

Delegating learning

JUAN F. ESCOBAR

Department of Industrial Engineering, University of Chile

QIAOXI ZHANG

School of Economics, Xiamen University and Departamento de Economía, Universidad Diego Portales

Learning is crucial to organizational decision making but often needs to be delegated. We examine a dynamic delegation problem where a principal decides on a project with uncertain profitability. A biased agent, who is initially as uninformed as the principal, privately learns the profitability over time and communicates to the principal. We formulate learning delegation as a dynamic mechanism design problem and characterize the optimal delegation scheme. We show that private learning gives rise to the trade-off between how much information to acquire and how promptly it is reflected in the decision. We discuss implications on learning delegation for distinct organizations.

KEYWORDS. Private learning, delegation, delays, deadlines, commitment, cheap talk.

JEL CLASSIFICATION. D82, D83.

1. INTRODUCTION

Suppose that two people need to decide whether to invest in a project. If they invest, they could receive a gain or suffer a loss. If they do not invest, they wait, obtain new information, and may invest in the future. Now suppose that given the information so far, one of them—the principal—prefers to learn more, while the other—the agent—prefers to invest. If learning has to be delegated to the agent and the principal cannot observe the learning outcome, can the agent convey it truthfully? If so, what should be the optimal delegation scheme? How does it change over time given what the agent has learned so far? For how long should learning take place?

Delegating learning is a common occurrence. For example, suppose the board of directors of a company is deliberating whether to acquire another company. Apart from the financial value of the acquisition, its strategic value—e.g., its impact on the price and competition, the current employees, the bargaining power with suppliers—is also relevant. Since much of this information is hard to observe directly, the board needs to rely

Juan F. Escobar: jescobar@dii.uchile.cl

Qiaoxi Zhang: jdzhangqx@gmail.com

We are extremely grateful for comments from Nageeb Ali, Alessandro Bonatti, Rahul Deb, Elton Dusha, Federico Echenique, Andrew McClellan, Asher Wolinsky, and participants of various conferences and seminars. We acknowledge financial support from the Institute for Research in Market Imperfections and Public Policy (MIPP) ICM IS130002. Zhang also acknowledges financial support from CONICYT-FONDECYT postdoctoral award (Project 3170783).

on learning by the manager, who has direct access to the parties involved. For another example, the Food and Drug Administration (FDA) relies on pharmaceutical companies to develop drugs and test their efficacy. Although the companies are required to submit clinical trial results, the trials themselves cannot be fully monitored and, therefore, the results can be manipulated.¹

We study how a principal should delegate an investment decision to an agent who privately learns about the investment over time. Our analysis extends the traditional static delegation approach (Holmström 1984, Melumad and Shibano 1991, Alonso and Matouschek 2008, Amador and Bagwell 2013) to allow for evolving private information. We formulate delegation as a dynamic mechanism design problem and show that the optimal delegation scheme should have delays in response to information acquired so as to incentivize the agent's learning. Moreover, we uncover a unique trade-off arising in learning delegation, which is that between the amount of information acquired and how promptly it is reflected in the decision. The optimal duration of learning solves this trade-off. We also discuss some implications on organization decisions.

In our model, a principal and an agent face a project that never expires. The project's quality, which can be good or bad, is initially uncertain. Players share the same belief about the project's quality. The principal needs to decide when, if ever, to invest in the project. The project generates a signal whose arrival time is random. As long as no investment has happened, the agent privately observes the signal, or the absence thereof, without cost. Hence, investing and learning are two sides of the same coin in that as long as investment has not happened, learning continues. A signal perfectly reveals the project quality. At each point in time before investment happens, the agent sends a cheap talk message about the information learned so far to the principal. The principal commits to a delegation rule that specifies, for each point in time and each possible message history at that time, whether to invest. Once the investment happens, the game ends.

No one receives a payoff if no investment has happened. Once the investment happens, each player receives a time-discounted payoff determined by project quality and the player's identity. A good investment brings gains to both players, while a bad investment brings losses to both players. Consequently, each player will want to invest if they are optimistic enough about project quality and will want to wait for more information otherwise. However, the common initial belief is high enough for the agent such that he prefers to invest immediately and low enough for the principal such that she prefers to wait and invest only when a good signal arrives. The challenge for the principal is then to incentivize the agent to tell the truth when he has not received any signal, while trying to invest as soon as possible after a good signal arrives.

First, we note that investment must follow a good signal with delay. If a good signal triggers immediate investment, the agent would like to pretend to be informed when

¹Several authors have documented frauds in clinical trials (George and Buyse 2015). Seife (2015) shows that the FDA has found substantive evidence of fraudulent data in biomedical research on humans. For a summary, see Seife, "Are your medications safe," *Slate*, February 9, 2015. However, we think the more common cases are probably less extreme and that some results are verifiable while others are not. We focus on the residual unverifiability.

in fact no signal has arrived, and no learning would take place. To see how the delay should evolve over time to incentivize learning, we need to understand the driving forces behind learning. On the one hand, learning benefits the agent because if a bad signal arrives, he would then learn that the project is bad and avoid the loss from investing. On the other hand, learning costs the agent in that it takes time. Suppose that no signal has arrived and the agent is still optimistic enough to prefer to invest right away. At this point, the cost of learning outweighs the benefit. To encourage learning, the principal needs to decrease the cost by making investment respond faster to the good signal throughout time. Suppose that the principal wants to encourage the agent to learn for one more day. If a claim of good signal leads the principal to invest immediately, then the cost of learning is a one day delay. However, if a good signal that arrives today leads to investment 5 days later while a good signal that arrives tomorrow leads to investment 4.5 days after tomorrow, the cost of learning is only a half day delay. Delays in investment that decrease in the arrival times of the good signal allow the principal to balance the cost and benefit of learning for the agent, hence his truthful revelation.

It is natural to think that if no good signal has arrived, investment should not happen. This is not always the case. In fact, as a result of the trade-off between the amount of information acquired and how effectively it is used, the principal may prefer to invest at a deadline even if no good signal has been claimed. Suppose that at some time T , even if no signal has arrived, learning stops and investment happens. Since no incentives for learning is required from T on, the delay decreases gradually to 0 at T . If instead the principal decides to incentivize learning after T , the delay at T must be positive. Accordingly, delays of investment for claims of the good signal at each time before T must also be increased. Therefore, the longer learning takes place, the better information the principal receives and the more accurate but less prompt her decision to invest is. If no investment happens unless a good signal arrives, learning could take place for an arbitrarily long period of time. Consequently, the decision to invest is 100% accurate because the principal only invests if she is completely sure that the project is good. However the downside is that she has to provide incentives for learning for a long time. The resulting long delays in investment when a good signal arrives can, therefore, be a prohibitive cost for the principal. Therefore, due to the trade-off between the amount of information acquired and how effectively it is used, the principal may find it optimal to commit to investing with a deadline even if no good signal has been claimed.

The optimal duration of learning that balances the trade-off depends on the players' preferences for learning. Given that the principal prefers to wait while the agent prefers to invest initially, if the players' gain-loss ratios are sufficiently high, maintaining incentives for learning becomes very costly for the principal and, therefore, the duration of learning is short. Alternatively, when the players' gain-loss ratios are sufficiently low, it is in the principal's interest to maintain a longer learning phase. Our results speak to how distinct organizations should use distinct protocols to delegate learning. In the FDA example, losses from approving a damaging drug are substantial. This maps to our model when the principal has a low gain-loss ratio and, therefore, is strongly inclined to learn. The optimal course of action for the FDA is then to be prudent by establishing long revision processes. Not only do they ensure that a damaging drug never gets

approved, the ensuing long delays in approval also guarantee truthful revelation from pharmaceutical companies. Alternatively, if a manager's career concern is strong and he has a high gain–loss ratio, the optimal action for the board when it comes to acquisition decisions is to set short learning phases and then acquire as long as no negative news has arrived.

Our paper contributes to the delegation literature initiated by [Holmström \(1984\)](#) and extended by [Melumad and Shibano \(1991\)](#), [Alonso and Matouschek \(2008\)](#), [Armstrong and Vickers \(2010\)](#), [Amador and Bagwell \(2013\)](#), and [Ambrus and Egorov \(2017\)](#), among others. As in all these papers, in our model, the principal may grant flexibility to the agent so that he can use his information, but granting too much flexibility may open up room for opportunistic behavior. However, these papers study static models and do not address the issue of how to provide incentives to an agent with evolving private information.² In particular, our work emphasizes how the dynamic provision of incentives determines how information is used and for how long learning takes place.

[Grenadier et al. \(2016\)](#) and [Guo \(2016\)](#) explore delegation models in dynamic contexts. In [Grenadier et al. \(2016\)](#), a timing decision needs to be made and an agent who is informed at time 0 communicates with the principal throughout time. Whereas [Grenadier et al. \(2016\)](#) explore how the value of commitment for the principal depends on the sign of the agent's bias, we take commitment for granted but explore how to delegate with evolving private information. As [Grenadier et al. \(2016\)](#) point out, their full commitment case is similar to standard static delegation problems and, as a result, interval delegation is optimal. In [Guo \(2016\)](#), the principal delegates the decision to experiment over time to an agent who has private information about its profitability at time 0.³ Once experimentation starts, however, all signals are public. A comparison between our paper and [Guo \(2016\)](#) highlights the differences between private and public learning, which have important implications for the design of incentive schemes. In her model with a continuum of types, since signals are public, once a good signal arrives, the risky project is publicly known to be optimal and is fully implemented. In our model, however, investment decisions commonly known to be optimal are nonetheless delayed. This is the principal's response to the problem of providing incentives to an agent with evolving private information.

Our paper is also related to the study of optimal delegation decisions when information acquisition is endogenous. In [Aghion and Tirole \(1997\)](#), [Szalay \(2005\)](#), and [Deimen and Szalay \(2019\)](#), information acquisition is a one-time decision; therefore, the trade-off between extracting information and using information efficiently is different from ours. [Lewis and Ottaviani \(2008\)](#) study a setting where the agent searches for the best alternative over time and money transfers are used, which we rule out.

[Frankel \(2016\)](#), [Li et al. \(2017\)](#), [Lipnowski and Ramos \(2020\)](#), [Guo and Hörner \(2020\)](#), and [Chen \(2018\)](#) study repeated delegation models in which parties face a stream of

²While in our model the principal dynamically screens the agent's information, we depart from the growing dynamic mechanism design literature ([Pavan et al. 2014](#), [Bergemann and Välimäki 2019](#), [Madsen 2018](#)) by assuming transfers are infeasible.

³[Guo \(2016\)](#) focuses on the full commitment case, but she also shows that the sign of the agent's bias determines the value of commitment.

decisions. In these models, incentives can be provided by linking the different decisions. In contrast, we study situations in which a single, irreversible decision is to be made and, therefore, linking decisions is infeasible.⁴

Finally, our work is related to dynamic persuasion models such as [Ely \(2017\)](#), [McClellan \(2017\)](#), [Henry and Ottaviani \(2019\)](#), and [Orlov et al. \(2020\)](#). These papers explore how to design approval rules when learning is costly, signals are public, and incentives are misaligned ex post. In contrast, we mainly focus on the case where learning is costless, signals are private, and incentives are misaligned ex ante. Our result is reminiscent of [Ely \(2017\)](#), where the delay of information is used to influence the receiver's beliefs so that he is exactly indifferent between working and stopping. In our model, the delay of implementation is used to distort the consequences of the sender's reports so that he is exactly indifferent between truth-telling and lying.

The rest of the paper is organized as follows. [Section 2](#) presents the model. [Section 3](#) formalizes the dynamic delegation problem. [Section 4](#) presents our main results. [Section 5](#) concludes.

