is q . This value is privately observed and independent across periods. More specifically, at the beginning of each period, the value is drawn and the agent is informed of it.

We write $v_n = h, l$ for the realized value in period n , as well as $\bar{v}$ for the expected value of the good, that is, $\bar{v} = qh + (1 - q)l$.

Players are impatient and share a common discount factor $\delta \in [0, 1)$. To rule out trivial cases, we assume throughout $\delta > l/\bar{v}$ as well as $\delta > 1/2$.

Let $x_n \in \{0, 1\}$ refer to the production decision in period n , *e.g.*, $x_n = 1$ means that the good is produced in period n . The principal's realized payoff

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

where $\delta \in [0, 1)$ is a discount factor. Our purpose is to derive the *optimal* policy, namely, the one that achieves the *value*, or maximum expected payoff, given that the agent seeks to maximize the expectation of his realized utility, namely,

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

We assume full commitment by the principal. Hence, it is without loss that we focus on policies in which the agent truthfully reports his type in every period, and the principal commits to a (possibly random) production decision as a function of this last report, as well as of the entire history of reports and production decisions.

Following standard arguments (see Spear and Srivastava, 1987), this problem can be described as a Markov decision process, in which the state variable is the promised utility to the agent. Let $W(U) \in \mathbf{R}$ denote the value given that the agent is promised an expected utility of U . The promised utility U is restricted to the range of $[0, \bar{v}]$, corresponding to the two possible extreme courses of actions –never or always producing the good. Conversely, for any value in the range $[0, \bar{v}]$, it is easy to construct an allocation that delivers it in expectation (for instance, an *ex ante* randomization over the two extremes).

Hence, we can directly define a policy as a map from $[0, \bar{v}] \times \{l, h\}$ into two pairs $\{p_l, u_l\}, \{p_h, u_h\} \in [0, 1] \times [0, \bar{v}]$, mapping each promised utility U and report $v = l, h$ into a probability of producing the good (p_l, p_h) and a (continuation) promised utility (u_l, u_h), subject to incentive compatibility conditions, and such that U is indeed the expected utility delivered to the agent (*promise-keeping*).

Given that payoffs are discounted, the Bellman equation characterizes both the value and the (set of) optimal policies. For any fixed $U \in [0, \bar{v}]$, the optimality equation states that

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

subject to the incentive compatibility constraints, the promise keeping constraint and the

feasibility constraint:

$$(1 - \delta)p_h h + \delta u_h \geq (1 - \delta)p_l h + \delta u_l, \quad (\text{ICH})$$

$$(1 - \delta)p_l l + \delta u_l \geq (1 - \delta)p_h l + \delta u_h, \quad (\text{ICL})$$

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

$$(p_h, p_l, u_h, u_l) \in [0, 1] \times [0, 1] \times [0, \bar{v}] \times [0, \bar{v}].$$

The dependence of the optimal policy on U is omitted whenever no confusion arises. Our objective is to calculate the payoff $W(U)$ as well as the optimal policies for any $U \in [0, \bar{v}]$. Incentive compatibility conditions and promise-keeping conditions will be referred to as IC (or ICH, ICL) and PK.

Whenever we refer to the optimal policy in the sequel, we mean the map that maximizes the principal's payoff for any given (feasible) promise U . Obviously, not the entire map might be relevant once we take into the specific choice of the initial promise. We refer to the latter as the optimal choice of the initial promise.

2.2 First-best

We start with the first-best (or symmetric-information) scenario in which the principal observes the agent's value as well. The optimization problem is the same as before except that IC constraints are dropped. Since the agent's value is i.i.d., it is without loss that we restrict attention to stationary allocation rules. We only need to determine the probability that the agent obtains the unit as a function of his realized value. For any fixed U , the principal chooses p_h and p_l to maximize

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

subject to the PK constraint $qp_h h + (1 - q)p_l l = U$. We state the results in the following lemma.

Lemma 1 *The first-best scenario admits a stationary optimal policy*

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

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

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

Hence, the optimal choice of utility in the initial period is $U_0 = qh$.

While this policy is the only Markovian policy achieving the first-best, there are many other ones that do so as well. While they are not Markovian, their structure can nonetheless be intuitive, and we will encounter some in the sequel.

2.3 Optimal Mechanism

It turns out that the solution depends on whether or not the utility is below $\underline{U} := q(h - l)$. The following lemma states that we can achieve the first-best value function if U is below $\underline{U}$. That is, when the initial value U lies in the range $[0, \underline{U}]$, $\bar{W}$ is a solution to the Bellman equation with the incentive constraints. The trick is that we can pick continuation values u_h and u_l (satisfying all the constraints) that lie in the range $[0, \underline{U}]$ as well. It then suffices to exhibit a feasible $(p_l, p_h, u_l, u_h) \in [0, 1]^2 \times [0, \underline{U}]^2$ for which $\bar{W}$ solves the Bellman equation with the incentive constraints.

While there is considerable leeway in the specification of the optimal policy in this range, we pick here

$$p_l = 0, \quad p_h = \min \left\{ 1, \frac{U}{(1 - \delta)\bar{v}} \right\},$$

with promises

$$u_h = \frac{U - (1 - \delta)p_h\bar{v}}{\delta}, \quad u_l = \frac{U - (1 - \delta)p_h\underline{U}}{\delta}.$$

Two remarks are in order. First, note that for $U = \underline{U}$, $p_h = 1$, as our assumption that $\delta \geq l/\bar{v}$ is equivalent to $\underline{U} \geq (1 - \delta)\bar{v}$. Second, we note that $u_h \leq u_l \leq U \leq \underline{U}$, where the inequality $u_l \leq \underline{U}$ follows from the definition of $\underline{U}$. It is easy to verify that $\bar{W}$ (alongside the specification of (p_l, p_h, u_l, u_h)) satisfies all the constraints, including incentive compatibility.

Also, if $U = \bar{v}$, the first-best value $\bar{W}$ can be achieved by always giving the agent the unit. This mechanism is incentive compatible. Therefore, we have $W(\bar{v}) = \bar{W}(\bar{v})$. To summarize:

Lemma 2 *For all $U \in [0, \underline{U}]$ or $U = \bar{v}$, it holds that $W(U) = \bar{W}(U)$.*

This result is simple enough, but it is rather surprising: for low enough utility levels (a range that does not vanish with discounting) first-best can be achieved. We will see in Section 3 that this only holds when the Markov chain is not too persistent. To understand how the first best is possible, note that, when the utility level is low enough, *even the same promised expected utility* tomorrow is worth more to a low type than this utility level today, despite the discounting. This is because his value might be high tomorrow. Cashing in on this expected

utility would be a bad calculation today, if this comes at the expense of this promised utility, even if foregoing the unit does not get rewarded.

We now turn to $U \in (\underline{U}, \bar{v})$. Plainly, $W \leq \bar{W}$, and W is (weakly) concave, hence continuous on $(0, \bar{v})$ and differentiable almost everywhere, with a decreasing one-sided derivative denoted W' .

Lemma 3 *For all $U \in [0, \bar{v}]$, it holds that $W \leq \bar{W}$.*

Combining Lemma 1 and 3 with the concavity of W , it follows that W' must be in the interval $[1 - c/l, 1 - c/h]$.

The IC constraints can be written as

$$(1 - \delta)(p_h - p_l)h \geq \delta(u_l - u_h) \geq (1 - \delta)(p_h - p_l)l.$$

By single-crossing, $p_h - p_l$ and $u_l - u_h$ must be weakly positive. Given that W is (weakly) concave, it is (weakly) better to decrease $u_l - u_h$ while keeping $qu_h + (1 - q)u_l$ constant. Therefore, it is without loss to assume that (ICL) binds. Based on (PK) and the binding (ICL), we solve for u_h, u_l as a function of p_h, p_l and U :

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

$$u_l = \frac{U - (1 - \delta)(p_h q(h - l) + p_l l)}{\delta}. \quad (2)$$

Lemma 4 *For all $U \in (\underline{U}, \bar{v})$, an optimal policy is such that (i) either u_h as defined in (1) equals 0 or $p_h = 1$; and (ii) either u_l as defined in (2) equals $\bar{v}$ or $p_l = 0$.*

Proof. 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 (1)–(2) (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. ■

We are almost ready to prove the main theorem of this section. Recall that for $U \leq \underline{U}$, the characterization is already achieved. We introduce $\bar{U} := q(h - l) + \delta l \in [\underline{U}, \bar{v})$.

Theorem 1 *On the range $U \in [\underline{U}, \bar{v}]$, an optimal policy is given by*

$$\begin{cases} p_h = 1, & p_l = 0 & \text{if } U \in [\underline{U}, \bar{U}], \\ p_h = 1, & p_l = 1 - \frac{\bar{v} - U}{(1 - \delta)l} & \text{if } U \in [\bar{U}, \bar{v}] \end{cases}$$

The continuation value is given by (1) and (2) respectively (given p_h and p_l), that is,

$$\begin{cases} u_h = \frac{U - (1 - \delta)\bar{v}}{\delta}, & u_l = \frac{U - (1 - \delta)\underline{U}}{\delta} & \text{if } U \in [\underline{U}, \bar{U}], \\ u_h = \frac{U - (1 - \delta)\bar{v}}{\delta}, & u_l = \bar{v} & \text{if } U \in [\bar{U}, \bar{v}] \end{cases}$$

Proof. Immediate given previous lemmata. ■

To compare the optimal policy with the first-best policy, it is useful to note that the first-best policy can be achieved in many other ways than the one described in Lemma 1. For instance, it is also optimal to give the unit if and only if the value is h as long as it is possible given the promised utility U , and update the promised utility U to $(U - (1 - \delta)qh)/\delta$; if the utility drops below $(1 - \delta)qh$, the next unit is produced with probability $U/((1 - \delta)qh)$ if the value is high, and it is not produced if it is low and it is never produced again; if the utility goes above $\delta\bar{v} + (1 - \delta)qh$, it is produced for sure if the value is high, and with probability $(U - (\delta\bar{v} + (1 - \delta)qh))/((1 - \delta)(1 - q))$ even if it is low, and forever after. In short, efficient provision is chosen, as long as doing so is possible given the promised utility. Note that, according to this policy, promised utility evolves deterministically; if $U > qh$, it increases ineluctably until absorption at $\bar{v}$; similarly, if $U < qh$, it decreases until absorption at 0. In both cases, efficiency is eventually sacrificed. If $U = qh$, promised utility remains constant, efficient provision is guaranteed indefinitely, and maximum social welfare results.

This first-best policy is almost the same as the optimal policy in the incentive-constraint problem. In both problems, efficiency is maintained as long as possible. Given promise-keeping, this also implies that expected continuation utility is the same in both cases. But in the second best, incentives require a wedge between the promised utilities after a low and a high report, of size $((1 - \delta)/\delta)(\bar{v} - \underline{U}) = (1 - \delta)l/\delta$. Promised utility in the second best

is a mean-preserving spread of promised utility in the first best. Higher volatility in the process of realized utility results, and this comes at the cost, as eventual absorption cannot be prevented.

Lemma 5 *The function W is strictly concave on $[\underline{U}, \bar{v}]$ (and so strictly below $\bar{W}$ on $(\underline{U}, \bar{v})$). It is also continuously differentiable.*

The proof can be found in Appendix A.

Lemma 6 *It holds that*

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

Proof. These limits exist as W is concave. We only deal with the first case, the second being analogous. Fix $U_0 \in [\underline{U}, \bar{U}]$ such that u_h is below $\underline{U}$. Consider the sequence $u_l, u_{ll}, \dots$; as is immediate to verify, its k -th term, denoted u_k , is $\underline{U} + \delta^{-k}x$, where $x := U_0 - \underline{U}$. Let n be the last term of this sequence for which it holds that a h report leads to a utility below $\underline{U}$. Let u_k^h denote the promised utility after k low reports followed by one high report. We have $n = \sup\{k : u_k^h \leq \underline{U}\}$. Note that $\delta^{(n+1)}/x \geq \frac{\delta}{1-\delta} \frac{1}{l}$ and $n \rightarrow \infty$ as $U_0 \downarrow \underline{U}$.

