

Contracting over persistent information

WEI ZHAO

School of Economics, Renmin University of China

CLAUDIO MEZZETTI

School of Economics, The University of Queensland

LUDOVIC RENOU

School of Economics and Finance, Queen Mary University of London, CEPR, and Department of Economics, University of Adelaide

TRISTAN TOMALA

HEC Paris and GREGHEC-CNRS

We consider a dynamic principal-agent problem, where the *sole* instrument the principal has to incentivize the agent is the disclosure of information. The principal aims at maximizing the (discounted) number of times the agent chooses the principal's preferred action. We show that there exists an optimal policy, where the principal recommends its most preferred action and discloses information as a reward in the next period, until either this action becomes statically optimal for the agent or the agent *perfectly* learns the state.

KEYWORDS. Dynamic, contract, information, revelation, disclosure, sender, receiver, persuasion.

JEL CLASSIFICATION. C73, D82.

1. INTRODUCTION

We consider a dynamic “principal-agent” model, where the sole instrument the principal has is information.¹ Principal and agent are engaged in a long-term relationship. The principal aims at inducing the agent to choose an action—the principal's most preferred action—as often as possible, and can only do so by disclosing information about

Wei Zhao: wei.zhao@outlook.fr

Claudio Mezzetti: c.mezzetti@uq.edu.au

Ludovic Renou: lrenou.econ@gmail.com

Tristan Tomala: tomala@hec.fr

Wei Zhao gratefully acknowledges the support of the HEC Foundation. Claudio Mezzetti thankfully acknowledges financial support from the Australian Research Council Discovery grant DP190102904. Ludovic Renou gratefully acknowledges the support of the Agence Nationale pour la Recherche under grant ANR CIGNE (ANR-15-CE38-0007-01) and through the ORA Project “Ambiguity in Dynamic Environments” (ANR-18-ORAR-0005). Tristan Tomala gratefully acknowledges the support of the HEC Foundation and ANR/Investissements d'Avenir under grant ANR-11-IDEX-0003/Labex Ecodec/ANR-11-LABX-0047.

¹That is, the principal cannot make transfers, terminate the relationship, choose allocations, or constrain the agent's choices.

an unknown state. To give examples, the principal is: (i) an external consultant with a clear agenda about what a company (the agent) should do, (ii) a department in a corporation aiming to maintain a central role while advising the CEO, (iii) a technology leading, multinational firm in a joint venture with a local firm in a less developed country; (iv) a lobbyist attempting to influence a politician.

We assume that the principal commits to a disclosure policy, which we refer to as the offer of a “contract.” The dynamic contracting problem we study is, therefore, a *dynamic persuasion problem*.

The standard approach in the study of dynamic contracting models (e.g., [Spear and Srivastava \(1987\)](#)) is to use the agent’s continuation value, or promised utility, as a state variable. The principal’s Bellman equation is then the fixed point of an operator, which satisfies a promise-keeping constraint in addition to incentive constraints. However, in dynamic persuasion models, there are additional complications.

First, since the belief of the agent changes over time due to information disclosure, we must treat it as an additional state variable. This increases the dimensionality of the principal’s problem. Second, any information disclosure policy, to which the principal commits, generates a *martingale* of beliefs. We must therefore impose the constraint that the belief process is a martingale. To the best of our knowledge, we are the first to be able to provide a complete characterization of an optimal contract by solving for the fixed point of a Bellman equation with two state variables tracking the evolution of the agent’s beliefs and of his promised utility.

We now illustrate the general properties of our optimal policy. First, the principal uses information disclosure as a “carrot” to motivate the agent to take the principal’s most preferred action until either the agent perfectly learns the state, or choosing the principal’s most preferred action becomes statically optimal. Moreover, if the agent learns the state, he learns it in finite time. After the agent has learned the state, he will take his optimal action in that state. Alternatively, as long as the agent keeps getting pieces of information from the principal (and thus, has not learned the state yet), he will take the principal’s preferred action. By trickling down bits of information, the principal is able to induce the agent to delay moving away from his favorite course of action. In some instances, the principal will promise eventual full disclosure of the state with probability one. In other instances, the principal will be able to stir the agent’s beliefs so that, with positive probability, the agent will take the principal’s favorite action forever. We provide a characterization of when this occurs.

Define the agent’s opportunity cost at a state as the difference between the agent’s stage payoff at his optimal action and the stage payoff when taking the principal’s preferred action. Generically, the agent’s opportunity cost, relative to the principal’s benefit from his preferred action, is different in different states, and our optimal policy exploits these differences. The second property of our optimal policy is that, along the paths at which the agent plays the principal’s most preferred action, his belief about the likelihood of the “high opportunity cost” state is decreasing. Intuitively, the optimal contract

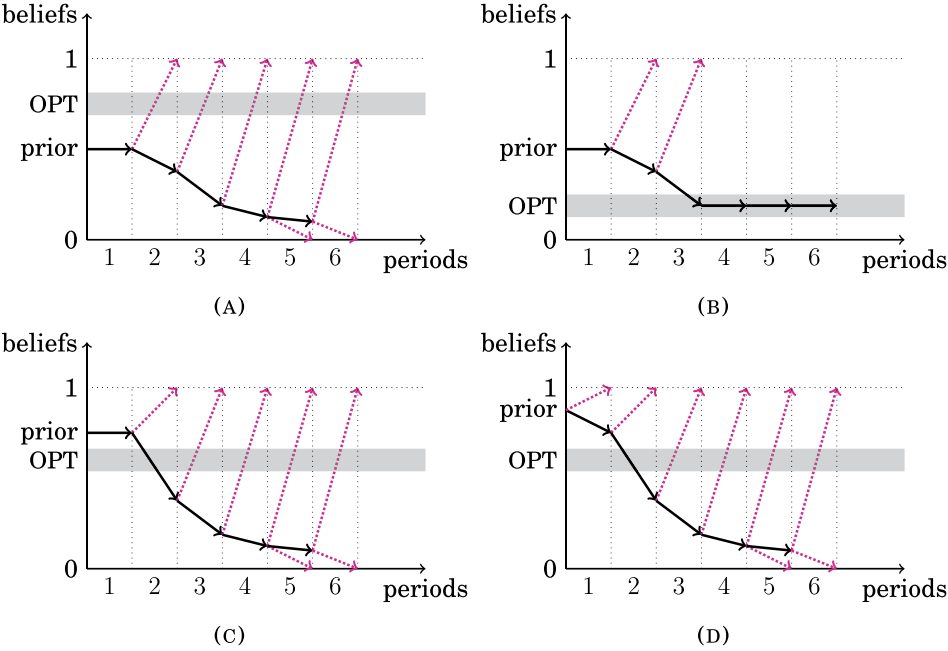


FIGURE 1. Evolution of actions and beliefs over time.

exploits the asymmetry in opportunity costs and lowers the agent’s expected opportunity cost—hence making it easier to incentivize the agent—by biasing information disclosure in the direction of informing him when the opportunity cost is high.²

Figure 1 plots four representative evolutions of the agent’s belief about the high opportunity cost state. In each panel, the grey region “OPT” indicates the region at which choosing the principal’s most preferred action is optimal for the agent. An arrow pointing from one belief to another indicates how the agent revises his belief within the period following a signal’s realization. Multiple arrows originating from the same point thus represent the information disclosed by the policy. Within a period, the agent takes a decision after having revised his beliefs. Arrows have different colors/patterns. At all beliefs at the end of continuous black arrows, the agent chooses the principal’s most preferred action. At all beliefs at the end of dotted magenta arrows, he chooses what is best given his current belief.

Third, in panels (A), (B), and (C), the policy does not disclose information to the agent at the first period. Starting from the second period, the policy discloses just enough information to compensate the agent for the opportunity cost of choosing the principal’s preferred action; no rent is left to the agent. However, as panel (D) illustrates, in some cases the policy discloses information in the first period, which may leave a strictly positive rent to the agent. For instance, it does so if the promise of full information disclosure at the next period would not incentivize the agent to choose the principal’s preferred

²To be precise, under our policy, upon receiving the signal “the opportunity cost is high,” the agent learns that this is indeed true. However, the signal is not sent with probability one. This corresponds to the (magenta/dotted) arrows pointing at 1 in Figure 1.

action. Disclosing information at the first period may also be necessary to reduce the agent's expected opportunity cost of following the principal's recommendation.

Finally, with the exception of panel (B), the policy does not induce the agent to believe that playing the principal's most preferred action is optimal. This is markedly different from what we would expect from the static analysis of [Kamenica and Gentzkow \(2011\)](#). Intuitively, the "static" persuasion policy is suboptimal because it does not extract all the information surplus it creates. Even in panel (B), the beliefs do not jump immediately to the "OPT" region. In fact, the belief process may approach the "OPT" region only asymptotically.

These properties highlight that in our dynamic environment, information is used as a compensation tool for creating intertemporal incentives, more than as a persuasion tool to affect the agent's myopic incentives.

Related literature The paper is part of the literature on Bayesian persuasion, pioneered by [Kamenica and Gentzkow \(2011\)](#), and recently surveyed by [Kamenica \(2019\)](#). The three most closely related papers are [Ball \(2023\)](#), [Ely and Szydlowski \(2020\)](#), and [Orlov, Skrzypacz, and Zryumov \(2020\)](#). In common with our paper, these papers study the optimal disclosure of information in dynamic games and show how the disclosure of information can be used as an incentive tool. The observation that information can be used to incentivize agents is not new and dates back to the literature on repeated games with incomplete information, for example, [Aumann, Maschler, and Stearns \(1995\)](#). See [Garicano and Rayo \(2017\)](#) and [Fudenberg and Rayo \(2019\)](#) for some more recent papers exploring the role of information provision as an incentive tool.

The classes of dynamic games studied differ considerably from one paper to another, and this makes comparisons difficult. In [Ely and Szydlowski \(2020\)](#), the agent has to repeatedly decide whether to continue working on a project or to quit (i.e., unlike our paper, there are only two actions); quitting ends the game. The principal aims at maximizing the number of periods the agent works on the project and can only do so by disclosing information about its complexity, modeled as the number of periods required to complete the project. Thus, their dynamic game is a quitting game, while ours is a repeated game. When the project is either easy or difficult (i.e., when there are two states), the optimal disclosure policy initially persuades the agent that the task is easy, so that he starts working. (Naturally, if the agent is sufficiently convinced that the project is easy, there is no need to persuade him initially.) If the project is in fact difficult, the policy then discloses it at a later date, when completing the project is now within reach. A main difference with our optimal disclosure policy is that information comes in lumps in [Ely and Szydlowski \(2020\)](#), that is, information is disclosed only at the initial period and at a later period, while information is repeatedly disclosed in our model.³ Another main difference is as follows. In [Ely and Szydlowski](#), only when the promise of full information disclosure at a later date is not enough to incentivize the agent to start working

³When there are more than two states, the optimal policy discloses information more frequently in [Ely and Szydlowski \(2020\)](#). The frequency of disclosure is thus a consequence of the dimensionality of the state space in their model, while it is not so in our model.

does the principal persuade the agent initially. This is not so with our policy: the principal persuades the agent in a larger set of circumstances. This initial persuasion reduces the cost of incentivizing the agent in future periods.

Orlov, Skrzypacz, and Zryumov (2020) also consider a quitting game, where the principal aims at delaying the quitting time as far as possible.⁴ The quitting time is when the agent decides to exercise an option, which has different values to the principal and the agent. The principal chooses a disclosure policy informing the agent about the option's value. When the principal is able to commit to a long-run policy, it is optimal to fully reveal the state with some delay. This policy is not optimal in Ely and Szydlowski (2020), or in our paper. See Au (2015), Bizzotto, Rüdiger, and Vigier (2021), Che, Kim, and Mierendorff (2023), Henry and Ottaviani (2019), and Smolin (2021) for other papers on information disclosure in quitting games, where the agent either waits and obtains additional information, or takes an irreversible action and stops the game.

Ball (2023) studies a continuous time model of information provision, where the state changes over time and payoffs are the ones of the quadratic example of Crawford and Sobel (1982). Ball shows that the optimal disclosure policy requires the sender to disclose the current state at a later date, with the delay shrinking over time. The main difference between his work and ours is the persistence of the state (also, we consider two different classes of games). When the state is fully persistent, as in Ely and Szydlowski (2020) and our model, full information disclosure with delay is not optimal in general. (See the discussion of Example 1 in Section 3.)

Finally, there are a few papers on dynamic persuasion, where the agent takes an action repeatedly. However, either the agent is myopic, for example, Ely (2017) and Renault, Solan, and Vieille (2017), or the principal cannot commit, for example, Escude and Sinander (2023).

2. THE PROBLEM

2.1 *The model*

A principal and an agent interact over an infinite number of periods, indexed by $t \in \{1, 2, \dots\}$. At the first period, the principal learns a payoff-relevant state $\omega \in \Omega = \{\omega_0, \omega_1\}$, while the agent remains uninformed. The prior probability of ω is $p_0(\omega) > 0$. At each period t , the principal sends a signal $s \in S$ and, upon observing s , the agent takes decision $a \in A$. The sets A and S are finite. The cardinality of S is as large as necessary for the principal to be unconstrained in his information disclosure policy.⁵ Throughout, we interchangeably use the words “period” and “stage.”

We assume that there exists $a^* \in A$ such that the principal's stage payoff is strictly positive whenever a^* is chosen, and zero otherwise. The principal's stage payoff function is thus $v : A \times \Omega \rightarrow \mathbb{R}$, with $v(a^*, \omega_0) > 0$, $v(a^*, \omega_1) > 0$, and $v(a, \omega_0) = v(a, \omega_1) = 0$ for all $a \in A \setminus \{a^*\}$. The agent's stage payoff function is $u : A \times \Omega \rightarrow \mathbb{R}$. The (common) discount factor is $\delta \in (0, 1)$.

⁴We refer to what Orlov, Skrzypacz, and Zryumov (2020) call the agent as the principal, and vice versa.

⁵From Makris and Renou (2023), it is enough to have the cardinality of S as large as the cardinality of A .

We write A^{t-1} for $\underbrace{A \times \cdots \times A}_{t-1 \text{ times}}$ and S^{t-1} for $\underbrace{S \times \cdots \times S}_{t-1 \text{ times}}$, with generic elements a^t and s^t , respectively. A behavioral strategy for the agent is a collection of maps $\sigma = (\sigma_t)_{t=1}^\infty$ with $\sigma_t : A^{t-1} \times S^t \rightarrow \Delta(A)$.

Before learning the state, the principal commits to a strategy, or *contract*, specifying, as a function of the state, the information to be disclosed (i.e., the statistical experiment to be conducted) at each history of realized signals and actions. Formally, the principal commits to a collection of maps (a contract) $\tau = (\tau_t)_{t=1}^\infty$, with $\tau_t : A^{t-1} \times S^{t-1} \times \Omega \rightarrow \Delta(S)$. The contract enables the principal to use information disclosures to reward or punish the agent for choosing the “right” or the “wrong” action.

We denote by $\mathbf{V}(\tau, \sigma)$ and $\mathbf{U}(\tau, \sigma)$ the principal’s and the agent’s overall expected payoff under the profile (τ, σ) . Let $\mathbb{P}_{\tau, \sigma}(\cdot | \omega)$ be the distribution over sequences of signals and actions induced by (τ, σ) conditional on ω . The principal’s expected payoff $\mathbf{V}(\tau, \sigma)$ is

$$\sum_{\omega} p_0(\omega) \left(\sum_t \sum_{s^t, a^{t-1}} (1 - \delta) \delta^{t-1} \mathbb{P}_{\sigma, \tau}(s^{t-1}, a^{t-1} | \omega) \tau_t(s_t | s^{t-1}, a^{t-1}, \omega) \sigma_t(a^* | s^t, a^{t-1}) \right) \times v(a^*, \omega). \quad (1)$$

The agent’s expected payoff is defined similarly. The objective is to characterize the maximal expected payoff $V^{\max}$ the principal can achieve by committing to a contract τ before learning the state, that is,

$$V^{\max} = \begin{cases} \sup_{(\tau, \sigma)} \mathbf{V}(\tau, \sigma) \\ \text{subject to } \mathbf{U}(\tau, \sigma) \geq \mathbf{U}(\tau, \sigma') \text{ for all } \sigma'. \end{cases}$$

Several comments are worth making. First, an alternative interpretation of our model is that neither the principal nor the agent know the state, but the principal has the ability to conduct statistical experiments contingent on the state and past signals and actions. Second, the only additional information the agent obtains each period is the outcome of the statistical experiment. Third, the state is fully persistent and the principal perfectly monitors the action of the agent. Finally, the only instrument available to the principal is information. The principal can neither remunerate the agent nor terminate the relationship nor allocate different tasks to the agent. We purposefully make all these assumptions to address our main question of interest: What is the optimal way to incentivize the agent with information only?

2.2 An example

Throughout the paper, we illustrate our results with the help of the following example.

EXAMPLE 1. The agent has three possible actions a^0 , a^1 and a^* , with a^0 (resp., a^1) the agent’s optimal action when the state is ω_0 (resp., ω_1). The prior probability of ω_1 is

TABLE 1. Payoff table.

	a^0	a^1	a^*
ω_0	0, 1	0, 0	1, 1/2
ω_1	0, 0	0, 2	1, 1/2

1/3 and the discount factor is 1/2. The payoffs are in Table 1, with the first coordinate corresponding to the principal payoff.

We start with few preliminary observations. First, regardless of the agent’s belief, action a^* is never optimal. Second, the opportunity cost of playing a^* is higher when the state is ω_1 than ω_0 , that is, $u(a^1, \omega_1) - u(a^*, \omega_1) > u(a^0, \omega_0) - u(a^*, \omega_0)$. It is, therefore, harder to incentivize the agent to play a^* when he is more confident that the state is ω_1 . As we shall see, the optimal policy exploits this asymmetry.

We now consider some simple strategies the principal may commit to. To start with, assume that the principal commits to disclose information at the initial stage only. We call it the KG policy, in reference to [Kamenica and Gentzkow \(2011\)](#). Clearly, since a^* is never optimal, the principal’s payoff is 0. To obtain a positive payoff, the principal must condition his information disclosure on the agent’s actions.

The simplest such policy is to “reward” the agent with full disclosure of the state for playing a^* at the beginning of the relationship, say up to period T^* . If the agent deviates, the harshest punishment the principal can impose is to reveal no information in subsequent periods, inducing a normalized expected payoff of 2/3. We are thus looking for the largest T^* such that

$$(1 - \delta) \left(\frac{1}{2} (\delta^0 + \delta^1 + \dots + \delta^{T^*-1}) + \left(\frac{1}{3} \cdot 2 + \frac{2}{3} \cdot 1 \right) (\delta^{T^*} + \dots) \right) \geq \frac{2}{3},$$

which is $T^* = \lfloor \ln(5)/\ln(2) \rfloor = 2$, yielding the principal a payoff of $(1 - \frac{1}{2}) \cdot (1 + \frac{1}{2}) = \frac{3}{4}$.

Another simple strategy the principal can commit to is a “random full-disclosure policy,” where he fully discloses the state with probability α at period t (and withholds all information with the complementary probability) if the agent plays a^* at period $t - 1$.⁶ (Again if the agent deviates, the harshest punishment is to withhold all information in all subsequent periods.) Thus, if we write V (resp., U) for the principal (resp., agent) payoff, the best recursive policy is to choose α so as to maximize

$$\begin{aligned} V &= \frac{1}{2} 1 + \frac{1}{2} (1 - \alpha) V, \quad \text{subject to:} \\ U &= \frac{1}{2} \left(\frac{1}{2} \right) + \frac{1}{2} \left[(1 - \alpha) U + \alpha \frac{4}{3} \right] \geq \frac{2}{3}. \end{aligned}$$

The principal’s best payoff is $V = 4/5$ with $\alpha = 1/4$. The random full-disclosure policy does better than the policy of fully disclosing the state with delay since it circumvents

⁶Full information with delay plays an important role in the work of [Ball \(2023\)](#) and [Orlov, Skrzypacz, and Zryumov \(2020\)](#).

the integer constraint on T . Intuitively, it makes it possible to incentivize the agent to play a^* a discounted number of periods slightly larger than 2, namely $\ln(5)/\ln(2)$.

As we will see in Section 3.5, the random full-disclosure policy is still suboptimal since it does not exploit the asymmetry in the agent's opportunity cost of choosing a^* in the two states. The optimal policy exploits such asymmetry by disclosing no information in the first period and then either revealing that the state is ω_1 , the high opportunity cost state, or lowering the agent's belief that the state is ω_1 . By doing so, the policy incentivizes the agent to take action a^* for a longer expected time. $\diamond$

3. OPTIMAL CONTRACTS

This section characterizes optimal contracts and discusses their most salient properties.

3.1 A recursive formulation

The first step toward characterizing optimal contracts is to reformulate the principal's problem as a recursive problem. To do so, we introduce two state variables. The first state variable is promised payoff. It is well known that classical dynamic contracting problems admit recursive formulations if we introduce promised payoff as a state variable and impose promise-keeping constraints, for example, [Spear and Srivastava \(1987\)](#). The second state variable we introduce is beliefs. We now turn to the formal reformulation of the problem.

We first need some additional notation. We denote by $p \in [0, 1]$ a generic belief, with p the probability of ω_1 . We let $u(a, p) := p[u(a, \omega_1) - u(a, \omega_0)] + u(a, \omega_0)$ be the agent's expected stage payoff of choosing a when his belief is p . We define $m(p) := \max_{a \in A} u(a, p)$ as the agent's optimal stage payoff when his belief is p , and $M(p) := p[m(1) - m(0)] + m(0)$ as the agent's expected stage payoff if he learns the state prior to choosing an action. Note that m is a piecewise linear convex function that M is linear and that $m(p) \leq M(p)$ for all p . Similarly, we let $v(a, p)$ be the principal's expected stage payoff when the agent chooses a and the principal's belief is p . Finally, let $P := \{p \in [0, 1] : m(p) = u(a^*, p)\}$ be the set of beliefs at which a^* is optimal. If nonempty, the set P is the closed interval $[\underline{p}, \bar{p}]$.

Let $\mathcal{W} \subseteq [0, 1] \times \mathbb{R}$ be such that $(p, w) \in \mathcal{W}$ if and only if $w \in [m(p), M(p)]$. Throughout, we consider the complete metric space of bounded, continuous functions $V : \mathcal{W} \rightarrow \mathbb{R}$, with the interpretation that $V(p, w)$ is the principal's payoff if he promises a payoff of w to the agent when the agent's current belief is p . Consider the following maximization program:

$$T(V)(p, w) := \begin{cases} \max_{((\lambda_s, (p_s, w_s), a_s) \in [0, 1] \times \mathcal{W} \times A)_{s \in S}} \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta V(p_s, w_s)], \\ \text{subject to:} \\ (1 - \delta)u(a_s, p_s) + \delta w_s \geq m(p_s) \quad \text{for all } s \text{ such that } \lambda_s > 0, \\ \sum_{s \in S} \lambda_s [(1 - \delta)u(a_s, p_s) + \delta w_s] \geq w, \\ \sum_{s \in S} \lambda_s p_s = p, \sum_{s \in S} \lambda_s = 1. \end{cases}$$

The program maximizes the principal's expected payoff over *policies*, that is, maps from $\mathcal{W}$ to $([0, 1] \times \mathcal{W} \times \mathcal{A})^{|\mathcal{S}|}$. At each (p, w) , a policy prescribes the probability λ_s that the realized signal is s and conditional on s , the belief p_s , the promised continuation utility w_s , and the recommended action a_s . The first constraint is the incentive-compatibility condition that the agent prefers to obey the recommendation a_s , when w_s is the promised continuation payoff and p_s is the agent's belief. To understand the right-hand side, observe that the agent can always play a static best-reply to any belief, so that his expected payoff must be at least $m(p_s)$ when his current belief is p_s .⁷ Conversely, if the contract recommends action a_s and the agent does not obey, the contract can specify no further information revelation, in which case the agent's payoff is at most $m(p_s)$. Therefore, $m(p_s)$ is the agent's min-max payoff. The second constraint is the promise-keeping constraint: if the principal promises the payoff w at a period, the contract must honor that promise in subsequent periods. The third constraint states that the policy selects a splitting of p , that is, a distribution over posteriors with expectation p .