2. THE MODEL

We consider an infinite-horizon continuous-time game played by a principal and an agent. There is an initially unknown state $\theta \in \{0, 1\}$. We call $\theta = 1$ the good state and $\theta = 0$ the bad state. At time 0, the agent and the principal are symmetrically uninformed about the state θ , with $\mathbb{P}[\theta = 1] = p_0$ being the initial prior.

The agent privately learns about the state without cost. A signal is generated according to an exponential distribution with arrival rate λ^θ , which depends on θ . Specifically, conditional on θ , over an interval $[t, t + dt]$, a signal $s_t = \theta$ is realized with probability $\lambda^\theta dt$. The arrival of the signal is privately observed by the agent. Thus, the arrival of a signal perfectly reveals the state to the agent. We say that the agent is uninformed if he has not observed a signal. The agent's private history up to period t is denoted h^t . We use $\emptyset$ to denote the history with no signal.

The private belief process $p_t = \mathbb{P}[\theta = 1 | h^t]$ is formed according to the initial prior p_0 and the agent's private history h^t up to period t . The law of motion for the agent's private belief p_t can be derived as follows. If a signal $s_t = 1$ arrives during the interval $[t, t + dt]$, the belief jumps to 1; if a signal $s_t = 0$ arrives during the interval, the belief jumps to 0. If no signal arrives, Bayes's rule can be used to deduce that the posterior at the end of $t + dt$ is

$$p_t + dp_t = \frac{p_t(1 - \lambda^1 dt)}{(1 - p_t)(1 - \lambda^0 dt) + p_t(1 - \lambda^1 dt)}.$$

That is, when no signal arrives, the evolution of the belief is governed by the differential equation⁵

$$\frac{dp_t}{dt} = -(\lambda^1 - \lambda^0)p_t(1 - p_t).$$

⁴Another difference is that, with the exception of [Guo and Hörner \(2020\)](#), the repeated delegation literature has focused on serially uncorrelated incomplete information.

⁵See [Liptser and Shiryaev \(2013\)](#) for details.

We assume that $\lambda^0 < \lambda^1$ and thus no news is bad news. In other words, the belief decreases in the absence of a signal. We show that our results extend to the case $\lambda^0 \geq \lambda^1$ in [Appendix E](#) (also see [Remark 1](#)).

The principal chooses $y_t \in \{0, 1\}$ at each $t \geq 0$, where $y_t = 1$ means to invest and $y_t = 0$ means not to invest. The decision to invest is irreversible: if $y_t = 1$ for some t , then $y_\tau = 1$ for all $\tau > t$ and the interaction ends.

Players' preferences over investment coincide conditional on θ . During each interval $[t, t + dt)$ for which $y_t = 0$, both players receive zero payoff. Conditional on θ , if the principal invests at time t , she gets total discounted payoffs equal to

$$e^{-rt}V \quad \text{if } \theta = 1 \quad \text{and} \quad e^{-rt}(-\nu) \quad \text{if } \theta = 0,$$

whereas the agent gets total discounted payoffs equal to

$$e^{-rt}W \quad \text{if } \theta = 1 \quad \text{and} \quad e^{-rt}(-\omega) \quad \text{if } \theta = 0,$$

where V , ν , W , and ω are strictly positive, and $r > 0$ is the common discount rate.⁶ For the rest of the paper, we normalize $\nu = \omega = 1$, and, therefore, V and W are the gain–loss ratios of the principal and the agent, respectively.

To state our assumption on the conflict of interest, it is useful to describe the one-person benchmark. Suppose that the agent not only perfectly observes the arrival of the signal, but also has the right to invest. The optimal policy for the agent is characterized by a cutoff $p^* := (\lambda^1 + r)/(rW + \lambda^1 + r)$ ([Keller et al. 2005](#)). The agent finds it optimal to invest given the current belief p if and only if he is optimistic enough about the state; that is, $p \geq p^*$. Intuitively, the optimal policy must be a cutoff policy because if the uninformed agent does not find it attractive to invest at t , then neither does the uninformed agent at $t + dt$, who is more pessimistic about the value of the investment than at t . Similarly, suppose that the principal not only controls decisions, but also observes the signal. Given the current belief p , the principal would find it optimal to invest if and only if $p \geq q^* = (\lambda^1 + r)/(rV + \lambda^1 + r)$.

We can now state the assumption on the conflict of interest, which is maintained throughout the paper.

ASSUMPTION 1. *The agent's cutoff belief is lower than the principal's, and the common prior is inbetween. That is, $p^* < p_0 < q^*$.*

This assumption implies that at $t = 0$, the agent wants to invest immediately, whereas the principal wants to invest only after observing a good signal. An equivalent formulation for [Assumption 1](#) is

$$W \frac{r}{\lambda^1 + r} \frac{p_0}{1 - p_0} > 1 > V \frac{r}{\lambda^1 + r} \frac{p_0}{1 - p_0}.$$

⁶One can allow for heterogeneous discount rates. As long as the principal is weakly more patient than the agent, all the results hold qualitatively. When the agent is strictly more patient, the optimal deterministic mechanism is characterized in the same way, but the optimal mechanism is random.

The above inequality means that the gain–loss ratio for the agent, W , is sufficiently high while the gain–loss ratio for the principal, V , is sufficiently low. Note that when [Assumption 1](#) does not hold, the principal can easily align the agent’s incentives.⁷

Since $\lambda^1 > \lambda^0$, as time goes on and no signal is received, the agent gets more pessimistic. At some point, the agent would prefer to wait and invest only after observing the good signal. Let t^* be the time at which the agent becomes indifferent between investing and waiting for a good signal. Formally, for $\lambda^1 > \lambda^0$,

$$t^* = \frac{1}{\lambda^1 - \lambda^0} \ln \left(\frac{p_0}{1 - p_0} W \frac{r}{(\lambda^1 + r)} \right).$$

For $t < t^*$, the principal’s and the agent’s interests are not aligned when no signal has arrived. For $t > t^*$, the principal’s and the agent’s interests coincide for all private histories. We can thus interpret t^* as a measure of how long it takes for the incentives to be aligned. Note that t^* increases as the agent becomes more willing to invest without any information (i.e., when W becomes larger) and as the absence of signal becomes less informative (i.e., when $\lambda^1 - \lambda^0$ becomes smaller so that learning becomes slower).

3. THE DYNAMIC DELEGATION PROBLEM

We set up the principal’s problem of eliciting the agent’s evolving private information to maximize her expected profits. Following the delegation literature ([Holmström 1984](#)), we focus on incentive provision through the design of control rights in the absence of transfers. To do this when learning is private, we formulate a dynamic mechanism design problem with commitment. At each $t \in [0, \infty)$ the agent sends a costless message $m_t \in \{0, 1, \emptyset\}$ given the private history h^t . The principal commits to an action $y_t \in \{0, 1\}$ as a function of the message history up to t .

Given our single-agent setting, it is without loss to restrict to direct mechanisms (see [Sugaya and Wolitzky 2020](#) for details) and restrict the message space so that once the agent announces a signal, the future message space becomes a singleton and the game essentially ends.⁸ A contract is, therefore, a function mapping, for each t , $(\{m_\tau\}_{0 \leq \tau \leq t}, \{y_\tau\}_{0 \leq \tau < t})$ to $y_t \in \{0, 1\}$ with the following irreversibility property: for any t , if $y_\tau = 1$ for some $0 \leq \tau < t$, then $y_t = 1$. In [Appendix A](#), we show that it is without loss to restrict to contracts that specify an investment time for each time at which a good signal is claimed, as well as a deadline at which investment happens even if no signal is claimed. In particular, the principal never invests after the agent reports a bad signal. Therefore, for the rest of the paper, we define a contract as a tuple $\langle T, \tau \rangle$, with $T \in \mathbb{R}_+ \cup \{\infty\}$ and

⁷To see this, note that if $q^* < p_0$, then the principal would like to invest at $t = 0$ and would not need the agent. If $p_0 < p^*$ and $p_0 < q^*$, both the principal and the agent would like to invest only after observing a good signal. In this case, both parties’ preferences are perfectly aligned throughout the game and the first-best can be achieved even without commitment.

⁸Alternatively, one may think that the principal could benefit from allowing the agent to withdraw after announcing a fake good news and receive payoff 0 before investment occurs. In this case, incentive compatibility implies that the optimal contract features a fixed investment time given good news, which makes the principal worse off. Since we assume that the principal has full commitment power, it is without loss to rule out the ability to withdraw.

$\tau: [0, T] \rightarrow \mathbb{R}_+$ if $T < \infty$ while $\tau: \mathbb{R}_+ \rightarrow \mathbb{R}_+$ if $T = \infty$. We use $\text{dom}(\tau)$ to denote the domain of τ . If the agent has reported $m_t = \emptyset$ for all $t \in \text{dom}(\tau)$, the principal invests at time T . The function τ is the time at which the investment is made when the agent reports that he has received the good signal at t ($m_t = 1$).

We now describe the feasibility and incentive constraints. Since time is irreversible, $\tau(t) \geq t$ for all $t \in \text{dom}(\tau)$. To ensure the agent truthfully reveals when he is informed that the state is good at t instead of delaying the report, it must be that $\tau(t)$ is nondecreasing. Otherwise, take $\tau(t_1) > \tau(t_2)$ with $t_1 < t_2$ and note that the agent who receives the good signal at t_1 could wait and report the good signal at $t_2 > t_1$. The principal also needs to ensure that the informed agent at t reveals truthfully instead of pretending to be uninformed during the rest of the game. Formally, $\tau(t) \leq T$ for all $t \in \text{dom}(\tau)$.

A key incentive constraint is to ensure the uninformed agent at t does not want to claim that he is informed and has received a good signal. To ensure truthful revelation of the uninformed agent at t , $\langle T, \tau \rangle$ must satisfy

$$\begin{aligned} & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau(s)} W ds + \left(p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right) \\ & \geq \max\{e^{-r\tau(t)}(p_t W + p_t - 1), 0\} \end{aligned}$$

for all $t \in \text{dom}(\tau)$. Note that the agent can always claim that the state is bad and ensure a payoff equal to 0. The right-hand side is the maximum between 0 and the expected payoff of an uninformed agent at t (who has belief p_t) if he claims the state is good and induces investment at $\tau(t)$. The left-hand side is the agent's expected payoff if he claims to be uninformed and his continuation policy is to report truthfully. In this case, he could receive the good signal at $s < T$ and get the payoff $e^{-r\tau(s)} W$ with conditional probability $p_t \lambda^1 e^{-\lambda^1(s-t)} ds$, or receive no signal before T and induce an uninformed investment decision at T .⁹

The dynamic delegation problem can be formulated as

$$\max_{T \in \mathbb{R}_+ \cup \{\infty\}, \tau(\cdot)} \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} V ds + [p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT}] \quad (1)$$

subject to

$$\tau(t) \geq t \quad \forall t \in \text{dom}(\tau) \quad (2)$$

$$\tau \text{ is nondecreasing} \quad (3)$$

$$\tau(t) \leq T \quad \forall t \in \text{dom}(\tau) \quad (4)$$

$$\begin{aligned} & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau(s)} W ds + \left[p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right] \\ & \geq \max\{e^{-r\tau(t)}(p_t W + p_t - 1), 0\}, \quad t \in \text{dom}(\tau). \end{aligned} \quad (5)$$

⁹This incentive constraint could be considered insufficient as the agent could find it optimal to be truthful in some interval $[t, t + \epsilon]$ and lie after $t + \epsilon$. As we show in [Appendix A](#), this is not the case.

This problem maximizes the principal's expected payoffs (1) over all contracts subject to the feasibility constraint (2) and the dynamic incentive constraints (3)–(5). The dynamic incentive constraints ensure that at any private history, the agent has incentives to truthfully reveal his information.

We can also allow the principal to choose a random timing of investment. In particular, apart from delaying the decision to invest, now the principal can also commit to never invest with some probability. By varying this probability with the time at which the agent claims a good news, the principal can use the threat of no investment to incentivize learning. We show that for any given deadline T , the optimal deterministic mechanism remains optimal. Therefore, it is without loss to restrict to deterministic mechanisms.