We consider a lower bound to the value of the optimal mechanism starting from U_0 , using $\tilde{W}$ instead of $W(U)$ for all $U \geq u_{n+1}$. Recall that $W(U) = (1 - c/h)U$ for $U \leq \underline{U}$. The sequence $W(u_k)$ solves the recursion

$$W(u_k) = (1 - \delta)q(h - c) + \delta q \left(1 - \frac{c}{h}\right) \frac{u_k - (1 - \delta)\bar{v}}{\delta} + \delta(1 - q)W(u_{k+1}), \quad W(u_{n+1}) = \tilde{W},$$

whose solution gives

$$\begin{aligned} \frac{W(U_0) - \left(1 - \frac{c}{h}\right)\underline{U}}{x} &= \left(1 - \frac{c}{h}\right) (1 - (1 - q)^{n+1}) + \frac{\delta^{n+1}}{x} (1 - q)^{n+1} \left(\tilde{W} - \left(1 - \frac{c}{h}\right)\underline{U}\right) \\ &\geq \left(1 - \frac{c}{h}\right) (1 - (1 - q)^{n+1}) + \frac{\delta}{1 - \delta} \frac{1}{l} (1 - q)^{n+1} \left(\tilde{W} - \left(1 - \frac{c}{h}\right)\underline{U}\right). \end{aligned}$$

Because $(1 - q)^{n+1} \rightarrow 0$ as $U_0 \downarrow \underline{U}$, it follows that $\frac{W(U_0) - \left(1 - \frac{c}{h}\right)\underline{U}}{x} \rightarrow 1 - \frac{c}{h}$, or $\lim_{U \downarrow \underline{U}} W'(U) \rightarrow 1 - \frac{c}{h}$. ■

Comparative statics:

Lemma 7 *It holds that*

1. W converges (monotonically) to $\bar{W}$ as $\delta \rightarrow 1$.

2. Fix $U_0 = U$. The process $\{U_n\}$ converges a.s. to 0 or $\bar{v}$.

If $U_0 > qh$, then $\forall \varepsilon > 0$, there exists $\bar{\delta} < 1$, $\forall \delta \in [\bar{\delta}, 1)$, $\mathbf{P}[\lim_{n \rightarrow \infty} U_n = \bar{v}] > 1 - \varepsilon$.

Similarly, if $U_0 < qh$, then $\forall \varepsilon > 0$, there exists $\bar{\delta} < 1$, $\forall \delta \in [\bar{\delta}, 1)$, $\mathbf{P}[\lim_{n \rightarrow \infty} U_n = 0] > 1 - \varepsilon$.

3. W has a unique maximizer, $U^* = \arg \max_{U \in [0, \bar{v}]} W(U)$. U^* is non-increasing in c .

Proof. Part 1. Convergence follows from standard results (e.g., Jackson and Sonnenschein, 2007), and monotonicity from the first-order stochastic dominance of the distribution of the time at which inefficiency occurs as a function of the discount factor. As for part 2, note that $|u_l - U| > |U - u_h| \Leftrightarrow U > qh$, and that $|u_l - U|, |U - u_h| \rightarrow 0$ as $\delta \rightarrow 1$, so that the result follows from Hoeffding's inequality. Part 3: Uniqueness of U^* follows from strict concavity of W ; Consider now (27)–(28), replacing the function W that appears on the r.h.s. with a function W_n , defining the left-hand side W_{n+1} iteratively, and setting $W_0 = 0$ identically. By induction, the function W_n admits a cross-partial in (U, c) a.e., which is non-positive (note that the “flow payoff” in (28) involves a non-zero cross-partial term $-(1-q)/l$). It follows that by convergence of W_n to W that W has a non-positive cross-partial (a.e.) as well, implying that U^* is non-increasing in c . ■

We note that U^* can be higher or lower than qh (and clearly it is 0 if $c = h$ and $\bar{v}$ if $c = l$). We also note that the drift of $\{U_n\}$ is the same as for the first-best (in its alternative implementation described above), which is not surprising given promise-keeping. More formally, $\mathbf{E}[U_{n+1} | U_n] > U_n$ if and only if $U_n > qh$.² However, unlike in the first-best, dynamics are not deterministic, so that it might happen that the random walk gets absorbed at 0, say, despite starting above qh . However, the probability of such an event vanishes as $\delta \rightarrow 1$. The mean time to absorption as $\delta \rightarrow 1$ follows from Lemma 7.1–2.: since the payoff until absorption is $(1 - \delta)q(h - c) = (1 - \delta)\bar{W}(qh)$, and the payoff at absorption is $W(0)$ or $W(\bar{v})$, and $W \rightarrow \bar{W}$ which is affine on $[0, qh]$ and $[qh, \bar{v}]$, it follows that as $\delta \rightarrow 1$, $\delta^\tau \rightarrow U/(qh)$ a.s. for $U \in [0, qh)$, and $\delta^\tau \rightarrow (U - (\bar{v} - c))/(\bar{v} - qh)$ a.s. for $U \in (qh, \bar{v}]$, where τ denotes the random time of absorption.

2.4 Implementation

Let $f := (1 - \delta)\underline{U}$, and $p := (1 - \delta)\bar{v} - f = (1 - \delta)l$. The obvious scheme to implement the first-best is as follows. Give the agent a budget of U^* in the initial period. At the beginning of every round, charge him a fixed fee equal to f ; if he asks for the item, produce it and

²As for the social welfare, it is clear that it is a supermartingale, since efficient provision is front loaded.

charge a fixed price of p for it; increment his budget by the interest rate $r = \frac{1}{\delta} - 1$ per period—at least, do so as long as it is feasible.

It might be infeasible for two reasons: his budget might no longer allow him to pay p for a unit that he asks for; give him whatever fraction his budget can buy (at unit price p); or his budget might be so close to $\bar{v}$ that it is no longer possible to pay him an interest rate r on his budget; give him the excess back, independently of his report, at a conversion rate given by the price p as well. It is immediate that this scheme induces truth-telling and implements the first best.

For budgets below $\underline{U}$, the agent is “in the red,” and even if he does not buy a unit, his budget will shrink. If his budget is above $\underline{U}$, he is “in the black,” and forfeiting a unit will lead to a higher budget. When the budget is above $\bar{U}$, the agent “breaks the bank” and gets to $\bar{v}$ which is an absorbing state.

This structure is somewhat reminiscent of results in the literature on optimal financial contracting (see, for instance, Biais, Mariotti, Plantin and Rochet, 2007), a literature that assumes transfers:³ in their analysis as well, one obtains (at least for some parameters) an upper absorbing boundary (where the agent gets first-best), and a lower absorbing boundary (where 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, so efficiency is achieved in this intermediate region.

It is also reminiscent of the literature on immiseration (see Thomas and Worrall, 1990, among others), but note that in our environment both immiseration and its exact opposite, ultimate affluence, are possible long-run outcomes.

3 Markovian Types

We now drop the assumption of independence. Here instead, we assume that types follow a Markov chain, with

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

where $\rho_l, \rho_h \in (0, 1)$. The (invariant) probability of h is $q := \rho_l / (\rho_h + \rho_l)$. We define $\kappa := 1 - \rho_h - \rho_l$, a useful measure of persistence. We assume that $\kappa \geq 0$, or equivalently

³There are other important differences in the set-up: they allow two instruments: downsizing the firm and payments; and the problem is of the moral hazard type, as the agent can divert resources from a risky project, reducing the chances it succeeds in a given period.

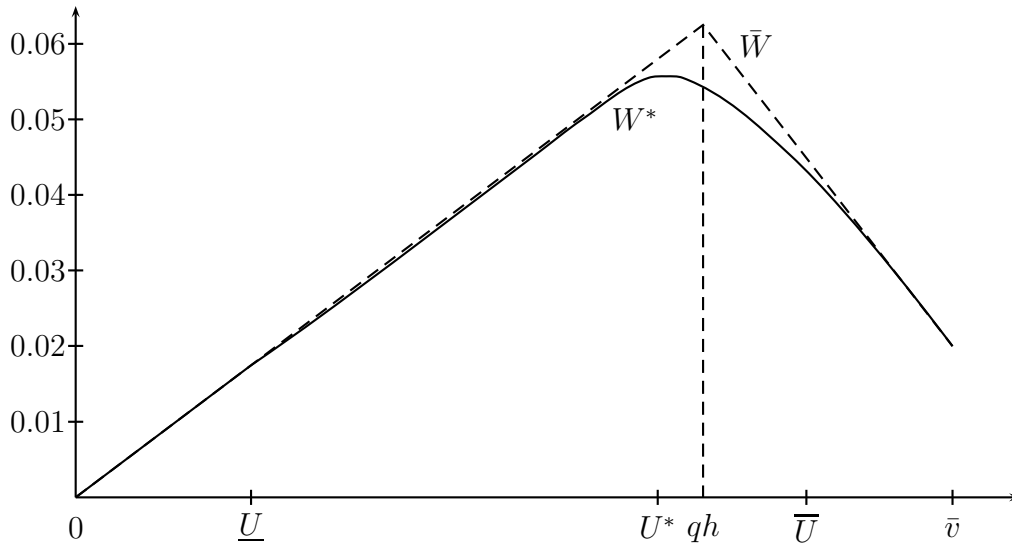


Figure 1: Optimal value, as a function of utility. Parameters are $(\delta, l, h, q, c) = (.95, .40, .60, .60, .50)$

$1 - \rho_h \geq \rho_l$: 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 . The special case $1 - \rho_h = \rho_l$ is the i.i.d. case.

Everything else in the model remains as in the baseline model, and the reader is referred to Section 2.1 for details. To fix ideas, assume that the probability that $v_0 = h$ in period 0 is q as well (although nothing depends on this).

Note that we rule out the case of perfectly persistent types, that is, the case $\rho_h = \rho_l = 0$. This case is trivial: if types never change, there is simply no possibility for the principal to use the future allocations as an instrument to elicit truth-telling. We are back to the static problem, whose solution is either to always provide the good (if $qh + (1 - q)l \geq c$, where q is the prior on the high type), or never to do so.

This suggests that persistence plays an ambiguous role *a priori*. Because current types assign different probabilities of being (say) high types tomorrow, one might hope that tying the promised utility in the future to the current reports might facilitate truth-telling. On the other hand, the case of perfectly persistent types makes clear that correlation also diminishes the scope for using future allocations as a “transfer”: utilities might still be separable between the contribution from the current and future allocations, but the current type affects both terms.

It is no longer possible to summarize continuation play (that is, the dynamic allocation rule that will be followed given the history of announcements so far) by a single utility that one could use as a state variable. This is because a given dynamic allocation rule that will be followed from the next period onward is valued differently depending on the agent's current type, as his current type is correlated with his type tomorrow.