In most dynamic contracting papers, the promise-keeping constraint holds as an equality everywhere in the state space. Here, on the contrary, we will show that under the contract solving the recursive problem, there are two regions: A region where the promise-keeping constraint holds as an equality, and a region where it holds as a strict inequality. Importantly, we will also show (see Corollary 4) that under this contract the second region can be visited only in the very first period, and hence the contract is feasible.

Throughout, we slightly abuse notation and write τ for a policy (i.e., a map from $\mathcal{W}$ to $([0, 1] \times \mathcal{W} \times \mathcal{A})^{|\mathcal{S}|}$). A policy is *feasible* if it specifies a feasible tuple $(\lambda_s, (p_s, w_s), a_s)_{s \in \mathcal{S}}$ for each (p, w) , that is, a tuple satisfying the constraints of the maximization problem $T(V)(p, w)$.

Three important observations, implying Proposition 1, are worth making. First, for any function V , it is easy to show that $T(V)$ is a concave function of the pair (p, w) . This is because, if τ is feasible at state variables (p, w) and τ' is feasible at state variables (p', w') , then a policy which follows τ with probability α and follows τ' with probability $1 - \alpha$ is feasible at state variables $(\alpha p + (1 - \alpha)p', \alpha w + (1 - \alpha)w')$.⁸ Second, for any function V , the mapping $T(V)$ is weakly decreasing in w , since a policy that is feasible at state variables (p, w) is also feasible at state variables (p, w') for any $w' \leq w$. The more the principal promises to the agent, the harder it is to incentivize the agent to play a^* . Third, the operator T is a contraction. Indeed, T is monotone, that is, $T(V) \geq T(V')$ for all $V \geq V'$, and satisfies $T(V + c) \leq T(V) + \delta c$ for all positive constant $c \geq 0$, for all V . Hence, T is a contraction by Blackwell's theorem. Let V^* be its unique fixed point.

PROPOSITION 1. *The value function V^* is concave in both arguments and weakly decreasing in w .*

⁷More precisely, if the agent's belief at period t is p_t , he obtains the payoff $m(p_t)$ by playing a static best-reply. Since the function m is convex and beliefs follow a martingale, his expected payoff is therefore at least $(1 - \delta) \sum_{t' \geq t} \delta^{t'-t} \mathbb{E}[m(\mathbf{p}_{t'}) | \mathcal{F}_t] \geq m(p_t)$, where $\mathcal{F}_t$ is the agent's filtration at period t .

⁸Note that $(\alpha p + (1 - \alpha)p', \alpha w + (1 - \alpha)w') \in \mathcal{W}$ since $\alpha w + (1 - \alpha)w' \geq \alpha m(p) + (1 - \alpha)m(p') \geq m(\alpha p + (1 - \alpha)p')$, by the convexity of m .

With this recursive formulation, the principal's maximal payoff $V^{\max}$ is $V^*(p_0, m(p_0))$, and the solutions to the optimization problem $T(V^*)(p_0, m(p_0))$ define the optimal contracts. Characterizing the optimal solutions is the objective of the rest of the paper.

In a working paper, Ely (2015) discusses the extension of his model in Ely (2017) to the interaction between a long-run principal and a long-run agent and derives a recursive reformulation.⁹ The main difference with our formulation is that the promise-keeping constraint is an equality in Ely (2015). In Appendix C.1, we prove that $V^{\max} = \max_{w \in [m(p_0), M(p_0)]} \widehat{V}^*(p_0, w)$, where $\widehat{V}^*$ is the fixed point of the “Ely's operator.”

Proposition 1, together with the recursive formulation, has a number of implications, which are summarized in Proposition 2. First, if the principal induces the posterior p_s while recommending the action a_s and promising the continuation payoff w_s , he should not have an incentive to disclose more information in that period, that is, we cannot have $V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s) > (1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)$. In other words, the tuple $(1, p_s, w_s, a_s)$ must be optimal at state variables $(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)$. To see this, observe that $V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s) \geq (1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)$ for all (p_s, w_s, a_s) . Indeed, the tuple $(1, p_s, w_s, a_s)$ is feasible at state variables $(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)$ and gives a payoff of $(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)$ to the principal. Therefore, if the inequality were strict at some s , the principal would strictly benefit from releasing further information at p_s so as to get $V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)$ —he would do so by following the optimal policy at $(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)$.

Second, if the principal does not recommend a^* at a period, then he never recommends a^* at any subsequent periods, that is, the principal's continuation value is zero. In other words, as soon as an action other than a^* is played, the principal stops incentivizing the agent to play a^* . The intuition is simple. Suppose to the contrary that the principal were to recommend $a_s \neq a^*$ after the signal s at period t and a^* at the next period. Consider the policy change where the principal anticipates the disclosure of information: what incentivizes the agent to play a^* at period $t + 1$ is disclosed at period t . This policy change is feasible and increases the principal's payoff, a contradiction. This property justifies thinking of the principal's preferred action a^* as a status quo, which the principal tries to induce the agent to maintain as long as possible. Note that, unlike quitting games, the irreversibility is endogenous here—this explains why our solution differs from the ones found in previous works.

Third, under any optimal policy, if a^* is recommended at signal s , that is, $a_s = a^*$, then $V^*(p_s, w_s) > V^*(p_s, w'_s)$ for all $w'_s > w_s$. This implies that the agent's expected continuation payoff is the promised continuation utility w_s , that is, the promise-keeping constraint binds at state variables (p_s, w_s) . (If the constraint were not binding, there would exist some $w'_s > w_s$ such that $V^*(p_s, w'_s) = V^*(p_s, w_s)$), a contradiction.)

PROPOSITION 2. *For all (p, w) and all solutions $(\lambda_s, p_s, w_s, a_s)_{s \in S}$ to $T(V^*)(p, w)$, we have:*

⁹Ely (2017) analyzes the interaction between a long-run principal and a sequence of short-run agents (see also Renault, Solan, and Vieille (2017)).

(i) For all $s \in S$ such that $\lambda_s > 0$,

$$(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s) = V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s).$$

(ii) For all $s \in S$ such that $\lambda_s > 0$ and $a_s \neq a^*$, $V^*(p_s, w_s) = 0$.

(iii) If $a_s = a^*$, then $V^*(p_s, w_s) > V^*(p_s, w'_s)$ for all $w'_s \in (w_s, M(p_s)]$.

While the principal's value function is unique, there might be several, payoff-equivalent, optimal policies. One such optimal policy—the one we focus on—has the following two additional properties, summarized in Proposition 3. First, there is at most one signal s^* at which the principal recommends the agent to play a^* . Intuitively, if two signals recommended a^* , the principal would not lose from merging them into one. In addition, upon receiving s^* , the agent is made indifferent between obeying the recommendation and deviating to his outside option. Second, when the principal does not recommend a^* , the principal perfectly informs the agent of the payoff-relevant state. This follows from the principal's indifference over actions $a \neq a^*$. Since the principal's continuation value is zero when a^* is not recommended (Proposition 2(ii)), full information revelation does not hurt the principal while increasing the agent's payoff, which relaxes the promise-keeping constraint.

PROPOSITION 3. *For all (p, w) , there exists a solution $(\lambda_s, p_s, w_s, a_s)_{s \in S}$ to $T(V^*)(p, w)$ such that*

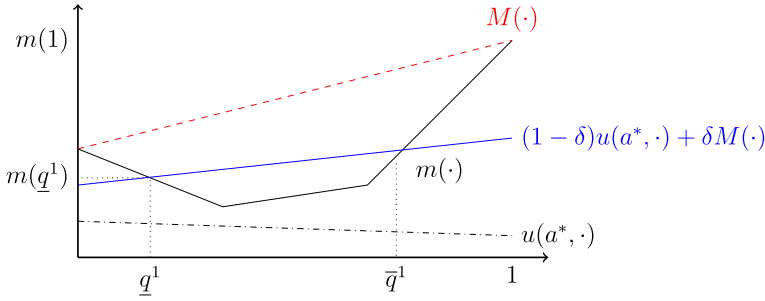
(i) *There exists at most one signal $s^* \in S$ such that $\lambda_{s^*} > 0$ and $a_{s^*} = a^*$. Moreover, $(1 - \delta)u(a_{s^*}, p_{s^*}) + \delta w_{s^*} = m(p_{s^*})$.*

(ii) *If $a_s \neq a^*$, then $p_s = 1$ or $p_s = 0$.*

The main implication of Proposition 3 is that we can restrict our attention to policies with at most three messages s^* , s_0 , and s_1 in its support. At s^* , the policy recommends a^* and the agent is made indifferent between obeying and disobeying (if $\lambda_{s^*} > 0$). At s_0 (resp., s_1), the agent knows that the state is ω_0 (resp., ω_1) and his payoff is $m(0)$ (resp., $m(1)$). It is also worth noting that, in Example 1, the policy of fully disclosing the state with delay satisfies neither property (iii) of Proposition 2 nor property (i) of Proposition 3. The “random full-disclosure” policy satisfies these properties, but its induced value function is not concave.¹⁰

Several important questions remain. What are the beliefs at which the agent plays a^* ? How does the principal compensate the agent for playing a^* ? Does the principal need to reveal information at the prior belief? Does the agent learn the state? If so, does he learn it in finite time? Before formally answering these questions, we build some further intuition on the optimal policies.

¹⁰We remark that Propositions 2 and 3 remain true with more than two states.

FIGURE 2. The set Q^1 .

3.2 Optimal policy: Building intuition

Let Q^1 be the set of beliefs at which the agent has an incentive to play a^* when promised full information disclosure at the next period. That is,

$$Q^1 := \{p \in [0, 1] : (1 - \delta)u(a^*, p) + \delta M(p) \geq m(p)\}.$$

If Q^1 is empty, then all policies are optimal, as the principal can never incentivize the agent to play a^* . If Q^1 is nonempty, then it is a closed interval $[\underline{q}^1, \bar{q}^1]$. Note that $\underline{q}^1 = 0$ if, and only if, a^* is optimal at $p = 0$. See Figure 2 for an illustration.

For all $p \in Q^1$, we write $\mathbf{w}(p) \in [m(p), M(p)]$ for the continuation payoff that makes the agent indifferent between playing action a^* and receiving the continuation payoff $\mathbf{w}(p)$ in the future, and playing a best-reply to the belief p forever. That is, $\mathbf{w}(p)$ solves

$$(1 - \delta)u(a^*, p) + \delta \mathbf{w}(p) = m(p).$$

An important feature of our model is that the agent's opportunity cost of choosing a^* rather than his best action, relative to the principal's benefit, differs across states. When the state is ω_0 (resp., ω_1) the opportunity cost relative to the benefit is $[m(0) - u(a^*, 0)]/v(a^*, 0)$ (resp., $[m(1) - u(a^*, 1)]/v(a^*, 1)$). Without loss of generality, we assume the following.

ASSUMPTION 1.

$$\frac{m(1) - u(a^*, 1)}{v(a^*, 1)} \geq \frac{m(0) - u(a^*, 0)}{v(a^*, 0)}.$$

As we shall see, our optimal policy heavily exploits this asymmetry. It also follows from Assumption 1 that if a^* is optimal for the agent at $p = 1$, that is, $m(1) = u(a^*, 1)$, then a^* is also optimal at $p = 0$. Consequently, a^* is optimal at all beliefs, that is, $P = [0, 1]$. (Recall that P is the set of beliefs at which a^* is optimal.) In what follows, we exclude this trivial case and assume that $1 \notin P$.

To strengthen our intuition, we briefly return to the original (nonrecursive) description of the problem. Let (τ, σ) be a profile of strategies. We can rewrite the principal's

expected payoff $V(\tau, \sigma)$ in equation (1) as $V(\tau, \sigma) = \lambda^* v^*(a^*, p^*)$, with

$$\lambda^* := (1 - \delta) \sum_{\omega} p_0(\omega) \left(\sum_t \sum_{s^t, a^{t-1}} \delta^{t-1} \mathbb{P}_{\sigma, \tau}(s^t, a^{t-1} | \omega) \sigma_t(a^* | s^t, a^{t-1}) \right)$$

the discounted probability of recommending action a^* , and

$$p^* := \frac{(1 - \delta) p_0(\omega_1) \left(\sum_t \sum_{s^t, a^{t-1}} \delta^{t-1} \mathbb{P}_{\sigma, \tau}(s^t, a^{t-1} | \omega_1) \sigma_t(a^* | s^t, a^{t-1}) \right)}{\lambda^*},$$

the average discounted probability of ω_1 when a^* is played.¹¹ As expected, the principal's payoff only depends on how often a^* is played, and the average belief at which it is played.

We now make two observations, which will enable us to rewrite the principal's expected payoff and get important insights on optimal policies. First, if we let $p^\dagger$ be the average discounted probability of ω_1 when a^* is *not* recommended, we have that $\lambda^* p^* + (1 - \lambda^*) p^\dagger = p_0$ since the belief process is a martingale. Second, recall from Proposition 3(ii) that the agent's belief is either 0 or 1, when he does not play a^* . Therefore, conditional on not playing a^* , his expected payoff is $M(p^\dagger)$, and his overall expected payoff is

$$\lambda^* u(a^*, p^*) + (1 - \lambda^*) M(p^\dagger) = \lambda^* [u(a^*, p^*) - M(p^*)] + M(p_0). \quad (2)$$

Since the agent's ex ante payoff must be at least $m(p_0)$, the agent's ex ante rent is

$$c := \lambda^* [u(a^*, p^*) - M(p^*)] + M(p_0) - m(p_0) \geq 0. \quad (3)$$

With the help of these two observations, we can rewrite the principal's expected payoff as

$$\frac{v(a^*, p^*)}{M(p^*) - u(a^*, p^*)} \times (M(p_0) - m(p_0) - c). \quad (4)$$

The first term captures the benefit of incentivizing the agent to play a^* relative to the cost. Since $\frac{v(a^*, 0)}{v(a^*, 1)} \geq \frac{m(0) - u(a^*, 0)}{m(1) - u(a^*, 1)}$, it is decreasing in p^* .¹² Ceteris paribus, the lower the average belief at which the agent plays a^* , the higher the principal's expected payoff.

The second term captures how the principal rewards the agent for playing a^* with his only instrument: information. The term $M(p_0) - m(p_0)$ is the maximal value of information the principal can create. Ceteris paribus, the principal's payoff is decreasing in c , that is, the best is to leave no rents to the agent and to create as much information as necessary to repay the agent.

The above discussion, along with our previous results, thus suggests some guiding principles in constructing an optimal policy. First, the policy must recommend a^* at the

¹¹ Note that p^* cannot be lower than $\underline{q}^1$ since the agent would never play a^* at beliefs lower than $\underline{q}^1$.

¹² This follows from the observation that $M(p^*) - u(a^*, p^*) = p^* [(m(1) - m(0)) - (u(a^*, 1) - u(a^*, 0))] + m(0) - u(a^*, 0)$, $v(a^*, p^*) = p^* (v(a^*, 1) - v(a^*, 0)) + v(a^*, 0)$, and simple algebra.

lowest possible beliefs p_s^* . Second, the policy should leave as little rent as possible to the agent. Naturally, it is not always possible to leave no rents. For example, when the prior belief $p_0 \notin Q^1$, the agent must be given some strictly positive rent if he is to ever play a^* . In the next subsection, we will construct an optimal policy with all these features.

3.3 Benchmark scenarios

Before defining our optimal policy, we discuss two benchmark scenarios. In the first benchmark, the agent only has two actions, $A = \{a^*, a^\dagger\}$. The constraint that the agent's ex ante expected payoff must be at least as high as his outside option is

$$\lambda^* u(a^*, p^*) + (1 - \lambda^*) u(a^\dagger, p^\dagger) \geq m(p_0).$$

Observe that if either $\lambda^*(u(a^*, p^*) - u(a^\dagger, p^\dagger)) < 0$ or $(1 - \lambda^*)(u(a^\dagger, p^\dagger) - u(a^*, p^*)) < 0$, then the constraint cannot be satisfied.¹³ In words, if an action is recommended with strictly positive probability, the agent must find that action optimal at the corresponding belief; the same is true in the static problem. Therefore, the KG policy is an optimal policy, when the agents has only two actions.

In the second benchmark, the agent's opportunity cost relative to the principal's benefit of inducing a^* is the same across states, that is, $\frac{m(1) - u(a^*, 1)}{v(a^*, 1)} = \frac{m(0) - u(a^*, 0)}{v(a^*, 0)}$. Consider the best random full-disclosure policy. It requires that the state be fully revealed with probability α at each period, where α is such that the agent's payoff equals his outside option payoff, that is,

$$(1 - \delta)u(a^*, p_0) + \delta[\alpha M(p_0) + (1 - \alpha)m(p_0)] = m(p_0).$$

It follows that

$$\alpha = \frac{1 - \delta}{\delta} \frac{m(p_0) - u(a^*, p_0)}{M(p_0) - m(p_0)}.$$

Since the principal's payoff satisfies $V = (1 - \delta)v(a^*, p_0) + \delta(1 - \alpha)V$, it follows that

$$V = \frac{M(p_0) - m(p_0)}{M(p_0) - u(a^*, p_0)} v(a^*, p_0).$$

Now, from the above *relaxed* version of the principal's maximization problem, where only the (ex ante) participation constraint needs to be satisfied, an upper bound on the principal's payoff is given by equation (4). Since $\frac{v(a^*, 0)}{v(a^*, 1)} = \frac{m(0) - u(a^*, 0)}{m(1) - u(a^*, 1)}$, the first term is constant in p^* (see footnote 12), and thus, equals to $\frac{v(a^*, p_0)}{M(p_0) - u(a^*, p_0)}$. Since $c \geq 0$, the second term is at most $M(p_0) - m(p_0)$. Therefore, $\frac{M(p_0) - m(p_0)}{M(p_0) - u(a^*, p_0)} v(a^*, p_0)$ is an upper bound of the relaxed problem. Since the random full-disclosure policy achieves this upper bound, it is an optimal policy.

COROLLARY 1. (i) *When $|A| = 2$, the KG policy is optimal.*

¹³In the former case, the left-hand side would be strictly less than $u(a^\dagger, p_0) \leq m(p_0)$, while it would be strictly less than $u(a^*, p_0) \leq m(p_0)$ in the latter case, a contradiction in both cases.

(ii) When $\frac{v(a^*, 0)}{v(a^*, 1)} = \frac{m(0) - u(a^*, 0)}{m(1) - u(a^*, 1)}$, the random full-disclosure policy is optimal.

We hasten to stress that both the KG and random full-disclosure policies are not optimal in general, as Example 1 demonstrates. See Section 3.5.

3.4 Optimal policy: A formal description

We define a family of policies $(\tau_q)_{q \in [\underline{q}^1, \bar{q}^1]}$ indexed by a belief q , and prove later the existence of $q^* \in [\underline{q}^1, \bar{q}^1]$ such that the policy τ_{q^*} is optimal. At each $(p, w) \in \mathcal{W}$, a policy prescribes a feasible tuple $(\lambda_s, (p_s, w_s), a_s)_{s \in S}$, that is, a splitting $(\lambda_s, p_s)_{s \in S}$, a profile of recommendations $(a_s)_{s \in S}$ and a profile of continuation payoffs $(w_s)_{s \in S}$. There are four different types of prescription, depending on which of four regions the state variables (p, w) belong to; the belief q parameterizes these regions. The four regions are

$$\begin{aligned} \mathcal{W}_q^1 &:= \left\{ (p, w) : p \in \left[0, \underline{q}^1 \right), w \leq \frac{\underline{q}^1 - p}{\underline{q}^1} m(0) + \frac{p}{\underline{q}^1} m(\underline{q}^1) \right\}, \\ \mathcal{W}_q^2 &:= \left\{ (p, w) : p \in (q, 1], \frac{1-p}{1-q} m(q) + \frac{p-q}{1-q} m(1) < w \leq \frac{1-p}{1-\underline{q}^1} m(\underline{q}^1) + \frac{p-\underline{q}^1}{1-\underline{q}^1} m(1) \right\} \\ &\quad \cup \left\{ (p, w) : p \in [\underline{q}^1, q], w \leq \frac{1-p}{1-\underline{q}^1} m(\underline{q}^1) + \frac{p-\underline{q}^1}{1-\underline{q}^1} m(1) \right\}, \\ \mathcal{W}_q^3 &:= \left\{ (p, w) : p \in (q, 1], w \leq \frac{1-p}{1-q} m(q) + \frac{p-q}{1-q} m(1) \right\}, \\ \mathcal{W}_q^4 &:= \mathcal{W} \setminus (\mathcal{W}_q^1 \cup \mathcal{W}_q^2 \cup \mathcal{W}_q^3). \end{aligned}$$

Figure 3 illustrates the four regions, with $\mathcal{W}_q^1$ the black region, $\mathcal{W}_q^2$ the region with vertical lines, $\mathcal{W}_q^3$ the gray region, and $\mathcal{W}_q^4$ the region with slanted lines. Observe that regions $\mathcal{W}_q^1$ and $\mathcal{W}_q^4$ do not depend on the parameter q , while the other two do.

We begin with an informal overview of our optimal policy. In region $\mathcal{W}_q^2$, the principal recommends a^* and discloses information in the next period as a reward. The belief p decreases over time, until either it reaches a point at which the agent will choose a^* forever, or it enters region $\mathcal{W}_q^4$.

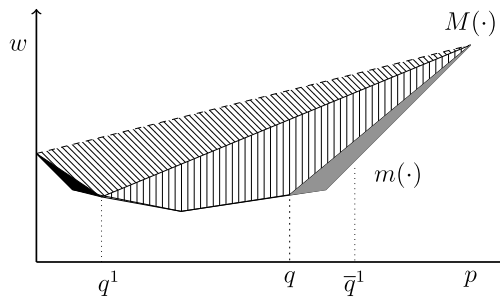


FIGURE 3. Regions $\mathcal{W}_q^1$ (black), $\mathcal{W}_q^2$ (vertical lines), $\mathcal{W}_q^3$ (gray), and $\mathcal{W}_q^4$ (slanted lines).

In region $\mathcal{W}_q^4$, the principal discloses the state with sufficiently high probability so that, when disclosure does not occur, the agent's belief is $\underline{q}^1$. At the belief $\underline{q}^1$, the agent plays a^* one final time. (Recall that at $\underline{q}^1$, the agent plays a^* if rewarded with full information disclosure at the next period.)

In region $\mathcal{W}_q^1$, the belief p is so low that even the promise of full information disclosure at the next period does not incentivize the agent to play a^* , not even once. In this region, the principal sends either the signal s^* or the signal s_0 . The signal s_0 perfectly informs the agent that the state is ω_0 , while the signal s^* induces the belief $\underline{q}^1$, at which the agent plays a^* one final time.

In region $\mathcal{W}_q^3$, the belief p is higher than q and, even possibly, higher than $\bar{q}^1$. In this region, the principal sends either the signal s^* or the signal s_1 . The signal s_1 perfectly informs the agent that the state is ω_1 , while the signal s^* induces the belief $q \leq \bar{q}^1$, at which the agent plays a^* . It is almost the mirror image of what the policy does in region of $\mathcal{W}_q^1$; the only conceptual difference is that the policy induces the belief q rather than $\bar{q}^1$, the natural counterpart of $\underline{q}^1$. This asymmetry is a consequence of trying to induce a^* at the lowest possible average belief.

We now formally define the policy τ_q , starting with region $\mathcal{W}_q^2$. Define the functions $\lambda : \mathcal{W} \rightarrow [0, 1]$ and $\varphi : \mathcal{W} \rightarrow [0, 1]$ so that $(\lambda(p, w), \varphi(p, w))$ is the unique solution of

$$\begin{pmatrix} p \\ w \end{pmatrix} = \lambda(p, w) \begin{pmatrix} \varphi(p, w) \\ m(\varphi(p, w)) \end{pmatrix} + (1 - \lambda(p, w)) \begin{pmatrix} 1 \\ m(1) \end{pmatrix} \quad (5)$$

for all $w > m(p)$ and $(\lambda(p, m(p)), \varphi(p, m(p))) = (1, p)$. When (p, w) is in region $\mathcal{W}_q^2$, the policy splits p into two beliefs $\varphi(p, w)$ and 1, with probability $\lambda(p, w)$ and $1 - \lambda(p, w)$, respectively. When the posterior belief is $\varphi(p, w)$, the policy recommends a^* and promises the continuation payoff $\mathbf{w}(\varphi(p, w))$ if the recommendation is followed. Therefore, if the agent follows the recommendation, his discounted expected payoff is $m(\varphi(p, w)) = (1 - \delta)u(a^*, \varphi(p, w)) + \delta\mathbf{w}(\varphi(p, w))$. When the posterior belief is 1, the policy recommends a^1 and promises the continuation payoff $m(1)$, with a^1 an optimal action at state ω_1 . Therefore, if the agent follows the recommendation, he achieves the discounted expected payoff $m(1)$. Note that when $w = m(p)$, the principal recommends a^* with probability one, and promises the continuation payoff $\mathbf{w}(p)$ in the future. Upon following the recommendation, the agent achieves the discounted expected payoff $m(p)$.