As it turns out, the discount factor $e^{-r\tau(t)}$ resulting from the delay function in a deterministic contract can act as probabilities in a random contract. As a result, the optimal random mechanism can be characterized in essentially the same way. After drawing a connection between the “discount factors” from a deterministic and random mechanism, we show that the ability to randomize at T does not benefit the principal. Therefore, we have the following lemma (see [Appendix B](#) for the proof).

LEMMA 1. *The optimal deterministic contract with deadline is weakly better than all random contracts.*

4. ANALYSIS

In this section, we characterize the solution to the dynamic delegation problem.

4.1 Delays

This subsection characterizes the delay with which an investment commonly known to be profitable is implemented. The proofs are relegated to [Appendix C](#). Our first result shows that optimal investments are delayed in any contract that satisfies the dynamic incentive constraints.

LEMMA 2. *Let $\langle T, \tau \rangle$ satisfy (2) and (5). Then $\tau(t) > t$, for all $t < \min\{t^*, T\}$.*

Conditional on the project being revealed profitable at $t < \min\{t^*, T\}$, the implementation time is inefficient (from both the principal's and the agent's perspectives). This distortion arises precisely due to the fact that learning is private: if the implementation time were not distorted and $\tau(t) = t$ for some $t < \min\{t^*, T\}$, the uninformed agent at t would claim he learned that the state is good so as to induce immediate investment.

To solve our dynamic delegation problem, it is useful to find a solution τ to (1) keeping $T \in \mathbb{R}_+ \cup \{\infty\}$ fixed. Solving the dynamic delegation problem for a fixed T is analytically useful and allows us to illustrate the trade-offs involved when delegating to an agent who privately learns over time. The dynamic delegation problem keeping T fixed can be analyzed by finding solutions to the following *relaxed problem* (6). It is obtained

by ignoring constraints (3) and (4), and by imposing the feasibility constraint (2) and the dynamic incentive constraint (5) over subsets of $\text{dom}(\tau)$:

$$\max_{\tau(\cdot)} \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} V ds + [p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT}] \quad (6)$$

subject to

$$\tau(t) \geq t \quad \forall t \geq \min\{t^*, T\} \quad (7)$$

$$\begin{aligned} & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau(s)} W ds + \left[p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right] \\ & \geq \max\{e^{-r\tau(t)}(p_t W + p_t - 1), 0\}, \quad \forall t \leq \min\{t^*, T\}. \end{aligned} \quad (8)$$

The following result establishes a necessary and sufficient optimality condition for the relaxed problem (6).

LEMMA 3. *Let τ satisfy (7) and (8).*

- (a) *Suppose $T \leq t^*$. Then τ solves the relaxed problem if and only if (8) binds for almost every $t \in [0, T]$.*
- (b) *Suppose $T > t^*$. Then τ solves the relaxed problem if and only if (7) binds for almost every $t \geq t^*$ and (8) binds for almost every $t \leq t^*$.*

In the optimal solution to the relaxed problem, the uninformed agent is indifferent between truthful revelation and claiming to know that the state is good for almost all $t \leq \min\{t^*, T\}$. To see this, suppose that τ is optimal and there is a set $A \subseteq [0, \min\{t^*, T\}]$ of positive measure such that for any $t' \in A$, the uninformed agent strictly prefers to reveal the truth. The principal could construct a new function τ' that coincides with τ outside of A but is slightly smaller than τ inside A . Then τ' results in higher expected payoffs for the principal than τ , and it satisfies (7) and (8). Thus, τ cannot be optimal. Moreover, Lemma 3 also shows that for $t > \min\{t^*, T\}$, there is no need to distort investment. Since after t^* the incentives are aligned, delaying investments only makes it harder to provide incentives before t^* .

We now further explore an important consequence of the binding incentive constraint (8) over $[0, \min\{t^*, T\}]$.

LEMMA 4. *Fix T and $\tau(\cdot)$ such that (8) binds for all $t < \min\{t^*, T\}$. Then the derivative of τ with respect to t is given by*

$$\dot{\tau}(t) = \left(\frac{\lambda^0}{r} \right) \frac{1}{W \frac{p_t}{1 - p_t} - 1}$$

for all $t < \min\{t^, T\}$. In particular, over $t < \min\{t^*, T\}$, τ is strictly increasing and convex, and its slope is strictly less than 1.*

Outcomes	$s^t = 1$	$s^t = 0$	$s^t = \emptyset, \theta = 1$	$s^t = \emptyset, \theta = 0$
Lie at t	$e^{-r\tau(t)}W$	$-e^{-r\tau(t)}$	$e^{-r\tau(t)}W$	$-e^{-r\tau(t)}$
Lie at $t + dt$	$e^{-r\tau(t+dt)}W$	0	$e^{-r\tau(t+dt)}W$	$-e^{-r\tau(t+dt)}$
Probabilities	$p_t \lambda^1 dt$	$(1 - p_t) \lambda^0 dt$	$p_t(1 - \lambda^1 dt)$	$(1 - p_t)(1 - \lambda^0 dt)$

Note: Under the first policy (lie at t), the uninformed agent claims that the state is good at t . Under the second policy (lie at $t + dt$), the uninformed agent claims to be uninformed at t , but lies at $t + dt$ if he remains uninformed.

TABLE 1. Payoffs from two different policies.

This lemma characterizes the slope of a timing policy τ when (8) is binding. It can be intuitively derived as follows. Since (8) is binding everywhere in $[0, \min\{t^*, T\})$, the uninformed agent at t is indifferent between claiming he has received the good signal and truth-telling for all $t' \geq t$. The expected payoff the agent gets from truth-telling for $t' \geq t$ can be decomposed into the current and continuation payoffs. Current payoffs are 0, as by declaring truthfully no investment is made at t . For continuation payoffs, note that since the incentive constraint (8) is also binding at $t + dt$, the uninformed agent at $t + dt$ gets the same expected payoff from truth-telling for all $t' \geq t + dt$ and from pretending to have observed the good signal at $t + dt$. Combining these two remarks, the payoff the uninformed agent gets at t from being truthful for all $t' \geq t$ is the same as what he gets from truth-telling at t and lying at $t + dt$. As a result, the uninformed agent at t is indifferent between (i) claiming to have observed the good signal at t (lie at t) and (ii) being truthful at t but lying at $t + dt$ if he is still uninformed (lie at $t + dt$). Table 1 shows the agent's payoffs from both policies for all possible outcomes.

Since the expected payoffs from both policies coincide,

$$e^{-r\tau(t)}(p_t W + p_t - 1) = e^{-r\tau(t+dt)}(p_t W - (1 - p_t)(1 - \lambda^0 dt) + 0 \cdot (1 - p_t)\lambda^0 dt).$$

Equivalently,

$$(1 - p_t)\lambda^0 dt e^{-r\tau(t)} = (e^{-r\tau(t)} - e^{-r\tau(t+dt)})(p_t W - (1 - p_t)(1 - \lambda^0 dt)).$$

Rearranging terms, dividing by dt , and taking $dt \rightarrow 0$, we deduce that

$$(1 - p_t)\lambda^0 = r\dot{\tau}(t)(p_t W + p_t - 1), \quad (9)$$

which provides the characterization in Lemma 4.

Equation (9) illustrates how τ balances the costs and benefits of learning for the agent. The left-hand side in (9) is the benefit from learning, as the agent could avoid investment when the project is bad. The right-hand side in (9) is the cost of learning, as when no signal arrives the investment is just delayed. An important implication from this characterization is that $\dot{\tau}(t) < 1$ and, thus, the delay with which investment decisions are made, $\tau(t) - t$, is decreasing in t . Intuitively, to motivate the agent to learn, the agent's cost of learning has to be lower than that in the single-player benchmark for the agent and, therefore, the principal sets $\dot{\tau}(t) < 1$. In Appendix E, we show that this feature of decreasing delays is robust when $\lambda^0 \geq \lambda^1$.

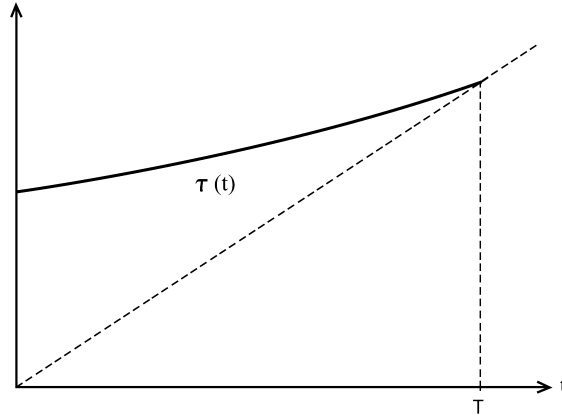


FIGURE 1. The dark line shows the time at which the investment is made as a function of the time at which the good signal is received. For $t < T$, the investment decision is delayed and the delay, $\tau^T(t) - t$, is decreasing. Parameter values: $\lambda^1 = 5$, $\lambda^0 = 4.8$, $r = 0.03$, $p^0 = 0.8$, $v = 120$, $V = 480$, $W = 180$, $T = 2.4$.

4.2 Optimal dynamic delegation

This subsection characterizes the solutions to the optimal delegation problem and establishes the trade-off between the amount of information acquired and how effectively it is used.

We first find a solution τ^T to the relaxed problem when $T \leq t^*$. We impose (8) binding everywhere in $[0, T]$. Note that (8) binding at T gives us

$$\tau^T(T) = T.$$

By Lemma 4, this together with (8) binding in $[0, T)$ gives us

$$\tau^T(t) = T - \frac{\lambda^0}{r} \int_t^T \frac{1}{W \frac{p_s}{1-p_s} - 1} ds, \quad t \leq T. \quad (10)$$

Figure 1 illustrates the solution.

Since $\tau^T(\cdot)$ satisfies the conditions in Lemma 3, it solves the relaxed problem. We now verify that it actually solves the original dynamic delegation problem (1) for a given T . First note that τ^T satisfies (2). Indeed, $\tau^T(t) = \tau^T(T) - \int_t^T \dot{\tau}^T(s) ds$ and, since $\tau^T(T) = T$ and $\dot{\tau}^T(t) < 1$, $\tau^T(t) \geq \tau^T(T) - (T - t) = t$ for all $t \in [0, T]$. Second, τ^T satisfies (5) because it holds with equality over $t \in [0, T]$. Finally, since $\tau^T(t)$ is increasing over $[0, T]$ and $\tau^T(T) = T$, τ^T also satisfies the incentive constraints (3) and (4). As a result, τ^T indeed solves the dynamic delegation problem (1) for a given T . As can be seen, the incentive constraint for the uninformed agent is the key to pinning down the optimal contract when $T \leq t^*$.

The following result provides a key insight for solving for the optimal $T < t^*$.

PROPOSITION 1. *Let $t < T < \hat{T} < t^*$. Then $\tau^T(t) < \tau^{\hat{T}}(t)$.*

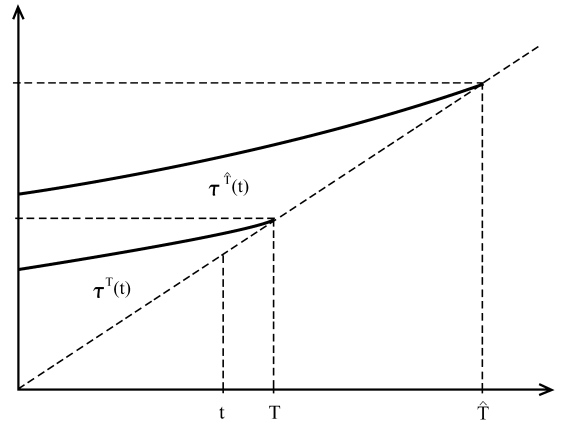


FIGURE 2. For $t \leq T < \hat{T}$, $\tau^T(t) < \tau^{\hat{T}}(t)$. Parameter values: $\lambda^1 = 5$, $\lambda^0 = 4.8$, $r = 0.03$, $p^0 = 0.8$, $v = 120$, $V = 480$, $W = 180$, $T = 2.4$, $\hat{T} = 4.3$.