On the other hand, conditional on the agent's type tomorrow, his type today carries no information about future types, by the Markovian assumption. Hence, we can summarize continuation play by two values: his utility conditional on his type tomorrow being high or low. Of course, his type tomorrow is not observable, so we must use instead the utility he gets from reporting his type tomorrow, conditional on truthful reporting. This creates no difficulty, as on path, the agent has incentive to report truthfully his type tomorrow. Hence, he does so as well after having lied in the previous period (conditional on his current type and his previous report, his previous type does not affect the decision problem that 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, so that the promised utility pair that we use corresponds to his actual possible continuation utilities if he plays optimally in the continuation, whether or not he reports truthfully today.⁴

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

Hence, we must now carry three state variables. The belief of the principal, $\mu = \mathbf{P}[v = h] \in [0, 1]$, and the pair of promised utilities that the principal promises as a function of the current report, U_h, U_l . We note that the highest utility $\bar{v}_h$ (resp. $\bar{v}_l$) that can be promised to a player whose type is high (resp. low) must solve⁵

$$\bar{v}_h = (1 - \delta)h + \delta(1 - \rho_h)\bar{v}_h + \delta\rho_h\bar{v}_l, \quad \bar{v}_l = (1 - \delta)l + \delta(1 - \rho_l)\bar{v}_l + \delta\rho_l\bar{v}_h,$$

that is,

$$\bar{v}_h = h - \frac{\delta\rho_h(h - l)}{1 - \delta + \delta(\rho_h + \rho_l)}, \quad \bar{v}_l = l + \frac{\delta\rho_l(h - l)}{1 - \delta + \delta(\rho_h + \rho_l)}.$$

We note that

$$\bar{v}_h - \bar{v}_l = \frac{1 - \delta}{1 - \delta + \delta(\rho_h + \rho_l)}(h - l).$$

⁴Of course, we are not the first ones to point out the necessity to use as state variable the vector of promised utilities, as opposed to the expected promised utility, in case of serial correlation. See in particular Townsend (1982), Fernandes and Phelan (2000), Cole and Kocherlakota (2001), Doepke and Townsend (2006) and Zhang and Zenios (2008).

⁵Clearly, the corresponding policy is to produce the unit in all periods independently of the reports.

Hence, as one would expect, the gap between the maximum utilities that can be promised as a function of the types decreases in the discount factor, and vanishes when $\delta \rightarrow 1$.

In addition to the state variables, we must define the choice variables. A *policy* is a pair of maps $p = (p_h, p_l) : \mathbf{R}^2 \rightarrow [0, 1]^2$, mapping the current utility vector to the probability with which the good is produced as a function of the report, and a pair of maps $(U(h), U(l)) : \mathbf{R}^2 \rightarrow \mathbf{R}^2$, mapping the current utility vector $U = (U_h, U_l)$ into the promised utilities $(U_h(h), U_l(h))$ in case the current report is h , and $(U_h(l), U_l(l))$ in case the current report is l . These definitions abuse notation, since the domain of p and $(U(h), U(l))$ should be those utility vectors that are feasible and incentive-compatible (clearly, a subset of $[0, \bar{v}_h] \times [0, \bar{v}_l]$). We postpone derivation of the domain to the next subsection.

So we define the function $W : [0, \bar{v}_h] \times [0, \bar{v}_l] \times [0, 1] \rightarrow \mathbf{R} \cup \{-\infty\}$, that solves the following program, for all $U_h \in [0, \bar{v}_h]$, $U_l \in [0, \bar{v}_l]$, and $\mu \in [0, 1]$,

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

over $(p_l, p_h) \in [0, 1]^2$, and $U_h(h), U_h(l) \in [0, \bar{v}_h]$, $U_l(h), U_l(l) \in [0, \bar{v}_l]$ subject to promise-keeping and incentive compatibility, namely,

$$U_h = (1 - \delta)p_h h + \delta(1 - \rho_h)U_h(h) + \delta\rho_h U_l(h) \quad (3)$$

$$\geq (1 - \delta)p_l h + \delta(1 - \rho_h)U_h(l) + \delta\rho_h U_l(l), \quad (4)$$

and

$$U_l = (1 - \delta)p_l l + \delta(1 - \rho_l)U_l(l) + \delta\rho_l U_h(l) \quad (5)$$

$$\geq (1 - \delta)p_h l + \delta(1 - \rho_l)U_l(h) + \delta\rho_l U_h(h), \quad (6)$$

with the convention that $\sup W = -\infty$ whenever the feasible set is empty. Note that W is concave on its domain (by linearity of the constraints in the promised utilities). The optimal policy is any map from triples (U_h, U_l, μ) into $(p_h, p_l, U_h(h), U_l(h), U_h(l), U_l(l))$ that achieves the supremum.

3.1 First-best

The first-best mechanism obtains by considering the same program, except that constraints (4) and (6) are ignored. Write $\bar{W}$ for the resulting value function. The optimal allocation rule (ignoring promises) from the principal's point of view is to produce the good if the agent's

type is h and not to produce if l . Let v_h^* (resp. v_l^*) denote the utility that a high (resp. low) type agent obtains under this optimal rule. The pair (v_h^*, v_l^*) satisfies the recursive equations

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

which give

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

Given that there are no IC constraints, the set of the promised utility (U_h, U_l) is simply $[0, \bar{v}_h] \times [0, \bar{v}_l]$. When a high type's promised utility U_h is in $[0, v_h^*]$, the principal produces the good only if the agent's type is high. Therefore, the principal's payoff is $U_h(1 - c/h)$. When $U_h \in (v_h^*, \bar{v}_h]$, the principal always produces the good if the agent's type is high. To fulfill the promised utility, the principal also produces the good when the agent's type is low. The principal's payoff is $v_h^*(1 - c/h) + (U_h - v_h^*)(1 - c/l)$. Analogously, we can calculate the principal's payoff facing a low type who is promised U_l . To sum up, $\bar{W}(U_h, U_l, \mu)$ is given by

$$\begin{cases} \mu \frac{U_h(h-c)}{h} + (1 - \mu) \frac{U_l(h-c)}{h} & \text{if } (U_h, U_l) \in [0, v_h^*] \times [0, v_l^*] \\ \mu \frac{U_h(h-c)}{h} + (1 - \mu) \left(\frac{v_l^*(h-c)}{h} + \frac{(U_l - v_l^*)(l-c)}{l} \right) & \text{if } (U_h, U_l) \in [0, v_h^*] \times [v_l^*, \bar{v}_l] \\ \mu \left(\frac{v_h^*(h-c)}{h} + \frac{(U_h - v_h^*)(l-c)}{l} \right) + (1 - \mu) \frac{U_l(h-c)}{h} & \text{if } (U_h, U_l) \in [v_h^*, \bar{v}_h] \times [0, v_l^*] \\ \mu \left(\frac{v_h^*(h-c)}{h} + \frac{(U_h - v_h^*)(l-c)}{l} \right) + (1 - \mu) \left(\frac{v_l^*(h-c)}{h} + \frac{(U_l - v_l^*)(l-c)}{l} \right) & \text{if } (U_h, U_l) \in [v_h^*, \bar{v}_h] \times [v_l^*, \bar{v}_l]. \end{cases}$$

For future purposes, it is useful to note that the slope of W (differentiable except when either $U_h = v_h^*$ or $U_l = v_l^*$) is in the interval $[1 - c/l, 1 - c/h]$, as should be expected: the latter corresponds to the most efficient way of allocating utility, the former to the most inefficient one.

3.2 Feasible and Incentive-Feasible Payoffs

As defined, the value function W might take the value $-\infty$. This may occur because the constraint set (3)–(6) might be empty. Promising to give all future units to the agent in case his current report is high, while giving none if this report is low is simply not incentive-compatible. This issue also arose with types that are independent across periods, but because it is then possible to work with promised *expected* utility, instead of the conditional utility for each possible report, there was no need to determine precisely the domain of these conditional utility pairs.

The set of *feasible* utility pairs (that is, the subset of $[0, \bar{v}_h] \times [0, \bar{v}_l]$ that solves (3) and (5)) is easy to describe. Because the two equations are uncoupled, it is simply the set

$[0, \bar{v}_h] \times [0, \bar{v}_l]$ itself. It follows that the first-best policy is independent of μ (although the expected value $\bar{W}$ obviously does).

What is challenging is to solve for the largest (bounded) subset of values (U_h, U_l) for which there exists a pair $(p_h, p_l) \in [0, 1]^2$ and two pairs $U(h) := (U_h(h), U_l(h))$ and $U(l) := (U_h(l), U_l(l))$ *within that set itself* solving (3)–(6). These are the pairs of utilities for which there exists an allocation rule and pairs of promised utilities tomorrow that are feasible and incentive-compatible (for short, incentive-feasible), and such that the promised utilities are themselves incentive-feasible, etc. Formally, we define V as follows. Given an arbitrary $A \subset [0, \bar{v}_h] \times [0, \bar{v}_l]$, let

$$\mathcal{B}(A) := \{(U_h, U_l) \in [0, \bar{v}_h] \times [0, \bar{v}_l] : \exists (p_h, p_l) \in [0, 1]^2, U(h) \in A, U(l) \in A \text{ solving (3)–(6)}\},$$

a possible empty set. We let V be the largest bounded fixed point of $\mathcal{B}$ (this operator being monotone, it is well-defined). Clearly, this is very close to the notion of self-generation in repeated games (see Abreu, Pearce and Stacchetti, 1990), though in the context of different types of a given agent, as opposed to the different players in the game.⁶

Our first step towards solving for the optimal mechanism is to solve for V . Clearly, V must contain the points $(0, 0)$ and $(\bar{v}_h, \bar{v}_l)$, as $(0, 0)$ can be obtained by the choices $(p_h, p_l) = (0, 0)$ and $U(l) = U(h) = (0, 0)$ (in an abuse of notation, we also write 0 for the $(0, 0)$ vector), and similarly for $\bar{v} := (\bar{v}_h, \bar{v}_l)$. Clearly, the segment connecting those two points can be obtained as well, by picking some $p_h = p_l$ and promises $U(h), U(l)$ on that segment. What is challenging is to derive the boundary of this (clearly, compact convex) set V , especially since our interest does not only lie in the limit as $\delta \rightarrow 1$.

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

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

$$\bar{u}_l^\nu = \delta^\nu \bar{v}_l + \delta^\nu q(\bar{v}_h - \bar{v}_l)(1 - \kappa^\nu), \quad (8)$$

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

$$\underline{u}_h^\nu = (1 - \delta^\nu) \bar{v}_h + \delta^\nu (1 - q)(\bar{v}_h - \bar{v}_l)(1 - \kappa^\nu), \quad (9)$$

$$\underline{u}_l^\nu = (1 - \delta^\nu) \bar{v}_l - \delta^\nu q(\bar{v}_h - \bar{v}_l)(1 - \kappa^\nu), \quad (10)$$

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 $\bar{u}^0 = \bar{v}$, with $\lim_{\nu \rightarrow \infty} \bar{u}^\nu = 0$. Similarly, $\underline{u}^\nu$ is increasing, with $\underline{u}^0 = 0$ and $\lim_{\nu \rightarrow \infty} \underline{u}^\nu = \bar{v}$.

⁶The distinction is not merely a matter of interpretation, as a high type can become a low type and vice-versa, for which there is no analogue in repeated games.

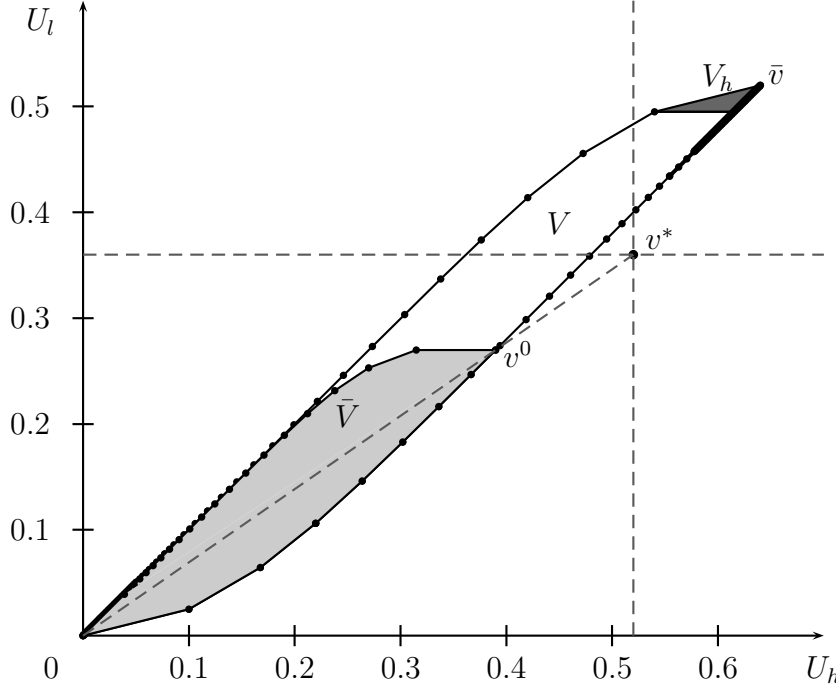


Figure 2: The set V for parameters $(\delta, \rho_h, \rho_l, l, h) = (9/10, 1/3, 1/4, 1/4, 1)$

The main result of this subsection is the following. The proof of this lemma and several of the next ones are gathered in Appendix B.