The key feature of the policy in region $\mathcal{W}_q^2$ is to disclose, with some probability, that the state is ω_1 . As we already suggested, the rationale for disclosing when the state is ω_1 is two-fold. First, the lower the agent's belief, the lower the cost of incentivizing the agent to play a^* relative to the principal's benefit. Second, to satisfy the promise-keeping constraint, the policy needs to compensate the agent for playing a^* . Since the principal's payoff is zero when the agent takes any action different from a^* , the best is to choose a compensation, which guarantees the highest probability of playing a^* . Putting these two observations together, at (p, w) , policy $\tau_q(p, w)$ finds two beliefs (p', p'') such that (i) the agent is asked to play a^* at p' , (ii) $p' < p$ since the agent should play a^* at the lowest belief, and (iii) the probability of p' is as high as possible. The best splitting is to

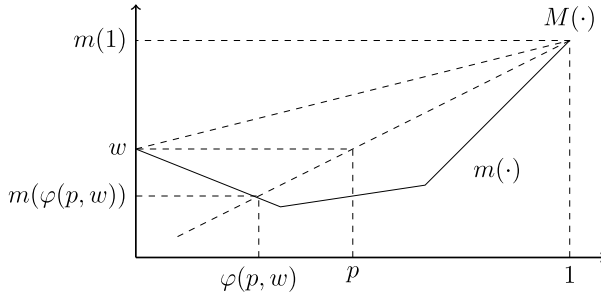


FIGURE 4. Construction of λ and φ : $p = \lambda\varphi + (1 - \lambda)1$; $w = \lambda m(\varphi) + (1 - \lambda)m(1)$.

have p' as close as possible to p and p'' as far as possible, that is, equal to 1. Observe that since $(1 - \lambda(p, w))m(1) + \lambda(p, w)m(\varphi(p, w)) = w$, the promise-keeping constraint binds in region $\mathcal{W}_q^2$. See Figure 4 for an illustration.

Note that starting with $(p, w) \in \mathcal{W}_q^2$, the decreasing sequence of beliefs $(\varphi(p, w), \varphi^2(p, w), \dots)$ (and corresponding payoffs) reaches either region $\mathcal{W}_q^4$ —as in panels (A) and (C) of Figure 1—or a belief in P at which it is statically optimal for the agent to play a^* —as in panel (B) of Figure 1.¹⁴ In the latter case, the policy recommends a^* and stops disclosing information (i.e., the belief stays constant).

When (p, w) is in region $\mathcal{W}_q^4$, the agent cannot be incentivized to play a^* at (p, w) .¹⁵ In that case, the policy splits p into posteriors 0, $\underline{q}^1$, and 1 with respective probabilities λ_0 , $\lambda_{\underline{q}^1}$, and λ_1 . Conditional on 0 (resp., 1), the policy recommends an action optimal at 0 (resp., an action optimal at 1), and promises a continuation payoff of $m(0)$ (resp., $m(1)$). Conditional on $\underline{q}^1$, the policy recommends action a^* and promises a continuation payoff of $\mathbf{w}(\underline{q}^1)$. Doing so, the principal ensures that the agent plays a^* one more time. The probabilities $(\lambda_0, \lambda_{\underline{q}^1}, \lambda_1) \in \mathbb{R}_+ \times \mathbb{R}_+ \times \mathbb{R}_+$ are the unique solution to

$$\lambda_0 \begin{pmatrix} 0 \\ m(0) \\ 1 \end{pmatrix} + \lambda_{\underline{q}^1} \begin{pmatrix} \underline{q}^1 \\ m(\underline{q}^1) \\ 1 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ m(1) \\ 1 \end{pmatrix} = \begin{pmatrix} p \\ w \\ 1 \end{pmatrix}.$$

A solution exists since $\mathcal{W}_q^4$ is the convex hull of $(0, m(0))$, $(\underline{q}^1, m(\underline{q}^1))$, and $(1, m(1))$. In this region, the promise-keeping constraint is also binding.

When (p, w) is in region $\mathcal{W}_q^1$, the policy splits p into 0 (i.e., discloses that the state is ω_0) and $\underline{q}^1$ with respective probabilities $\frac{\underline{q}^1 - p}{\underline{q}^1}$ and $\frac{p}{\underline{q}^1}$. If the realized belief is 0, the policy recommends an action optimal at 0 and promises a continuation payoff of $m(0)$. If the realized belief is $\underline{q}^1$, the policy recommends a^* and promises a continuation payoff of $\mathbf{w}(\underline{q}^1)$. The agent is thus made indifferent between playing a^* and receiving $\mathbf{w}(\underline{q}^1)$ in the future, and playing a best reply to the belief $\underline{q}^1$ forever. Intuitively, in region $\mathcal{W}_q^1$, the principal cannot incentivize the agent to take action a^* by promising future information

¹⁴We write $\varphi^2(p, w)$ for $\varphi(\varphi(p, w), m(\varphi(p, w)))$.

¹⁵Recall that $\underline{q}^1$ is the lowest belief at which the agent can be incentivized to play a^* .

disclosure (since $p < \underline{q}^1$). Hence, the principal must first persuade the agent by disclosing some information. Note that the promise-keeping constraint is slack in this region whenever (p, w) satisfies $\frac{q^1-p}{q^1}m(0) + \frac{p}{q^1}m(q^1) > w$.

When (p, w) is in region $\mathcal{W}_q^3$, the policy splits p into q and 1 with respective probabilities $\frac{1-p}{1-q}$ and $\frac{p-q}{1-q}$. Conditional on 1, the policy recommends an action optimal at 1 and promises a continuation payoff of $m(1)$. Conditional on q , the policy recommends a^* and promises a continuation payoff of $\mathbf{w}(q)$. The agent is thus made indifferent between playing a^* and receiving $\mathbf{w}(q)$ in the future, and playing a best-reply to the belief q forever. The policy in this region is analogous to the one in region $\mathcal{W}_q^1$ —the policy starts by disclosing some information. When $q = \bar{q}^1$, the reason for the analogy is immediate, as $\bar{q}^1$ is the highest belief at which the agent is willing to take action a^* at the current period in exchange for full information at the next period. As we shall see later, the optimal policy τ_{q^*} may require $q^* < \bar{q}^1$, in order to guarantee that the principal's value function is concave, a necessary requirement to minimize the cost of incentivizing the agent relative to the benefit to the principal. As in region $\mathcal{W}_q^1$, the promise-keeping constraint is also slack in this region whenever (p, w) satisfies $w < \frac{1-p}{1-q}m(q) + \frac{p-q}{1-q}m(1)$. This completes the description of the policy τ_q .

Before moving on, we first verify that our policy τ_{q^*} is optimal under the two benchmark scenarios discussed in Section 3.3. Given the value function, we just need to check whether τ_{q^*} solves the Bellman equation.

COROLLARY 2. *The policy τ_{q^*} is optimal both when $|A| = 2$ and when $\frac{v(a^*, 0)}{v(a^*, 1)} = \frac{m(0) - u(a^*, 0)}{m(1) - u(a^*, 1)}$.*

We now illustrate our construction by revisiting Example 1.

3.5 Example 1 revisited

We have that $M(p) = 1 + p$, $m(p) = \max(1 - p, 2p)$ and $\mathbf{w}(p) = 2 \max(2p, 1 - p) - (1/2)$. Therefore, $Q^1 = [1/6, 1/2]$. Assume that $q = 1/3$ (we will show that this choice is the optimal's one). Remember that the prior probability of ω_1 is $1/3$ and the discount factor is $1/2$. Let us start with the pair $(p, m(p)) = (1/3, 2/3)$, which is in region $\mathcal{W}_{1/3}^2$. The policy recommends a^* to the agent and promises a continuation payoff of $\mathbf{w}(1/3) = 5/6$. The next value of the state variables is therefore $(1/3, 5/6)$, which is again in $\mathcal{W}_{1/3}^2$. If the agent had been obedient, the policy then splits the prior probability $1/3$ into $3/11$ and 1 with probability $22/24$ and $2/24$, respectively. Indeed, we have

$$\begin{pmatrix} 1 \\ 3 \\ 5 \\ 6 \end{pmatrix} = \frac{22}{24} \begin{pmatrix} 3 \\ 11 \\ 3 \\ 11 \end{pmatrix} + \frac{2}{24} \begin{pmatrix} 1 \\ m(1) \end{pmatrix}.$$

Conditional on the posterior $3/11$, the policy recommends a^* to the agent and promises a continuation payoff of $\mathbf{w}(3/11) = 21/22$. Conditional on the posterior 1, the

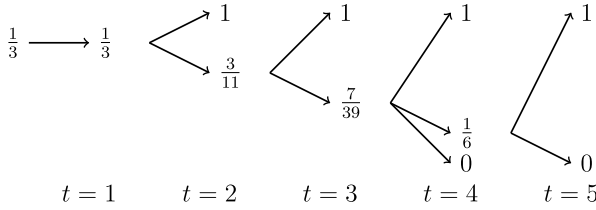


FIGURE 5. Evolution of the beliefs.

policy recommends a_1 and promises a continuation payoff of $m(1) = 2$. Therefore, the next value of the state variables is either $(3/11, 21/22)$ or $(1, 2)$, with the former again in $\mathcal{W}_{1/3}^2$.

If the value of the state variables is $(1, 2)$, the policy yet again recommends a_1 and a continuation payoff of 2. If the value of the state variables is $(3/11, 21/22)$, the policy splits $3/11$ into $7/39$ and 1 , with probability $39/44$ and $5/44$, respectively. Conditional on the posterior $7/39$, the policy recommends a^* to the agent and promises a continuation payoff of $\mathbf{w}(7/39) = 89/78$. Conditional on the posterior 1 , the policy recommends a_1 and promises a continuation payoff of $m(1) = 2$.

Finally, at the state variables value of $(7/39, 89/78)$, which is in region $\mathcal{W}_{1/3}^4$, the policy does a penultimate split of $7/39$ into 0 , $1/6$ and 1 with probability $113/156$, $18/156$ and $25/156$, respectively. Conditional on the posterior $1/6$, the policy recommends a^* and promises a continuation payoff of $7/6$, that is, full information disclosure at the next period. The policy fully discloses the state in finite time to the agent. See Figure 5 for the evolution of the beliefs at the beginning of each period. At all beliefs other than 0 and 1 , the agent is recommended to play a^* . The principal's expected payoff is $1285/1536$, that is, about 0.83 .

Our optimal policy performs strictly better than the random full-disclosure policy because it exploits the asymmetry in the agent's opportunity cost of choosing a^* in the two states. At each period in which information is disclosed and a^* is played, our policy decreases the belief at which a^* is played; the average discounted beliefs is $p^* \approx 0.197 < 1/3$.

On the contrary, the random full-disclosure policy does not alter the belief that the state is ω_1 when a^* is played; the belief stays fixed at the prior $p_0 = 1/3$.

It remains to explain how to choose the parameter q^* to guarantee the optimality of τ_{q^*} .

3.6 Construction of q^* and optimality

For all $q \in [\underline{q}^1, \bar{q}^1]$, let $V_q : \mathcal{W} \rightarrow \mathbb{R}$ be the value function induced by the policy τ_q . For all q , note that $V_q(1, m(1)) = 0$ since a^* is not optimal at $p = 1$, and $V_q(0, m(0)) = 0$ if a^* is not optimal at $p = 0$ (resp., $V_q(0, m(0)) = v(a^*, 0)$ if a^* is optimal at $p = 0$). Also, $V_q(\underline{q}^1, m(\underline{q}^1)) = (1 - \delta)v(a^*, \underline{q}^1)$ if $\underline{q}^1 > 0$ (resp., $V_q(0, m(0)) = v(a^*, 0)$ if $\underline{q}^1 = 0$, since a^* is then optimal at $p = 0$). Therefore, any two policies τ_q and $\tau_{q'}$ induce the same values at all $(p, w) \in \mathcal{W}_q^1 \cup \mathcal{W}_q^4 = \mathcal{W}_{q'}^1 \cup \mathcal{W}_{q'}^4$. (Remember that the regions $\mathcal{W}_q^1$ and $\mathcal{W}_q^4$ do not vary with q ; see Figure 3.)

Similarly, any two policies τ_q and $\tau_{q'}$ induce the same values at all $(p, w) \in \mathcal{W}_{\min(q, q')}^2$. Thus, in particular, τ_q and $\tau_{\bar{q}^1}$ induce the same values at all $(p, w) \in \mathcal{W} \setminus \mathcal{W}_q^3$. Finally, at all $(p, w) \in \mathcal{W}_q^3$, $V_q(p, w) = \frac{1-p}{1-q} V_q(q, m(q)) = \frac{1-p}{1-q} V_{\bar{q}^1}(q, m(q))$. Hence, characterizing $V_{\bar{q}^1}$ is enough to characterize V_q . (See Appendix B for more details.)

Recall that V^* is the unique solution to the fixed-point problem—to be optimal, a policy must therefore induce the value function V^* . Let

$$q^* = \sup\{p \in [\underline{q}^1, \bar{q}^1] : V_{\bar{q}^1}(p, m(p)) \geq V_{\bar{q}^1}(p, w) \text{ for all } w\}.$$

We are now ready to state our main result.

THEOREM 1. *The policy τ_{q^*} is optimal: $V_{q^*} = V^*$.*

To understand the role of q^* , recall that for all $p \in [q^*, 1]$, the policy leaves rents to the agent.¹⁶ To minimize these rents, the principal therefore would like to have q^* as high as possible, that is, equal to $\bar{q}^1$, the highest belief at which the agent is willing to play a^* in exchange for full information disclosure at the next period. However, $V_{\bar{q}^1}(\cdot, m(\cdot))$ is not guaranteed to be concave in p , a necessary condition for optimality. To see that $V^*(\cdot, m(\cdot))$ must be concave in p , consider any pair $(p, p') \in [0, 1] \times [0, 1]$ and $\alpha \in [0, 1]$. We have

$$\begin{aligned} \alpha V^*(p, m(p)) + (1 - \alpha) V^*(p', m(p')) &\leq V^*(\alpha p + (1 - \alpha) p', \alpha m(p) + (1 - \alpha) m(p')) \\ &\leq V^*(\alpha p + (1 - \alpha) p', m(\alpha p + (1 - \alpha) p')), \end{aligned}$$

where the first inequality follows from the concavity of V^* in both arguments and the second from V^* decreasing in w and the convexity of m . The optimal choice of q^* is thus the largest q , which guarantees $V_q(\cdot, m(\cdot))$ to be concave.

More precisely, as we show in Appendix A.5, the definition of q^* guarantees that V_{q^*} is concave in both arguments and decreasing in w , so that $V_{q^*}(\cdot, m(\cdot))$ is a concave function of p . We also prove that $V_{q^*}(p, m(p)) \geq V_{\bar{q}^1}(p, m(p))$ for all p . Since it is clearly the smallest such function, V_{q^*} is the concavification of $V_{\bar{q}^1}$. In particular, $q^* = \bar{q}^1$ if $V_{\bar{q}^1}(\cdot, m(\cdot))$ is already concave. Figure 6 illustrates the concavification for Example 1. In dashed red is the value function of policy $\tau_{\bar{q}^1}$; in solid blue its concavification—the value function of policy τ_{q^*} , with $q^* = \frac{1}{3}$.

The policy τ_{q^*} leaves rents to the agent, that is, the (ex ante) participation constraint does not bind, for all priors in $[0, \underline{q}^1] \cup (q^*, 1]$. This is quite natural for all priors in $[0, 1] \setminus Q^1$ since the agent cannot be incentivized to play a^* even once. In the language of Ely and Szydlowski (2020), “the goalposts need to move,” that is, one needs to disclose information at the ex ante stage to persuade the agent to play a^* . However, our policy also leaves rents for all priors in $(q^*, \bar{q}^1]$. The intuitive reason is that the initial information disclosure reduces the cost of incentivizing the agent in subsequent periods sufficiently enough to compensate for the initial loss. (When the realized posterior is 1, the agent never plays a^* , thus creating the loss.)

¹⁶That is, the agent is promised a payoff of $\frac{1-p}{1-q^*} m(q^*) + \frac{p-q^*}{1-q^*} m(1) > m(p)$.

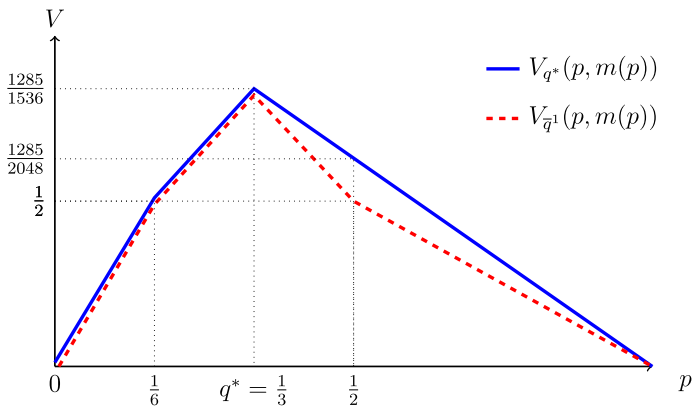


FIGURE 6. The concavification of $V_{q^{-1}}(\cdot, m(\cdot))$ in Example 1.

4. EVOLUTION OF BELIEFS IN THE OPTIMAL POLICY

The optimal policy discloses information gradually over time, with beliefs evolving until either the agent learns the state or believes that a^* is statically optimal. We can be more specific. First, we consider the instances when the policy converges with positive probability to a belief $p \in P = [\underline{p}, \overline{p}]$, the set of beliefs at which a^* is optimal. Let $Q^\infty = [\underline{p}, \overline{q}^\infty]$, with $\overline{q}^\infty$ the solution to

$$m(\overline{q}^\infty) = (1 - \delta)u(a^*, \overline{q}^\infty) + \delta\left(\frac{1 - \overline{q}^\infty}{1 - \underline{p}}m(\underline{p}) + \frac{\overline{q}^\infty - \underline{p}}{1 - \underline{p}}m(1)\right),$$

if P is nonempty, and $Q^\infty = \emptyset$, otherwise. Note that $P \subseteq Q^\infty$. See Figure 7 for a graphical illustration.

Intuitively, the set Q^∞ has the “fixed-point property,” that is, if one starts with a belief $p \in Q^\infty$ and promised utility $\mathbf{w}(p)$, then the belief $\varphi(p, \mathbf{w}(p)) \in Q^\infty$. To see this, note that the pair $(p, \mathbf{w}(p))$ is in region $\mathcal{W}_q^2$. Since $\varphi(p, \mathbf{w}(p)) \leq p$ (with a strict inequality if $p \notin P$), we then have a decreasing sequence of beliefs converging to an element in P . This is because, at all beliefs $p \in Q^\infty$, the policy splits p into $p' = \varphi(p, \mathbf{w}(p))$ and 1,

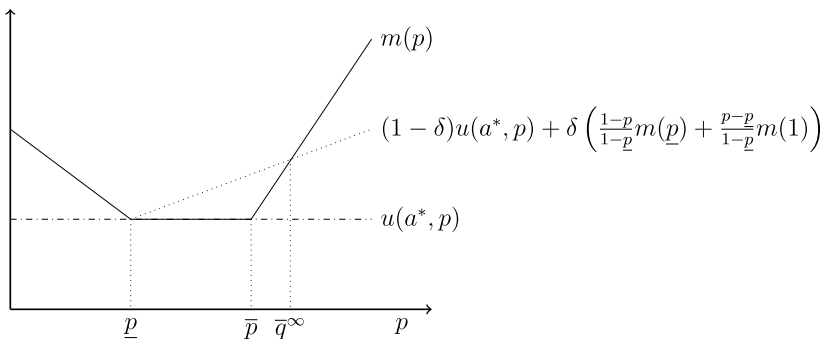


FIGURE 7. Construction of $\overline{q}^\infty$.

then splits p' into $p'' = \varphi(p', \mathbf{w}(p'))$ and 1, etc. The decreasing sequence $(p, p', p'', \dots)$ converges, either in finite time or asymptotically, to a belief in P , at which no further splitting occurs and the agent plays a^* forever. See panel (B) of Figure 1 for an illustration.

Recall that if the prior p_0 is larger than q^* , the policy first splits p_0 into q^* and 1. Hence, if $q^* \leq \bar{q}^\infty$, the agent's belief enters the set Q^∞ with strictly positive probability.¹⁷ Therefore, if the agent's prior belief is in the set $Q_{q^*}^\infty$, then there is a strictly positive probability that the agent chooses action a^* forever, where

$$Q_{q^*}^\infty := \begin{cases} Q^\infty & \text{if } q^* > \bar{q}^\infty, \\ [\underline{p}, 1) & \text{otherwise.} \end{cases}$$

Second, at all priors in $[0, 1] \setminus Q_{q^*}^\infty$, there exists $T_\delta < \infty$ such that the belief process is absorbed in the degenerate beliefs 0 or 1 after at most T_δ periods. In other words, the agent learns the state for sure in finite time. The number of periods T_δ corresponds to the maximal number of periods the agent can be incentivized to play a^* . We provide an explicit computation in Appendix B. In Example 1, $T_\delta = 3$. Moreover, the number T_δ is increasing in δ and converges to $+\infty$ as δ converges to 1. (Note that the convergence is uniform in that it does not depend on $p_0 \in [0, 1] \setminus Q_{q^*}^\infty$.) Thus, we have the following corollary.

COROLLARY 3. *Under the optimal disclosure policy τ_{q^*} , there is a strictly positive probability that the agent chooses action a^* forever if, and only if, $p_0 \in Q_{q^*}^\infty$. Alternatively, if $p_0 \notin Q_{q^*}^\infty$, then there exists T_δ such that the agent perfectly learns the state (i.e., p reaches either 0 or 1) with probability 1 after at most T_δ periods.*

The interval $Q_{q^*}^\infty$ includes the subinterval $[\underline{p}, \bar{p}]$, where the agent takes action a^* with probability one. In the complementary set $Q_{q^*}^\infty \setminus [\underline{p}, \bar{p}]$, the probability that the agent takes action a^* forever is strictly less than 1. That is, the principal discloses the state with positive probability, and with the complementary probability he lowers the agent's belief so that it converges to the region where taking action a^* is statically optimal. Convergence may be asymptotic or may happen in finite time.

As already mentioned, the promise-keeping constraint binds in regions $\mathcal{W}_{q^*}^2$ and $\mathcal{W}_{q^*}^4$, but may not bind in the other two regions. We now argue that under our policy τ_{q^*} , the promise-keeping constraint can only be slack in the first period. In other words, the promised-keeping constraint binds from period two onwards. To see this, suppose that $(p_0, m(p_0))$ is in region $\mathcal{W}_{q^*}^3$, hence the prior belief $p_0 \in (q^*, 1)$. What the policy τ_{q^*} does is to split p_0 into q^* and 1, so that the state variable transit to either $(q^*, m(q^*))$ or $(1, m(1))$. In the latter case, the promise-keeping constraint clearly binds and will continue to bind in all subsequent periods, since the agent has learned that the state is ω_1 . In the former case, since $(q^*, m(q^*)) \in \mathcal{W}_{q^*}^2$, the promise-keeping constraint binds and will continue to bind in all subsequent periods since the subsequent state variables will

¹⁷From the definition of q^* , we have that $q^* \geq \bar{p}$ since $V_{\bar{q}^1}(p, m(p)) = u(a^*, p)$ for all $p \in P$.

either be in regions $\mathcal{W}_{q^*}^2$ or $\mathcal{W}_{q^*}^4$ or equal to $(1, m(1))$. A symmetric argument holds when $(p_0, m(p_0))$ is in region $\mathcal{W}_{q^*}^1$.