Figure 2 illustrates Proposition 1. Increasing the deadline is beneficial for the principal in that more information is acquired and, thus, investment in the bad state is less likely to happen. Proposition 1 shows that more learning imposes a nontrivial incentive cost on the principal because when T increases, $\tau^T(t)$ must increase too. This means that when T increases, investments are delayed more when the good signal is received.

Formally, Proposition 1 follows immediately from (10). To better understand Proposition 1, take $t < T < \hat{T}$ and assume for the moment that t is close to T . When the uninformed agent at t faces the contract $\langle T, \tau^T \rangle$, he knows that by declaring truthfully, the investment will be made at T (unless a bad signal is received in the meanwhile). Now, when the uninformed agent at t faces the contract $\langle \hat{T}, \tau^{\hat{T}} \rangle$, the earliest time at which the investment could be made is $\tau^{\hat{T}}(T) > T$. As a result, the expected continuation payoff that the uninformed agent gets at t by being truthful is lower when he faces $\langle \hat{T}, \tau^{\hat{T}} \rangle$ than when he faces $\langle T, \tau^T \rangle$. Therefore, to provide incentives for truthful revelation at t , contract $\langle \hat{T}, \tau^{\hat{T}} \rangle$ must punish the agent even more when he claims a good signal. In other words, $\tau^{\hat{T}}(t) > \tau^T(t)$. This intuition can be iteratively applied backward to render this property for all $t < T$.

We now solve the relaxed problem given $T > t^*$ by imposing (8) binding everywhere in $[0, t^*)$ and (7) binding everywhere in $[t^*, T]$. By Lemma 4, (8) binding in $[0, t^*)$ combined with (7) binding for $t \geq t^*$ gives us

$$\tau^T(t) = \begin{cases} t^* - \frac{\lambda^0}{r} \int_t^{t^*} \frac{1}{W \frac{p_s}{1-p_s} - 1} ds & t \leq t^* \\ t & t > t^*. \end{cases}$$

Last, to make sure that τ^T satisfies (8) at t^* and, therefore, solves the relaxed problem, we need T to be infinity. To see this, notice that at t^* , by revealing truthfully that he has not received a signal, the agent receives the payoff from the policy “invest as soon as a good signal arrives before T and invest at T if no signal arrives before T ,” which is weakly

less preferred to the policy “invest as soon as a good signal arrives and do not invest if no signal arrives.” Since at t^* the agent is indifferent between the latter policy and the policy “invest right away,” we need $T = \infty$ to ensure incentive compatibility. We have shown that $\tau^\infty(\cdot)$ is a solution to the relaxed problem. Moreover, since $\tau^\infty(\cdot)$ is increasing and (5) is satisfied everywhere in $[0, \infty)$, it solves the original dynamic delegation problem (1) given that $T > t^*$.

The following theorem summarizes our characterization.

THEOREM 1. *The optimal contract takes one of the following two forms:*

- (a) *There is a deadline $T < t^*$. If a good signal arrives before T , investment happens with a delay. If no signal arrives before T , investment happens at T .*
- (b) *There is no deadline. If a good signal arrives before t^* , investment happens with a delay. If a good signal arrives after t^* , investment happens with no delay.*

To find the optimal contract $\langle T^*, \tau^{T^*} \rangle$, it suffices to compare the optimal solution when $T \in [0, t^*]$ to the case in which $T = \infty$. It is thus enough to compare the expected payoff for the principal from the optimal τ^T when $T \leq t^*$ to that from τ^∞ .

The optimal contract can be implemented by setting time-dependent delegation sets. At any $t < \min\{t^*, T^*\}$, the agent is allowed to commit to invest in $[\tau^{T^*}(t), \infty)$ or just wait and commit later. For $t \geq \min\{t^*, T^*\}$, the agent is granted full freedom.

REMARK 1. When $\lambda^0 = \lambda^1$, the agent’s belief remains constant given no news. Therefore, the uninformed agent is never indifferent between investing and waiting; that is, $t^* = \infty$. In this case, we show that the optimal contract always features a finite deadline and $\tau(\cdot)$ is linear. When $\lambda^0 > \lambda^1$, the agent’s belief drifts up as time goes on. In this case, there is a t^* after which even the principal would like to invest. The deadline is also finite in this case, and $\tau(\cdot)$ is concave. In both cases, the decreasing delay feature remains (see [Appendix E](#) for details).

4.3 Comparative statics

We now derive some comparative statics results. These results assume that parameters satisfy [Assumption 1](#), that is, $(Wr p_0)/((\lambda^1 + r)(1 - p_0)) > 1 > (Vr p_0)/((\lambda^1 + r)(1 - p_0))$.

PROPOSITION 2. (a) *Fix all parameters except W . There exist cutoffs $0 < \underline{\kappa} < \bar{\kappa}$ such that for all $W < \underline{\kappa}$, the optimal contract sets no deadline, whereas for $W > \bar{\kappa}$, the optimal contract sets a deadline $T^* < t^*$.*

- (b) *Fix all parameters except V . There exist cutoffs $0 < \underline{\eta} < \bar{\eta}$ such that for all $V < \underline{\eta}$, the optimal contract sets no deadline, whereas for $V > \bar{\eta}$, the optimal contract sets a deadline $T^* < t^*$.*

Part (a) shows that when W is sufficiently small, it is optimal to invest only after the principal has perfectly learned that $\theta = 1$. In this case, t^* is small, so the incentives become aligned rapidly and there is no need to significantly delay investments for $t < t^*$. In

contrast, when W is large, t^* is large and the conflict of interest is severe. So decisions to invest need to be significantly distorted for $t < \min\{T^*, t^*\}$. To economize on distortions early in the game, the principal commits to invest at a deadline $T^* < t^*$ even when this means learning stops early.

Part (b) characterizes the solutions as we vary the principal's payoff. When V is small, it is relatively costly for the principal to invest when the state is bad. To avoid the costs of a failure, the principal prefers to perfectly learn the state even when this entails significant delays for $t < t^*$. In contrast, when V is large, the cost of a failure is relatively small, and the principal fixes a deadline $T^* < t^*$ that stops learning and reduces distortions for $t < T^*$.¹⁰

Proposition 2 sheds light on how different organizations should provide incentives for learning. For example, the FDA incurs significant costs when approving bad drugs. Our results suggest that the FDA should set lengthy revision processes to ensure pharmaceutical companies learn the value of the drugs even if this entails substantial delays between the drugs' discovery and the FDA's final approval. In contrast, the board of a company that is contemplating a partially reversible acquisition or that cannot align the manager's career incentives should set a deadline $T^* < t^*$ that facilitates truthful communication even at the possible cost of an incorrect decision.

5. CONCLUDING REMARKS

We study a dynamic delegation model in which learning is private. Evolving private information shapes the optimal contract in distinctive ways. Indeed, we show that to ensure truthful revelation from the agent, the principal needs to delay investment commonly known to be optimal. As time goes on, the principal grants more flexibility to the agent and eventually the agent is free to make any decision. Our analysis uncovers a new trade-off between the amount of information acquired and how promptly it is used.

One may be interested in what happens if transfers are allowed. First, one can easily see that with unlimited transfers, the principal can perfectly align the incentives through charging the agent for investment as well as at the beginning of the relationship. With limited liability, however, one can show that transfers rewarding good news exacerbate the incentive problem and, therefore, are never optimal.

Our model is stylized. The learning process is assumed to be Poisson,¹¹ the principal has full commitment power, investment is irreversible, and the agent has little freedom to decide how to learn.¹² The model could also be extended to allow for money burning.¹³ Our dynamic delegation model with private learning can also be used as a workhorse to explore applied issues in political economy, finance, and organizational economics. We leave these research projects for future work.

¹⁰As W has a direct impact on τ as well as an indirect impact on τ through t^* , it is not obvious that the payoff difference between the optimal finite deadline policy and the policy with no deadline is monotonic in W . Therefore, it is not clear how to show that the cutoffs coincide.

¹¹Exploring a model with Brownian learning would be interesting, but evolving private information makes the problem hard to analyze. When learning is Brownian, delays and deadlines are likely to play a role, but the contract may have additional features.

¹²At the other extreme, the agent could decide any experiment that reveals information about the state.

¹³This means that the agent can spend resources that have no value for the principal.

APPENDIX

This Appendix consists of five parts. [Appendix A](#) justifies our mechanism design formulation. [Appendix B](#) provides the proof for [Lemma 1](#). [Appendix C](#) provides proofs for [Section 4.1](#). [Appendix D](#) provides proofs for [Section 4.3](#). [Appendix E](#) extends the information structure to no news is no news and no news is good news.

APPENDIX A: FORMULATION OF MECHANISM DESIGN PROBLEM

Our objective in this section is to show that it is without loss to represent a mechanism as $\langle T, \tau \rangle$ and, therefore, the principal's problem can be written as in (1)–(5).

First we define a contract. To do this, we first need some terminology. A *public history at t* is $h_t^p := \{(m_\tau, y_\tau)\}_{0 \leq \tau < t} \in H_t^p$. It contains the sequence of messages and investment decisions strictly before t . A *private history at t* is $h_t^\alpha := \{ \{(m_\tau, y_\tau)\}_{0 \leq \tau < t}, \{h_\tau\}_{0 \leq \tau \leq t} \} \in H_t^\alpha$. It contains the sequence of signal, messages, and investment decisions strictly before t as well as the signal observation at t . The message space at time t , $M_t : H_t^p \rightarrow 2^{\{0,1,\emptyset\}}$, is defined $\forall t$ as

$$M_t(h_t^p) = \begin{cases} \{m_\tau\} & \text{if } \exists \tau < t \text{ s.t. } m_\tau \neq \emptyset \text{ or } y_\tau = 1 \\ \{0, 1, \emptyset\} & \text{otherwise.} \end{cases}$$

We use m^t to denote the message sequence up to and including t : $m^t = \{m_\tau\}_{0 \leq \tau \leq t}$. Through abuse of notation, we say that $m^t = \emptyset$ if $m_\tau = \emptyset$ for all $0 \leq \tau \leq t$.

A contract Γ is a function mapping $\{ \{m_\tau\}_{0 \leq \tau \leq t}, \{y_\tau\}_{0 \leq \tau < t} \}$ to $y_t \in \{0, 1\}$ with the following irreversibility property: for any t , if $y_\tau = 1$ for some $0 \leq \tau < t$, then $y_t = 1$. From now on, we keep in mind this property and omit the dependence of y_t on $\{y_\tau\}_{0 \leq \tau < t}$ and simply write $y_t = y(m^t)$.

Our next goal is to simplify the principal's problem. To do this, we first show that any contract can be represented by three components: a “deadline” $\bar{T} \in \mathbb{R}_+ \cup \{\infty\}$, a function $\tau_0(\cdot)$ that maps the arriving time of the first 0 message to an investment time $\mathbb{R}_+ \cup \{\infty\}$, and a function $\tau_1(\cdot)$ that maps the arriving time of the first 1 message to an investment time $\mathbb{R}_+ \cup \{\infty\}$.