Lemma 8 *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 (but nevertheless only two accumulation points). From now on, we may restrict attention to utility pairs in V , as this is the domain of W over which $W > -\infty$. We note that this does not mean that the probability pair $p := (p_h, p_l)$ is unrestricted: plainly, the utility pair $\bar{v}$, for instance, is only incentive-feasible if the principal sets $p = (1, 1)$. See Figure 3.2 for an illustration. The subsets $\bar{V}$ (defined in Lemma 9) and V_h (defined after Lemma 9) are also depicted.

It is easily checked that

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

In particular, the slopes of the upper and lower loci are less than 1. Because $(\bar{v}_l - v_l^*)/(\bar{v}_h - v_h^*) > 1$, it follows that the vector v^* is outside V .

The vertices of V admit a simple interpretation. The utility vector $\bar{u}^\nu$ corresponds to *backloading* of good provision: reports are ignored, and the good is not provided for ν

consecutive periods, and then provided forever. (Utility vectors between two edges obtain by randomizing between ν and $\nu + 1$ periods). Similarly, the utility vector $\underline{u}^\nu$ corresponds to *frontloading* of good provision: here as well, reports are ignored, and the good is provided for ν consecutive periods, and then never again. This is intuitive: if one's type is low today, and adjusting for discounting, it is best to get a given promised unit as late as possible, as the probability that the type is high is increasing over time, given that the current type is low; conversely, a high type prefers to get it as soon as possible, as the probability that the type is high is decreasing over time, given that the current type is high. As a result, for a given utility promise to the high type, it is worst for the low type if it corresponds to frontloading, and best if it corresponds to backloading.

This explains why the shape of V is particularly simple in the i.i.d. case: the lower type prefers to get a larger fraction (or probability) of the good tomorrow rather than today (adjusting for discounting), but has no preference over later days; and similarly for the high type. As a result, all the vertices $\{\bar{u}^\nu\}_{\nu=1}^\infty$ (resp., $\{\underline{u}^\nu\}_{\nu=1}^\infty$) are perfectly aligned, and V is a quadrangle (in fact, as easily checked, a parallelogram) whose vertices are 0 , $\bar{v}$, $\bar{u}^1$ and $\underline{u}^1$.

While the lower and upper boundary are most easily understood in terms of these extreme policies (front- and backloading), these are not the only policies achieving those boundaries. It is not hard to see that the lower locus corresponds to those policies that (starting from this locus) assign as high a probability as possible to the good being produced whenever the report is h , while promising continuation utilities that make ICL bind in every period. Similarly, the upper boundary corresponds to those policies that (starting from this locus) assign as low a probability as possible to the good being produced whenever the report is l , while promising continuation utilities that make ICH bind in every period. Front- and backloading as much as possible are representative examples in each class.

3.3 The Optimal Mechanism

As in the i.i.d. case, it is actually possible to achieve the first-best allocation, for a given subset of promised utilities (and aside from the trivial promises 0 and $\bar{v}$). But not necessarily: persistence matters.

To describe this subset, we define the point $v^0 := (v_h^0, v_l^0)$ as the intersection in $\mathbf{R}_{++}$, if any, of the (lower) boundary of V with the line

$$U_l = \frac{\delta \rho_l}{1 - \delta(1 - \rho_l)} U_h.$$

There might be no (strictly positive) intersection, because the line $U_l = \frac{\delta \rho_l}{1 - \delta(1 - \rho_l)} U_h$ that

defines v^0 might be flatter than the flattest segment of the lower boundary of V (namely, the segment that connects the vector 0 to $\underline{u}^1$). An immediate computation gives that this intersection exists if and only if⁷

$$\frac{h-l}{l} > \frac{1-\delta}{\delta\rho_l}. \quad (11)$$

Otherwise, we set $\bar{V} = \{0\}$. When (11) holds, the point v^0 lies along the segment $[0, v^*]$. We then define the sequence $\{v^\nu\}_{\nu \geq 1}$ by

$$v_h^\nu = \delta^\nu ((1-q)v_l^0 + qv_h^0 + (1-q)\kappa^\nu(v_h^0 - v_l^0)), \quad v_l^\nu = \delta^\nu ((1-q)v_l^0 + qv_h^0 - q\kappa^\nu(v_h^0 - v_l^0)),$$

and define

$$\bar{V} = \text{co}\{\{(0,0)\} \cup \{v^\nu\}_{\nu \geq 0}\}. \quad (12)$$

Note that $\bar{V}$ has non-empty interior if and only if ρ_l is sufficiently large, see (11).

We have the following result. Recall that $\bar{W}$ is the maximum value when ignoring incentive compatibility.

Lemma 9 *For all $U = (U_h, U_l) \in \bar{V}$, or $U = (\bar{v}_h, \bar{v}_l)$, it holds that, for all μ ,*

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

Conversely, if $U \in V \setminus \bar{V}$, $U \neq (\bar{v}_h, \bar{v}_l)$, then $W(U, \mu) < \bar{W}(U, \mu)$ for all μ .

Let V_h be $\{(U_h, U_l) : (U_h, U_l) \in V, U_l \geq \bar{u}_l^1\}$ 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 enforceable at U if and only if $U \in V \setminus (V_h \cup V_l)$.

Finally, we start examining the optimal policy in the domain $V \setminus \bar{V}$. We shall introduce one more sequence, Namely, we define $\hat{u}^\nu := (\hat{u}_h^\nu, \hat{u}_l^\nu)$, $\nu \geq 0$, as follows:

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

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

⁷This condition is trivially satisfied in our analysis of the i.i.d. case because of our maintained assumption that $\delta > l/\bar{v}$, see Section 2.

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

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

Definition 1 For all $U = (U_h, U_l) \in V$, set

$$p_l = \max \left\{ 0, 1 - \frac{\bar{v}_l - U_l}{(1 - \delta)l} \right\}, \quad p_h = \min \left\{ 1, \frac{U_h}{(1 - \delta)h} \right\}, \quad (13)$$

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 so that ICH binds.

For each continuation utility vector $U(h)$ or $U(l)$, this gives one constraint (either an incentive constraint, or the constraint that the utility vector lies on one of the boundaries). In addition to the two promise-keeping equations, this gives four constraints, which uniquely define the pair of points $(U(h), U(l))$. It is readily checked that the policy as well as the choices of $U(l), U(h)$ also imply that ICL binds for all $U \in P_b$.

A very surprising property of this policy is its independence of the principal's belief μ . That is, the principal's belief about the agent's value is entirely irrelevant, given the promised utility. However, we will see that the initial choice of promised utility depends on this belief.

Figure 3 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 it is constrained in $\bar{V} \cup V_h \cup V_l$), and promised utilities depend on the report: a report of l takes the utility to the right (towards higher utilities), while a report of h takes it to the left and to the lower boundary. Below the polygonal chain, the l report also takes us to the lower boundary (and ICL binds), while above it, it does not, and it is ICH which is binding. In fact, note that if the utility vector is on the upper boundary, the continuation utility after l remains there.

Theorem 2 Fix $U_0 \in V$; given U_0 , the policy stated above is optimal. The optimal choice of U_0 is in $P_b \cap (V \setminus \bar{V})$, with U_0 increasing in the principal's initial belief of the high type.

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

Given that $U_0 \in P_b$, and given the structure of the optimal policy, the promised utility vector actually never leaves P_b . It is also simple to check that, as in the i.i.d. case (and with the

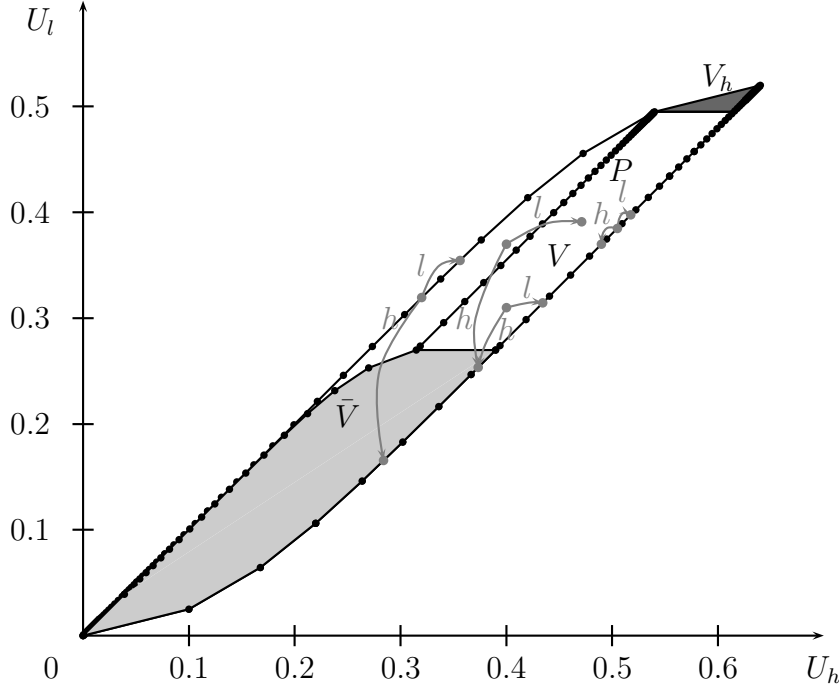


Figure 3: Optimal continuation utilities for $(\delta, \rho_h, \rho_l, l, h) = (9/10, 1/3, 1/4, 1/4, 1)$

same arguments), the (one-sided) derivative of W approaches the derivative of W^* as either U approaches $\bar{v}$ or the set $\bar{V}$. As a result, the optimal choice of U_0 is strictly interior.

Proof. 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 $\mu \in \{\rho_l, 1 - \rho_h\}$,

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

over $(p_l, p_h) \in [0, 1]^2$, and $U(h), U(l) \in V$ subject to PKH, PKL, ICL. Note that ICH 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 $U \in \bar{V}$. So, we focus on the subset $V \setminus \bar{V}$.