COROLLARY 4. *Under the optimal policy τ_q^* , the promise-keeping constraint can only be slack in the first period.*

All in all, information disclosure plays two roles in our optimal policy. First, the promise of future information disclosure motivates the agent to take action a^* in early periods. The intertemporal incentives make it possible to motivate to play a^* at beliefs outside P . Second, information disclosure decreases the discounted average belief that the state is the high opportunity cost state ω_1 and, therefore, makes it easier to incentivize the agent to take action a^* for a longer expected time.

APPENDIX A: PROOFS

A.1 Mathematical preliminaries

We collect without proofs some useful results about concave functions. Let $f : [a, b] \rightarrow \mathbb{R}$ be a concave function and $a \leq x < y < z \leq b$. The following properties hold:

- (a) $\frac{f(y)-f(x)}{y-x} \geq \frac{f(z)-f(y)}{z-y}$.
- (b) $\frac{f(y)-f(a)}{y-a} \geq \frac{f(z)-f(a)}{z-a}$.
- (c) $\frac{f(b)-f(x)}{b-x} \geq \frac{f(b)-f(y)}{b-y}$.
- (d) $\frac{f(y)-f(x)}{y-x} \geq \frac{f(y+\Delta)-f(x+\Delta)}{y-x}$ for all $\Delta \geq 0$ such that $y + \Delta \leq b$.

Note that property (a) implies (d) and is true irrespective of whether $x + \Delta \leq y$. We will repeatedly use these properties in most of the following proofs.

To prove Lemma 3, we will use the following property: if $f : [a, b] \rightarrow \mathbb{R}$ satisfies $\frac{f(x)-f(a)}{x-a} \geq \frac{f(y)-f(a)}{y-a}$ for all $a < x \leq y \leq b$, then f is concave.

A.2 Proposition 2

PROOF OF PROPOSITION 2(I). By contradiction, assume that there exists $s' \in S$ such that $\lambda_{s'} > 0$ and

$$(1 - \delta)v(a_{s'}, p_{s'}) + \delta V^*(p_{s'}, w_{s'}) < V^*(p_{s'}, (1 - \delta)u(a_{s'}, p_{s'}) + \delta w_{s'}).$$

Let $(\lambda_s^*, p_s^*, w_s^*, a_s^*)_{s \in S}$ be the policy, which achieves $V^*(p_{s'}, (1 - \delta)u(a_{s'}, p_{s'}) + \delta w_{s'})$, and consider the new policy

$$((\lambda_s, p_s, w_s, a_s)_{s \in S \setminus \{s'\}}, (\lambda_{s'} \lambda_s^*, p_s^*, w_s^*, a_s^*)_{s \in S}).$$

By construction, the new policy is feasible. Moreover, we have that

$$\begin{aligned}
 & \sum_{s \in S \setminus \{s'\}} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)] + \lambda_{s'} \sum_{s \in S} \lambda_s^* [(1 - \delta)v(a_s^*, p_s^*) + \delta V^*(p_s^*, w_s^*)] \\
 &= \sum_{s \in S \setminus \{s'\}} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)] + \lambda_{s'} V^*(p_{s'}, (1 - \delta)u(a_{s'}, p_{s'}) + \delta w_{s'}) \\
 &> \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)],
 \end{aligned}$$

a contradiction with the optimality of $(\lambda_s, p_s, w_s, a_s)_{s \in S}$. Thus, we must have $(1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s) \geq V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)$ for all s such that $\lambda_s > 0$.

Since the fixed point satisfies $V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s) \geq (1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s)$, we have the desired result. $\square$

PROOF OF PROPOSITION 2(II). Let $s \in S$ such that $\lambda_s > 0$ and $a_s \neq a^*$. We have

$$\begin{aligned}
 (1 - \delta)v(a_s, p_s) + \delta V^*(p_s, w_s) &= \delta V^*(p_s, w_s) \geq V^*(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s) \\
 &\geq V^*(p_s, w_s),
 \end{aligned}$$

where the first inequality follows from Proposition 2(i) and the second follows from V^* decreasing in w and $w_s \geq u(a_s, p_s)$ for

$$(1 - \delta)u(a_s, p_s) + \delta w_s \geq m(p_s),$$

to hold. It follows that $V^*(p_s, w_s) = 0$. $\square$

PROOF OF PROPOSITION 2(III). The proof is by contradiction. Suppose to the contrary that $V^*(p_s, w_s) = V^*(p_s, w'_s)$ for some $w'_s \in (w_s, M(p_s)]$ and $a_s = a^*$. By Proposition 2(i), we have

$$\begin{aligned}
 V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) &= (1 - \delta)v(a^*, p_s) + \delta V^*(p_s, w_s) \\
 &= (1 - \delta)v(a^*, p_s) + \delta V^*(p_s, w'_s) \\
 &\leq V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w'_s).
 \end{aligned}$$

Since V^* is decreasing in w , the inequality cannot be strict, hence

$$V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) = V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w'_s). \quad (6)$$

We now show that

$$V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) = V^*(p_s, w_s), \quad (7)$$

hence

$$V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w'_s) = V^*(p_s, w'_s) = v(a^*, p_s),$$

where the last equality follows from $a_s = a^*$ and Proposition 2(i). This means that after signal s , action a^* is taken with probability one in all periods. This is the required contradiction, since $u(a^*, p_s) \leq (1 - \delta)u(a^*, p_s) + \delta w_s < (1 - \delta)u(a^*, p_s) + \delta w'_s$: No feasible policy promising utility w'_s guarantees that a^* is chosen with probability one in all periods.

It remains to prove that equation (7) is true. Recall that $(1 - \delta)u(a^*, p_s) + \delta w_s \leq w_s < w'_s$. We consider two cases. First, assume that $(1 - \delta)u(a^*, p_s) + \delta w'_s < w_s$. Equation (7) then follows from

$$\begin{aligned} V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) &\geq V^*(p_s, w_s) \\ &\geq \alpha V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w'_s) + (1 - \alpha)V^*(p_s, w'_s) \\ &= \alpha V^*(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) + (1 - \alpha)V^*(p_s, w_s), \end{aligned}$$

where α is the weight on $(1 - \delta)u(a^*, p_s) + \delta w'_s$ such that the convex combination of $(1 - \delta)u(a^*, p_s) + \delta w'_s$ and w'_s is equal to w_s .

Second, if $(1 - \delta)u(a^*, p_s) + \delta w'_s \geq w_s$, then equation (7) follows from a similar argument using a convex combination of $(1 - \delta)u(a^*, p_s) + \delta w_s$ and $(1 - \delta)u(a^*, p_s) + \delta w'_s$. $\square$

A.3 Proposition 3

PROOF OF PROPOSITION 3(1), PART A. We show that we can restrict attention to contracts where $a_s = a^*$ for at most one signal s such that $\lambda_s > 0$. Let $(\lambda'_s, p'_s, w'_s, a'_s)_{s \in S'}$ be a solution to the maximization program $T(V^*)(p, w)$. Let $S^* \subseteq S'$ be the set of signals such that $a_s = a^*$ and $\lambda_s > 0$. If S^* is empty, there is nothing to prove. If S^* is nonempty, define p^* as

$$\sum_{s \in S^*} \left(\frac{\lambda'_s}{\sum_{s \in S^*} \lambda'_s} \right) p_s = p^*,$$

and $\sum_{s \in S^*} \lambda'_s = \lambda^*$. From the concavity of V^* , we have that

$$\begin{aligned} \sum_{s \in S^*} \lambda'_s (v(a^*, p'_s)(1 - \delta) + \delta V^*(p'_s, w'_s)) &= \lambda^* \left(v(a^*, p^*)(1 - \delta) + \delta \sum_{s \in S^*} \left(\frac{\lambda'_s}{\lambda^*} \right) V^*(p'_s, w'_s) \right) \\ &\leq \lambda^* (v(a^*, p^*)(1 - \delta) + \delta V^*(p^*, w^*)), \end{aligned}$$

where

$$w^* = \sum_{s \in S^*} \left(\frac{\lambda'_s}{\sum_{s \in S^*} \lambda'_s} \right) w'_s.$$

Notice that $w^* \in [m(p^*), M(p^*)]$ since the convexity of m implies

$$M(p^*) = \sum_{s \in S^*} \left(\frac{\lambda'_s}{\sum_{s \in S^*} \lambda'_s} \right) M(p'_s) \geq \sum_{s \in S^*} \left(\frac{\lambda'_s}{\sum_{s \in S^*} \lambda'_s} \right) w_s \geq \sum_{s \in S^*} \left(\frac{\lambda'_s}{\sum_{s \in S^*} \lambda'_s} \right) m(p'_s) \geq m(p^*).$$

It is routine to verify that the new contract

$$((\lambda'_s, p'_s, w'_s, a'_s)_{s \in S \setminus S^*}, (\lambda^*, p^*, a^*, w^*))$$

is feasible and, therefore, also optimal. $\square$

PROOF OF PROPOSITION 3(II). Let $(\lambda'_s, p'_s, w'_s, a'_s)_{s \in S}$ be a solution to the maximization program $T(V^*)(p, w)$. From Proposition 3(i), part A, we can assume that there exists a unique signal s^* such that $a'_{s^*} = a^*$. We construct a new contract with three messages, s^* , s_0 , and s_1 , as follows. First, $(\lambda_{s^*}, p_{s^*}, w_{s^*}, a_{s^*}) = (\lambda'_{s^*}, p'_{s^*}, w'_{s^*}, a'_{s^*})$. Second, $\lambda_{s_0} = \sum_{s \in S \setminus \{s^*\}} \lambda'_s (1 - p'_s)$, $p_{s_0} = 0$, $w_{s_0} = m(0)$, and $a_{s_0} \in \arg \max_{a \in A} u(a, 0)$. Third, $\lambda_{s_1} = \sum_{s \in S \setminus \{s^*\}} \lambda'_s p'_s$, $p_{s_1} = 1$, $w_{s_1} = m(1)$, and $a_{s_1} \in \arg \max_{a \in A} u(a, 1)$. It is routine to verify that this new contract is feasible. (To check that the promise-keeping constraint is satisfied, we simply need to observe that $(1 - \delta)u(a'_s, p'_s) + \delta w'_s \leq M(p'_s) = (1 - p'_s)m(0) + p'_s m(1)$.) Since the new contract gives the same payoff to the principal, it is optimal. $\square$

PROOF OF PROPOSITION 3(I), PART B. From part A and (ii), we can restrict attention to contracts with three messages s^* , s_0 , and s_1 , such that $p_{s_0} = 0$, $p_{s_1} = 1$, and a^* is recommended at s^* . To ease notation, we denote such a contract by $(\lambda'_{s^*}, p'_{s^*}, w'_{s^*}, \lambda'_{s_0}, \lambda'_{s_1})$. In words, the contract induces the beliefs p'_{s^*} , 0, and 1, with probability λ'_{s^*} , λ'_{s_0} , and λ'_{s_1} , respectively. At s^* , the contract recommends a^* and promises a continuation payoff of w'_{s^*} . Throughout the proof, we refer to such a contract as a *simple* contract.

Among all optimal simple contracts at (p, w) , fix one that minimizes the probability of recommending a^* . Denote it $(\lambda_{s^*}, p_{s^*}, w_{s^*}, \lambda_{s_0}, \lambda_{s_1})$. The existence of such a contract follows from standard arguments. (See Appendix C.2 for details.) We want to show that $(1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*} = m(p_{s^*})$ if $\lambda_{s^*} > 0$.

Under this contract, the principal's payoff is

$$V^*(p, w) = \lambda_{s^*} [(1 - \delta)v(a^*, p_{s^*}) + \delta V^*(p_{s^*}, w_{s^*})] = \lambda_{s^*} V^*(p_{s^*}, (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*}),$$

where the second equality follows from Proposition 2(i). $\square$

We complete the proof by contradiction. Suppose that $\lambda_{s^*} > 0$, but $(1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*} > m(p_{s^*})$. We will construct another simple contract, which is also optimal and has a strictly lower probability of recommending a^* , thus contradicting the hypothesis that $(\lambda_{s^*}, p_{s^*}, w_{s^*}, \lambda_{s_0}, \lambda_{s_1})$ minimizes the probability of recommending a^* . We need the following lemma.

LEMMA 1. *For any $(p, w) \in \mathcal{W}$, such that $w > m(p)$ and*

$$V^*(p, w) = (1 - \delta)v(a^*, p) + \delta V^*\left(p, \frac{w - (1 - \delta)u(a^*, p)}{\delta}\right),$$

there exist $(\hat{p}_{s^}, \hat{w}_{s^*}) \in \mathcal{W}$ and $(\hat{\lambda}_{s^*}, \hat{\lambda}_{s_0}, \hat{\lambda}_{s_1}) \in [0, 1]^3$ such that $\hat{\lambda}_{s^*} + \hat{\lambda}_{s_0} + \hat{\lambda}_{s_1} = 1$,*

$$\hat{\lambda}_{s^*} \begin{pmatrix} \hat{p}_{s^*} \\ \hat{w}_{s^*} \end{pmatrix} + \hat{\lambda}_{s_0} \begin{pmatrix} 0 \\ m(0) \end{pmatrix} + \hat{\lambda}_{s_1} \begin{pmatrix} 1 \\ m(1) \end{pmatrix} = \begin{pmatrix} p \\ w \end{pmatrix},$$

$V^(p, w) \leq \hat{\lambda}_{s^*} V^*(\hat{p}_{s^*}, \hat{w}_{s^*})$, and $\hat{\lambda}_{s^*} < 1$.*

Since

$$V^*(p_{s^*}, (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*}) = (1 - \delta)v(a^*, p_{s^*}) + \delta V^*(p_{s^*}, w_{s^*}),$$

an application of Lemma 1 at $(p_{s^*}, (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*})$ guarantees the existence of $(\hat{p}_{s^*}, \hat{w}_{s^*}) \in \mathcal{W}$ and $(\hat{\lambda}_{s^*}\hat{\lambda}_{s_0}, \hat{\lambda}_{s_1}) \in [0, 1]^3$ such that $\hat{\lambda}_{s^*} + \hat{\lambda}_{s_0} + \hat{\lambda}_{s_1} = 1$, $\hat{\lambda}_{s^*} < 1$,

$$\hat{\lambda}_{s^*}\hat{p}_{s^*} + \hat{\lambda}_{s_0}0 + \hat{\lambda}_{s_1}1 = p_{s^*},$$

$$\hat{\lambda}_{s^*}\hat{w}_{s^*} + \hat{\lambda}_{s_0}m(0) + \hat{\lambda}_{s_1}m(1) = (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*},$$

$$V^*(p_{s^*}, (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*}) \leq \hat{\lambda}_{s^*}V^*(\hat{p}_{s^*}, \hat{w}_{s^*}).$$

Consider the following simple contract:

$$\left(\lambda_{s^*}\hat{\lambda}_{s^*}, \hat{p}_{s^*}, \frac{\hat{w}_{s^*} - (1 - \delta)u(a^*, \hat{p}_{s^*})}{\delta}, \lambda_{s^*}\hat{\lambda}_{s_0} + \lambda_{s_0}, \lambda_{s^*}\hat{\lambda}_{s_1} + \lambda_{s_1} \right).$$

We first argue that the contract is feasible at (p, w) . Since $\hat{w}_{s^*} \geq m(\hat{p}_{s^*})$, the contract satisfies

$$(1 - \delta)u(a^*, \hat{p}_{s^*}) + \delta \frac{\hat{w}_{s^*} - (1 - \delta)u(a^*, \hat{p}_{s^*})}{\delta} = \hat{w}_{s^*} \geq m(\hat{p}_{s^*}),$$

that is, obedience is guaranteed at s^* . We also have that

$$\begin{aligned} & \lambda_{s^*}\hat{\lambda}_{s^*} \left[(1 - \delta)u(a^*, \hat{p}_{s^*}) + \delta \frac{\hat{w}_{s^*} - (1 - \delta)u(a^*, \hat{p}_{s^*})}{\delta} \right] \\ & \quad + (\lambda_{s^*}\hat{\lambda}_{s_0} + \lambda_{s_0})m(0) + (\lambda_{s^*}\hat{\lambda}_{s_1} + \lambda_{s_1})m(1) \\ & = \lambda_{s^*} [\hat{\lambda}_{s^*}\hat{w}_{s^*} + \hat{\lambda}_{s_0}m(0) + \hat{\lambda}_{s_1}m(1)] + \lambda_{s_0}m(0) + \lambda_{s_1}m(1) \\ & = \lambda_{s^*} [(1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*}] + \lambda_{s_0}m(0) + \lambda_{s_1}m(1) \geq w, \end{aligned}$$

that is, the contract satisfies the promise-keeping constraint. Finally, the splitting is feasible since

$$\lambda_{s^*}\hat{\lambda}_{s^*} \times \hat{p}_{s^*} + (\lambda_{s^*}\hat{\lambda}_{s_0} + \lambda_{s_0}) \times 0 + (\lambda_{s^*}\hat{\lambda}_{s_1} + \lambda_{s_1}) \times 1 = p.$$

The contract is therefore feasible.

We next argue that the new contract is optimal, since under it the principal's payoff is

$$\lambda_{s^*}\hat{\lambda}_{s^*}V^*(\hat{p}_{s^*}, \hat{w}_{s^*}) \geq \lambda_{s^*}V^*(p_{s^*}, (1 - \delta)u(a^*, p_{s^*}) + \delta w_{s^*}).$$

Finally, since $\hat{\lambda}_{s^*} < 1$, the new contract recommends a^* with probability $\lambda_{s^*}\hat{\lambda}_{s^*} < \lambda_{s^*}$, the required contradiction. It remains to prove Lemma 1.

PROOF OF LEMMA 1. We organize the proof around two claims.

CLAIM 1. For any $w' \in [m(p), w]$, we have that

$$V^*(p, w') = (1 - \delta)v(a^*, p) + \delta V^*\left(p, \frac{w' - (1 - \delta)u(a^*, p)}{\delta}\right).$$

PROOF OF CLAIM 1. Fix any $w' \in [m(p), w]$. Since $u(a^*, p) \leq m(p)$, we have

$$\begin{aligned} w' &< \min \left\{ w, \frac{w' - (1 - \delta)u(a^*, p)}{\delta} \right\} \leq \max \left\{ w, \frac{w' - (1 - \delta)u(a^*, p)}{\delta} \right\} \\ &< \frac{w - (1 - \delta)u(a^*, p)}{\delta}, \end{aligned}$$

and

$$\frac{1}{1 + \delta}w' + \frac{\delta}{1 + \delta} \frac{w - (1 - \delta)u(a^*, p)}{\delta} = \frac{1}{1 + \delta}w + \frac{\delta}{1 + \delta} \frac{w' - (1 - \delta)u(a^*, p)}{\delta}.$$

The concavity of V^* implies that

$$\begin{aligned} \frac{1}{1 + \delta}V^*(p, w') + \frac{\delta}{1 + \delta}V^*\left(p, \frac{w - (1 - \delta)u(a^*, p)}{\delta}\right) \\ \leq \frac{1}{1 + \delta}V^*(p, w) + \frac{\delta}{1 + \delta}V^*\left(p, \frac{w' - (1 - \delta)u(a^*, p)}{\delta}\right). \end{aligned}$$

Rearranging, we obtain

$$\begin{aligned} V^*(p, w') - \delta V^*\left(p, \frac{w' - (1 - \delta)u(a^*, p)}{\delta}\right) &\leq V^*(p, w) - \delta V^*\left(p, \frac{w - (1 - \delta)u(a^*, p)}{\delta}\right) \\ &= (1 - \delta)v(a^*, p). \end{aligned}$$

Since, by definition,

$$V^*(p, w') \geq (1 - \delta)v(a^*, p) + \delta V^*\left(p, \frac{w' - (1 - \delta)u(a^*, p)}{\delta}\right),$$

it follows that

$$V^*(p, w') = (1 - \delta)v(a^*, p) + \delta V^*\left(p, \frac{w' - (1 - \delta)u(a^*, p)}{\delta}\right). \quad \square$$

CLAIM 2. On the domain $[m(p), \frac{w - (1 - \delta)u(a^*, p)}{\delta}]$, the map $w' \mapsto V(p, w')$ is linear.

PROOF OF CLAIM 2. By contradiction, suppose that $w' \mapsto V(p, w')$ is not linear, that is, suppose that there exists $w'' \in [m(p), \frac{w - (1 - \delta)u(a^*, p)}{\delta}]$ such that $\forall \alpha \in (0, 1)$:

$$V^*(p, \alpha m(p) + (1 - \alpha)w'') > \alpha V^*(p, m(p)) + (1 - \alpha)V^*(p, w''). \quad (8)$$

Without loss of generality, we can assume that $w'' \in (\frac{m(p) - (1 - \delta)u(a^*, p)}{\delta}, \frac{w - (1 - \delta)u(a^*, p)}{\delta}]$. (If $w'' \leq \frac{m(p) - (1 - \delta)u(a^*, p)}{\delta}$, then the concavity of V^* implies that (8) is also satisfied for any larger w''' .)

Note that

$$\begin{aligned} m(p) &< \min \left\{ (1-\delta)u(a^*, p) + \delta w'', \frac{m(p) - (1-\delta)u(a^*, p)}{\delta} \right\} \\ &\leq \max \left\{ (1-\delta)u(a^*, p) + \delta w'', \frac{m(p) - (1-\delta)u(a^*, p)}{\delta} \right\} \\ &\leq w'', \end{aligned}$$

and, therefore, there exist a unique $\beta \in (0, 1)$ and $\gamma \in (0, 1]$ such that

$$\begin{cases} (1-\delta)u(a^*, p) + \delta w'' = \beta m(p) + (1-\beta)w'' \\ \frac{m(p) - (1-\delta)u(a^*, p)}{\delta} = \gamma m(p) + (1-\gamma)w'', \end{cases}$$

which implies that $\beta + \delta\gamma = 1$. In addition, observe that

$$\frac{1}{1+\delta}m(p) + \frac{\delta}{1+\delta}w'' = \frac{1}{1+\delta}[(1-\delta)u(a^*, p) + \delta w''] + \frac{\delta}{1+\delta} \frac{m(p) - (1-\delta)u(a^*, p)}{\delta}.$$

From Claim 1, we have that

$$\begin{aligned} V^*(p, m(p)) &= (1-\delta)v(a^*, p) + \delta V^*\left(p, \frac{m(p) - (1-\delta)u(a^*, p)}{\delta}\right) \\ \delta V^*(p, w'') &= -(1-\delta)v(a^*, p) + V^*(p, (1-\delta)u(a^*, p) + \delta w'') \end{aligned}$$

Together with the concavity of V^* , we therefore have

$$\begin{aligned} &\frac{1}{1+\delta}V^*(p, m(p)) + \frac{\delta}{1+\delta}V^*(p, w'') \\ &= \frac{1}{1+\delta}V^*(p, (1-\delta)u(a^*, p) + \delta w'') + \frac{\delta}{1+\delta}V^*\left(p, \frac{m(p) - (1-\delta)u(a^*, p)}{\delta}\right) \\ &\geq \frac{1}{1+\delta}[\beta V^*(p, m(p)) + (1-\beta)V^*(p, w'')] \\ &\quad + \frac{\delta}{1+\delta}[\gamma V^*(p, m(p)) + (1-\gamma)V^*(p, w'')] \\ &= \frac{1}{1+\delta}V^*(p, m(p)) + \frac{\delta}{1+\delta}V^*(p, w''), \end{aligned}$$

which contradicts (8), by setting $\alpha = \frac{1}{1+\delta}$. □

We now complete the proof of Lemma 1. Define the set

$$W := \left\{ w' \in (m(p), M(p)) : V^*(p, w') = (1-\delta)v(a^*, p) + \delta V^*\left(p, \frac{w' - (1-\delta)u(a^*, p)}{\delta}\right) \right\}.$$

The set W is nonempty since $w \in W$. Let $\bar{w} := \sup W$. From Claims 1 and 2, we have that $[m(p), \bar{w}] \subseteq W$ and $w' \mapsto V^*(p, w')$ is linear on the domain $[m(p), \frac{\bar{w} - (1-\delta)u(a^*, p)}{\delta}]$. (Claims 1 and 2 are valid for any $w' \in W$.)