Let us define $\bar{T} := \inf\{t : m^t = \emptyset, y(m^t) = 1\}$. It follows that for any t such that $m^t = \emptyset$, if $t < \bar{T}$, then $y(m^t) = 0$; otherwise $y(m^t) = 1$. In other words, $\bar{T}$ pins down the principal's action for an empty message history of any length. Now let us consider $m^t \neq \emptyset$. We define $\gamma(m^t) := \min\{\tau : m_\tau \neq \emptyset\}$, the time that a nonempty message history jumps from $\emptyset$ to 0 or 1. Then by the definition of M_t , any $m^t \neq \emptyset$ is completely characterized by $\gamma(m^t)$, the value of $m_t \in \{0, 1\}$, and the value of t . For each $x \geq 0$, let us define $\tau_0(x) = \inf\{t : y(m^t) = 1, m_t = 0, \gamma(m^t) = x\}$. Therefore, for each m^t for which $m_t = 0$ and $t < \tau_0(\gamma(m^t))$, we have $y(m^t) = 0$; for each m^t for which $m_t = 0$ and $t \geq \tau_0(\gamma(m^t))$, $y(m^t) = 1$. Similarly, $\tau_1(\cdot) := \inf\{t : y(m^t) = 1, m_t = 1, \gamma(m^t) = x\}$ pins down $y(m^t)$ for all $m^t \neq \emptyset$ and $m_t = 1$. Therefore, $\tau_0(\cdot)$ and $\tau_1(\cdot)$ pin down the principal's action at any nonempty message history. Note that since the infimum of an empty set is $+\infty$, we allow the case that $\bar{T} = \infty$ or $\tau_i(x) = \infty$ for $i = 0, 1$ for some x .

Right now the domains of $\tau_0(\cdot)$ and $\tau_1(\cdot)$ are both $[0, \infty)$. We argue that it suffices to restrict them to $[0, \bar{T}]$ whenever $\bar{T} < \infty$. In other words, it is redundant to define the investment time for a history $m^t \neq \emptyset$ for which $\gamma(m^t) > \bar{T}$. The argument is simple: if for some m^t we have $\gamma(m^t) > \bar{T}$, then it must be the case that $m^{\bar{T}} = \emptyset$ and $y(m^{\bar{T}}) = y(m^t) = 1$. We sum up our discussion in the following proposition.

PROPOSITION 3. *A contract belongs to one of the groups*

- (a) $\bar{T} < \infty$, $\tau_0, \tau_1 : [0, \bar{T}] \rightarrow [0, \infty]$
- (b) $\bar{T} = \infty$, $\tau_0, \tau_1 : [0, \bar{T}) \rightarrow [0, \infty]$.

Now we have demonstrated that a contract consists of three components $\bar{T}$, $\tau_0(\cdot)$, and $\tau_1(\cdot)$. We are now ready to state the principal's objective function:

$$\begin{aligned} & \int_0^{\bar{T}} [p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_1(s)} V - (1 - p_0) \lambda^0 e^{-\lambda^0 s} e^{-r\tau_0(s)}] ds \\ & + [p_0 e^{-\lambda^1 \bar{T}} + (1 - p_0) e^{-\lambda^0 \bar{T}}] e^{-r\bar{T}} (p_{\bar{T}} V + p_{\bar{T}} - 1). \end{aligned}$$

For the constraints faced by the principal, first note that the principal's actions must be feasible; therefore, $\tau_i(x) \geq x$ for all x . For the incentive compatibility constraints of the agent, we require that at any on- or off-path history, the agent prefers to tell the truth from then on. Hence, we discuss the possible histories h_t^α faced by the agent at which M_t is not a singleton:

Case 1. Suppose that h^t contains signal 1. Then choosing $m_t = 1$ is preferred by the agent to

- (a) choosing $m_t = 0$: $e^{-r\tau_1(t)} W \geq e^{-r\tau_0(t)} W$
- (b) choosing $m_t = \emptyset$ and $m_s = 1$ for some $s > t$: $e^{-r\tau_1(t)} W \geq e^{-r\tau_1(s)} W \forall s > t$
- (c) choosing $m_t = \emptyset$ and $m_s = 0$ for some $s > t$: $e^{-r\tau_1(t)} W \geq e^{-r\tau_0(s)} W \forall s > t$
- (d) choosing $m_s = \emptyset$ for all $s \geq t$.

That is,

- (a) $\tau_1(t) \leq \tau_0(t) \forall t$
- (b) $\tau_1(t) \leq \tau_1(s) \forall s > t$
- (c) $\tau_1(t) \leq \tau_0(s) \forall s > t$
- (d) $\tau_1(t) \leq \bar{T} \forall t$.

Case 2. Suppose that h^t contains signal 0. Then choosing $m_t = 0$ is preferred by the agent to

- (a) choosing $m_t = 1$: $-e^{-r\tau_0(t)} \geq -e^{-r\tau_1(t)}$

- (b) choosing $m_t = \emptyset$ and $m_s = 0$ for some $s > t$: $-e^{-r\tau_0(t)} \geq -e^{-r\tau_0(s)} \forall s > t$
- (c) choosing $m_t = \emptyset$ and $m_s = 1$ for some $s > t$: $-e^{-r\tau_0(t)} \geq -e^{-r\tau_1(s)} \forall s > t$
- (d) choosing $m_s = \emptyset$ for all $s \geq t$.

That is,

- (a) $\tau_0(t) \geq \tau_1(t) \forall t$
- (b) $\tau_0(t) \geq \tau_0(s) \forall s > t$
- (c) $\tau_0(t) \geq \tau_1(s) \forall s > t$
- (d) $\tau_0(t) \geq \bar{T} \forall t$.

Case 3. Suppose that $h^t = \emptyset$. Then truth-telling forever from now on maximizes the agent's expected payoff. That is, the agents expected payoff from using the truth-telling strategy given that $h^t = \emptyset$,

$$\begin{aligned} V(t) = & \left(\int_t^{\bar{T}} e^{-\lambda^1 s - r\tau_1(s)} ds \right) p_t \lambda^1 e^{(r+\lambda^1)t} W \\ & - \left(\int_t^{\bar{T}} e^{-\lambda^0 s - r\tau_0(s)} ds \right) (1 - p_t) \lambda^0 e^{(r+\lambda^0)t} \\ & + p_t e^{(r+\lambda^1)(t-\bar{T})} W - (1 - p_t) e^{(r+\lambda^0)(t-\bar{T})}, \end{aligned}$$

satisfies

$$\begin{aligned} U(t) = & \max \{ e^{-r[\tau_1(t)-t]} (p_t W + p_t - 1), e^{-r[\tau_0(t)-t]} (p_t W + p_t - 1), \\ & e^{-r dt} p_t \lambda^1 dt e^{-r[\tau_1(t+dt)-t-dt]} W - e^{-r dt} (1 - p_t) \lambda^0 dt e^{-r[\tau_0(t+dt)-t-dt]} \\ & + e^{-r dt} [1 - p_t \lambda^1 dt - (1 - p_t) \lambda^0 dt] U(t + dt) \}. \end{aligned}$$

The first and second terms denote the agent's expected payoff if he chooses $m_t = 1$ and $m_t = 0$, respectively. Both actions essentially end the game and there is no need to specify future actions. The last term denotes the agent's expected payoff if he chooses $m_t = \emptyset$ and the optimal action at $t + dt$. The first component is the agent's payoff if he gets a 1 signal during $(t, t + dt)$. The incentive-compatibility (IC) conditions in Case 1 ensure that the optimal action is to choose $m_{t+dt} = 1$ in this case, which leads to an investment time $\tau_1(t + dt)$. The second component is the agent's payoff if he gets a 0 signal during $(t, t + dt)$. The third component is the agent's payoff if he receives no signal during $(t, t + dt)$.

The next lemma simplifies the incentive condition at $h^t = \emptyset$.

LEMMA 5. Suppose that a contract $\langle \bar{T}, \tau_0, \tau_1 \rangle$ satisfies IC at any history h_t^α for which $h^t \neq \emptyset$. Moreover, suppose that at h_t^α for which $h^t = \emptyset$, the following relationship holds:

$$e^{-rt} V(t) \geq \max \{ e^{-r\tau_0(t)} (p_t W + p_t - 1), e^{-r\tau_1(t)} (p_t W + p_t - 1) \}.$$

Then the strategy of truth-telling at every history also maximizes the agent's expected payoff at any h_t^α for which $h^t = \emptyset$.

PROOF. Fix an arbitrary h_t^α for which $M_t(h_t^\rho) = \{0, 1, \emptyset\}$ and $h^t = \emptyset$. Let σ^* denote the strategy of truth-telling at every history, whenever doing so is possible. Let σ denote an alternative strategy such that either $\sigma(h_t^\alpha) \neq \sigma^*(h_t^\alpha)$ or there exists a concatenation history of h_t^α, h_s^α , such that $s > t$ and $\sigma(h_s^\alpha) \neq \sigma^*(h_s^\alpha)$.

If $\sigma(h_t^\alpha) \neq \emptyset$, then by inequality (5), σ^* renders higher payoff than σ .

If $\sigma(h_t^\alpha) = \emptyset$, then take a concatenation history of h_t^α for which $\sigma^*(h_s^\alpha) \neq \sigma(h_s^\alpha)$. Note that at h_s^α , $m_\tau = \emptyset$ for all $\tau < s$; otherwise, $M_s(h_s^\rho)$ is a singleton. Moreover, $y_\tau = 0$ for all $\tau < s$; otherwise, a decision is already made at $\bar{T}$ and $M_s(h_s^\rho)$ is again a singleton. Therefore, the agent's cumulative payoff during $[0, s]$ equals 0 for both σ^* and σ . Now, if $h^s \neq \emptyset$, σ^* renders higher payoff since the contract is incentive compatible at such a history. If $h^s = \emptyset$, then $\sigma^*(h_s^\alpha) = \emptyset$ while $\sigma(h_s^\alpha) \in \{0, 1\}$. By inequality (5), σ^* still renders higher payoff. We have thus shown that σ^* renders higher payoff at any future (on- or off-path) history h_s^α for which σ^* and σ differ. Therefore, σ^* renders higher expected payoff than σ at the information set h_t^α . $\square$

Now we characterize the optimal contract. First we notice that any incentive-compatible optimal contract $\langle \bar{T}^*, \tau_0^*, \tau_1^* \rangle$ must have $\tau_0^* = \infty$ almost surely.

PROPOSITION 4. *Given an incentive-compatible optimal contract $\langle \bar{T}^*, \tau_0^*, \tau_1^* \rangle$, let us define $A := \{t : \tau_0^*(t) < \infty\}$. Then A has measure 0.*

PROOF. First notice that if $\bar{T}^* = \infty$, then IC requires that $\tau_0^*(t) = \infty$ for all t . So for the rest of the proof, let us assume that $\bar{T}^* < \infty$. By way of contradiction, suppose that A has positive measure. Then the part of the principal's payoff involving $\tau_0^*(\cdot)$ can be rewritten as

$$\begin{aligned} & \int_0^{\bar{T}^*} -e^{-r\tau_0^*(s)}(1-p_0)\lambda^0 e^{-\lambda^0 s} ds \\ &= \int_A -e^{-r\tau_0^*(s)}(1-p_0)\lambda^0 e^{-\lambda^0 s} ds + \int_{[0, \bar{T}^*] \setminus A} -e^{-r\tau_0^*(s)}(1-p_0)\lambda^0 e^{-\lambda^0 s} ds \\ &= \int_A -e^{-r\tau_0^*(s)}(1-p_0)\lambda^0 e^{-\lambda^0 s} ds \\ &< 0. \end{aligned}$$

Let $\bar{t}$ be such that $p_{\bar{t}}W - (1-p_{\bar{t}}) = 0$. We propose an alternative contract depending on whether $\bar{T}^*$ is greater or less than $\bar{t}$.

Option 1: $\bar{T}^* \leq \bar{t}$. Consider an alternative contract $\langle \bar{T}^*, \tilde{\tau}_0, \tau_1^* \rangle$, where

$$\tilde{\tau}_0(s) = \infty \quad \forall s.$$

Since the principal's payoff involving $\tilde{\tau}_0$ is 0, this contract strictly increases the principal's payoff. Now we show that $\langle \bar{T}^*, \tilde{\tau}_0, \tau_1^* \rangle$ is incentive compatible, contradicting that $\langle \bar{T}^*, \tau_0^*, \tau_1^* \rangle$ is a solution.