1. We want to show that for any U , it is optimal to set p_h, p_l as in (13) 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 ICL and PKH lies above P_b . It is also feasible to choose such a $U(l)$ as ICH 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 $[0, (1 - \rho_h)\bar{v}_h + \rho_h \bar{v}_l]$. 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 $[0, \rho_l \bar{v}_h + (1 - \rho_l)\bar{v}_l]$. 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

$$w'_h(X(U)) \geq w'_l(Y(U)). \quad (14)$$

Note that we have $w'_h(X(U)) = w'_l(Y(U)) = 1 - c/h$ when $U \in \bar{V}$. For any fixed $U \in P_b \setminus (\bar{V} \cup V_h)$, 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 $U \in P_b \setminus (\bar{V} \cup V_h)$. 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 (14) 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 $\bar{W}$ on $\bar{V}$;

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 (13) 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,

$$W_1(U_h + \varepsilon, U_l, \mu) - W_1(U_h, U_l, \mu) \geq W_1(U_h, U_l + \frac{1 - \rho_h}{\rho_h}\varepsilon, \mu) - W_1(U_h, U_l, \mu). \quad (15)$$

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

$$\delta(1 - \rho_h) \left(W_0(\tilde{U}_h(h), \tilde{U}_l(h), 1 - \rho_h) - W_0(U_h(h), U_l(h), 1 - \rho_h) \right), \quad (16)$$

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 $U \in V \setminus (\bar{V} \cup V_h)$, the right-hand side equals

$$\delta \rho_h \left(W_0(\tilde{U}_h(l), \tilde{U}_l(l), \rho_l) - W_0(U_h(l), U_l(l), \rho_l) \right), \quad (17)$$

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 (23) is greater than (24). 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 (15) is equivalent to

$$w'_h \left(\frac{U_h - (1 - \delta)h}{\delta} \right) \geq w'_l \left(\frac{U_l}{\delta} \right), \quad \forall (U_h, U_l) \in V \setminus (\bar{V} \cup V_h \cup V_l). \quad (18)$$

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 (18) holds when (U_h, U_l) are on P_b . This is true given that (i) w_h, w_l are concave; (ii) inequality (14) 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 (15) is given by $(1 - \rho_h)\varepsilon(1 - c/l)$. Inequality (15) 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 ICH, 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 PKH and ICL; (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 PKL and ICH. 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 ICH binds for $U \in V_t$. Next, we show that the conjecture policy is the solution to the second sub-problem.

For a fixed $U \in V_t$, PKL and ICH 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 PKL and ICH, 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

$$\gamma_h \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_h} + \gamma_l \frac{\partial W(U_h, U_l, \rho_l)}{\partial U_l} \leq \frac{(1 - \delta)(c - l)}{\delta}, \quad (19)$$

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 (19) 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 (19) follows. Using similar arguments, we can show that (19) 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 , (19) holds within V .

For any $x \in [0, (1 - \rho_h)\bar{v}_h + \rho_h\bar{v}_l]$, 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 (14), 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 . ■

Utility Dynamics: One of the striking results of Battaglini (2005)'s analysis is that the utility process of the agent always goes up—in fact, first best is achieved almost surely along every history. Here instead, a simple computation yields that, along the optimal path (so, on P_b), for any $U = (U_h, U_l)$, it holds that

$$\begin{aligned} q(\rho_h U_l(h) + (1 - \rho_h)U_h(h) - U_h) + (1 - q)(\rho_l U_h(l) + (1 - \rho_l)U_l(l) - U_l) \\ = \frac{1 - \delta}{\delta} (q(U_h - h) + (1 - q)U_l), \end{aligned}$$

implying that utility goes up on average if and only if

$$\rho_h U_l + \rho_l U_h \geq \rho_l h.$$

As is immediate, this condition is satisfied at $U = (0, 0)$, and violated at $U = (\bar{v}_h, \bar{v}_l)$. Hence, there exists a critical value of U_h —say, $\hat{U}_h$ (or equivalently $\hat{U}_l$) such that utility goes up (in expectations) if and only $U_h \geq \hat{U}_h$. The critical value is obtained by intersecting the lower boundary of V with the line $\rho_h U_l + \rho_l U_h = \rho_l h$.

It is also immediate that, given any initial choice of $U_0 \notin \bar{V} \cup \{\bar{v}\}$, finitely many consecutive reports of l (or h) suffice for the promised utility to reach $\bar{v}$ (or 0). As a result, both long-run outcomes have strictly positive probability under the optimal policy, for any optimal initial choice. Furthermore, absorption occurs with probability 1 a.s..

3.4 Implementation

While the analysis has required keeping track of the two-dimensional state variable, it leads to a policy that can be described with a one-dimensional state variable, because expected utilities never leave the lower locus of the feasible set V . As in the i.i.d. case, the low type's incentive constraint dictates the dynamics, as he is tempted to pretend being a high type and get the unit. To understand these dynamics, it is best to think of the utility vectors on the lower locus as the payoffs that would result, given the initial type, from giving the unit to the agent for a certain number of periods, irrespective of his sequence of messages, and then never again. Because of the discreteness of time, we represent such a policy by a pair $(n, \lambda) \in \mathbf{N}_0 \times [0, 1)$ (or by $n = \infty$) with the interpretation that the good is awarded for n periods with probability λ , and $n + 1$ periods with probability $1 - \lambda$ (or forever when $n = \infty$).

Each such policy that maximally frontloads the provision of the good leads to an expected utility conditional on the agent's initial type, which we write $U_h(n, \lambda), U_l(n, \lambda)$. If $n = \infty$, then the unit is promised forever, with corresponding utilities $\bar{v}_h, \bar{v}_l$.

We may think of the optimal policy as follows. In a given period t , the agent is promised (n_t, λ_t) . If the agent asks for the unit (and this is feasible, that is, $n_t \geq 1$), the next promise (n_{t+1}, λ_{t+1}) , is then the solution to

$$\frac{U_l(n_t, \lambda_t) - (1 - \delta)l}{\delta} = \mathbf{E}_l [U_{\theta_{t+1}}(n_{t+1}, \lambda_{t+1})], \quad (20)$$

where $\mathbf{E}_l [U_{\theta_{t+1}}(n_{t+1}, \lambda_{t+1})] = (1 - \rho_l)U_l(n_{t+1}, \lambda_{t+1}) + \rho_l U_h(n_{t+1}, \lambda_{t+1})$ is the expected utility from tomorrow's promise (n_{t+1}, λ_{t+1}) given that today's type is low. If $n_t < 1$ and the agent

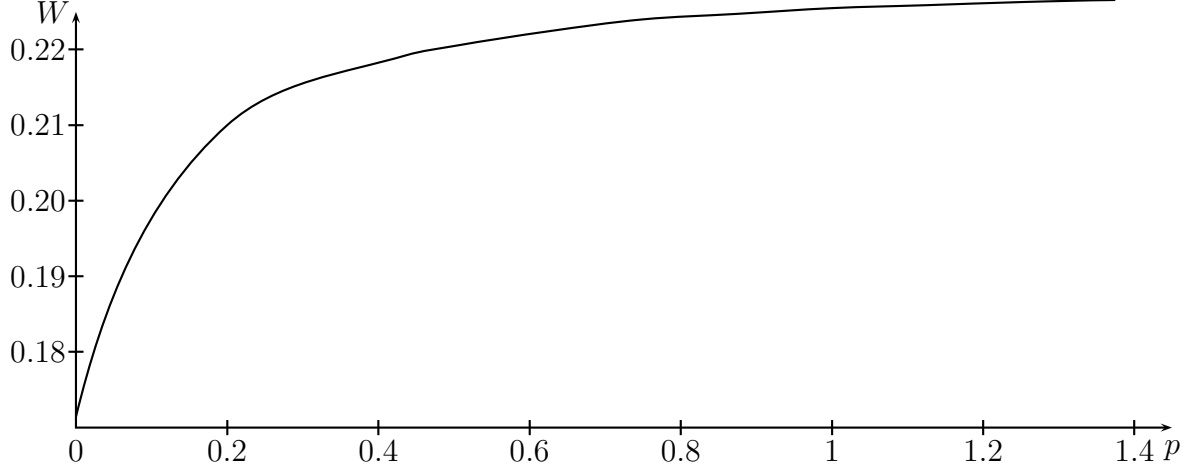


Figure 4: Payoff as a function of persistence $(\delta, \rho_h, \rho_l, l, h) = (9/10, p/3, p/4, 1/4, 1)$ (The initial belief is set at q).

claims to be high, he then gets with the probability $\tilde{q}$ that solves $U_l(n_t, \lambda_t) - \tilde{q}(1 - \delta)l = 0$.) On the other hand, claiming to being low simply leads to the revised promise

$$\frac{U_l(n_t, \lambda_t)}{\delta} = \mathbf{E}_l [U_{\theta_{t+1}}(n_{t+1}, \lambda_{t+1})], \quad (21)$$

provided that there exists a (finite) n_{t+1} and $\lambda_{t+1} \in [0, 1]$ that solve this equation.⁸ While it was perhaps more natural in the i.i.d. case to use the expected utility as opposed to the utility of a low type to describe the optimal policy, the policy described by (20)–(21) reduces to the one described in Section 2.4 in that case (a special case of the Markovian one). The policy described in the i.i.d. case obtains by taking expectations of these dynamics with respect to today's type.

It is perhaps surprising that the optimal policy can be solved for. Less surprising is that comparative statics are difficult to obtain by other means than numerical simulations. Figure 4 illustrates the impact of increasing persistence. 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 so not either the value $\bar{v}$. As can be seen, a decrease in persistence (increase in p) leads to a higher payoff. When $p = 0$, types never change and we are left with

⁸This is impossible if the promise (n_t, λ_t) is already too large (formally, if the corresponding payoff vector $(U_h(n, \lambda), U_l(n, \lambda)) \in V_h$), in which case the good is given even in that event with the probability that solves $\frac{U_l(n_t, \lambda_t) - \tilde{q}(1 - \delta)l}{\delta} = \mathbf{E}_l [U_{\theta_{t+1}}(\infty)]$.

a static problem (for the parameters chosen here, it is then best not to provide the good). When p increases, types change more rapidly, so that promised utility becomes a frictionless currency.

As mentioned, this comparative statics is merely suggested by (numerous) simulations. Given that promised utility varies as a random walk with unequal step size, on a grid that is itself a polygonal chain, there is a little hope to establish this result more formally here. To derive sharper analytic insights, we turn to a tractable limiting case.

3.5 The Continuous Limit

In this subsection, we examine the limiting stochastic process of utility and payoff as transitions are scaled according to the usual Poisson limit, when variable period 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 , so that, by convention, $v_t = l$ on any interval $[T_{2k}, T_{2k+1})$).

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, which takes values in $\{0, 1\}$. Namely, $\mu_t = 0$ over any interval $[T_{2k}, T_{2k+1})$, and $\mu_t = 1$ otherwise. Second, the utilities of the agent $(U_{l,t}, U_{h,t})_{t \geq 0}$, as a function of his type. Finally, the expected payoff of the principal, $(W_t)_{t \geq 0}$, 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

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

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

$$(\bar{v}_h, \bar{v}_l) = \left(h - \frac{\lambda_h}{\lambda_h + \lambda_l + r}(h - l), l + \frac{\lambda_l}{\lambda_h + \lambda_l + r}(h - l) \right)$$

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

given by

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

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

$$\begin{aligned}v_h(\tau) &= e^{-r\tau}((1-q)l_0 + qh_0) + e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau})(1-q)(h_0 - l_0), \\ v_l(\tau) &= e^{-r\tau}((1-q)l_0 + qh_0) - e^{-r\tau}(1 - e^{-(\lambda_h+\lambda_l)\tau})q(h_0 - l_0),\end{aligned}$$

where (h_0, l_0) are the coordinates of the 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 immediate to check that $\bar{V}$ has nonempty interior iff (cf.(11))

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

Hence, first-best cannot be achieved for any utility level (aside from 0 and $\bar{v}$) whenever the low state is too persistent. On the other hand, $\bar{V}$ is always non-empty when the agent is sufficiently patient.

Figure 5 illustrates this construction. Note that the boundary of V is smooth, except at 0 and $\bar{v}$. It is also easy to check 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 payoffs can be explicitly solved for. 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 of time over which h is continuously reported, it evolves deterministically, with increments

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

On the other hand, when l is reported, the evolution is more complicated. Algebra gives that

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

where

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

and $\bar{v} = qh + (1-q)l$, as before.