Fix $\tilde{w} \in (\bar{w}, \frac{\bar{w} - (1-\delta)u(a^*, p)}{\delta})$. From the linearity of $w' \mapsto V^*(p, w')$, there exists $\zeta \in (0, 1)$ such that $\zeta m(p) + (1 - \zeta)\tilde{w} = w$, and

$$\zeta V^*(p, m(p)) + (1 - \zeta)V^*(p, \tilde{w}) = V^*(p, w).$$

Moreover, since $\tilde{w} > \bar{w}$, by the definition of $\bar{w}$, we have that

$$V^*(p, \tilde{w}) > (1 - \delta)v(a^*, p) + \delta V^*\left(p, \frac{\tilde{w} - (1 - \delta)u(a^*, p)}{\delta}\right).$$

From part A and (ii), there exists a simple contract $(\tilde{\lambda}_{s^*}, \tilde{p}_{s^*}, \frac{\tilde{w}_{s^*} - (1-\delta)u(a^*, p)}{\delta}, \tilde{\lambda}_{s_0}, \tilde{\lambda}_{s_1})$ at $(p, \tilde{w})$ such that $V^*(p, \tilde{w}) = \tilde{\lambda}_{s^*} V^*(\tilde{p}_{s^*}, \tilde{w}_{s^*})$. It follows that

$$\begin{aligned} V^*(p, w) &= \zeta V^*(p, m(p)) + (1 - \zeta)V^*(p, \tilde{w}) \\ &= \zeta V^*(p, m(p)) + (1 - \zeta)\tilde{\lambda}_{s^*} V^*(\tilde{p}_{s^*}, \tilde{w}_{s^*}) \\ &\leq [\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}] V^*\left(\frac{\zeta p + (1 - \zeta)\tilde{\lambda}_{s^*}\tilde{p}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}}, \frac{\zeta m(p) + (1 - \zeta)\tilde{\lambda}_{s^*}\tilde{w}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}}\right), \end{aligned}$$

where the last inequality follows from the concavity of V^* . To conclude the proof, let

$$\begin{aligned} (\hat{p}_{s^*}, \hat{w}_{s^*}) &= \left(\frac{\zeta p + (1 - \zeta)\tilde{\lambda}_{s^*}\tilde{p}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}}, \frac{\zeta m(p) + (1 - \zeta)\tilde{\lambda}_{s^*}\tilde{w}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}}\right), \\ (\hat{\lambda}_{s^*}, \hat{\lambda}_{s_0}, \hat{\lambda}_{s_1}) &= (\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}, (1 - \zeta)\tilde{\lambda}_{s_0}, \tilde{\lambda}_{s_1}). \end{aligned}$$

To verify that $(\hat{p}_{s^*}, \hat{w}_{s^*}) \in \mathcal{W}$, note that the convexity of m implies that

$$\begin{aligned} m(\hat{p}_{s^*}) &\leq \frac{\zeta}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} m(p) + \frac{(1 - \zeta)\tilde{\lambda}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} m(\tilde{p}_{s^*}) \\ &\leq \frac{\zeta}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} m(p) + \frac{(1 - \zeta)\tilde{\lambda}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} \tilde{w}_{s^*}. \end{aligned}$$

Similarly,

$$\begin{aligned} &\frac{\zeta}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} m(p) + \frac{(1 - \zeta)\tilde{\lambda}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} \tilde{w}_{s^*} \\ &\leq \frac{\zeta}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} M(p) + \frac{(1 - \zeta)\tilde{\lambda}_{s^*}}{\zeta + (1 - \zeta)\tilde{\lambda}_{s^*}} M(\tilde{p}_{s^*}) = M(\hat{p}_{s^*}) \end{aligned}$$

as required. It is routine to verify the other constraints. This completes the proof of Lemma 1.

A.4 Corollary 2

PROOF. Let $\mathcal{A} = \{a^*, a^\dagger\}$, and w.l.o.g. assume that the optimal action for the principal, a^* , is optimal for the agent in state ω_0 , and hence also in the interval $p \in [0, \bar{p}]$, where

$$\bar{p}u(a^*, 1) + (1 - \bar{p})u(a^*, 0) = u(a^*, \bar{p}) = u(a^\dagger, \bar{p}) = \bar{p}u(a^\dagger, 1) + (1 - \bar{p})u(a^\dagger, 0).$$

This implies that

$$\frac{\bar{p}}{1 - \bar{p}} = \frac{u(a^*, 0) - u(a^\dagger, 0)}{u(a^*, 1) - u(a^\dagger, 1)}.$$

Assuming $p_0 > \bar{p}$, under our policy, the principal recommends the agent to take a^* in the first period and promises to split p_0 between 1 and $\bar{p}$ with probability λ in the second period, where $(\lambda, \bar{p})$ solves

$$\begin{cases} \lambda \bar{p} + (1 - \lambda)1 = p_0 \\ \lambda u(a^*, \bar{p}) + (1 - \lambda)u(a^\dagger, 1) = w(p_0) = \frac{u(a^\dagger, p_0) - (1 - \delta)u(a^*, p_0)}{\delta}. \end{cases}$$

Replacing $\lambda \bar{p} = p_0 - (1 - p_0)$ and $\lambda(1 - \bar{p}) = (1 - p_0)$ into the second equation yields

$$\begin{aligned} & \lambda u(a^*, \bar{p}) + (1 - \lambda)u(a^\dagger, 1) \\ &= \lambda u(a^*, 1) - (1 - p_0)u(a^*, 1) + (1 - p_0)u(a^*, 0) + (1 - \lambda)u(a^\dagger, 1) \\ &= u(a^*, p_0) + (1 - \lambda)[u(a^\dagger, 1) - u(a^*, 1)] \\ &= \frac{u(a^\dagger, p_0) - (1 - \delta)u(a^*, p_0)}{\delta} \\ \implies & \lambda = 1 - \frac{u(a^\dagger, p_0) - u(a^*, p_0)}{\delta[u(a^\dagger, 1) - u(a^*, 1)]} = 1 - \frac{p_0}{\delta} - \frac{1 - p_0}{\delta} \frac{\bar{p}}{1 - \bar{p}}. \end{aligned}$$

Then it follows that the principal's payoff is

$$\begin{aligned} V &= (1 - \delta)v(a^*, p_0) + \delta \lambda v(a^*, \bar{p}) \\ &= [(1 - \delta)p_0 + \delta \lambda \bar{p}]v(a^*, 1) + [(1 - \delta)(1 - p_0) + \delta \lambda(1 - \bar{p})]v(a^*, 0) \\ &= [p_0 - \delta(1 - \lambda)]v(a^*, 1) + (1 - p_0)v(a^*, 0) \\ &= v(a^*, p_0) - \delta(1 - \lambda)v(a^*, 1) \\ &= v(a^*, p_0) - \left[p_0 + (1 - p_0) \frac{\bar{p}}{1 - \bar{p}} \right] v(a^*, 1) \\ &= \frac{1 - p_0}{1 - \bar{p}} v(a^*, \bar{p}), \end{aligned}$$

which is exactly the payoff under the KG policy, which splits the initial belief p_0 into $\bar{p}$ with probability $\frac{1 - p_0}{1 - \bar{p}}$ and 1 with the complementary probability.

Let V^R be the value function under the random full-disclosure policy. To show that our policy τ is also optimal when $\frac{v(a^*, 0)}{m(0) - u(a^*, 0)} = \frac{v(a^*, 1)}{m(1) - u(a^*, 1)}$, we need to verify that

$$V^R(p, w) = \sum_{p_s \in \text{supp}(\tau)} \tau(p_s) [(1 - \delta)v(a_s, p_s) + \delta V^R(p_s, w_s)], \quad \forall (p, w) \in \mathcal{W}.$$

Note that $V^R(p, w) = \frac{M(p)-w}{M(p)-u(a^*, p)}v(a^*, p)$, since the probability of full disclosure α satisfies $\alpha M(p) + (1 - \alpha)u(a^*, p) = w$. Hence,

$$\begin{aligned} & \sum_{p_s \in \text{supp}(\tau)} \tau(p_s) [(1 - \delta)v(a_s, p_s) + \delta V^R(p_s, w_s)] \\ &= \lambda \cdot \left[(1 - \delta)v(a^*, \hat{p}) + \delta \frac{M(\hat{p}) - w(\hat{p})}{M(\hat{p}) - u(a^*, \hat{p})} v(a^*, \hat{p}) \right] \\ &= \lambda \frac{M(\hat{p}) - m(\hat{p})}{M(\hat{p}) - u(a^*, \hat{p})} v(a^*, \hat{p}), \end{aligned}$$

where $(\lambda, \hat{p})$ solves

$$\begin{cases} \lambda \hat{p} + (1 - \lambda)1 = p \\ \lambda m(\hat{p}) + (1 - \lambda)m(1) = w. \end{cases}$$

Since $\frac{v(a^*, 0)}{m(0) - u(a^*, 0)} = \frac{v(a^*, 1)}{m(1) - u(a^*, 1)}$, we have $\frac{v(a^*, \hat{p})}{M(\hat{p}) - u(a^*, \hat{p})} = \frac{v(a^*, p)}{M(p) - u(a^*, p)} = \frac{v(a^*, 1)}{M(1) - u(a^*, 1)}$. Therefore, recalling that $v(a^*, 1) = 0$, we have

$$\begin{aligned} & \lambda \frac{M(\hat{p}) - m(\hat{p})}{M(\hat{p}) - u(a^*, \hat{p})} v(a^*, \hat{p}) \\ &= \lambda \frac{M(\hat{p}) - m(\hat{p})}{M(\hat{p}) - u(a^*, \hat{p})} v(a^*, \hat{p}) + (1 - \lambda) \frac{M(1) - m(1)}{M(1) - u(a^*, 1)} v(a^*, 1) \\ &= \frac{v(a^*, p)}{M(p) - u(a^*, p)} (\lambda [M(\hat{p}) - m(\hat{p})] + (1 - \lambda) [M(1) - m(1)]) \\ &= \frac{v(a^*, p)}{M(p) - u(a^*, p)} (M(p) - w) = V^R(p, w). \end{aligned} \quad \square$$

A.5 Theorem 1

To prove Theorem 1, we first introduce the following lemma.

LEMMA 2. *Consider any feasible policy inducing the value function $\tilde{V}$. If $\tilde{V}$ is concave in both arguments, decreasing in w and satisfies*

$$\tilde{V}(p, m(p)) \geq (1 - \delta)v(a^*, p) + \delta \tilde{V}(p, \mathbf{w}(p)),$$

for all $p \in Q^1$, then the policy is optimal.

PROOF. We argue that $\tilde{V}$ is the fixed point of the operator T , hence $\tilde{V} = V^*$. Let $(\lambda_s, p_s, w_s, a_s)_{s \in S}$ be a solution to the maximization problem $T(\tilde{V})(p, w)$. We start by the following observation. Consider any s such that $a_s \neq a^*$. We have

$$(1 - \delta)v(a_s, p_s) + \delta \tilde{V}(p_s, w_s) = \delta \tilde{V}(p_s, w_s) \leq \tilde{V}(p_s, w_s) \leq \tilde{V}(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s),$$

where the last inequality follows from the fact that $\tilde{V}$ is decreasing in w and $m(p_s) \leq (1 - \delta)u(a_s, p_s) + \delta w_s \leq (1 - \delta)m(p_s) + \delta w_s \leq w_s$.

Consider now any s such that $a_s = a^*$. Since $(\lambda_s, p_s, w_s, a_s)_{s \in S}$ is feasible, we have

$$(1 - \delta)u(a^*, p_s) + \delta w_s \geq m(p_s),$$

hence $p_s \in Q^1$ and, therefore,

$$\tilde{V}(p_s, m(p_s)) \geq (1 - \delta)v(a^*, p_s) + \delta \underbrace{\tilde{V}\left(p_s, \frac{-(1 - \delta)u(a^*, p_s) + m(p_s)}{\delta}\right)}_{\mathbf{w}(p_s)}.$$

The concavity of $\tilde{V}$ implies that

$$\tilde{V}(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) - \tilde{V}(p_s, m(p_s)) \geq \delta[\tilde{V}(p_s, w_s) - \tilde{V}(p_s, \mathbf{w}(p_s))],$$

where we use the identity $(1 - \delta)u(a^*, p_s) + \delta w_s - m(p_s) = \delta(w_s - \mathbf{w}(p_s))$ and observation (a) about concave functions in Section A.1.

Combining the above two inequalities implies

$$\tilde{V}(p_s, (1 - \delta)u(a^*, p_s) + \delta w_s) \geq (1 - \delta)v(a^*, p_s) + \delta \tilde{V}(p_s, w_s).$$

It follows that

$$\begin{aligned} T(\tilde{V})(p, w) &= \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta \tilde{V}(p_s, w_s)] \\ &\leq \sum_{s \in S} \lambda_s [\tilde{V}(p_s, (1 - \delta)u(a_s, p_s) + \delta w_s)] \\ &\leq \tilde{V}\left(\sum_{s \in S} \lambda_s p_s, \sum_{s \in S} \lambda_s ((1 - \delta)u(a_s, p_s) + \delta w_s)\right) \\ &\leq \tilde{V}(p, w), \end{aligned}$$

where the second inequality follows from the concavity of $\tilde{V}$ and the third inequality from $\tilde{V}$ being decreasing in w .

Conversely, since the policy inducing $\tilde{V}$ is feasible, we must have that $T(\tilde{V})(p, w) \geq \tilde{V}(p, w)$ for all (p, w) . This completes the proof. $\square$

Invoking Lemma 2, we only need to prove the following proposition to prove Theorem 1.

PROPOSITION 4. *Let V_{q^*} be the value function induced by the policy τ^* , with*

$$q^* = \sup\{p \in Q^1 : V_{q^1}(p, m(p)) \geq V_{q^1}(p, w) \text{ for all } w\}.$$

Then V_{q^} is concave in (p, w) , decreasing in w , and satisfies*

$$V_{q^*}(p, m(p)) \geq (1 - \delta)v(a^*, p) + \delta V_{q^*}(p^*, \mathbf{w}(p)),$$

for all $p \in Q^1$.

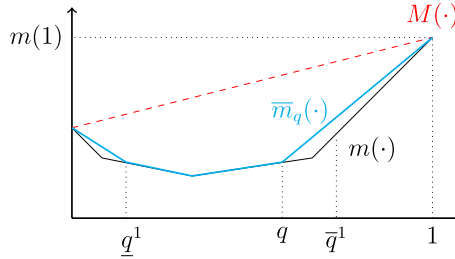


FIGURE 8. The function $\bar{m}_q$.

Proving Proposition 4 requires to construct the value function V_q induced by the policy τ_q . The construction is tedious, and we postpone it to Appendix B. In the rest of this section, we only report the properties we need to prove Proposition 4.

We start with an important identity, which we will use throughout. For any $q \in [\underline{q}^1, \bar{q}^1]$, define the function $\bar{m}_q : [0, 1] \rightarrow \mathbb{R}$ as

$$\bar{m}_q(p) = \begin{cases} \left(1 - \frac{p}{\underline{q}^1}\right)m(0) + \frac{p}{\underline{q}^1}m(\underline{q}^1) & \text{if } p \in [0, \underline{q}^1], \\ m(p) & \text{if } p \in (\underline{q}^1, q], \\ \frac{1-p}{1-q}m(q) + \frac{p-q}{1-q}m(1) & \text{if } p \in (q, 1]. \end{cases}$$

Note that $\bar{m}_q$ is convex, $\bar{m}_q(p) \geq m(p)$ for all $p \in [0, 1]$, $\bar{m}_q(0) = m(0)$, and $\bar{m}_q(1) = m(1)$. For a graphical illustration, see Figure 8.

It is straightforward to check that we have the following identity:

$$V_q(p, w) = \bar{\lambda}(p, w)V_q(\bar{\varphi}(p, w), \bar{m}_q(\bar{\varphi}(p, w))), \quad (9)$$

where the functions $\bar{\lambda}$ and $\bar{\varphi}$ are defined as in the main text, but with $\bar{m}_q$ instead of m ; see equation (5). This identity states that knowing V_q on the set $\{(p, w) \in \mathcal{W} : (p, w) = (p, \bar{m}_q(p))\}$ suffices to reconstruct V_q at all points on its domain. We now make two additional observations.

OBSERVATION A. For all $q \in [\underline{q}^1, \bar{q}^1]$, we have the following identity:

$$V_q(p, w) = \frac{1-p}{1-p'}V_q\left(p', \frac{1-p'}{1-p}w + \frac{p'-p}{1-p}\bar{m}_q(1)\right).$$

PROOF OF OBSERVATION A. Let $w' = \frac{1-p'}{1-p}w + \frac{p'-p}{1-p}\bar{m}_q(1)$.

Assume that $w' > \bar{m}_q(p')$. Since

$$\bar{\lambda}(p', w') \left(\frac{\bar{\varphi}(p', w')}{\bar{m}_q(\bar{\varphi}(p', w'))} \right) + (1 - \bar{\lambda}(p', w')) \left(\frac{1}{\bar{m}_q(1)} \right) = \left(\frac{p'}{w'} \right),$$

we have

$$\frac{1-p}{1-p'} \bar{\lambda}(p', w') \left(\frac{\bar{\varphi}(p', w')}{\bar{m}_q(\bar{\varphi}(p', w'))} \right) + \left(1 - \frac{1-p}{1-p'} \bar{\lambda}(p', w') \right) \left(\frac{1}{\bar{m}_q(1)} \right) = \left(\frac{p}{w} \right).$$

Therefore, $\bar{\lambda}(p, w) = \frac{1-p}{1-p'} \bar{\lambda}(p', w')$ and $\bar{\varphi}(p', w') = \bar{\varphi}(p, w)$ since the solution $(\bar{\lambda}(p', w'), \bar{\varphi}(p', w'))$ is unique when $w' > m_q(p')$. The statement then follows from equation (9).

Assume that $w' = \bar{m}_q(p')$. From the convexity of $\bar{m}_q$, this requires that $w = \bar{m}_q(p)$, so that $\bar{m}_q(p') = \frac{1-p'}{1-p} \bar{m}_q(p) + \frac{p'-p}{1-p} \bar{m}_q(1)$. The result follows from continuity as

$$\begin{aligned} V_q(p, \bar{m}_q(p)) &= \lim_{w \rightarrow \bar{m}_q(p)} V_q(p, w), \\ &= \lim_{w \rightarrow \bar{m}_q(p)} \frac{1-p}{1-p'} V_q\left(p', \frac{1-p'}{1-p} w + \frac{p'-p}{1-p} \bar{m}_q(1)\right), \\ &= \frac{1-p}{1-p'} V_q\left(p', \frac{1-p'}{1-p} \bar{m}_q(p) + \frac{p'-p}{1-p} \bar{m}_q(1)\right), \\ &= \frac{1-p}{1-p'} V_q(p', \bar{m}_q(p')). \end{aligned}$$

Note that this implies that

$$V_q(p, \mathbf{w}(p) + c) = \bar{\lambda}(p, \mathbf{w}(p)) V_q\left(\bar{\varphi}(p, \mathbf{w}(p)), \bar{m}_q(\bar{\varphi}(p, \mathbf{w}(p))) + \frac{c}{\bar{\lambda}(p, \mathbf{w}(p))}\right),$$

where c is a positive constant. □

OBSERVATION B. The value function $V_{\bar{q}^1}(p, \cdot) : [\bar{m}_{\bar{q}^1}(p), M(p)] \rightarrow \mathbb{R}$ is concave in w , for each p . See Lemma 3 in Section B.2.

A.5.1 Proposition 4(a) We prove that V_{q^*} is decreasing in w . To start with, fix $p \in [0, 1]$ and $(w, w') \in [\bar{m}_{q^*}(p), M(p)] \times [\bar{m}_{q^*}(p), M(p)]$, with $w' > w$.

First, assume that $p \leq q^*$. If $w = \bar{m}_{q^*}(p)$, then $V_{q^*}(p, w') \leq V_{q^*}(p, w)$ by construction of q^* . If $w > \bar{m}_{q^*}(p)$, we have that

$$\begin{aligned} \frac{V_{q^*}(p, w') - V_{q^*}(p, w)}{w' - w} &= \frac{V_{\bar{q}^1}(p, w') - V_{\bar{q}^1}(p, w)}{w' - w} \\ &\leq \frac{V_{\bar{q}^1}(p, w) - V_{\bar{q}^1}(p, \bar{m}_{q^*}(p))}{w - \bar{m}_{q^*}(p)} \\ &= \frac{V_{q^*}(p, w) - V_{q^*}(p, \bar{m}_{q^*}(p))}{w - \bar{m}_{q^*}(p)} \leq 0, \end{aligned}$$

where the inequality follows from the concavity of $V_{\bar{q}^1}$ with respect to w , for all $w \geq m_{\bar{q}^1}(p)$. (Recall that $m_{q^*}(p) = \bar{m}_{\bar{q}^1}(p)$ for all $p \leq q^*$.)

Second, assume that $p > q^*$. We show in detail how to make use of Observation A to deduce the result. We repeatedly use similar computations later on. We have

$$\begin{aligned} V_{q^*}(p, w') &= \bar{\lambda}(p, w') V_{q^*}(\bar{\varphi}(p, w'), \bar{m}_{q^*}(\bar{\varphi}(p, w'))) \\ &= \bar{\lambda}(p, w') \frac{1 - \bar{\varphi}(p, w')}{1 - \bar{\varphi}(p, w)} V_{q^*}\left(\bar{\varphi}(p, w), \frac{1 - \bar{\varphi}(p, w)}{1 - \bar{\varphi}(p, w')} \bar{m}_{q^*}(\bar{\varphi}(p, w'))\right) \end{aligned}$$

$$\begin{aligned}
 & + \left(1 - \frac{1 - \bar{\varphi}(p, w)}{1 - \bar{\varphi}(p, w')}\right) \bar{m}_{q^*}(1) \\
 & = \bar{\lambda}(p, w) V_{q^*} \left(\bar{\varphi}(p, w), \frac{\bar{\lambda}(p, w')}{\bar{\lambda}(p, w)} \bar{m}_{q^*}(\bar{\varphi}(p, w')) + \left(1 - \frac{\bar{\lambda}(p, w')}{\bar{\lambda}(p, w)}\right) \bar{m}_{q^*}(1) \right) \\
 & = \bar{\lambda}(p, w) V_{q^*} \left(\bar{\varphi}(p, w), \bar{m}_{q^*}(\bar{\varphi}(p, w)) + \frac{w' - w}{\bar{\lambda}(p, w)} \right),
 \end{aligned}$$

where the first line follows from the construction of V_{q^*} , the second line from Observation A, the third line from the definition of the functions $\bar{\lambda}$ and $\bar{\varphi}$, and the last line from the following computations:

$$\begin{aligned}
 & \frac{\bar{\lambda}(p, w')}{\bar{\lambda}(p, w)} \bar{m}_{q^*}(\bar{\varphi}(p, w')) + \left(1 - \frac{\bar{\lambda}(p, w')}{\bar{\lambda}(p, w)}\right) \bar{m}_{q^*}(1) \\
 & = \frac{1}{\bar{\lambda}(p, w)} w' + \left(1 - \frac{1}{\bar{\lambda}(p, w)}\right) \bar{m}_{q^*}(1) \\
 & = \frac{1}{\bar{\lambda}(p, w)} w' + \left(1 - \frac{1}{\bar{\lambda}(p, w)}\right) \left[\frac{w - \bar{\lambda}(p, w) \bar{m}_{q^*}(\bar{\varphi}(p, w))}{1 - \bar{\lambda}(p, w)} \right] \\
 & = \bar{m}_{q^*}(\bar{\varphi}(p, w)) + \frac{w' - w}{\bar{\lambda}(p, w)}.
 \end{aligned}$$

Thus, we are able to express $V_{q^*}(p, w')$ as $\bar{\lambda}(p, w) V_{q^*}(\bar{\varphi}(p, w), \tilde{w})$, with $\tilde{w}$ the above expression. Moreover, $\bar{\varphi}(p, w) \leq q^*$ as $w \geq \bar{m}_{q^*}(p)$. We can use the (already established) concavity of V_{q^*} in w for each $p \leq q^*$ to deduce the desired result. More precisely, we have that

$$\begin{aligned}
 & \frac{V_{q^*}(p, w') - V_{q^*}(p, w)}{w' - w} \\
 & = \frac{\bar{\lambda}(p, w) \left(V_{q^*} \left(\bar{\varphi}(p, w), \bar{m}_{q^*}(\bar{\varphi}(p, w)) + \frac{w' - w}{\bar{\lambda}(p, w)} \right) - V_{q^*}(\bar{\varphi}(p, w), \bar{m}_{q^*}(\bar{\varphi}(p, w))) \right)}{w' - w} \\
 & \leq 0,
 \end{aligned}$$

where the inequality follows from the concavity of V_{q^*} in w at all $p \leq q^*$.