It is obvious that the IC conditions when $h^t \neq \emptyset$ (i.e., Cases A and A) are still satisfied for the new contract. For the case when $h^t = \emptyset$, notice that under the new contract, the agent's payoff from truth-telling forever from time t on is

$$\begin{aligned}
 & [p_t e^{-\lambda^1(\bar{T}^*-t)} + (1-p_t) e^{-\lambda^0(\bar{T}^*-t)}] e^{-rT^*} [p_{\bar{T}^*} W - (1-p_{\bar{T}^*})] \\
 & + \int_t^{\bar{T}^*} (1-p_t) \lambda^0 e^{-\lambda^0(s-t)} e^{-r\tau_0^*(s)} \cdot 0 ds + \int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds \\
 & \geq [p_t e^{-\lambda^1(\bar{T}^*-t)} + (1-p_t) e^{-\lambda^0(\bar{T}^*-t)}] e^{-rT^*} [p_{\bar{T}^*} W - (1-p_{\bar{T}^*})] \\
 & - \int_t^{\bar{T}^*} (1-p_t) \lambda^0 e^{-\lambda^0(s-t)} e^{-r\tau_0^*(s)} ds + \int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds \\
 & \geq (p_t W + p_t - 1) e^{-r\tau_1^*(t)} \\
 & \geq 0.
 \end{aligned}$$

The second inequality follows from the fact that the old contract satisfies IC. Therefore, truth-telling forever from t on is preferred to lying that $m_t = 1$. The last inequality follows because for any $t \leq \bar{T}^* \leq \bar{t}$, $p_t W + p_t - 1 \geq 0$. Therefore, truth-telling forever from t on is preferred to lying that $m_t = 0$. We have just established that under the new contract, truth-telling forever from t on is preferred to lying at t . By Lemma 5, this ensures that the new contract is incentive compatible.

Option 2: $\bar{T}^* > \bar{t}$. Since p_t decreases in t , $p_{\bar{T}^*} W - (1-p_{\bar{T}^*}) < 0$. Consider an alternative contract $(\tilde{T}, \tilde{\tau}_0, \tilde{\tau}_1)$, where

$$\tilde{T} = \infty, \quad \tilde{\tau}_0(s) = \infty \quad \forall s$$

and

$$\tilde{\tau}_1(t) = \begin{cases} \tau_1^*(t) & \text{if } t \leq \bar{T}^* \\ t & \text{otherwise.} \end{cases}$$

The agent's payoff under this contract equals

$$\begin{aligned}
 & \int_0^{\bar{T}^*} (1-p_0) \lambda^0 e^{-\lambda^0 s} \cdot 0 ds + \int_0^{\bar{T}^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_1^*(s)} W ds \\
 & + [p_0 e^{-\lambda^1 \bar{T}^*} + (1-p_0) e^{-\lambda^0 \bar{T}^*}] \cdot \int_{\bar{T}^*}^{\infty} p_{\bar{T}^*} \lambda^1 e^{-\lambda^1(s-\bar{T}^*)} e^{-rs} W ds \\
 & > \int_0^{\bar{T}^*} -(1-p_0) \lambda^0 e^{-\lambda^0 s} e^{-r\tau_0^*(s)} ds + \int_0^{\bar{T}^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_1^*(s)} W ds \\
 & + [p_0 e^{-\lambda^1 \bar{T}^*} + (1-p_0) e^{-\lambda^0 \bar{T}^*}] \cdot e^{-r\bar{T}^*} [p_{\bar{T}^*} W - (1-p_{\bar{T}^*})].
 \end{aligned}$$

The inequality follows because

$$\int_{\bar{T}^*}^{\infty} p_{\bar{T}^*} \lambda^1 e^{-\lambda^1(s-\bar{T}^*)} e^{-rs} W ds > 0 > e^{-r\bar{T}^*} [p_{\bar{T}^*} W + (1-p_{\bar{T}^*})(-\omega)]$$

and A has positive measure by assumption. Now we show the new contract is incentive compatible. First, it is easy to see that at $h^t \neq \emptyset$ and $t \leq \bar{T}^*$, IC are satisfied. Second, at any $t > \bar{T}^* > t^*$, the interests of the principal and the agent are aligned. Therefore, the first-best action is incentive compatible. Last, at any $h^t = \emptyset$ and $t \leq \bar{T}^*$, the agent's payoff if he is truth-telling since then on equals

$$\int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds + \int_{\bar{T}^*}^{\infty} p_{\bar{T}^*} \lambda^1 e^{-\lambda^1(s-\bar{T}^*)} e^{-rs} W ds \geq 0.$$

Therefore, the payoff of truth-telling is greater than the payoff of lying that $m_t = 0$.

Alternatively, the payoff of lying that $m_t = 1$ is

$$\begin{aligned} & [p_t W - (1 - p_t)] e^{-r\tau_1^*(t)} \\ & \leq [p_t e^{-\lambda^1(\bar{T}^*-t)} + (1 - p_t) e^{-\lambda^0(\bar{T}^*-t)}] e^{-r\bar{T}^*} [p_{\bar{T}^*} W - (1 - p_{\bar{T}^*})] \\ & \quad + \int_t^{\bar{T}^*} -(1 - p_t) \lambda^0 e^{-\lambda^0(s-t)} e^{-r\tau_0^*(s)} ds + \int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds \\ & \leq \int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds \\ & \leq \int_t^{\bar{T}^*} p_t \lambda^1 e^{-\lambda^1(s-t)} e^{-r\tau_1^*(s)} W ds + \int_{\bar{T}^*}^{\infty} p_{\bar{T}^*} \lambda^1 e^{-\lambda^1(s-\bar{T}^*)} e^{-rs} W ds, \end{aligned}$$

which is the agent's payoff of truth-telling. The first inequality follows because the contract $\langle \bar{T}^*, \tau_0^*, \tau_1^* \rangle$ is incentive compatible. The second inequality follows because $p_{\bar{T}^*} W - (1 - p_{\bar{T}^*}) < 0$. Applying Lemma 5 again, we know that the new contract is incentive compatible. $\square$

It is easy to argue the following statement.

PROPOSITION 5. *Given an incentive-compatible optimal contract $\langle \bar{T}^*, \tau_0^*, \tau_1^* \rangle$ for which $\bar{T}^* < \infty$,*

$$\tau_0^*(t) = \infty \quad \forall t < \bar{T}^*.$$

PROOF. Suppose that $\tau_0^*(t) < \infty$ for some $t < \bar{T}^*$. Then at h_t^α , which includes a 0 signal, the agent could deviate to $m_\tau = \emptyset$ for all $t \leq \tau < s$ and $m_s = 0$ for some $s > t$. By Proposition 4, such s must exist. $\square$

We have shown that $\tau_0^*(t) = \infty$ for all t with the possible exception of $\tau_0^*(\bar{T}^*)$ when $\bar{T}^* < \infty$. We set $\tau_0(\bar{T}^*) = \bar{T}^*$ automatically whenever $\bar{T}^* < \infty$. This ensures that for any $\bar{T}^* < \infty$, truth-telling is incentive compatible at $\bar{T}^*$ and at $t < \bar{T}^*$, and $h^t \neq \emptyset$.

Now we are ready to rewrite the principal's constrained maximization problem as

$$\max_{\bar{T} \in \mathfrak{N}_+ \cup \{\infty\}, \tau(\cdot)} \int_0^{\bar{T}} p \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} V ds + [p e^{-\lambda^1 \bar{T}} + (1 - p) e^{-\lambda^0 \bar{T}}] e^{-r\bar{T}} (p_{\bar{T}} V + p_{\bar{T}} - 1)$$

subject to

$$\begin{aligned}
 \tau(t) &\geq t \quad \forall t \in [0, \bar{T}] \\
 \tau(s) &\geq \tau(t) \quad \forall s \geq t \\
 \tau(t) &\leq \bar{T} \quad \forall t \in [0, \bar{T}] \\
 \int_t^{\bar{T}} p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau(s)} W ds &+ \left[p_t \frac{e^{-\lambda^1 \bar{T}}}{e^{-\lambda^1 t}} e^{-r\bar{T}} W - (1 - p_t) \frac{e^{-\lambda^0 \bar{T}}}{e^{-\lambda^0 t}} e^{-r\bar{T}} \right] \\
 &\geq \max\{e^{-r\tau(t)}(p_t W + p_t - 1), 0\} \quad \forall t \in [0, \bar{T}].
 \end{aligned}$$

APPENDIX B: PROVING LEMMA 1

PROOF OF LEMMA 1. Following Pavan et al. (2014), we focus on random mechanisms that do not allow the agent to update his belief about the outcomes of the randomization until the game ends. In particular, we study the following family of contracts (which we call the *simple family*): let $T \in \mathbb{R}_+ \cup \{\infty\}$ be deterministic. At any $t \leq T$, if the agent announces $m_t = 1$, the principal chooses the investment time according to $\beta_\tau(\cdot | t)$, which is the probability measure from the space $([t, \infty], \mathcal{B}([t, \infty]), \beta_\tau(\cdot | t))$. In addition, if $T < \infty$ and the agent never announces $m_t = 1$ or $m_t = 0$ for any $t \leq T$, the principal chooses the investment time according to β_T , which is the probability measure from the space $([T, \infty], \mathcal{B}([T, \infty]), \beta_T(\cdot))$. Once the investment time has been determined by β_τ or β_T , the game ends.

The principal's constrained maximization problem is

$$\begin{aligned}
 \max_{T \in \mathbb{R}_+ \cup \{\infty\}, \beta_\tau(\cdot | t), \beta_T(\cdot)} &\int_0^T p_0 \lambda^1 e^{-\lambda^1 s} \left[\int_s^\infty e^{-r\tau} \beta_\tau(d\tau | s) \right] V ds \\
 &+ \int_T^\infty e^{-r\tau} \beta_T(d\tau) [p_0 e^{-\lambda^1 T} V - (1 - p_0) e^{-\lambda^0 T}]
 \end{aligned}$$

subject to

$$W \left[\int_t^\infty e^{-r\tau} \beta_\tau(d\tau | t) \right] \geq W \left[\int_s^\infty e^{-r\tau} \beta_\tau(d\tau | s) \right] \quad \forall s \geq t \quad (11)$$

$$W \left[\int_t^\infty e^{-r\tau} \beta_\tau(d\tau | t) \right] \geq W \int_T^\infty e^{-r\tau} \beta_T(d\tau) \quad \forall t \quad (12)$$

$$\begin{aligned}
 \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} W \left[\int_s^\infty e^{-r\tau} \beta_\tau(d\tau | s) \right] ds \\
 + \int_T^\infty e^{-r\tau} \beta_T(d\tau) [p_t e^{-\lambda^1 (T-t)} W - (1 - p_t) e^{-\lambda^0 (T-t)}] \\
 \geq \max \left\{ \left[\int_t^\infty e^{-r\tau} \beta_\tau(d\tau | t) \right] \cdot (p_t W + p_t - 1), 0 \right\} \quad \forall t. \quad (13)
 \end{aligned}$$

We fix T , ignore conditions (11) and (12), and argue that condition (13) should bind for almost all $t < t^*$. Suppose the strict inequality holds for all $t \in \mathcal{A}$, where $\mathcal{A}$ has positive

measure. Let us consider

$$\beta_\tau^\epsilon(\cdot | t) = \begin{cases} (1 - \epsilon)\beta_\tau(\cdot | t) + \epsilon \mathbb{1}_t & \text{if } t \in A \\ \beta_\tau(\cdot | t) & \text{otherwise.} \end{cases}$$

For ϵ sufficiently small, the strict inequality still holds at t . For $s < t$, the incentive is strengthened. For $s > t$, the incentive is unaffected. Since β_τ^ϵ first-order stochastically dominates β_τ and $e^{-r\tau}$ is decreasing in τ , the principal receives strictly higher payoff under β_τ^ϵ .