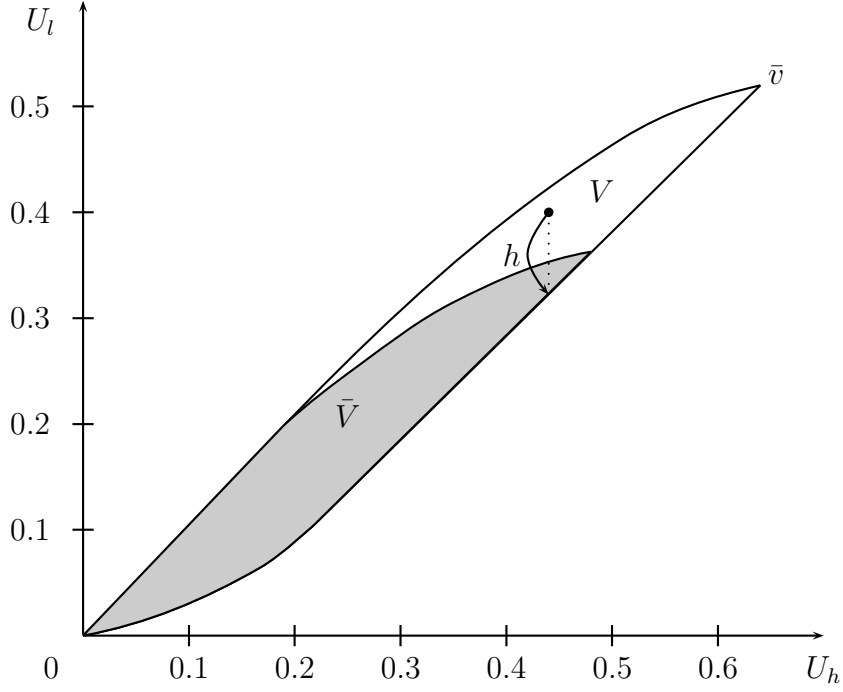


Figure 5: Incentive-feasible set for $(r, \lambda_h, \lambda_l, l, h) = (1, 10/4, 1/4, 1/4, 1)$

The increment $d\tau_l$ is positive or negative, depending upon whether τ maps into a utility vector in $\bar{V}$ or not. If $\bar{V}$ has nonempty interior, we can identify the value of τ that is the intersection of the critical line and the boundary; call it $\hat{\tau}$, which is simply the positive root (if any) of g . Otherwise, set $\hat{\tau} = 0$.

It might be worth pointing out that the evolution of utility is not continuous for utilities that are not on the lower boundary. A high report leads to a vertical jump in the utility of the low type, down to the lower boundary. See Figure 5. This is intuitive, as by promise-keeping the utility of the high type agent cannot jump, as such an instantaneous report has only a minute impact on his flow utility. A low report, on the other leads to a drift in the type's utility.

Our goal is to derive the principal's payoff (or value) functions. 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 (resp., low). By definition of the policy that is followed, the value functions solve the 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

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

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

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

and

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

subject to the following boundary conditions.⁹ First, at $\tau = \hat{\tau}$, the value must coincide with the one given by the first-best payoff $\bar{W}$ on 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 $\bar{v}$ be approached. Hence,

$$\lim_{\tau \rightarrow \infty} W_h(\tau) = \bar{v}_h, \lim_{\tau \rightarrow \infty} W_l(\tau) = \bar{v}_l.$$

Despite having variable coefficients, it turns out that this system can be solved. 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

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

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) that¹⁰

Proposition 1 *The value function of the principal is given by*

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

⁹To be clear, these are not HJB equations, as there is no need to verify the optimality of the policy that is being followed. This has already been established. These are simple recursive equations that these functions must satisfy.

¹⁰As $\tau \rightarrow t_0$, 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 $W_1(t_0)$.

where

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

It is straightforward to derive the closed-form expressions for first-best payoff, which we omit here. Figure 6 illustrates the value function for two levels of persistence, and compares it to the first-best payoff evaluated along the lower locus, $\bar{W}$ (the lower envelope of three curves).

Lemma 10 *The value $W(\tau)$ decreases pointwise in persistence $1/p$, where $\lambda_h = p\bar{\lambda}_h$, $\lambda_l = p\bar{\lambda}_l$, for some fixed $\bar{\lambda}_h, \bar{\lambda}_l$.*

The proof is in Appendix. Hence, persistence hurts the principal’s payoff, as is intuitive: with independent types, the agent’s preferences are quasilinear in promised utility, so that the only source of inefficiency derives from the bounds on this currency. When types are correlated, promised utility no longer enters independently of today’s types in the agent’s preferences, reducing the degree to which this can be used to provide incentives efficiently. With perfectly persistent types, there is no leeway anymore, and we are back to the inefficient static outcome.

How about the agent’s utility? We note that the utility of both types is increasing in τ . Indeed, since 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 get the good if he continuously claimed to be of the high type: this is precisely the definition of τ . But persistence plays an ambiguous role on the agent’s utility: indeed, perfect persistence is his favorite outcome if $\bar{v} > c$, so that always providing the good is best in the static game. Conversely, perfect persistence is worse if $\bar{v} < c$. Hence, persistence tends to improve the agent’s situation when $\bar{v} > c$.¹¹ As $r \rightarrow 0$, the principal’s value converges to the first-best payoff $q(h - c)$. Jackson and Sonnenschein (2007) shows that the rate of convergence is (at least) polynomial (in the number of identical copies). Cohn (2010) strengthens this result by exhibiting a refinement of Jackson and Sonnenschein’s mechanism that achieves convergence at an exponential rate (see also Eilat and Pauzner, 2011, for an exactly optimal mechanism in a simpler static setting). However, his mechanism is cast in the static version of the Bayesian decision problem, in which the agent knows the value of all units ahead of time. While it is unclear *a priori* whether this makes it easier or harder to overcome the incentive constraints, we show that his result carries over to the optimal mechanism in our environment. Namely, we show that the value converges to the first-best payoff at a rate that is linear in the discount rate

¹¹However, this convergence isn’t necessarily monotone, as is easy to check via examples.

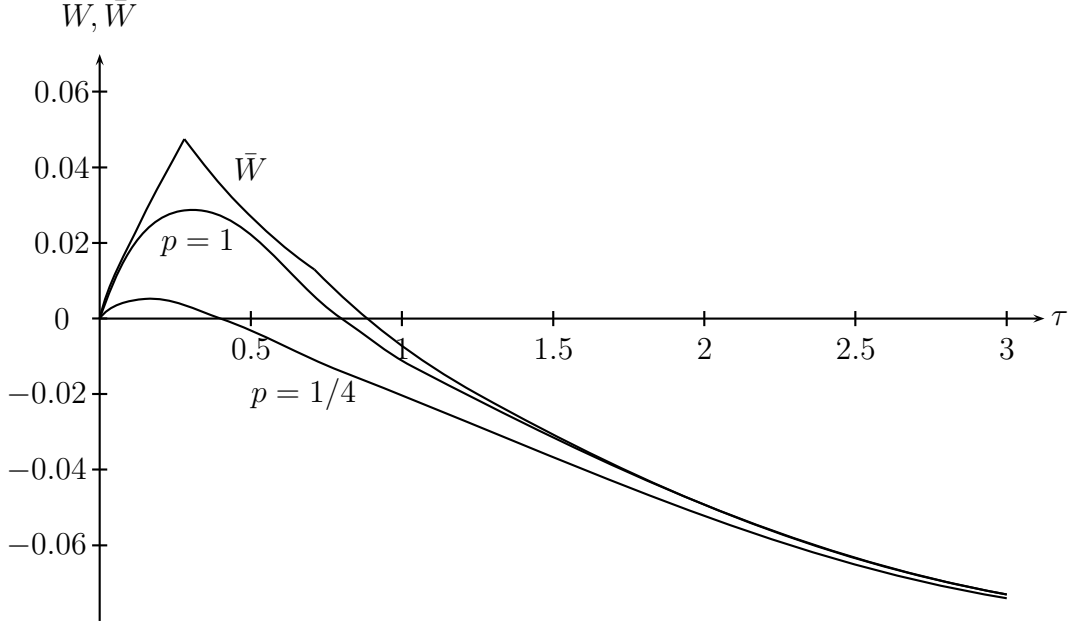


Figure 6: Equilibrium and first-best payoffs, as a function of τ (here, $(\lambda_l, \lambda_h, r, l, h, c) = (p/4, 10p/4, 1, 1/4, 1, 2/5)$ and $p = 1, 1/4$)

(which we can interpret as a hazard rate, as a way of connecting this result to the other ones that involve a finite number of identical copies).

Lemma 11 *It holds that*

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

3.6 A Comparison with the Case with Transfers

As mentioned, our model can be viewed as a no-transfer counterpart of Battaglini (2005).

At first sight, the difference in results is striking. One of the main findings of Battaglini, “no distortion at the top,” has no counterpart. With transfers, efficient provision occurs *forever* as soon as the agent reveals to be of the high type. Also, as noted, with transfers, even along the one history in which efficiency is not achieved in finite time, namely an uninterrupted string of low reports, efficiency is asymptotically approached. Here instead, as explained, we necessarily end up (with probability one) with an inefficient outcome, which can be implemented without using further reports. And both such outcomes (providing the good forever or never again) can arise. In summary, inefficiencies are frontloaded as much as possible with transfers, while 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 to achieve: simply charge c whenever the good has to 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 mimick 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 the binding one. Ideally, the principal would like to charge for these rents *before* the agent has private information, when the expected value of these rents to the agent are still common knowledge. When types are i.i.d., this poses no difficulty, and these rents can be expropriated one period ahead of time; with correlation, however, different types of the agent value these rents differently, as their likelihood of being high in the future depends on their current type. However, when considering information rents far enough in the future, the initial type hardly affects the conditional expectation of the value of these rents, so that they can be “almost” extracted. As a result, it is in the principal’s best interest to maximize the surplus and so offer a nearly efficient contract at all dates sufficiently far away.

We see that money plays two roles. First, because it is an instrument that allows to “clear” promises 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 this was the objective. Another role of money, as highlighted by Battaglini, is that it allows transferring value from the agent to the principal before private information realizes, so that information rents no longer stand in the way of efficiency, at least, as far as the remote future is concerned. Hence, these future inefficiencies can be eliminated, so that inefficiencies only arise in the short run.

4 More types

It is important to understand the role played by the assumption of two types only. Obviously, this is restrictive. As we know (see for instance, Battaglini and Lamba, 2014), identifying the binding incentive constraints and thus solving the problem becomes hard with more types, even with transfers. The situation is unlikely to improve without transfers. Nonetheless, such an analysis is called for in order to evaluate the robustness of our findings.

4.1 Independent Types

Suppose here that types are drawn i.i.d. from some atomless distribution F with support $[\underline{v}, 1]$, $\underline{v} \in [0, 1)$, and density $f > 0$ on $[\underline{v}, 1]$. Let $\bar{v} = \mathbf{E}[v]$ be the expected value of the type, and so the highest promised utility. Assume that the inverse hazard rate $\lambda(v) = \frac{1-F(v)}{f(v)}$ is differentiable and such that $v \mapsto \lambda(v) - v$ is monotone.

We start with the statement regarding the first-best policy.

Lemma 12 *The first-best policy is unique. The first-best value function W is strictly concave. The optimal policy is of the threshold type, with threshold v^* that is continuously decreasing from 1 to 0 as U goes from 0 to $\bar{v}$. Furthermore, given the initial promised utility U , future promised utility is equal to U .*

That is, given promised utility $U \in [0, \bar{v}]$, 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.