Lastly, since $V_{q^*}(p, w) = V_{q^*}(p, \bar{m}_{q^*}(p))$ for all $w \in [m(p), \bar{m}_{q^*}(p)]$, the result immediately follows for all (w, w') , with $w \in [m(p), \bar{m}_{q^*}(p)]$.

A.5.2 Proposition 4(b) We prove the concavity of V_{q^*} with respect to both arguments (p, w) .

Let $\bar{\mathcal{W}} = \{(p, w) : w \geq \bar{m}_{q^*}(p)\}$. Let $(p, w) \in \bar{\mathcal{W}}$, $(p', w') \in \bar{\mathcal{W}}$ and $\alpha \in [0, 1]$. Write (p_α, w_α) for

$$\alpha \begin{pmatrix} p \\ w \end{pmatrix} + (1 - \alpha) \begin{pmatrix} p' \\ w' \end{pmatrix}.$$

Without loss of generality, assume that $p \leq p'$. We have that

$$\begin{aligned}
 & \alpha V_{q^*}(p, w) + (1 - \alpha) V_{q^*}(p', w') \\
 &= \alpha \frac{1-p}{1-p'} V_{q^*} \left(p', \underbrace{\frac{1-p'}{1-p} w + \frac{p'-p}{1-p} \bar{m}_{q^*}(1)}_{\geq \bar{m}_{q^*}(p')} \right) + (1 - \alpha) V_{q^*}(p', w') \\
 &\leq \left(\alpha \frac{1-p}{1-p'} + (1 - \alpha) \right) V_{q^*} \left(p', \frac{\alpha \frac{1-p}{1-p'} \left(\frac{1-p'}{1-p} w + \frac{p'-p}{1-p} \bar{m}_{q^*}(1) \right) + (1 - \alpha) w'}{\alpha \frac{1-p}{1-p'} + (1 - \alpha)} \right) \\
 &= \frac{1-p_\alpha}{1-p'} V_{q^*} \left(p', \frac{1-p'}{1-p_\alpha} w_\alpha + \frac{p'-p_\alpha}{1-p_\alpha} \bar{m}_{q^*}(1) \right) \\
 &= V_{q^*}(p_\alpha, w_\alpha),
 \end{aligned}$$

where the inequality follows from the concavity of $V_{\bar{q}^1}$ with respect to w for each p and the property that $V_{q^*}(p, w) = V_{\bar{q}^1}(p, w)$ for all (p, w) such that $w \geq \bar{m}_{q^*}(p)$. Notice that we use twice Observation A.

Finally, for all $(p, w) \in \mathcal{W}$, for all $(p', w') \in \mathcal{W}$ and for all α , we have that

$$\begin{aligned}
 & \alpha V_{q^*}(p, w) + (1 - \alpha) V_{q^*}(p', w') \\
 &= \alpha V_{q^*}(p, \max(w, \bar{m}_{q^*}(p))) + (1 - \alpha) V_{q^*}(p', \max(w', \bar{m}_{q^*}(p'))) \\
 &\leq V_{q^*}(p_\alpha, \alpha \max(w, \bar{m}_{q^*}(p)) + (1 - \alpha) \max(w, \bar{m}_{q^*}(p'))) \\
 &\leq V_{q^*}(p_\alpha, w_\alpha),
 \end{aligned}$$

since $\alpha \max(w, \bar{m}_{q^*}(p)) + (1 - \alpha) \max(w, \bar{m}_{q^*}(p')) \geq w_\alpha$ and the fact that V_{q^*} is decreasing in w for all p . This completes the proof of concavity.

A.5.3 Proposition 4(c) We prove that $V_{q^*}(p, m(p)) \geq (1 - \delta)v(a^*, p) + \delta V_{q^*}(p, \mathbf{w}(p))$ for all $p \in Q^1$.

The statement is true for all $p \leq q^*$ by definition since $V_{q^*}(p, w) = V_{\bar{q}^1}(p, w)$ for all w .

Assume that $p > q^*$. From Lemma 4, there exists $\bar{q}$ such that $\varphi(p, \mathbf{w}(p)) \geq \varphi(p', \mathbf{w}(p'))$ for all $p' \geq p \geq \bar{q}$. Moreover, it follows from A.6.3 and A.6.4 that $V(p, m(p)) \geq V(p, w)$ for all w , for all $p \leq \bar{q}$. Therefore, we must have that $q^* \geq \bar{q}$. It follows that $\varphi(p, \mathbf{w}(p)) < \varphi(q^*, \mathbf{w}(q^*)) \leq q^*$, hence $\mathbf{w}(p) \geq \bar{m}_{q^*}(p)$. We therefore have that $V_{q^*}(p, \mathbf{w}(p)) = V_{\bar{q}^1}(p, \mathbf{w}(p))$.

Since $V_{\bar{q}^1}(p, m(p)) = (1 - \delta)v(a^*, p) + \delta V_{\bar{q}^1}(p, \mathbf{w}(p))$ for all $p \in Q^1$ and $V_{q^*}(p, m(p)) = V_{q^*}(p, \bar{m}_{q^*}(p)) = V_{\bar{q}^1}(p, \bar{m}_{q^*}(p))$, it is enough to prove that $V_{\bar{q}^1}(p, \bar{m}_{q^*}(p)) \geq V_{\bar{q}^1}(p, m(p))$.

Clearly, there is nothing to prove if $\bar{m}_{q^*}(p) = m(p)$ for all $p \in Q^1$, that is, if $q^* = \bar{q}^1$ (remember that $\bar{m}_{\bar{q}^1}(p) = m(p)$ for all $p \in Q^1$).

So, assume that $\bar{m}_{q^*}(p) > m(p)$ for some $p \in (q^*, \bar{q}^1)$, hence $\bar{m}_{q^*}(p) > m(p)$ for all $p \in (q^*, \bar{q}^1)$. We now argue that if $V_{\bar{q}^1}(p, w) > V_{\bar{q}^1}(p, m(p))$ for some $w \geq \bar{m}_{q^*}(p)$, then

$$V_{\bar{q}^1}(p', m(p')) < \frac{1-p'}{1-p} V_{\bar{q}^1}(p, w),$$

for all $p' > p$. To see this, observe that $w > m(p)$ and, accordingly,

$$\frac{1-p'}{1-p}w + \frac{p'-p}{1-p}m(1) - m(p') > 0,$$

since m is convex. Hence,

$$\begin{aligned} 0 &< \frac{V_{\bar{q}^1}(p, w) - V_{\bar{q}^1}(p, m(p))}{w - m(p)} \\ &= \frac{\frac{1-p}{1-p'} \left[V_{\bar{q}^1} \left(p', \frac{1-p'}{1-p}w + \frac{p'-p}{1-p}m(1) \right) - V_{\bar{q}^1} \left(p', \frac{1-p'}{1-p}m(p) + \frac{p'-p}{1-p}m(1) \right) \right]}{w - m(p)} \\ &\leq \frac{V_{\bar{q}^1} \left(p', \frac{1-p'}{1-p}w + \frac{p'-p}{1-p}m(1) \right) - V_{\bar{q}^1}(p', m(p'))}{\frac{1-p'}{1-p}w + \frac{p'-p}{1-p}m(1) - m(p')}, \end{aligned}$$

where the equality follows Observation A and the inequality from the concavity of $V_{\bar{q}^1}$ in w for each p . Since

$$V_{\bar{q}^1}(p, w) = \frac{1-p}{1-p'} V_{\bar{q}^1} \left(p', \frac{1-p'}{1-p}w + \frac{p'-p}{1-p}m(1) \right),$$

we have the desired result.

Finally, from the definition of q^* , for all $n > 0$, there exist $p_n \in (q^*, \min(q^* + \frac{1}{n}, \bar{q}^1)]$ and $w_n \geq m(p_n)$ such that $V_{\bar{q}^1}(p_n, m(p_n)) < V_{\bar{q}^1}(p_n, w_n)$. From the concavity of $V_{\bar{q}^1}$ in w for all p , $V_{\bar{q}^1}(p_n, m(p_n)) < V_{\bar{q}^1}(p_n, \bar{m}_{q^*}(p_n))$ for all n .

From the above argument, for all p , for all n sufficiently large, that is, such that $p_n < p$, we have that

$$V_{\bar{q}^1}(p, m(p)) < \frac{1-p}{1-p_n} V_{\bar{q}^1}(p_n, \bar{m}_{q^*}(p_n)).$$

Taking the limit as $n \rightarrow \infty$, we obtain that

$$V_{\bar{q}^1}(p, m(p)) < \frac{1-p}{1-q^*} V_{\bar{q}^1}(q^*, \bar{m}_{q^*}(q^*)) = V_{\bar{q}^1}(p, \bar{m}_{q^*}(p)),$$

which completes the proof.

APPENDIX B: CONSTRUCTING THE VALUE FUNCTION

This section characterizes the value function V_q induced by the policy τ_q . As explained in the text, it suffices to characterize $V_{\bar{q}^1}$ since $V_q(p, w) = V_{\bar{q}^1}(p, w)$ for all $(p, w) \in \mathcal{W} \setminus \mathcal{W}_q^3$

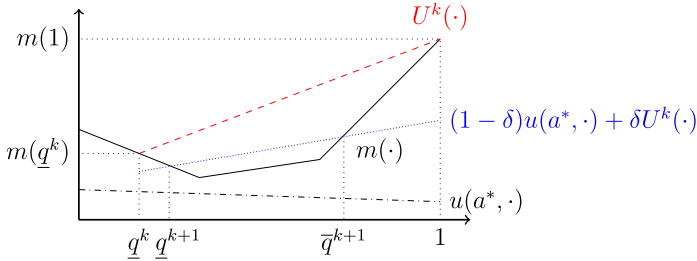


FIGURE 9. Construction of the thresholds.

and $V_q(p, w) = \frac{1-p}{1-q} V_{\bar{q}^1}(q, m(q))$ for all $(p, w) \in \mathcal{W}_q^3$. We first start with the definition of important subsets of $[0, 1]$.

B.1 Construction of the sets Q^k

Let $Q^0 := [0, 1]$. We define inductively the set $Q^k \subseteq [0, 1]$, $k \geq 0$. We write $\underline{q}^k$ (resp., $\bar{q}^k$) for $\inf Q^k$ (resp., $\sup Q^k$). For any $k \geq 0$, define the function $U^k : [\underline{q}^k, 1] \rightarrow \mathbb{R}$:

$$U^k(q) := \frac{1-q}{1-\underline{q}^k} m(\underline{q}^k) + \frac{q-\underline{q}^k}{1-\underline{q}^k} m(1),$$

with the convention that $U^k \equiv m(1)$ if $\underline{q}^k = 1$. Note that $U^0(q) = M(q)$ and $U^k(q) \geq m(q)$ for all k . We define Q^{k+1} as follows:

$$Q^{k+1} = \{q \in Q^k : (1-\delta)u(a^*, q) + \delta U^k(q) \geq m(q)\}.$$

For a graphical illustration, see Figure 9.

Few observations are worth making. First, we have that $P \subseteq Q^k$ for all k . Second, we have a decreasing sequence, that is, $Q^{k+1} \subseteq Q^k$ for all k . Third, if Q^k and P are nonempty, then they are closed intervals. Fourth, the limit $Q^\infty = \lim_{k \rightarrow \infty} Q^k = \bigcap_k Q^k$ exists and includes P . Moreover, if $P \neq \emptyset$, then $\underline{q}^\infty = \underline{p}$, where $\underline{p} := \inf P$. If $P = \emptyset$, then $Q^\infty = \emptyset$. Consequently, there exists $k^* < \infty$ such that $\emptyset = Q^{k^*+1} \subset Q^{k^*} \neq \emptyset$.

The first to the third observations are readily proved, so we concentrate on the proof of the fourth observation. The limit exists as we have a decreasing sequence of sets.

We prove that if $P = \emptyset$, then $Q^\infty = \emptyset$. So, assume that $P = \emptyset$. We first argue that it cannot be that $Q^k = Q^{k-1} \neq \emptyset$ for some $k \geq 0$. To the contrary, assume that $Q^k = Q^{k-1} \neq \emptyset$ for some $k \geq 0$, hence $Q^{k'} = Q^{k-1}$ for all $k' \geq k$. From the convexity and continuity of m and the linearity of u , Q^{k-1} is the closed interval $[\underline{q}^{k-1}, \bar{q}^{k-1}]$, with the two boundary points solution to

$$(1-\delta)u(a^*, q) + \delta U^{k-2}(q) = m(q).$$

Therefore, if $(\underline{q}^k, \bar{q}^k) = (\underline{q}^{k-1}, \bar{q}^{k-1})$, we have that

$$m(\underline{q}^{k-1}) = (1-\delta)u(a^*, \underline{q}^{k-1}) + \delta m(\underline{q}^{k-1}),$$

$$\begin{aligned}
m(\bar{q}^{k-1}) &= (1 - \delta)u(a^*, \bar{q}^{k-1}) + \delta \left[\frac{1 - \bar{q}^{k-1}}{1 - \underline{q}^{k-1}} m(\underline{q}^{k-1}) + \frac{\bar{q}^{k-1} - \underline{q}^{k-1}}{1 - \underline{q}^{k-1}} m(1) \right], \\
&\leq (1 - \delta)u(a^*, \bar{q}^{k-1}) + \delta m(\bar{q}^{k-1}).
\end{aligned}$$

This implies that $u(a^*, \underline{q}^{k-1}) = m(\underline{q}^{k-1})$ and $u(a^*, \bar{q}^{k-1}) = m(\bar{q}^{k-1})$ and, therefore, $\emptyset \neq Q^{k-1} \subseteq P$, a contradiction.

We thus have an infinite sequence of strictly decreasing nonempty closed intervals. Let $\varepsilon := \min_{p \in [0, 1]} m(p) - u(a^*, p)$. Since $P = \emptyset$, we have that $\varepsilon > 0$. For all $p \in Q^\infty$, for all k ,

$$\begin{aligned}
m(p) &\leq (1 - \delta)u(a^*, p) + \delta U^k(p) \\
&\leq (1 - \delta)(m(p) - \varepsilon) + \delta U^k(p).
\end{aligned}$$

Assume that Q^∞ is nonempty and let $\underline{q}^\infty$ its greatest lower bound. Since $\underline{q}^\infty \in Q^k$ for all k , we have that $U^k(\underline{q}^\infty) \geq m(\underline{q}^\infty) + \varepsilon(1 - \delta)/\delta$ for all k . Since $\lim_{k \rightarrow \infty} U^k(\underline{q}^\infty) = m(\underline{q}^\infty)$, we have that $m(\underline{q}^\infty) \geq m(\underline{q}^\infty) + \varepsilon(1 - \delta)/\delta$, a contradiction.

We now prove that if $P \neq \emptyset$, then $\underline{q}^\infty = \underline{p}$. From above, we have that if $Q^k = Q^{k-1} \neq \emptyset$ for some $k \geq 0$, hence $Q^{k'} = Q^{k-1}$ for all $k' \geq k$, then $P = Q^k$ since $P \subseteq Q^k$. If we have an infinite sequence of strictly decreasing sets, for all $q \in Q^\infty$,

$$(1 - \delta)u(a^*, q) + \delta \left[\frac{1 - q}{1 - \underline{q}^\infty} m(\underline{q}^\infty) + \frac{q - \underline{q}^\infty}{1 - \underline{q}^\infty} m(1) \right] \geq m(q).$$

Taking the limit $q \downarrow \underline{q}^\infty$, we obtain that $u(a^*, \underline{q}^\infty) = m(\underline{q}^\infty)$, that is, $\underline{q}^\infty \in P$. Hence, $\underline{q}^\infty = \underline{p}$.

B.1.1 Derivation of $V_{\bar{q}^1}$ We first derive $V_{\bar{q}^1}$ for all $(p, w) \in \mathcal{W} \setminus \mathcal{W}_{\bar{q}^1}^2$.

To start with, $V_{\bar{q}^1}(1, m(1)) = 0$ since a^* is not optimal at $p = 1$. Similarly, $V_{\bar{q}^1}(0, m(0)) = 0$ if a^* is not optimal at $p = 0$, while $V_{\bar{q}^1}(0, m(0)) = v(a^*, 0)$ if a^* is optimal at $p = 0$. Also, $V_{\bar{q}^1}(\underline{q}^1, m(\underline{q}^1)) = (1 - \delta)v(a^*, \underline{q}^1)$ if $\underline{q}^1 > 0$; while $V_{\bar{q}^1}(0, m(0)) = v(a^*, 0)$ if $\underline{q}^1 = 0$, since a^* is then optimal at $p = 0$.

With the function $V_{\bar{q}^1}$ defined at these three points, it is then defined at all points (p, w) in $\mathcal{W}_{\bar{q}^1}^1 \cup \mathcal{W}_{\bar{q}^1}^4$. In particular, it is easy to show that

$$V_{\bar{q}^1}(\underline{q}^1, w) = \frac{M(\underline{q}^1) - w}{M(\underline{q}^1) - m(\underline{q}^1)} (1 - \delta)v(a^*, \underline{q}^1) = \frac{M(\underline{q}^1) - w}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} v(a^*, \underline{q}^1),$$

for all $w \in [m(\underline{q}^1), M(\underline{q}^1)]$.

At all points $(p, w) \in \mathcal{W}_{\bar{q}^1}^3$,

$$V_{\bar{q}^1}(p, w) = \frac{1 - p}{1 - \bar{q}^1} V_{\bar{q}^1}(\bar{q}^1, m(\bar{q}^1)).$$

Therefore, $V_{\bar{q}^1}$ is well-defined at all $(p, w) \in \mathcal{W} \setminus \mathcal{W}_{\bar{q}^1}^2$.

At all points $(p, w) \in \mathcal{W}_{\bar{q}^1}^2$, $V_{\bar{q}^1}(p, w)$ is defined via the recursive equation:

$$\begin{aligned} V_{\bar{q}^1}(p, w) &= \lambda(p, w)[(1 - \delta)v(a^*, \varphi(p, w)) + \delta V_{\bar{q}^1}(\varphi(p, w), \mathbf{w}(\varphi(p, w)))] \\ &= \lambda(p, w)V_{\bar{q}^1}(\varphi(p, w), m(\varphi(p, w))). \end{aligned}$$

Since $V_{\bar{q}^1}(p, w) = \lambda(p, w)V_{\bar{q}^1}(\varphi(p, w), m(\varphi(p, w)))$, the value function is well-defined at all (p, w) if it is well-defined at all $(p, m(p))$, which we now prove.

By construction of the sets Q^k , observe that if $p \in Q^k \setminus Q^{k+1}$, then $\mathbf{w}(p) \in (U^k(p), U^{k+1}(p))$ and, therefore, $\varphi(p, \mathbf{w}(p)) \in [\underline{q}^{k-1}, \underline{q}^k] \subset Q^{k-1} \setminus Q^k$. Moreover, $\varphi(\bar{q}^k, \mathbf{w}(\bar{q}^k)) = \underline{q}^k$. We now use these observations to complete the derivation of $V_{\bar{q}^1}$.

For all $p \in Q^1 \setminus Q^2$, we have that $\mathbf{w}(p) \in Q^0 \setminus Q^1$, so that $(p, \mathbf{w}(p)) \in \mathcal{W}_{\bar{q}^1}^4$. Since

$$V_{\bar{q}^1}(p, m(p)) = (1 - \delta)v(a^*, p) + \delta V_{\bar{q}^1}(p, \mathbf{w}(p)),$$

$V_{\bar{q}^1}(p, m(p))$ is well-defined for all $p \in Q^1 \setminus Q^2$. By induction, assume that it is well-defined for all $p \in \bigcup_{\ell < k} Q^\ell \setminus Q^{\ell+1}$. We argue that it is well-defined for all $p \in Q^k \setminus Q^{k+1}$. Fix any $p \in Q^k \setminus Q^{k+1}$. From our initial observation, $\varphi(p, \mathbf{w}(p)) \in [\underline{q}^{k-1}, \underline{q}^k]$ and, therefore, $V_{\bar{q}^1}(p, m(p))$ is well-defined since

$$\begin{aligned} V_{\bar{q}^1}(p, m(p)) &= (1 - \delta)v(a^*, p) + \delta V_{\bar{q}^1}(p, \mathbf{w}(p)) \\ &= (1 - \delta)v(a^*, p) + \lambda(p, \mathbf{w}(p)) \underbrace{V_{\bar{q}^1}(\varphi(p, \mathbf{w}(p)), m(\varphi(p, \mathbf{w}(p))))}_{\text{defined by the induction step}}. \end{aligned}$$

Therefore, $V_{\bar{q}^1}(p, m(p))$ is well-defined for all $p \in \bigcup_{\ell} Q^\ell \setminus Q^{\ell+1} = Q^1 \setminus Q^\infty$. It remains to argue that it is well-defined for all $p \in Q^\infty$.

From the definition of Q^∞ , we have that $\mathbf{w}(p) \leq \frac{1-p}{1-\underline{q}^\infty}m(\underline{q}^\infty) + \frac{p-\underline{q}^\infty}{1-\underline{q}^\infty}m(1)$ and, therefore, $\varphi(p, \mathbf{w}(p)) \in Q^\infty$. In other words, if $p \in Q^\infty$, then $\varphi(p, \mathbf{w}(p)) \in Q^\infty$, so that the restriction of $V_{\bar{q}^1}(\cdot, m(\cdot))$ to Q^∞ is entirely defined by its value on Q^∞ via the contraction:

$$V_{\bar{q}^1}(p, m(p)) = (1 - \delta)v(a^*, p) + \delta \lambda(p, \mathbf{w}(p))V_{\bar{q}^1}(\varphi(p, \mathbf{w}(p)), m(\varphi(p, \mathbf{w}(p)))).$$

The unique solution to this fixed point problem is given by

$$V_{\bar{q}^1}(p, m(p)) = v(a^*, p) - \frac{m(p) - u(a^*, p)}{m(1) - u(a^*, 1)}v(a^*, 1),$$

for all $p \in Q^\infty$. To see this, with a slight abuse of notation, write (λ, φ) for $(\lambda(p, w), \varphi(p, \mathbf{w}(p)))$, and note that

$$\begin{aligned} &(1 - \delta)v(a^*, p) + \delta \lambda \left[v(a^*, \varphi) - \frac{m(\varphi) - u(a^*, \varphi)}{m(1) - u(a^*, 1)}v(a^*, 1) \right] \\ &= (1 - \delta)v(a^*, p) + \delta [v(a^*, p) - (1 - \lambda)v(a^*, 1)] \end{aligned}$$

$$\begin{aligned}
 & - \frac{m(p) - (1 - \lambda)m(1) - u(a^*, p)(1 - \delta)}{m(1) - u(a^*, 1)} v(a^*, 1) \\
 & + \delta \frac{u(a^*, p) - (1 - \lambda)u(a^*, 1)}{m(1) - u(a^*, 1)} v(a^*, 1) \\
 & = v(a^*, p) - \frac{m(p) - u(a^*, p)}{m(1) - u(a^*, 1)} v(a^*, 1),
 \end{aligned}$$

where we use the identities $\lambda\varphi + (1 - \lambda)1 = p$, $\lambda m(\varphi) + (1 - \lambda)m(1) = \mathbf{w}(p)$, and $\delta\mathbf{w}(p) = m(p) - (1 - \delta)u(a^*, p)$.

This completes the characterization of $V_{\bar{q}^1}$. Note that $V_{\bar{q}^1}$ and, therefore, all value functions $V_{\bar{q}}$, are continuous functions.

B.2 Concavity of $V_{\bar{q}^1}$ with respect to w for each p

LEMMA 3. For all p , the function $V_{\bar{q}^1}(p, \cdot) : [\bar{m}_{\bar{q}^1}(p), M(p)] \rightarrow \mathbb{R}$ is concave in w .

We must prove that

$$\begin{aligned}
 & \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p) + \eta(\bar{m}_{\bar{q}^1}(1) - u(a^*, 1))) - V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p))}{\eta} \\
 & \geq \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p) + \eta'(\bar{m}_{\bar{q}^1}(1) - u(a^*, 1))) - V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p))}{\eta'},
 \end{aligned}$$

for all (η, η') such that $\eta' \geq \eta$. (See the observations on concave functions.) We start with some preliminary results.