Now let us impose (13) binding for all $t < t^*$ while letting (12) and (13) hold for $t = T$ simultaneously. Moreover, for any $t \geq 0$, let $\tau(t)$ be such that $e^{-r\tau(t)} = \int_t^\infty e^{-r\tau} \beta_\tau(d\tau | t)$ and $e^{-rS} = \int_T^\infty e^{-r\tau} \beta_T(d\tau)$. We then have

$$\begin{aligned} & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} W e^{-r\tau(s)} ds + e^{-rS} [p_t e^{-\lambda^1(T-t)} W - (1 - p_t) e^{-\lambda^0(T-t)}] \\ &= e^{-r\tau(t)} \cdot (p_t W + p_t - 1) \quad \forall t \in [0, t^*] \end{aligned} \quad (14)$$

$$e^{-rS} = e^{-r\tau(T)}. \quad (15)$$

Equations (14) and (15) combined give us

$$\int_t^\infty e^{-r\tau} \beta_\tau^*(d\tau | t) = e^{-r\tau(t)} = e^{-r(S-T)} e^{-r\tau^{T*}(t)} \quad \forall t < t^*, \quad (16)$$

where τ^{T*} is the optimal deterministic contract given deadline T . It is then easy to see that β_τ^* satisfies the other conditions as well and, therefore, is feasible. Therefore, it solves the original problem.

Note that any β_τ^* that satisfies (16) must not assign probability 1 to $\{t\}$ for $t < \min\{t^*, T\}$. If this is the case, then

$$\int_0^\infty e^{-r\tau} \beta_\tau^*(d\tau | t) = e^{-rt} > e^{-r\tau^{T*}(t)} \geq e^{-r(S-T)} e^{-r\tau^{T*}(t)},$$

a contradiction.

The principal's payoff equals

$$\begin{aligned} & \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} \left[\int_s^\infty e^{-r\tau} \beta_\tau^*(d\tau | s) \right] V ds + \int_T^\infty e^{-r\tau} \beta_T(d\tau) [p_0 e^{-\lambda^1 T} V - (1 - p_0) e^{-\lambda^0 T}] \\ &= \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r(S-T)} e^{-r\tau^{T*}(s)} V ds + e^{-rS} [p_0 e^{-\lambda^1 T} V - (1 - p_0) e^{-\lambda^0 T}] \\ &= \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau^{T*}(s)} e^{-r(S-T)} V ds + e^{-rS} [p_0 e^{-\lambda^1 T} V - (1 - p_0) e^{-\lambda^0 T}] \\ &\leq \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau^{T*}(s)} V ds + e^{-rT} [p_0 e^{-\lambda^1 T} V - (1 - p_0) e^{-\lambda^0 T}]. \end{aligned}$$

The first equality comes from (16). We have thus shown that for any given T , the optimal random contract is at most as good as the optimal deterministic contract with the same T . $\square$

APPENDIX C: PROOFS FOR SECTION 4.1

PROOF OF LEMMA 2. We prove that $\tau(t) > t$ for $t < \min\{t^*, T\}$. For simplicity, take $t = 0$. By contradiction, assume that $\tau(0) = 0$. The left-hand side of (5) can be written as

$$\begin{aligned} & \int_0^T (p_0 \lambda^1 \exp(-\lambda^1 s) \exp(-r\tau(s)) W) ds \\ & \quad + (p_0 \exp(-\lambda^1 T) e^{-rT} W - (1 - p_0) \exp(-\lambda^0 T) e^{-rT}) \\ & \leq \int_0^T (\lambda^1 \exp(-\lambda^1 s) \exp(-rs) p_0 W) ds \\ & \quad + (p_0 \exp(-\lambda^1 T) e^{-rT} W - (1 - p_0) \exp(-\lambda^0 T) e^{-rT}). \end{aligned}$$

The inequality follows since $\tau(s) \geq s$. The term on the right-hand side of the inequality above is the expected payoff that the agent would get following the policy of investing if any good signal is revealed before T and investing at T if no signal is revealed before T . Since $p_0 > p^*$, this policy must result in strictly lower payoffs than the expected payoff from investing at $t = 0$. So

$$\begin{aligned} & \int_0^T (\lambda^1 \exp(-\lambda^1 s) \exp(-rs) p_0 W) ds \\ & \quad + (p_0 \exp(-\lambda^1 T) e^{-rT} W - (1 - p_0) \exp(-\lambda^0 T) e^{-rT}) \\ & < p_0 W - (1 - p_0). \end{aligned}$$

Combining these inequalities, we deduce that (5) is violated at $t = 0$ when $\tau(0) = 0$. It follows that $\tau(0) > 0$. $\square$

PROOF OF LEMMA 3. Let τ^* solve the relaxed problem. By way of contradiction, assume that for some $A \subseteq [0, \min\{t^*, T\})$ with positive Lebesgue measure and for all $t \in A$, the constraint (8) is slack. For $t \in \text{dom}(\tau)$, define

$$\varphi_t = \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau^*(s)} W ds + \left(p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right).$$

Now define τ' as follows. For $t \notin A$, $\tau'(t) = \tau^*(t)$, while for $t \in A$,

$$e^{-r\tau'(t)} (p_t W + p_t - 1) = \varphi_t.$$

For $t \in A$, $e^{-r\tau'(t)} > e^{-r\tau^*(t)}$. Therefore, $\tau'(t) \leq \tau^*(t)$ for all $t \in [0, \min\{t^*, T\}]$, with strict inequality for $t \in A$. We claim that τ' is feasible. To see this, note that for all

$t \in [0, \min\{t^*, T\}]$,

$$\begin{aligned}
 & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau'(s)} W ds + \left(p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right) \\
 & \geq \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau^*(s)} W ds + \left(p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right) \\
 & = \varphi_t \\
 & \geq e^{-r\tau'(t)} (p_t W + p_t - 1) \\
 & \geq \max\{0, e^{-r\tau'(t)} (p_t W + p_t - 1)\}.
 \end{aligned}$$

The first inequality follows since τ' is below τ^* and the equality is by definition of φ_t . The second inequality follows with equality when $t \in A$ (by definition of τ'), and for $t \notin A$ follows since τ' and τ^* coincide and τ^* satisfies (8). The third inequality follows since $t < t^*$. It follows that τ' satisfies (7) and (8) and results in higher expected payoffs than τ^* . This contradicts the optimality of τ^* for the relaxed problem.

We now argue that when $T > t^*$ (case (b) in the statement of the proposition), $\tau^*(t) = t$ for almost every $t \in [t^*, T]$. Otherwise, there is a set $A \subseteq [t^*, T]$ of positive measure such that for all $t \in A$, $\tau^*(t) > t$. Construct τ' that coincides with τ^* outside A , but $\tau'(t) = t$ for $t \in A$. It is clear that τ' satisfies (8) since for $t < t^*$, τ' does not change the payoff from lying, but does increase the payoff from truth-telling. It follows that τ' is feasible for the relaxed problem and results in higher expected payoffs for the principal than τ^* . This is a contradiction.

Now, to prove the converse, we assume that $T > t^*$. The proof of the converse when $T \leq t^*$ is analogous. Take τ^* such that (7) and (8) bind almost everywhere. Take τ' that solves the relaxed problem (6). From the first part of this proof, τ' and τ^* coincide for almost every $t \in [t^*, T]$. The previous step also shows that τ' is such that (8) binds for almost every $t \in [0, \min\{t^*, T\}]$. Define

$$u(t) = \int_t^T e^{-\lambda^1 s} (e^{-r\tau^*(s)} - e^{-r\tau'(s)}) ds$$

for $t \in [0, \min\{t^*, T\}]$. Note that $u(t)$ is absolutely continuous, and its derivative is defined almost everywhere and equals $-e^{-\lambda^1 t} (e^{-r\tau^*(t)} - e^{-r\tau'(t)})$. Now using the fact that the constraint binds almost everywhere for both τ' and τ^* , we deduce that for almost every $t \in [0, \min\{t^*, T\}]$,

$$-p_t \lambda^1 W u(t) = u'(t) (p_t W + p_t - 1)$$

and $u(\min\{t^*, T\}) = 0$. It follows that for almost every $t \in [0, \min\{t^*, T\}]$, $d(u(t)e^{\int_0^t H(s) ds})/dt = 0$, where H is a continuous function. Since $u(\min\{t^*, T\}) = 0$, $u(t) = 0$ for all $t \in [0, \min\{t^*, T\}]$. In particular, $0 = u'(t) = -e^{-\lambda^1 t} (e^{-r\tau^*(t)} - e^{-r\tau'(t)})$ almost everywhere, and, therefore, τ' and τ^* coincide for almost every $t \in [0, \min\{t^*, T\}]$. Since τ^* satisfies (7) and (8), τ^* solves the relaxed problem. $\square$

PROOF OF LEMMA 4. Since (5) is binding for all $t \in [0, \min\{t^*, T\})$,

$$\begin{aligned} & \int_t^T \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} W ds + \left(e^{-\lambda^1 T} e^{-rT} W - \frac{1-p_t}{p_t} e^{(\lambda^0 - \lambda^1)t} e^{-\lambda^0 T} e^{-rT} \right) \\ &= e^{-\lambda^1 t} e^{-r\tau(t)} \left(W - \frac{1-p_t}{p_t} \right), \end{aligned}$$

where we use the fact that $t \leq t^*$. Since the left-hand side of this equation and p_t are differentiable, so is τ . Taking derivatives and using the fact that $d(\frac{1-p_t}{p_t})/dt = (\lambda^1 - \lambda^0)(1 - p_t)/p_t$, we deduce that

$$-\lambda^1 e^{-\lambda^1 t} e^{-r\tau(t)} W = -(\lambda^1 + r\dot{\tau}(t)) e^{-\lambda^1 t - r\tau(t)} \left(W - \frac{1-p_t}{p_t} \right) - e^{-\lambda^1 t - r\tau(t)} (\lambda^1 - \lambda^0) \frac{1-p_t}{p_t}.$$

Solving for $\dot{\tau}(t)$, we deduce that

$$\dot{\tau}(t) = \left(\frac{\lambda^0}{r} \right) \frac{1}{W \frac{p_t}{1-p_t} - 1}.$$

The slope of τ is nonnegative. To see that τ is convex, note that $p_t/(1-p_t)$ is nonincreasing and, thus, $\dot{\tau}$ is nondecreasing. To see that $\dot{\tau}$ is less than 1, note that

$$\dot{\tau} < 1 \quad \text{if and only if} \quad 1 < W \frac{r}{\lambda^0 + r} \frac{p_t}{1-p_t}.$$

To verify this last property, note that investing at t results in higher expected payoffs for the agent than learning at t and investing at $t + dt$ unless the bad state is revealed. That is,

$$p_t W - (1-p_t) \geq -(1-p_t)(1-\lambda^0 dt) e^{-r dt} + p_t e^{-r dt} W.$$

Reordering terms and taking $dt \rightarrow 0$, we deduce that

$$1 < W \frac{r}{\lambda^0 + r} \frac{p_t}{1-p_t}. \quad \square$$

APPENDIX D: PROOFS FOR SECTION 4.3

PROOF OF PROPOSITION 2. We first prove part (a). Note that the principal's expected payoff from setting $T = \infty$ equals

$$\varphi(W) = \int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^*(s)} V ds + \int_{t^*}^{\infty} p_0 \lambda^1 e^{-\lambda^1 s} e^{-rs} V ds,$$

where

$$\tau_W^*(s) = t^* - \left(\frac{\lambda^0}{r} \right) \int_s^{t^*} \frac{1}{W \frac{p_x}{1-p_x} - 1} dx.$$

We claim that for all $\epsilon > 0$, there exists L such that for all $W > L$, $\varphi(W) < \epsilon$.