How about second-best? In an additional appendix, we prove that

Theorem 3 *The second-best payoff is strictly concave in U , continuously differentiable, and strictly below first-best (except for $U = 0, \bar{v}$). Given $U \in (0, \bar{v})$, the optimal policy $q : [0, 1] \rightarrow [0, 1]$ is not a threshold policy.*

Once again, we see 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 q , so that, given some value of U , all types above a certain threshold get the good supplied, while those below get it with probability zero. However, without transfers, incentive compatibility requires continuation utility to 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 point of view, conditional on the agent being in this interval, the outcome is a lottery over $q = 0, 1$, and corresponding continuation payoffs. Replacing this lottery by its expected value would leave the agent virtually indifferent, but it would certainly help the principal, because 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 for the binary case. Nonetheless, it follows from the envelope theorem that the marginal cost C' (where $C(U) := W(U) - U$) is a bounded martingale, that is, U -a.e.,

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

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

Lemma 13 *Given any initial level U^0 , the utility process U_n converges to $\{0, \bar{v}\}$, with both limits having strictly positive probability if $\underline{v} > 0$ (If $\underline{v} = 0$, 0 occurs a.s.).*

In the rest of this sub-section, we show that the optimal policy may be found using control theory. Let $x_1(v) = q(v)$ and $x_2(v) = \mathcal{U}(U, v)$. The optimal policy x_1, x_2 is the solution to the control problem,

$$\max \int_0^1 (1 - \delta)x_1(v)(v - c) + \delta W(x_2(v))dF$$

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

$$x_3(v) = \delta \mathcal{U}(U, v) + (1 - \delta) \int_v^1 q(s) \bar{F}(s) ds.$$

The law of motion of x_3 is $x'_3 = -(1 - \delta)(vu + x_1 \bar{F})$. The constraints are

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

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

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

The first-order conditions are

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

The transversality conditions are

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

The first observation is that $\gamma_3(v)$ is constant, denoted γ_3 . Moreover, $\dot{\gamma}_1$ involves no endogenous variables. Therefore, for a fixed $\gamma_1(1)$, the trajectory of γ_1 is determined. Whenever $u > 0$, we have $\mu = 0$. The first-order condition $\frac{\partial \mathcal{L}}{\partial u} = 0$ implies that

$$\gamma_2 = \delta \left(\frac{\gamma_1}{(1-\delta)v} - \gamma_3 \right) \quad \text{and} \quad \dot{\gamma}_2 = \frac{\delta(\gamma_1 - v\dot{\gamma}_1)}{(\delta-1)v^2}.$$

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

$$x_2 = (W')^{-1} \left(\frac{v\dot{\gamma}_1 - \gamma_1}{(\delta-1)fv^2} \right). \quad (22)$$

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

Case one occurs when U is sufficiently large: There exists $0 \leq v_1 \leq v_2 < 1$ 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 correspondingly

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

The continuity of x_1 at v_2 requires that

$$-\frac{\delta}{1-\delta} \int_{v_1}^{v_2} \frac{\dot{x}_2}{s} ds = 1. \quad (23)$$

The trajectory of γ_2 is given by

$$\gamma_2 = \begin{cases} \delta \left(\frac{\gamma_1(v_1)}{(1-\delta)v_1} - \gamma_3 \right) + \delta(F(v_1) - F(v)) \frac{v_1 \dot{\gamma}_1(v_1) - \gamma_1(v_1)}{(\delta-1)f(v_1)v_1^2} & \text{if } v < v_1 \\ \delta \left(\frac{\gamma_1(v)}{(1-\delta)v} - \gamma_3 \right) & \text{if } v_1 \leq v \leq v_2 \\ \delta \left(\frac{\gamma_1(v_2)}{(1-\delta)v_2} - \gamma_3 \right) - \delta(F(v) - F(v_2)) \frac{v_2 \dot{\gamma}_1(v_2) - \gamma_1(v_2)}{(\delta-1)f(v_2)v_2^2} & \text{if } v > v_2. \end{cases}$$

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

$$\delta \left(\frac{\gamma_1(v_1)}{(1-\delta)v_1} - \gamma_3 \right) + \delta F(v_1) \frac{v_1 \dot{\gamma}_1(v_1) - \gamma_1(v_1)}{(\delta-1)f(v_1)v_1^2} = 0 \quad (24)$$

$$\delta \left(\frac{\gamma_1(v_2)}{(1-\delta)v_2} - \gamma_3 \right) - \delta(1 - F(v_2)) \frac{v_2 \dot{\gamma}_1(v_2) - \gamma_1(v_2)}{(\delta-1)f(v_2)v_2^2} = -\delta\gamma_3. \quad (25)$$

We have four unknowns $v_1, v_2, \gamma_3, \gamma_1(1)$ and four equations, (23)-(25) and the promise-keeping constraint. Alternatively, for a fixed v_1 , (23)-(25) determine the three other unknowns $v_2, \gamma_3, \gamma_1(1)$. 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 $v \in (v_1, 1]$. The constraint $x_1(1) \leq 1$ does not bind. This implies that $\gamma_1(1) = 0$. When $v > v_1$, the state x_2 is pinned down by (22). From the condition that $\gamma_1(1) = 0$, we have that $W'(x_2(1)) = 1 - c$. Given strict concavity of W and $W'(0) = 1 - c$, we have $x_2(1) = 0$. The constraint $x_2(1) \geq 0$ binds, so (25) is replaced with

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

which is always satisfied given that $\gamma_1(1) \leq 0$. From (24), 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 $x_1(1) \leq 1$ does not bind. This requires that

$$-\frac{\delta}{1-\delta} \int_{v_1}^1 \frac{\dot{x}_2}{s} ds \leq 1. \quad (26)$$

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 $W'(U)$.

Proposition 2 *For power distribution $F(v) = v^a$ with $a \geq 1$, there exists $U^* \in (0, \bar{v})$ such that*

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

Proof. 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 (24), 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 (26) indeed increases as v_1 decreases.

We continue with case one. From (24) and (25), 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 (23) 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 the (23) is satisfied. ■

4.2 Correlated Types

Given the complexity of the problem, we see little hope for analytic results in this direction. 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 how about 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 get for any such pair, but the set of incentive-feasible payoffs has full dimension: the highest utility that he can get obtains when one of his incentive constraint binds, but there are two possibilities here, according to the incentive constraint. We obtain two hypersurfaces that do not seem to admit closed-form solutions. And the analysis of the i.i.d. case suggests that the optimal policy might well 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 a way.

5 Renegotiation-Proofness

The optimal policy, as described in Sections 2 and 3, is clearly not renegotiation-proof: after a history of reports such that the promised utility would be zero, both agent and principal would be better off by reneging and starting afresh.

There are many definitions of renegotiation-proofness with an infinite-horizon.¹² Here, we adopt the notion that offers the most favorable conditions to the agent, allowing him to renegotiate unilaterally to any contract that was offered in the past.¹³ That is, we require that the contract offers at least as much expected utility to the agent as it has ever done so far, and consider the contract that is best for the principal among those that satisfy this constraint. For simplicity, we consider the baseline model of Section 2, with two i.i.d. values.

Intuitively, because the utility of the agent can only increase (weakly), the principal must

¹²See Farrell and Maskin (1989).

¹³This is stronger than strong renegotiation-proofness in the sense of Farrell and Maskin (1989). Instead of requiring the promised utility to be a non-decreasing process, this would only require that the process be bounded below, with a strictly positive lower threshold that would act as a "reflecting barrier." Obviously, the dynamics would be somewhat different, but absorption would in that case as well necessarily occur at $\bar{v}$, unless utility remains at 0 forever.

compromise on the spread in continuation payoffs. Because continuation utilities are already larger than he would like to, the promise after a high report leads to an unchanged utility, while a low report leads to the minimum utility that makes the low type indifferent between both reports, given the probability with which a high report leads to the supply of the good. Together with this incentive constraint, promise-keeping then determines this utility. Unless it cannot be avoided because of promise-keeping, the good does not get supplied when a low report is sent.

Because of the nested structure of this policy, one can solve explicitly for the candidate value function, namely,

$$W(U) = \left(1 - \frac{qc}{\bar{v} - (1-q)l} + \left(\frac{qc}{\bar{v} - (1-q)l} - \frac{c}{l}\right) \left(\frac{(1-q)(\delta\bar{v} + (1-\delta)l)}{(1-\delta q)\bar{v}}\right)^{n+1}\right) U + \frac{\bar{v} - l}{l} c \left(\frac{\delta(1-q)}{1-\delta q}\right)^{n+1},$$

for $U \in [v/\beta^{n+1}, v/\beta^n]$, $\beta := 1 + \frac{1-\delta}{\delta} \frac{l}{\bar{v}}$, $n \in \mathbf{N}$, and the principle of optimality can be used to verify that the strategy is optimal, given U , on the domain over which W is decreasing (details omitted). Indeed, it is readily verified that the function W is single-peaked, with the maximum being achieved at some U^* . There is another candidate for the optimum, namely, setting U identically to 0. This cannot be optimal if $\bar{v} - c > 0$, as it would be better to always supply the good, but it cannot be ruled out in general. Hence, it holds that

Lemma 14 *The optimal policy involves*

$$U^0 = \begin{cases} 0 & \text{if } W(U^*) < 0, \\ U^* & \text{otherwise.} \end{cases}$$

If $U^0 > 0$, then given any $U \geq U^0$, it holds that (i) $U_h = U$, (ii) either $p_l = 0$ or $U_l = \bar{v}$, and (iii) ICL binds. Alongside promise-keeping, this uniquely determines p_h and U_l .

See Figure 7. As is clear, even the most extreme version of renegotiation-proofness does not alter the structure of the optimal contract significantly. To be sure, renegotiation-proofness puts a restriction on the continuation utility –and hence eliminates one of the possible long-run outcomes–, but this constraint tilts the structure of the contract to the minimum extent that is consistent with this additional constraint.

6 Concluding Comments

Here we discuss a few obvious extensions.

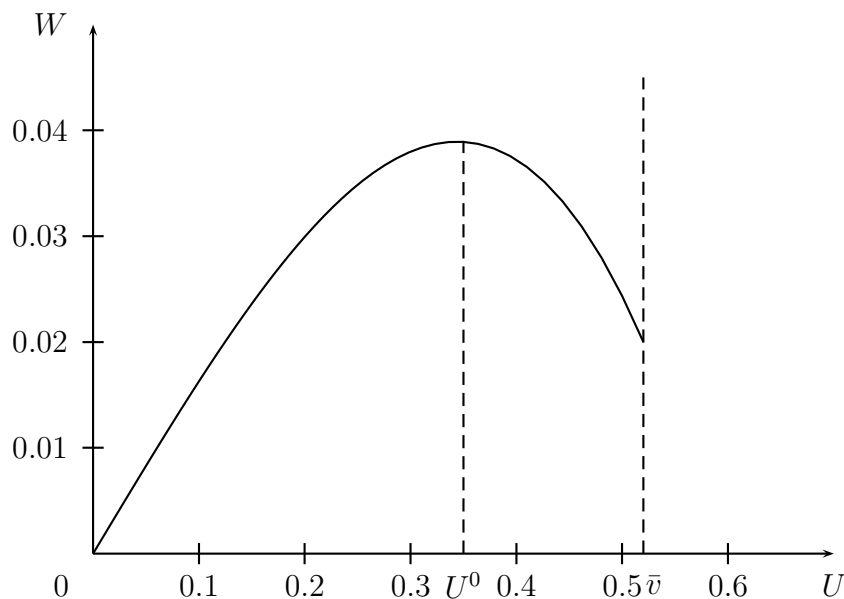


Figure 7: Renegotiation-proof payoff function W , as a function of U . Parameters are $(\delta, l, h, q, c) = (.95, .40, .60, .60, .50)$.

Public signals: While assuming no evidence whatsoever allows to clarify how the principal can take advantage of the repetition of the allocation decision to mitigate the inefficiency that goes along with private information, there are many applications for which some statistical evidence is available. This public signal depends on the current type, but also possibly on the action chosen by the principal. For instance, if we interpret the decision as filling a position (as in the labor example), we might get feedback on the quality of the applicant only if he is hired. If instead providing the good consists insuring the agent against a risk whose cost might be either high or low, it is when the principal fails to do so that he might find out that the agent's claim was warranted.

Incomplete Information regarding the process: So far, we have assumed that the agent's type is drawn from a distribution that is common knowledge. This 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 start, the revelation principle applies, and it is a matter of revisiting the analysis from Section 2, but with an incentive compatibility constraint at time 0.