B.2.1 Two preliminary results

LEMMA 4. There exists a nonempty interval $[q, \bar{q}]$ such that

- (1) For any $p' < p \leq q$ or $p' > p \geq \bar{q}$, $\varphi(p, \mathbf{w}(p)) \geq \varphi(p', \mathbf{w}(p'))$.
- (2) The ratio $\frac{m(1) - m(\varphi(p, \mathbf{w}(p)))}{1 - \varphi(p, \mathbf{w}(p))}$ is constant for all $p \in [q, \bar{q}]$.

PROOF. Observe that

$$\frac{m(1) - \mathbf{w}(p)}{1 - p} = \frac{m(1) - m(\varphi(p, \mathbf{w}(p)))}{1 - \varphi(p, \mathbf{w}(p))}.$$

Therefore, the convexity of m implies that if $\frac{m(1) - \mathbf{w}(p)}{1 - p} < \frac{m(1) - \mathbf{w}(p')}{1 - p'}$, then $\varphi(p, \mathbf{w}(p)) < \varphi(p', \mathbf{w}(p'))$.

Consider the function $h : [0, 1] \rightarrow \mathbb{R}$, defined by $h(p) = \frac{m(1) - \mathbf{w}(p)}{1-p}$. We argue that h is quasi-concave. For all (p, p') and $\alpha \in [0, 1]$, we have that

$$\begin{aligned} \frac{m(1) - \mathbf{w}(\alpha p + (1-\alpha)p')}{\alpha(1-p) + (1-\alpha)(1-p')} &\geq \frac{\alpha(m(1) - \mathbf{w}(p)) + (1-\alpha)(m(1) - \mathbf{w}(p'))}{\alpha(1-p) + (1-\alpha)(1-p')} \\ &= \frac{\alpha(1-p)}{\alpha(1-p) + (1-\alpha)(1-p')} \frac{m(1) - \mathbf{w}(p)}{1-p} \\ &\quad + \frac{(1-\alpha)(1-p')}{\alpha(1-p) + (1-\alpha)(1-p')} \frac{m(1) - \mathbf{w}(p')}{1-p'} \\ &\geq \min\left(\frac{m(1) - \mathbf{w}(p)}{1-p}, \frac{m(1) - \mathbf{w}(p')}{1-p'}\right), \end{aligned}$$

where the first inequality follows from the convexity of $\mathbf{w}$. (Note that the inequality is strict if $\mathbf{w}(\alpha p + (1-\alpha)p') < \alpha\mathbf{w}(p) + (1-\alpha)\mathbf{w}(p')$.)

It follows that if $h(p') \geq h(p)$, then it is also true for all $p'' \in (p, p')$. Since h is quasi-concave and continuous, the set of maxima is a nonempty convex set $[\underline{q}, \bar{q}]$, and the function is increasing for all $p < \underline{q}$ and decreasing for all $p > \bar{q}$. (Note that $m(1) - \mathbf{w}(1) = \frac{(1-\delta)(u(a^*, 1) - m(1))}{\delta} < 0$, hence the function is equal to $-\infty$ at $p = 1$.) $\square$

We can make few additional observations about the interval $[\underline{q}, \bar{q}]$. Let $k^* := \sup\{k : Q^k \neq \emptyset\}$. Since $\varphi(\bar{q}^k, \mathbf{w}(\bar{q}^k)) = \underline{q}^k$, the function h is decreasing for all $p \geq \bar{q}^{k^*}$. Similarly, since $\varphi(\underline{q}^k, \mathbf{w}(\underline{q}^k)) = \underline{q}^{k-1}$, the function h is increasing for all $p \leq \underline{q}^{k^*}$. Therefore, $[\underline{q}, \bar{q}] \subset Q^{k^*}$.

If $P \neq \emptyset$, so that $k^* = \infty$, then for all $p \in P$, the function h is increasing by convexity of m since $\mathbf{w}(p) = m(p)$. (This is clearly true since $\varphi(p, m(p)) = p$ in that region.) Therefore, $\bar{p} \leq \underline{q}$ if $P \neq \emptyset$.

Finally, let $\tilde{p} := \inf\{p : m(p) = u(a^1, p)\}$. By construction, m is linear from $\tilde{p}$ to 1, that is, $[\tilde{p}, 1]$ is the utmost right linear piece of m . We have that $\bar{q} < \tilde{p}$. To see this, observe that for all $p \geq \tilde{p}$,

$$\begin{aligned} \frac{m(1) - \mathbf{w}(p)}{1-p} &= \frac{(1-\delta)\overbrace{(u(a^*, 1) - u(a^1, 1))}^{<0}}{1-p} \\ &\quad + \frac{(u(a^1, 0) - u(a^1, 1)) - (1-\delta)(u(a^*, 0) - u(a^*, 1))}{\delta}, \end{aligned}$$

hence it is decreasing in p . (If there are multiple optimal actions at $p = 1$, the argument applies to all of them and, therefore, to the one that induces the smallest $\tilde{p}$.)

The second preliminary result is technical. For any $p \in (0, 1)$ and any $\eta \in [0, \frac{M(p) - \bar{m}_{\bar{q}^{-1}}(p)}{\bar{m}_{\bar{q}^{-1}}(1) - u(a^*, 1)}]$, define $w(p; \eta)$ as

$$\bar{m}_{\bar{q}^{-1}}(p) + \eta[\bar{m}_{\bar{q}^{-1}}(1) - u(a^*, 1)],$$

and write $(\lambda_\eta, \varphi_\eta)$ for $(\bar{\lambda}(p, w(p; \eta)), \bar{\varphi}(p, w(p; \eta)))$. To ease notation, we do not explicitly write the dependence of $(\lambda_\eta, \varphi_\eta)$ on p . We have the following.

LEMMA 5. φ_η , λ_η , and $\frac{1-\lambda_\eta}{\eta}$ are all decreasing in η .

The proof follows directly from the definition of $(\lambda_\eta, \varphi_\eta)$ and is omitted.

Finally, we conclude with the following implication of Observation A, which we will use throughout. For all (p, w, w') with $w \leq w'$, we have that

$$\begin{aligned} V_{\bar{q}^1}(p, w) - V_{\bar{q}^1}(p, w') \\ = \bar{\lambda}(p, w) \left[V_{\bar{q}^1}(\bar{\varphi}(p, w), \bar{m}_{\bar{q}^1}(p, w)) - V_{\bar{q}^1} \left(\bar{\varphi}(p, w), \bar{m}_{\bar{q}^1}(p, w) + \frac{w' - w}{\bar{\lambda}(p, w)} \right) \right]. \end{aligned}$$

We now turn to the proof of Lemma 3.

B.2.2 Proof of Lemma 3 We now prove that the gradient

$$\mathcal{G}(p; \eta) := \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta))}{\eta}$$

is increasing in $\eta \in [0, \frac{M(p) - \bar{m}_{\bar{q}^1}(p)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)}]$, for all p . We prove it on three separate intervals $\mathcal{I}_1$, $\mathcal{I}_2$, and $\mathcal{I}_3$. If $P = \emptyset$, the three intervals are $[0, \underline{q}]$, $(\underline{q}, \bar{q}]$ and $(\bar{q}, 1]$, respectively. If $P \neq \emptyset$, the three intervals are $[0, \underline{p}]$, $(\underline{p}, \bar{q}^\infty]$ and $(\bar{q}^\infty, 1]$, respectively.

Fact 1: For all $p \in \mathcal{I}_1$, $\mathcal{G}(p; \eta)$ is increasing in η . We limit attention to the case $P \neq \emptyset$. (The case $P = \emptyset$ is identical.) The proof is by induction. First, consider the interval $[0, \underline{q}^1]$. Remember that at $\underline{q}^1$, we have a closed-form solution for $V_{\bar{q}^1}(\underline{q}^1, w)$ for all w given by

$$V_{\bar{q}^1}(\underline{q}^1, w) = \frac{M(\underline{q}^1) - w}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} v(a^*, \underline{q}^1).$$

Therefore,

$$\begin{aligned} & \frac{V_{\bar{q}^1}(\underline{q}^1, \bar{m}_{\bar{q}^1}(\underline{q}^1)) - V_{\bar{q}^1}(\underline{q}^1, w(\underline{q}^1; \eta))}{\eta} \\ &= \frac{1}{\eta} \left[\frac{M(\underline{q}^1) - \bar{m}_{\bar{q}^1}(\underline{q}^1)}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} v(a^*, \underline{q}^1) - \frac{M(\underline{q}^1) - w(\underline{q}^1; \eta)}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} v(a^*, \underline{q}^1) \right] \\ &= \frac{v(a^*, \underline{q}^1)}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} \frac{[\bar{m}_{\bar{q}^1}(\underline{q}^1) + \eta(\bar{m}_{\bar{q}^1}(1) - u(a^*, 1))] - \bar{m}_{\bar{q}^1}(\underline{q}^1)}{\eta} \\ &= \frac{\underline{q}^1 v(a^*, 1) + (1 - \underline{q}^1) v(a^*, 0)}{\underline{q}^1 [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] + (1 - \underline{q}^1) [\bar{m}_{\bar{q}^1}(0) - u(a^*, 0)]} \frac{w(\underline{q}^1; \eta) - \bar{m}_{\bar{q}^1}(\underline{q}^1)}{\eta} \end{aligned}$$

$$\begin{aligned}
&= v(a^*, 1) \frac{\underline{q}^1 + (1 - \underline{q}^1) \frac{v(a^*, 0)}{v(a^*, 1)}}{\underbrace{\underline{q}^1 + (1 - \underline{q}^1) \frac{\bar{m}_{\bar{q}^1}(0) - u(a^*, 0)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)}}_{\geq 1 \text{ since } \frac{v(a^*, 0)}{v(a^*, 1)} \geq \frac{\bar{m}_{\bar{q}^1}(0) - u(a^*, 0)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)}}} \geq v(a^*, 1).
\end{aligned}$$

We now consider any $p \in [0, \underline{q}^1]$. From Observation A, we have that

$$\begin{cases} V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) = \frac{1-p}{1-\underline{q}^1} V_{\bar{q}^1}\left(\underline{q}^1, \frac{1-\underline{q}^1}{1-p} \bar{m}_{\bar{q}^1}(p) + \left(1 - \frac{1-\underline{q}^1}{1-p}\right) \bar{m}_{\bar{q}^1}(1)\right) \\ V_{\bar{q}^1}(p, w(p; \eta)) = \frac{1-p}{1-\underline{q}^1} V_{\bar{q}^1}\left(\underline{q}^1, \frac{1-\underline{q}^1}{1-p} \bar{m}_{\bar{q}^1}(p) + \left(1 - \frac{1-\underline{q}^1}{1-p}\right) \bar{m}_{\bar{q}^1}(1) \right. \\ \left. + \frac{1-\underline{q}^1}{1-p} \eta [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) \end{cases}$$

It follows that

$$\begin{aligned}
&\frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta))}{\eta} \\
&= \frac{1-p}{1-\underline{q}^1} \left(V_{\bar{q}^1}\left(\underline{q}^1, \frac{1-\underline{q}^1}{1-p} \bar{m}_{\bar{q}^1}(p) + \left(1 - \frac{1-\underline{q}^1}{1-p}\right) \bar{m}_{\bar{q}^1}(1)\right) \right. \\
&\quad \left. - V_{\bar{q}^1}\left(\underline{q}^1, \frac{1-\underline{q}^1}{1-p} \bar{m}_{\bar{q}^1}(p) + \left(1 - \frac{1-\underline{q}^1}{1-p}\right) \bar{m}_{\bar{q}^1}(1) \right. \right. \\
&\quad \left. \left. + \frac{1-\underline{q}^1}{1-p} \eta [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) \right) \eta^{-1} \\
&= \frac{1-p}{1-\underline{q}^1} \frac{1-\underline{q}^1}{1-p} \frac{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)}{M(\underline{q}^1) - u(a^*, \underline{q}^1)} v(a^*, \underline{q}^1) \\
&= \frac{1-p}{1-\underline{q}^1} \frac{1-\underline{q}^1}{1-p} v(a^*, 1) \frac{\underline{q}^1 + (1 - \underline{q}^1) \frac{v(a^*, 0)}{v(a^*, 1)}}{\underline{q}^1 + (1 - \underline{q}^1) \frac{\bar{m}_{\bar{q}^1}(0) - u(a^*, 0)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)}} \\
&\geq \frac{1-p}{1-\underline{q}^1} \frac{1-\underline{q}^1}{1-p} v(a^*, 1) = v(a^*, 1).
\end{aligned}$$

Therefore, $\mathcal{G}(p; \eta) \geq v(a^*, 1)$ for all η , for all $p \in [0, \underline{q}^1]$. Moreover, the gradient $\mathcal{G}(p; \eta)$ is independent of η for all $p \in [0, \underline{q}^1]$, hence is (weakly) increasing.

By induction, assume that $\mathcal{G}(p; \eta) \geq v(a^*, 1)$ for all $p \in [0, \underline{q}^k]$ and is increasing in η , and we want to prove that both properties also hold for all $p \in (\underline{q}^k, \underline{q}^{k+1}]$.

We rewrite $V_{\bar{q}^1}(p, w(p; \eta))$ as follows:

$$\begin{aligned}
 V_{\bar{q}^1}(p, w(p; \eta)) &= \lambda_\eta V_{\bar{q}^1}(\varphi_\eta, \bar{m}_{\bar{q}^1}(\varphi_\eta)) = \lambda_\eta [(1 - \delta)v(a^*, \varphi_\eta) + \delta V_{\bar{q}^1}(\varphi_\eta, \mathbf{w}(\varphi_\eta))] \\
 &= (1 - \delta)\lambda_\eta v(a^*, \varphi_\eta) + \delta \lambda_\eta V_{\bar{q}^1}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \\
 &= (1 - \delta)\lambda_\eta v(a^*, \varphi_\eta) + \delta V_{\bar{q}^1}(p, \lambda_\eta \mathbf{w}(\varphi_\eta) + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1)) \\
 &= (1 - \delta)\lambda_\eta v(a^*, \varphi_\eta) + \delta V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right).
 \end{aligned}$$

The second to last equality follows from Observation A, while the last equality follows from

$$\begin{aligned}
 &\lambda_\eta \mathbf{w}(\varphi_\eta) + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1) \\
 &= \lambda_\eta \frac{-(1 - \delta)u(a^*, \varphi_\eta) + \bar{m}_{\bar{q}^1}(\varphi_\eta)}{\delta} + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1) \\
 &= \frac{-(1 - \delta)}{\delta} \lambda_\eta u(a^*, \varphi_\eta) + \frac{1}{\delta} \lambda_\eta \bar{m}_{\bar{q}^1}(\varphi_\eta) + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1) \\
 &= \frac{-(1 - \delta)}{\delta} [u(a^*, p) - (1 - \lambda_\eta)u(a^*, 1)] \\
 &\quad + \frac{1}{\delta} [w(p; \eta) - (1 - \lambda_\eta)\bar{m}_{\bar{q}^1}(1)] + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1) \\
 &= \frac{-(1 - \delta)}{\delta} [u(a^*, p) - (1 - \lambda_\eta)u(a^*, 1)] \\
 &\quad + \frac{1}{\delta} [\bar{m}_{\bar{q}^1}(p) + \eta(\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)) - (1 - \lambda_\eta)\bar{m}_{\bar{q}^1}(1)] + [1 - \lambda_\eta]\bar{m}_{\bar{q}^1}(1) \\
 &= \left[\frac{-(1 - \delta)}{\delta} u(a^*, p) + \frac{1}{\delta} \bar{m}_{\bar{q}^1}(p) \right] + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)].
 \end{aligned}$$

For future reference, recall that

$$\begin{aligned}
 &\lambda_\eta \mathbf{w}(\varphi_\eta) + (1 - \lambda_\eta)\bar{m}_{\bar{q}^1}(1) \\
 &= \lambda_\eta [\bar{\lambda}(\varphi_\eta, \mathbf{w}(\varphi_\eta))\bar{m}_{\bar{q}^1}(\bar{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta))) + (1 - \bar{\lambda}(\varphi_\eta, \mathbf{w}(\varphi_\eta))\bar{m}_{\bar{q}^1}(1))] \\
 &\quad + (1 - \lambda_\eta)\bar{m}_{\bar{q}^1}(1),
 \end{aligned}$$

so that

$$\begin{aligned}
 \bar{\varphi}\left(p, \mathbf{w}(p) + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) &= \bar{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)), \quad \text{and} \\
 \bar{\lambda}\left(p, \mathbf{w}(p) + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) &= \lambda_\eta \bar{\lambda}(\varphi_\eta, \mathbf{w}(\varphi_\eta)).
 \end{aligned}$$

Since φ_η is decreasing in η , we have $\varphi_{\eta'} \leq \varphi_\eta$ when $\eta' > \eta$, and hence $\overline{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \leq \overline{\varphi}(\varphi_{\eta'}, \mathbf{w}(\varphi_{\eta'}))$, as $\varphi_{\eta'} \leq \varphi_\eta \leq p \leq \underline{q}$. Similarly, since $\varphi_\eta < p \leq \underline{q}$, we have that $\overline{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \leq \overline{\varphi}(p, \mathbf{w}(p))$ and, therefore, $\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} > 0$.

We now return to the computation of the gradient. We have

$$\begin{aligned}
 &= \left([(1-\delta)v(a^*, p) + \delta V_{\bar{q}^1}(p, \mathbf{w}(p))] \right. \\
 &\quad \left. - \left[(1-\delta)\lambda_\eta v(a^*, \varphi_\eta) + \delta V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right) \right] \right) \\
 &\quad \times \eta^{-1} \\
 &= \frac{(1-\delta)}{\eta} [v(a^*, p) - \lambda_\eta v(a^*, \varphi_\eta)] \\
 &\quad + \frac{\delta}{\eta} \left[V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right) \right] \\
 &= \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) \\
 &\quad + \frac{\delta}{\eta} \left[V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right) \right]. \quad (10)
 \end{aligned}$$

We further develop the above expression. To ease notation, we write $(\varphi(p), \lambda(p))$ for $(\overline{\varphi}(p, \mathbf{w}(p)), \overline{\lambda}(p, \mathbf{w}(p)))$. Note that $\varphi(p) \in (\underline{q}^{k-1}, \underline{q}^k]$, since $p \in (\underline{q}^k, \underline{q}^{k+1}]$. As $\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} > 0$, we have that

$$\begin{aligned}
 &= \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) \\
 &\quad + \frac{\delta}{\eta} \left[V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right) \right] \\
 &= \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) + \frac{\delta}{\eta} \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} \\
 &\quad \times \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right)}{\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta}} \\
 &= \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_\eta)}{\eta} \right] \\
 &\quad \times \left(\lambda(p) \left[V_{\bar{q}^1}(\varphi(p), \overline{m}_{\bar{q}^1}(\varphi(p))) \right. \right. \\
 &\quad \left. \left. - V_{\bar{q}^1}\left(\varphi(p), \overline{m}_{\bar{q}^1}(\varphi(p)) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta \lambda(p)} [m(1) - u(a^*, 1)]\right) \right] \right) \\
 &\quad \times \left(\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} \right)^{-1}
 \end{aligned}$$

$$\begin{aligned}
&= \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_\eta)}{\eta} \right] \\
&\quad \times \left(V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) \right. \\
&\quad \left. - V_{\bar{q}^1} \left(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p)) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta \lambda(p)} [m(1) - u(a^*, 1)] \right) \right) \\
&\quad \times \left(\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta \lambda(p)} \right)^{-1} \\
&\geq \frac{(1-\delta)}{\eta} (1-\lambda_\eta) v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_\eta)}{\eta} \right] v(a^*, 1) = v(a^*, 1),
\end{aligned}$$

where we use Observation A and the induction step.

We now show that the gradient is increasing in η . To start with, note that $\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta}$ is increasing in η since $\frac{1-\lambda_\eta}{\eta}$ is decreasing in η (see Lemma 5). For any $\eta > \eta'$, we have the following:

$$\begin{aligned}
&\frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1} \left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] \right)}{\frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta}} \\
&= \left(\lambda(p) V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) \right. \\
&\quad \left. - \lambda(p) V_{\bar{q}^1} \left(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p)) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta \lambda(p)} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] \right) \right) \\
&\quad \times \left(\frac{\eta - (1-\delta)(1-\lambda)}{\delta} \right)^{-1} \\
&= \left(V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) \right. \\
&\quad \left. - V_{\bar{q}^1} \left(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p)) + \frac{\eta - (1-\delta)(1-\lambda_\eta)}{\delta \lambda(p)} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] \right) \right) \\
&\quad \times \left(\frac{\eta - (1-\delta)(1-\lambda)}{\delta \lambda(p)} \right)^{-1} \\
&\geq \left(V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) \right. \\
&\quad \left. - V_{\bar{q}^1} \left(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p)) + \frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta \lambda(p)} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] \right) \right) \\
&\quad \times \left(\frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta \lambda(p)} \right)^{-1}
\end{aligned}$$

$$= \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right)}{\frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta}},$$

where the inequality follows from the fact that $\varphi(p) \in (q^{k-1}, q^k]$ and, therefore, the gradient $\mathcal{G}(\varphi(p); \eta)$ being increasing in η by the induction hypothesis.

Finally, we have that

$$\begin{aligned} & \frac{1}{\eta} [V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta))] \\ &= \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} \right] \\ & \quad \times \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1-\delta)(1-\lambda_{\eta})}{\delta} [m(1) - u(a^*, 1)]\right)}{\frac{\eta - (1-\delta)(1-\lambda_{\eta})}{\delta}} \\ & \geq \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} \right] \\ & \quad \times \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right)}{\frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta}} \\ &= \frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} \right] \\ & \quad \times \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right)}{\frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta}} \\ & \quad + \left[\frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} - \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} \right] \\ & \quad \times \left[\frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right)}{\frac{\eta' - (1-\delta)(1-\lambda_{\eta'})}{\delta}} \right. \\ & \quad \left. - v(a^*, 1) \right] \\ & \geq \frac{1}{\eta'} [V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta'))] \\ & \quad + \left[\frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} - \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} \right] \end{aligned}$$

$$\begin{aligned}
& \times \left[\left(V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) \right. \right. \\
& \left. \left. - V_{\bar{q}^1}(\varphi(p), \bar{m}_{\bar{q}^1}(\varphi(p))) + \frac{\eta' - (1 - \delta)(1 - \lambda_{\eta'})}{\delta \lambda(p)} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)] \right) \right) \\
& \times \left(\frac{\eta' - (1 - \delta)(1 - \lambda_{\eta'})}{\delta \lambda(p)} \right)^{-1} - v(a^*, 1) \Big] \\
& \geq \frac{1}{\eta'} [V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta'))].
\end{aligned}$$

The last inequality follows from the fact that the gradient in the second bracket is weakly larger than $v(a^*, 1)$ by the induction hypothesis and the fact that $\frac{1 - \lambda_{\eta}}{\eta} < \frac{1 - \lambda_{\eta'}}{\eta'}$ (Lemma 5).

Since $\lim_{k \rightarrow \infty} \underline{q}^k = \underline{p}$ when $P \neq \emptyset$, this completes the proof that the gradient is greater than $v(a^*, 1)$ for all $p \in [0, \underline{p}]$.

Fact 2: For all $p \in \mathcal{I}_2$, $\mathcal{G}(p; \eta)$ is increasing in η . We first treat the case $P \neq \emptyset$. Recall that for all $p \in (\underline{p}, \bar{q}^\infty]$, we have an explicit definition of the value function $V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p))$ as

$$v(a^*, p) - \frac{\bar{m}_{\bar{q}^1}(p) - u(a^*, p)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)} v(a^*, 1).$$

Define $\bar{\eta}(p)$ as the solution to $\varphi_{\bar{\eta}(p)} = \bar{\varphi}(p, w(p; \bar{\eta}(p))) = \underline{p}$. Note that for any $p \in (\underline{p}, \bar{q}^\infty]$, for any $\eta \leq \bar{\eta}$, $\varphi_\eta \in [\underline{p}, \bar{q}^\infty]$. Therefore,

$$\begin{aligned}
V_{\bar{q}^1}(p, w(p; \eta)) &= \lambda_\eta V_{\bar{q}^1}(\varphi_\eta, \bar{m}_{\bar{q}^1}(\varphi_\eta)) = \lambda_\eta \left[v(a^*, \varphi_\eta) - \frac{\bar{m}_{\bar{q}^1}(\varphi_\eta) - u(a^*, \varphi_\eta)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)} v(a^*, 1) \right] \\
&= v(a^*, p) - \frac{w(p; \eta) - u(a^*, p)}{\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)} v(a^*, 1).
\end{aligned}$$

It follows that the gradient is equal to $v(a^*, 1)$ for all $p \in (\underline{p}, p^*]$, for all $\eta \leq \bar{\eta}$.

Consider now $\eta > \bar{\eta}$. We rewrite the gradient $\mathcal{G}(p; \eta)$ as follows:

$$\begin{aligned}
& \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta))}{\eta} \\
&= \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta_1(p)))}{\eta} \\
& \quad + \frac{V_{\bar{q}^1}(p, w(p; \eta_1(p))) - V_{\bar{q}^1}(p, w(p; \eta))}{\eta} \\
&= \frac{\eta_1(p)}{\eta} \frac{V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p, \eta_1(p)))}{\eta_1(p)} \\
& \quad + \frac{\eta - \eta_1(p)}{\eta} \frac{V_{\bar{q}^1}(p, w(p; \eta_1(p))) - V_{\bar{q}^1}(p, w(p; \eta))}{\eta - \eta_1(p)}
\end{aligned}$$

$$\begin{aligned}
& \frac{1-p}{1-\underline{p}} \left[V_{\bar{q}^1}(\underline{p}, \bar{m}_{\bar{q}^1}(\underline{p})) - V_{\bar{q}^1}\left(\underline{p}, w\left(\underline{p}; \frac{\eta - \eta_1(p)}{1-\underline{p}}\right)\right) \right] \\
&= \frac{\eta_1(p)}{\eta} v(a^*, 1) + \frac{\eta - \eta_1(p)}{\eta} \frac{\eta - \eta_1(p)}{\eta - \eta_1(p)} \\
&= \frac{\eta_1(p)}{\eta} v(a^*, 1) + \frac{\eta - \eta_1(p)}{\eta} \mathcal{G}\left(\underline{p}; \frac{\eta - \eta_1(p)}{1-\underline{p}}\right).
\end{aligned}$$

Since we have already shown that $\mathcal{G}(\underline{p}; \eta)$ is increasing in η and weakly larger than $v(a^*, 1)$, we have that the gradient $\mathcal{G}(p; \eta)$ is also weakly increasing in η (and greater than $v(a^*, 1)$).

We now treat the case $P = \emptyset$. Define $\bar{\eta}(p)$ as the solution to $\varphi_{\bar{\eta}(p)} = \bar{\varphi}(p, w(p; \bar{\eta}(p))) = \underline{q}$. Note that for any $p \in [\underline{q}, \bar{q}]$, for any $\eta \leq \bar{\eta}$, $\varphi_\eta \in [\underline{q}, \bar{q}]$. Therefore, for all $\eta \leq \bar{\eta}$, $\eta = (1 - \delta)(1 - \lambda_\eta)$ since the ratio $\frac{\bar{m}_{\bar{q}^1}(1) - \mathbf{w}(\varphi_\eta)}{1 - \varphi_\eta}$ is constant in η and so is $\bar{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta))$. (Recall that we vary η at a fixed p .) It follows then from equation (10) that

$$\begin{aligned}
\mathcal{G}(p; \eta) &= \frac{(1 - \delta)}{\eta} (1 - \lambda_\eta) v(a^*, 1) \\
&\quad + \frac{\delta}{\eta} \left[V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [m(1) - u(a^*, 1)]\right) \right] \\
&= \frac{(1 - \delta)}{\eta} (1 - \lambda_\eta) v(a^*, 1) = v(a^*, 1).
\end{aligned}$$

We have that the gradient $\mathcal{G}(p; \eta)$ is equal to $v(a^*, 1)$ for all $p \in (\underline{q}, \bar{q}]$, for all $\eta \leq \bar{\eta}$. Finally, when $\eta > \bar{\eta}$, the same decomposition as in the case $P \neq \emptyset$ completes the proof.

Fact 3: For all $p \in \mathcal{I}_3$, the gradient $\mathcal{G}(p; \eta)$ is increasing in η . We only treat the case $P \neq \emptyset$. (The case $P = \emptyset$ is treated analogously.) Define $\bar{\eta}(p)$ as the solution to $\varphi_{\bar{\eta}(p)} = \bar{\varphi}(p, w(p; \bar{\eta}(p))) = \bar{q}^\infty$. By construction, for all $p \in (\bar{q}^\infty, 1]$, for all $\eta \leq \bar{\eta}(p)$, we have that $\varphi_\eta \in (\bar{q}^\infty, 1]$. Therefore, $\varphi_\eta > \bar{q}$.

Choose $\bar{\eta}(p) \leq \eta' \leq \eta$. We have that $\varphi_{\eta'} \geq \varphi_\eta \geq \bar{q}$ since $\bar{q}^\infty \geq \bar{q}$ and, therefore,

$$\begin{aligned}
& \bar{\varphi}\left(p, \mathbf{w}(p) + \frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) \\
&= \bar{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \\
&\geq \bar{\varphi}(\varphi_{\eta'}, \mathbf{w}(\varphi_{\eta'})) = \bar{\varphi}\left(p, \mathbf{w}(p) + \frac{\eta' - (1 - \delta)(1 - \lambda_{\eta'})}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right).
\end{aligned}$$

Also, since $\bar{q} \leq \varphi_\eta \leq p$, we have that $\bar{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \geq \bar{\varphi}(p, \mathbf{w}(p))$ and, therefore, $\frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} \leq 0$. The same applies to η' . Finally, as already shown,

$$\frac{\eta - (1 - \delta)(1 - \lambda_\eta)}{\delta} < \frac{\eta' - (1 - \delta)(1 - \lambda_{\eta'})}{\delta}.$$

To ease notation, define $(\tilde{\lambda}_\eta, \tilde{\varphi}_\eta)$ as follows:

$$\begin{cases} \tilde{\lambda}_\eta = \lambda \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} [m(1) - u(a^*, 1)] \right) \\ \tilde{\varphi}_\eta = \varphi \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} [m(1) - u(a^*, 1)] \right) \end{cases} \quad (11)$$

Notice that $\tilde{\varphi}_\eta = \overline{\varphi}(\varphi_\eta, \mathbf{w}(\varphi_\eta)) \in \mathcal{I}_1$ since $\varphi_\eta > \overline{q}^\infty$.

The rest of the proof is purely algebraic and mirrors the case $p \in \mathcal{I}_1$. First, we have the following:

$$\begin{aligned} & \frac{V_{\tilde{q}^1}(p, \mathbf{w}(p)) - V_{\tilde{q}^1} \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} [\overline{m}_{\tilde{q}^1}(1) - u(a^*, 1)] \right)}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta}} \\ &= \left(\tilde{\lambda}_\eta V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, \overline{m}_{\tilde{q}^1}(\tilde{\varphi}_\eta) + \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} [\overline{m}_{\tilde{q}^1}(1) - u(a^*, 1)] \right) \right. \\ & \quad \left. - \tilde{\lambda}_\eta V_{\tilde{q}^1}(\tilde{\varphi}_\eta, \overline{m}_{\tilde{q}^1}(\tilde{\varphi}_\eta)) \right) \left(\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} \right)^{-1} \\ &= \frac{V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, w \left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} \right) \right) - V_{\tilde{q}^1}(\tilde{\varphi}_\eta, \overline{m}_{\tilde{q}^1}(\tilde{\varphi}_\eta))}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta}}, \end{aligned}$$

where we again use Observation A. Similarly, we have

$$\begin{aligned} & \frac{V_{\tilde{q}^1}(p, \mathbf{w}(p)) - V_{\tilde{q}^1} \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\overline{m}_{\tilde{q}^1}(1) - u(a^*, 1)] \right)}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\ &= \left(\tilde{\lambda}_\eta V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, w \left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} \right) \right) \right. \\ & \quad \left. - \tilde{\lambda}_\eta V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, w \left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_\eta} \right) \right) \right) \\ & \quad \times \left(\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} \right)^{-1} \\ &= \left(V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, w \left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} \right) \right) \right. \\ & \quad \left. - V_{\tilde{q}^1} \left(\tilde{\varphi}_\eta, w \left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta} - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_\eta} \right) \right) \right) \end{aligned}$$

$$\times \left(\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_{\eta}} \right)^{-1},$$

where again we use Observation A and the fact

$$\frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta \tilde{\lambda}_{\eta}} > \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_{\eta}}.$$

Since $\tilde{\varphi}_{\eta} \in \mathcal{I}_1$, we have that

$$\begin{aligned} & \left(V_{\tilde{q}^1} \left(\tilde{\varphi}_{\eta}, w \left(\tilde{\varphi}_{\eta}; \frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta \tilde{\lambda}_{\eta}} \right) \right) \right. \\ & \quad \left. - V_{\tilde{q}^1} \left(\tilde{\varphi}_{\eta}, w \left(\tilde{\varphi}_{\eta}; \frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta \tilde{\lambda}_{\eta}} - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_{\eta}} \right) \right) \right) \\ & \quad \times \left(\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_{\eta}} \right)^{-1} \\ & \leq \frac{V_{\tilde{q}^1} \left(\tilde{\varphi}_{\eta}, w \left(\tilde{\varphi}_{\eta}; \frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta \tilde{\lambda}_{\eta}} \right) \right) - V_{\tilde{q}^1}(\tilde{\varphi}_{\eta}, \bar{m}_{\tilde{q}^1}(\tilde{\varphi}_{\eta}))}{\frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta \tilde{\lambda}_{\eta}}}, \end{aligned}$$

where the inequality follows from our previous argument on the interval $\mathcal{I}_1$.

It follows that

$$\begin{aligned} & \frac{V_{\tilde{q}^1}(p, \mathbf{w}(p)) - V_{\tilde{q}^1} \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\tilde{q}^1}(1) - u(a^*, 1)] \right)}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\ & \leq \frac{V_{\tilde{q}^1}(p, \mathbf{w}(p)) - V_{\tilde{q}^1} \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta} [\bar{m}_{\tilde{q}^1}(1) - u(a^*, 1)] \right)}{\frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta}}. \end{aligned}$$

From equation (10), we then have that

$$\begin{aligned} & \frac{1}{\eta} [V_{\tilde{q}^1}(p, \bar{m}_{\tilde{q}^1}(p)) - V_{\tilde{q}^1}(p, \mathbf{w}(p; \eta))] \\ & = \frac{(1-\delta)(1-\lambda_{\eta})}{\eta} v(a^*, 1) + \left[\frac{(1-\delta)(1-\lambda_{\eta})}{\eta} - 1 \right] \\ & \quad \times \frac{V_{\tilde{q}^1}(p, \mathbf{w}(p)) - V_{\tilde{q}^1} \left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta} [m(1) - u(a^*, 1)] \right)}{\frac{(1-\delta)(1-\lambda_{\eta}) - \eta}{\delta}} \end{aligned}$$

$$\begin{aligned}
 &\geq \frac{(1-\delta)(1-\lambda_\eta)}{\eta} v(a^*, 1) + \left[\frac{(1-\delta)(1-\lambda_\eta)}{\eta} - 1 \right] \\
 &\quad \times \frac{V_{\bar{q}^1}(p, \mathbf{w}(p)) - V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right)}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\
 &= \frac{(1-\delta)(1-\lambda_\eta)}{\eta} v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_\eta)}{\eta} \right] \\
 &\quad \times \frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\
 &= \frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} v(a^*, 1) + \left[1 - \frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} \right] \\
 &\quad \times \frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\
 &\quad + \left[\frac{(1-\delta)(1-\lambda_{\eta'})}{\eta'} - \frac{(1-\delta)(1-\lambda_\eta)}{\eta} \right] \\
 &\quad \times \left[\frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \right. \\
 &\quad \left. - v(a^*, 1) \right] \\
 &\geq \frac{1}{\eta'} [V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, w(p; \eta'))],
 \end{aligned}$$

where the last inequality follows from

$$\begin{aligned}
 &\frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta} [\bar{m}_{\bar{q}^1}(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \\
 &= \frac{\tilde{\lambda}_{\eta'} V_{\bar{q}^1}(\tilde{\varphi}_{\eta'}, \bar{m}_{\bar{q}^1}(\tilde{\varphi}_{\eta'})) - \tilde{\lambda}_{\eta'} V_{\bar{q}^1}\left(\tilde{\varphi}_{\eta'}, w\left(\tilde{\varphi}_{\eta'}; \frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta \tilde{\lambda}_{\eta'}}\right)\right)}{\frac{(1-\delta)(1-\lambda_{\eta'}) - \eta'}{\delta}} \geq v(a^*, 1).
 \end{aligned}$$

We now show that the the gradient $\mathcal{G}(p; \eta)$ is smaller than $v(a^*, 1)$ for any $\eta \leq \bar{\eta}(p)$.

From equation (10), we have that

$$\begin{aligned}
& \frac{1}{\eta} [V_{\bar{q}^1}(p, \bar{m}_{\bar{q}^1}(p)) - V_{\bar{q}^1}(p, \mathbf{w}(p; \eta))] \\
&= \frac{(1-\delta)(1-\lambda_\eta)}{\eta} v(a^*, 1) - \left[\frac{(1-\delta)(1-\lambda_\eta)}{\eta} - 1 \right] \\
& \quad \times \frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} [m(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta}} \\
&= v(a^*, 1) - \left[\frac{(1-\delta)(1-\lambda_\eta)}{\eta} - 1 \right] \\
& \quad \times \left[\frac{V_{\bar{q}^1}\left(p, \mathbf{w}(p) - \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta} [m(1) - u(a^*, 1)]\right) - V_{\bar{q}^1}(p, \mathbf{w}(p))}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta}} \right. \\
& \quad \left. - v(a^*, 1) \right] \\
&= v(a^*, 1) - \left[\frac{(1-\delta)(1-\lambda_\eta)}{\eta} - 1 \right] \\
& \quad \times \left[\frac{\tilde{\lambda}_\eta V_{\bar{q}^1}(\tilde{\varphi}_\eta, \bar{m}_{\bar{q}^1}(\tilde{\varphi}_\eta)) - \tilde{\lambda}_\eta V_{\bar{q}^1}\left(\tilde{\varphi}_\eta, w\left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta}\right)\right)}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta}} - v(a^*, 1) \right] \\
&= v(a^*, 1) - \underbrace{\left[\frac{(1-\delta)(1-\lambda_\eta)}{\eta} - 1 \right]}_{\geq 0} \\
& \quad \times \underbrace{\left[\frac{V_{\bar{q}^1}(\tilde{\varphi}_\eta, \bar{m}_{\bar{q}^1}(\tilde{\varphi}_\eta)) - V_{\bar{q}^1}\left(\tilde{\varphi}_\eta, w\left(\tilde{\varphi}_\eta; \frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta}\right)\right)}{\frac{(1-\delta)(1-\lambda_\eta) - \eta}{\delta \tilde{\lambda}_\eta}} - v(a^*, 1) \right]}_{\geq 0} \\
&\leq v(a^*, 1),
\end{aligned}$$

where the inequality follows from the fact that $\tilde{\varphi}_\eta \leq \underline{p}$ (therefore, from our arguments on the interval $\mathcal{I}_1$, where we show that the gradient is larger than $v(a^*, 1)$).

Finally, we can use a similar decomposition as in the case $p \in \mathcal{I}_2$ to prove that the gradient is increasing for all η .

APPENDIX C

C.1 Recursive formulation: A proof

Ely (2015) proves that the principal's maximal payoff is $\max_{w \in [m(p_0), M(p_0)]} \hat{V}^*(p_0, w)$, with $\hat{V}^*$ the unique fixed point of the contraction $\hat{T}$, with the operator $\hat{T}$ differing from the operator T in that the promise-keeping constraint is written as an equality in all maximization problems $\hat{T}(V)(p, w)$; all other constraints are the same. Note that, like T , the operator $\hat{T}$ is monotone.

For any $(p, w) \in \mathcal{W}$, let $\tilde{V}^*(p, w) := \max_{\tilde{w} \in [w, M(p)]} \hat{V}^*(p, \tilde{w})$ and $\tilde{w}^*(p, w)$ a maximizer. (If there are multiple maximizers, choose an arbitrary one.)

We prove that $V^* = \tilde{V}^*$. To do so, we prove that $T(\tilde{V}^*) = \tilde{V}^*$. Since T is a contraction, hence has a unique fixed point, it follows that $V^* = \tilde{V}^*$. (Note that we are not arguing that $T = \hat{T}$.)

We start with two simple observations: (i) $T(V)(p, w) \geq \hat{T}(V)(p, w)$ for all $(p, w) \in \mathcal{W}$, for all V , and (ii) $\tilde{V}^*(p, w) \geq \hat{V}^*(p, w)$ for all $(p, w) \in \mathcal{W}$. The first observation follows from the fact the promised-keeping constraint is an equality in $\hat{T}(V)(p, w)$, while it is an inequality in $T(V)(p, w)$. The second observation follows immediately from the definition of $\tilde{V}^*$.

We now prove that $T(\tilde{V}^*) \geq \tilde{V}^*$. For all $(p, w) \in \mathcal{W}$, we have

$$\begin{aligned} \tilde{V}^*(p, w) &= \hat{V}^*(p, \tilde{w}^*(p, w)) = \hat{T}(\hat{V}^*)(p, \tilde{w}^*(p, w)), \\ &\leq T(\hat{V}^*)(p, \tilde{w}^*(p, w)), \\ &\leq T(\hat{V}^*)(p, w), \\ &\leq T(\tilde{V}^*)(p, w), \end{aligned}$$

where the first line follows from the definitions of $\tilde{V}^*$, $\hat{V}^*$, and $\hat{T}$, and the fact that $\hat{V}^* = \hat{T}(\hat{V}^*)$; the second line from observation (i); the third line from the fact that $\tilde{w}^*(p, w) \geq w$, so that all feasible solutions to $T(\hat{V}^*)(p, \tilde{w}^*(p, w))$ are also feasible for $T(\hat{V}^*)(p, w)$; and the fourth line from observation (ii) and the definition of $T(V)(p, w)$, $V = \hat{V}^*$, $\tilde{V}^*$.

We next prove that $T(\tilde{V}^*) \leq \tilde{V}^*$. By contradiction, suppose that there exists $(p, w) \in \mathcal{W}$ and a feasible policy $(\lambda_s, p_s, a_s, w_s)_{s \in S}$ such that

$$\tilde{V}^*(p, w) < \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta \tilde{V}^*(p_s, w_s)].$$

Moreover, we have that

$$\begin{aligned} \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta \tilde{V}^*(p_s, w_s)] &= \sum_{s \in S} \lambda_s [(1 - \delta)v(a_s, p_s) + \delta \hat{V}^*(p_s, \tilde{w}^*(p_s, w_s))] \\ &\leq \hat{V}^*\left(p, \sum_{s \in S} [(1 - \delta)u(a_s, p_s) + \delta \tilde{w}^*(p_s, w_s)]\right) \\ &\leq \tilde{V}^*(p, w), \end{aligned}$$

where the first line follows from the definition of $\tilde{V}$; the second line from the observation that $(\lambda_s, p_s, a_s, \tilde{w}^*(p_s, w_s))_{s \in S}$ is feasible for the maximization problem $\hat{T}(\hat{V}^*)(p, \sum_{s \in S} [(1 - \delta)u(a_s, p_s) + \delta \tilde{w}^*(p_s, w_s)])$ and the fact that $\hat{T}(\hat{V}^*) = \hat{V}^*$; and the third line from the fact that

$$M(p) \geq \sum_{s \in S} [(1 - \delta)u(a_s, p_s) + \delta \tilde{w}^*(p_s, w_s)] \geq \sum_{s \in S} [(1 - \delta)u(a_s, p_s) + \delta w_s] \geq w,$$

since $\tilde{w}^*(p_s, w_s) \in [w_s, M(p_s)]$, for all $s \in S$, and the definition of $\tilde{V}^*$. We have the required contradiction, which completes the proof.

C.2 Proof of Proposition 3, part B

The existence of an optimal contract, which is simple and minimizes the probability of recommending a^* , follows from the compactness of the set of optimal contracts, and Proposition 3(i). To prove compactness, recall that $T(V^*)(p, w)$ is a constrained maximization problem, parameterized by (p, w) . Moreover, if V^* is continuous, so is $T(V^*)(p, w)$. The Berge maximum theorem implies the compactness of the set of optimal solutions of $T(V^*)(p, w)$ at all (p, w) and the continuity of $T(V^*)$. Thus, T is mapping continuous, concave and bounded functions into continuous, bounded, and bounded functions. Since the space of continuous, concave, and bounded functions is complete with respect to the sup-norm, the fixed point V^* is continuous, concave, and bounded (and its existence follows by Banach fixed-point theorem).

Next, since the set of optimal contracts is compact, there exists an optimal contract, which minimizes the probability of recommending a^* . Finally, the proof of Proposition 3(i) shows that there exists another optimal contract, which is simple and recommends a^* with the same probability.

REFERENCES

- Au, Pak Hung (2015), “Dynamic information disclosure.” *The RAND Journal of Economics*, 46, 791–823. [921]
- Aumann, Robert J., Michael B. Maschler, and Richard E. Stearns (1995), *Repeated Games With Incomplete Information*. MIT Press, Cambridge (Mass.) London. [920]
- Ball, Ian (2023), “Dynamic information provision: Rewarding the past and guiding the future.” *Econometrica*, 91, 1363–1391. [920, 921, 923]
- Bizzotto, Jacopo, Jesper Rüdiger, and Adrien Vigier (2021), “Dynamic persuasion with outside information.” *American Economic Journal: Microeconomics*, 13, 179–194. [921]
- Che, Yeon-Koo, Kyungmin Kim, and Konrad Mierendorff (2020), “Keeping the listener engaged: A dynamic model of Bayesian persuasion.” *Journal of Political Economy*, 131, 1619–1947. [921]
- Crawford, Vincent P. and Joel Sobel (1982), “Strategic information transmission.” *Econometrica*, 50, 1431–1451. [921]

- Ely, Jeffery (2015), “Beeps.” Northwestern University. [926, 972]
- Ely, Jeffrey (2017), “Beeps.” *American Economic Review*, 107, 31–53. [921, 926]
- Ely, Jeffrey and Martin Szydlowski (2020), “Moving the goalposts.” *Journal of Political Economy*, 128, 468–506. [920, 921, 936]
- Escude, Matteo and Ludvig Sinander (2023), “Slow persuasion.” *Theoretical Economics*, 18, 129–162. [921]
- Fudenberg, Drew and Luis Rayo (2019), “Training and effort dynamics in apprenticeship.” *American Economic Review*, 109, 3780–3812. [920]
- Garicano, Luis and Luis Rayo (2017), “Relational knowledge transfers.” *American Economic Review*, 107, 2695–2730. [920]
- Henry, Emeric and Marco Ottaviani (2019), “Research and the approval process: The organization of persuasion.” *American Economic Review*, 109, 911–955. [921]
- Kamenica, Emir (2019), “Bayesian persuasion and information design.” *Annual Review of Economics*, 11, 249–272. [920]
- Kamenica, Emir and Matthew Gentzkow (2011), “Bayesian persuasion.” *American Economic Review*, 101, 2590–2615. [920, 923]
- Makris, Miltos and Ludovic Renou (2023), “Information design in multi-stage games.” *Theoretical Economics*, 18, 1475–1509. [921]
- Orlov, Dmitry, Andrzej Skrzypacz, and Pavel Zryumov (2020), “Persuading the principal to wait.” *Journal of Political Economy*, 128, 2542–2578. [920, 921, 923]
- Renault, Jérôme, Eilon Solan, and Nicolas Vieille (2017), “Optimal dynamic information provision.” *Games and Economic Behavior*, 104, 329–349. [921, 926]
- Smolin, Alex (2021), “Dynamic evaluation design.” *American Economic Journal: Microeconomics*, 13, 300–331. [921]
- Spear, Stephen E. and Sanjay Srivastava (1987), “On repeated moral hazard with discounting.” *The Review of Economic Studies*, 54, 599–617. [918, 924]

Co-editor Marina Halac handled this manuscript.

Manuscript received 6 October, 2021; final version accepted 19 June, 2023; available online 3 July, 2023.