Or the agent might not have any such information either initially, but be able to learn

from successive arrivals what the underlying distribution is. This is the more challenging case in which the agent himself is learning about q (or more generally, the transition matrix) as time passes by. In that case, the agent's belief might be private (in case 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 (about the agent's belief), a report by the agent of this belief, denoted by ν , his report on his current type (h or l) into a decision to allocate the good or not, and the promised continuation utility.

References

- [1] Abdulkadiroğlu, A., and K. Bagwell (2012). "The Optimal Chips Mechanism in a Model of Favors," mimeo, Duke University.
- [2] Abdulkadiroğlu, A., and S. Loertscher (2007). "Dynamic House Allocations," mimeo, 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] Battaglini, M. (2005). "Long-Term Contracting with Markovian Consumers," *American Economic Review*, **95**, 637–658.
- [5] Battaglini, M. and R. Lamba (2014). "Optimal Dynamic Contracting: the First-Order Approach and Beyond," working paper, Princeton University.
- [6] Benveniste, L.M. and J.A. Scheinkman (1979). "On the Differentiability of the Value Function in Dynamic Models of Economics," *Econometrica*, **41**, 727–733.
- [7] 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.
- [8] Casella, A. (2005). "Storable Votes," *Games and Economic Behavior*, **51**, 391–419.
- [9] Cohn, Z. (2010). "A Note on Linked Bargaining," *Journal of Mathematical Economics*, **46**, 238–247.
- [10] Cole, H.L. and N. Kocherlakota (2001). "Efficient Allocations with Hidden Income and Hidden Storage," *Review of Economic Studies*, **68**, 523–542.

- [11] Derman, C. Lieberman, G.J. and S.M. Ross (1972). “A Sequential Stochastic Assignment Problem,” *Management Science*, **18**, 349–355.
- [12] Doepke, M. and R.M. Townsend (2006). “Dynamic Mechanism Design with Hidden Income and Hidden Actions,” *Journal of Economic Theory*, **126**, 235–285.
- [13] Fang, H. and P. Norman (2006). “To Bundle or Not To Bundle,” *RAND Journal of Economics*, **37**, 946–963.
- [14] Eilat, R. and A. Pauzner (2011). “Optimal Bilateral Trade of Multiple Objects,” *Games and Economic Behavior*, **71**, 503–512.
- [15] Farrell, J., and E. Maskin (1989). “Renegotiation in Repeated Games,” *Games and Economic Behavior*, **1**, 327–360.
- [16] Fernandes, A. and C. Phelan (2000). “A Recursive Formulation for Repeated Agency with History Dependence,” *Journal of Economic Theory*, **91**, 223–247.
- [17] Frankel, A. (2011). “Contracting over Actions,” PhD Dissertation, Stanford University.
- [18] Gershkov, A. and B. Moldovanu (2010). “Efficient Sequential Assignment with Incomplete Information,” *Games and Economic Behavior*, **68**, 144–154.
- [19] Hauser, C. and H. Hopenhayn (2008), “Trading Favors: Optimal Exchange and Forgiveness,” mimeo, UCLA.
- [20] Hortala-Vallve, R. (2010). “Inefficiencies on Linking Decisions,” *Social Choice and Welfare*, **34**, 471–486.
- [21] Hylland, A., and R. Zeckhauser (1979). “The Efficient Allocation of Individuals to Positions,” *Journal of Political Economy*, **87**, 293–313.
- [22] Jackson, M.O. and H.F. Sonnenschein (2007). “Overcoming Incentive Constraints by Linking Decision,” *Econometrica*, **75**, 241–258.
- [23] Johnson, T. (2013). “Dynamic Mechanism Without Transfers,” mimeo, Notre Dame.
- [24] 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.

- [25] Krishna, 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.
- [26] Miralles, A. (2012). “Cardinal Bayesian Allocation Mechanisms without Transfers,” *Journal of Economic Theory*, **147**, 179–206.
- [27] Möbius, M. (2001). “Trading Favors,” mimeo.
- [28] Renault, J., E. Solan and N. Vieille (2013). “Dynamic Sender-Receiver Games,” *Journal of Economic Theory*, **148**, 502–534.
- [29] Sendelbach, S. (2012). “Alarm Fatigue,” *Nursing Clinics of North America*, **47**, 375–382.
- [30] Spear, S.E. and S. Srivastava (1987). “On Repeated Moral Hazard with Discounting,” *Review of Economic Studies*, **54**, 599–617.
- [31] 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.
- [32] Townsend, R.M. (1982). “Optimal Multiperiod Contracts and the Gain from Enduring Relationships under Private Information,” *Journal of Political Economy*, **90**, 1166–1186.
- [33] 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.
- [34] 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 2

Proof of Lemma 5. We start the proof with some notation and preliminary remarks. First, given any interval $I \subset [0, \bar{v}]$, we write $I_h := \left[\frac{a-(1-\delta)\bar{v}}{\delta}, \frac{b-(1-\delta)\bar{v}}{\delta} \right] \cap [0, \bar{v}]$ and $I_l := \left[\frac{a-(1-\delta)\underline{U}}{\delta}, \frac{b-(1-\delta)\underline{U}}{\delta} \right] \cap [0, \bar{v}]$ 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

$$W(U) = (1 - \delta)(qh - c) + \delta qW\left(\frac{U - (1 - \delta)\bar{v}}{\delta}\right) + \delta(1 - q)W\left(\frac{U - (1 - \delta)\underline{U}}{\delta}\right), \quad (27)$$

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 $I \subset [\bar{U}, \bar{v}]$, identically, for $U \in I$,

$$W(U) = (1 - q)\left(U - c - (U - \bar{v})\frac{c}{l}\right) + (1 - \delta)q(\bar{v} - c) + \delta qW\left(\frac{U - (1 - \delta)\bar{v}}{\delta}\right), \quad (28)$$

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 $[\bar{U}, \bar{v}]$, with the convention that its slope at $u_l = \bar{v}$ 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 $\bar{v}$). We make the following observations.

1. For any $I \subseteq (\underline{U}, \bar{v})$ (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 $\bar{v}$). 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 $W(\bar{v}) = \bar{v} - c$.
2. It follows that there is no interval $I \subseteq [\underline{U}, \bar{v}]$ 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 $I \subseteq [\underline{U}, \bar{v}]$ on which W has slope $1 - c/l$.
3. It immediately follows from 2 that $W < \bar{W}$ on $(\underline{U}, \bar{v})$: if there is a $U \in (\underline{U}, \bar{v})$ 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 $[U, \bar{v}]$, both being impossible.

4. Next, suppose that there exists an interval $I \subset [\underline{U}, \bar{v})$ 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 $[\underline{U}, \bar{v}]$ 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 $I' \subset [\underline{U}, \bar{v})$ 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 $J \subset [\underline{U}, \bar{v})$ 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 $[\underline{U}, \bar{v})$ 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 $(\underline{U}, \bar{v})$ 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 $\epsilon > 0$, 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. ■

B Missing Proof For Section 3

Proof of Lemma 8. Let W denote the set $\text{co}\{\bar{u}^\nu, \underline{u}^\nu : \nu \geq 0\}$. The point $\bar{u}^0$ is supported by $(p_h, p_l) = (1, 1), U(h) = U(l) = (\bar{v}_h, \bar{v}_l)$. 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

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

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

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

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 $\bar{w}^0 = (0, \bar{v}_l)$ with $\lim_{\nu \rightarrow \infty} \bar{w}^\nu = 0$. Similarly, $\underline{w}^\nu$ starts at $\underline{w}^0 = (\bar{v}_h, 0)$ and $\lim_{\nu \rightarrow \infty} \underline{w}^\nu = \bar{v}$. 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 (3)–(6), $p_h, p_l \in [0, 1]$, 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:

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

It is easy to verify that

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

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

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

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

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

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

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

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:

$$W(\nu) = \text{co}(\{(0, 0), (\bar{v}_h, \bar{v}_l)\} \cup \overline{W}(\nu) \cup \underline{W}(\nu)).$$

Since $W(0)$ equals $[0, \bar{v}_h] \times [0, \bar{v}_l]$, 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 9. 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 $A \subset [0, \bar{v}_h] \times [0, \bar{v}_l]$, let

$$\mathcal{B}_{ij}(A) := \{(U_h, U_l) \in [0, \bar{v}_h] \times [0, \bar{v}_l] : U(h) \in A, U(l) \in A \text{ solving (3)–(6) for } (p_h, p_l) = (i, j)\},$$

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 $U_l \in (\delta \rho_l \bar{v}_h + \delta(1 - \rho_l) \bar{v}_l, \bar{v}_l]$ 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 ICH 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 ICL. In addition, we must argue that $U(h) \in \bar{V}$. We prove by construction. Pick any v^ν

such that $\nu \geq 1$. Once we fix a $p_h \in [0, 1]$, PKH 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 PKH and ICL, we want to show that there exists $U(h) \in \bar{V}$ such that both PKH and ICL are satisfied. It is easy to verify that the intersection of PKH and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ is below the intersection of the binding ICL and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$. Therefore, the intersection of PKH and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ satisfies both PKH and ICL. In addition, the constructed PKH goes through the boundary point v^ν , so the intersection of PKH and $U_l(h) = \frac{\delta\rho_l}{1-\delta(1-\rho_l)}U_h(h)$ is inside $\bar{V}$.

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

The proof that $V_0 \subset \bar{V}$ follows the same lines as determining the boundaries of V in the proof of Lemma 8: 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 (3)–(6), $p_h, p_l \in [0, 1]$, and $U(h), U(l) \in \hat{W}^{\nu-1}$. The set $\mathcal{B}(\hat{W}^{\nu-1})$ is given by

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

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

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

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

Next, we argue that this achieves the first-best payoff. First, note that $\bar{V} \subseteq V \cap \{U : U_l \leq v_l^*\}$. 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 first-best cannot be achieved. It follows from the definition of $\bar{V}$ as the largest fixed point of $\mathcal{B}_0$ that starting from any utility vector $U \in V \setminus \bar{V}$, $U \neq \bar{v}$, there is a positive probability that the unit is given (after some history that has positive probability) to the low type. This implies that first-best cannot be achieved in case $U \leq v^*$. For $U \geq v^*$, first-best 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 Lemma 10. 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

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

It is easily verified that the function W given in Proposition 1 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,

$$k(\tau)(w''(\tau) + rw'(\tau)) = r\bar{v}w(\tau) + e^{r\tau}w_0(\tau)w'(\tau), \quad (29)$$

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

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

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 $r\bar{v}w(\tau) + e^{r\tau}w_0(\tau)w'(\tau) < 0$, and so $w''(\tau) + rw'(\tau) < 0$ (see (29)). 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

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

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 1. Computing $r\bar{v}w(\tau) + e^{r\tau}w_0(\tau)w'(\tau)$ over the two intervals $(\hat{\tau}, \tau_0)$ and (τ_0, ∞) yields on both intervals, after simplification,

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

[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 first-best 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 (29), 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 (29) 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 11. The proof is divided into two steps. First we show that the difference in payoffs between $W(\tau)$ and the first-best 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 first-best payoff maximizing pair of utilities converges to 0 at a rate linear in r . Given that the first-best payoff is piecewise affine in utilities, the result follows from the triangle inequality.

1. We first note that the first-best 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 1 and

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

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 $[0, \bar{v}_h] \times [v_l^*, \bar{v}_l]$). 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 $le^z > \bar{v}$ (that is, such that $g(z/r) > 0$ and hence $z \geq r\hat{\tau}$ for small enough r). Algebra gives

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

and similarly

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

as well as

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

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 $r\tau_0 = \ln(\bar{v}/((1 - q)l))$, 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 1 that

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

[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 first-best 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

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

It is immediate to check that

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

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.

■