First, notice that since $t^* \rightarrow \infty$ as $W \rightarrow \infty$, there exists L_1 such that for all $W > L_1$,

$$\int_{t^*}^{\infty} p_0 \lambda^1 e^{-\lambda^1 s} e^{-rs} V ds < \epsilon/2.$$

Now we show that there exists L_2 such that for all $W > L_2$,

$$\int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^*(s)} V ds < \frac{\epsilon}{2}.$$

To show this, we first show that for any δ , there exists L_3 such that for all $W > L_3$,

$$\frac{1}{W \frac{p_x}{1-p_x} - 1} < \delta \quad \forall x \in [0, t^*].$$

Since $p_x/(1-p_x)$ decreases in x , it suffices to show

$$\frac{1}{W \frac{p_0}{1-p_0} - 1} < \delta.$$

This is done by letting

$$L_3 = \frac{\delta + 1}{\delta} \frac{2(1-p_0)}{p_0}.$$

Given this, we now show that for any η , there exists L_4 such that $W > L_4$ implies

$$e^{-r\tau(s)} < \eta \quad \forall s \in [0, t^*].$$

To show this, first we notice that $\tau(s)$ increases in s , so it suffices to show that there exists L_4 such that $W > L_4$ implies

$$e^{-r[t^* - \frac{\lambda^0}{r} \int_0^{t^*} \frac{1}{W \frac{p_x}{1-p_x} - 1} dx]} < \eta.$$

In other words,

$$t^* - \frac{\lambda^0}{r} \int_0^{t^*} \frac{1}{W \frac{p_s}{1-p_s} - 1} ds > \frac{\ln \eta}{-r}.$$

Given what we showed in the previous step, we can find L_3 such that

$$\frac{1}{W \frac{p_x}{1-p_x} - 1} < \frac{r}{2\lambda^0} \quad \forall x \in [0, t^*].$$

Therefore,

$$t^* - \frac{\lambda^0}{r} \int_0^{t^*} \frac{1}{W \frac{p_x}{1-p_x} - 1} dx > t^* - \frac{\lambda^0}{r} \int_0^{t^*} \frac{r}{2\lambda^0} dx$$

$$\begin{aligned}
&= t^* - \frac{\lambda^0}{r} \frac{r}{2\lambda^0} t^* \\
&= \frac{t^*}{2} \rightarrow \infty
\end{aligned}$$

as $W \rightarrow \infty$. We have, therefore, shown that there exists L_4 such that $W > L_4$ implies

$$e^{-r\tau(s)} < \eta \quad \forall s \in [0, t^*].$$

Now find L_4 such that

$$e^{-r\tau(s)} < \frac{\epsilon}{4Vp_0} \quad \forall s \in [0, t^*].$$

Therefore,

$$\begin{aligned}
\int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} V ds &< \int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} \frac{\epsilon}{4Vp_0} V ds \\
&= p_0 \lambda^1 \frac{\epsilon}{4Vp_0} V \int_0^{t^*} e^{-\lambda^1 s} ds \\
&= p_0 \lambda^1 \frac{\epsilon}{4Vp_0} V \cdot \frac{1}{\lambda^1} (1 - e^{-\lambda^1 t^*}) \\
&= p_0 \frac{\epsilon}{4Vp_0} V (1 - e^{-\lambda^1 t^*}) \\
&= \frac{\epsilon}{4} (1 - e^{-\lambda^1 t^*}) \\
&< \frac{\epsilon}{2}.
\end{aligned}$$

Therefore, for $W > L_2 := \max\{L_3, L_4\}$, we have

$$\int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^*(s)} V ds < \frac{\epsilon}{2}.$$

Last, letting $L := \max\{L_1, L_2\}$, we then have

$$\varphi(W) < \epsilon$$

for $W > L$.

Now note that by setting an optimal deadline $T \in [0, t^*]$, the principal's payoff equals

$$\Phi(W) = \max_{T \in [0, t^*]} \Phi(W, T),$$

where

$$\Phi(W, T) = \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^T(s)} V ds + (p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT}).$$

Note that

$$\begin{aligned}\Phi(W, T) &> p_0 e^{-rT} (1 - e^{-\lambda^1 T}) V + (p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT}) \\ &= e^{-rT} (p_0 V - (1 - p_0) e^{-\lambda^0 T}).\end{aligned}$$

Fix any T such that the expression above is strictly positive and equals $\eta > 0$. Let $\epsilon = \eta/2$, and take $W > L$ such that $\varphi(W) < \epsilon = \eta/2$ and $T < t^*$. In particular,

$$\Phi(W) \geq \eta > \eta/2 \geq \varphi(W),$$

which proves that there exists some $\bar{\kappa}$ such that for all $W > \bar{\kappa}$, $T \in [0, t^*]$ results in higher payoffs than $T = \infty$.

To complete the proof of part (a), note that as W goes to x (where $(xrp_0)/((\lambda^1 + r)(1 - p_0)) = 1$), $t^* \rightarrow 0$. In particular,

$$\varphi(W) \rightarrow \int_0^\infty p_0 \lambda^1 e^{-\lambda^1 s} e^{-Rs} V ds = p_0 V \frac{\lambda^1}{\lambda^1 + r},$$

whereas

$$\Phi(W) \rightarrow p_0 V + (1 - p_0)(-v).$$

Since $p_0 V \lambda^1 / (\lambda^1 + r) > p_0 V - (1 - p_0)$, there exists $\underline{\kappa}$ such that for all $W < \underline{\kappa}$, $\varphi(W) > \Phi(W)$.

To prove part (b), we note that

$$\varphi(V) = \int_0^{t^*} p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^*(s)} V ds + \int_{t^*}^\infty p_0 \lambda^1 e^{-\lambda^1 s} e^{-rs} V ds$$

for the principal's payoff when $T = \infty$ and

$$\Phi(V, T) = \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau_W^T(s)} V ds + (p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT})$$

for the principal payoff when setting $T < t^*$. Note that as $V \rightarrow 0$,

$$\varphi(V) \rightarrow 0, \quad \Phi(V, T) \rightarrow -(1 - p_0) e^{-\lambda^0 T} e^{-rT}.$$

Since Φ is continuous in (V, T) , there exists $\underline{\eta} > 0$ such that for all $V < \underline{\eta}$,

$$\varphi(V) > \max_{T \in [0, t^*]} \Phi(V, T)$$

and, thus, it is optimal for the principal to set $T = \infty$.

To complete part (b), define y such that $1 = yrp_0/((\lambda^1 + r)(1 - p_0))$. By definition,

$$\varphi(y) < \max_{T \in [0, t^*]} \Phi(y, T),$$

where the maximum on the right is attained at $T = 0$. By continuity, there exists $\bar{\eta} < y$ such that for all $V > \bar{\eta}$,

$$\varphi(V) < \max_{T \in [0, t^*]} \Phi(V, T)$$

and the principal sets a deadline $T < t^*$. $\square$

APPENDIX E: NO NEWS IS NO NEWS AND NO NEWS IS GOOD NEWS

In the absence of a signal, the evolution of the agent's belief $(p_t)_{t \geq 0}$ satisfies

$$\frac{dp_t}{dt} = -(\lambda^1 - \lambda^0)p_t(1 - p_t).$$

We first consider the case when $\lambda^0 = \lambda^1$; that is, $dp_t/dt = 0$. In this case, the agent's belief remains constant if no signal has arrived and jumps to 0 or 1 at the first signal. Therefore, the uninformed agent is never indifferent between investing and waiting to invest after a good signal, and we define $t^* = \infty$.

The single-player problem is solved identically as that in [Section 2](#). The dynamic delegation problem and the relaxed problem are set up in the same way. When solving the relaxed problem, since $t^* = \infty$, we need only to consider the $T \leq t^*$ case. Since it is infeasible to set $T = \infty$ for any combination of parameters that satisfy [Assumption 1](#), the optimal contract always features a deadline T and the corresponding contract τ^T is solved in the same way as in [Section 4.2](#).

When $\lambda^0 > \lambda^1$, the agent's belief drifts up as time goes on. Suppose that the agent observes the signal and decides whether to invest at each point in time. Following arguments similar to those in [Section 2](#), it is relatively simple to show that there exists p^* such that the agent invests if and only if $p_t \geq p^*$. Analogously, there exists q^* such that the principal would make the decision if and only if $p_t \geq q^*$.¹⁴ We assume that $p^* < p_0 < q^*$. This means that at time 0, the agent would like to invest, whereas the principal would like to wait for information. In contrast to [Section 2](#), the assumption $\lambda^0 > \lambda^1$ now implies that there exists t^* such that if no signal has been received, the principal would like to invest at any $t > t^*$. In particular, for $t > t^*$, the principal's and the agent's preferences surely coincide, as both would like to invest. This implies that there is always a deadline $T \leq t^*$.

We find the contract $\langle T, \tau \rangle$ that solves (1) under constraints (2)–(5). All these constraints remain relevant in this setup, as they capture feasibility and truth-telling incentives that need to be provided regardless of the direction followed by the belief path.

The solution method is similar to [Section 4](#). We now sketch and discuss the main steps.

LEMMA 6. *Let $\langle T, \tau \rangle$ satisfy (2) and (5). Then $\tau(t) > t$ for all $t \leq T$.*

¹⁴Note that the thresholds p^* and q^* in this subsection do not coincide with the thresholds derived in [Section 2](#).

This result is similar to [Lemma 2](#). The main difference is that now the uninformed agent prefers to invest for all $t \in \text{dom}(\tau)$ and, as a result, all the investment times need to be distorted.

We also solve the dynamic delegation problem for fixed T and, as in [Section 4](#), it is convenient to formulate the relaxed problem

$$\max_{\tau(\cdot)} \int_0^T p_0 \lambda^1 e^{-\lambda^1 s} e^{-r\tau(s)} V ds + (p_0 e^{-\lambda^1 T} e^{-rT} V - (1 - p_0) e^{-\lambda^0 T} e^{-rT}) \quad (17)$$

subject to

$$\tau(T) \geq T \quad (18)$$

$$\begin{aligned} & \int_t^T p_t \lambda^1 \frac{e^{-\lambda^1 s}}{e^{-\lambda^1 t}} e^{-r\tau(s)} W ds + \left[p_t \frac{e^{-\lambda^1 T}}{e^{-\lambda^1 t}} e^{-rT} W - (1 - p_t) \frac{e^{-\lambda^0 T}}{e^{-\lambda^0 t}} e^{-rT} \right] \\ & \geq \max\{e^{-r\tau(t)}(p_t W + p_t - 1), 0\}, \quad \forall t \leq T. \end{aligned} \quad (19)$$

LEMMA 7. Let τ^* satisfy (18) and (19). Then τ^* solves the relaxed problem (17) if and only if (19) binds for almost every $t \in [0, T]$.

This lemma is similar to [Lemma 3](#). Intuitively, if the constraint were slack, the principal could slightly reduce the investment time and improve her expected payoffs.

A solution to the relaxed problem is then found by imposing (19) binding over $[0, T]$. Since $p_t W + p_t - 1 > 0$ for all $t \geq 0$, (19) binding at T implies $\tau(T) = T$. Using [Lemma 4](#), we can solve for the binding constraint (19) by simply solving the system

$$\tau(T) = T, \quad \dot{\tau}(t) = \left(\frac{\lambda^0}{r} \right) \frac{1}{W \frac{p_t}{1 - p_t} - 1} \quad t < T.$$

The solution τ^T to this system is given by (10). This function is concave and its slope is less than 1. As it satisfies all the constraints of the dynamic delegation problem, τ^T actually solves the dynamic delegation problem for fixed T . As T increases, so does $\tau^T(t)$ and, thus, the principal needs to distort more investment decisions. The optimal T is chosen as follows. Over $T \geq t^*$, the principal should optimally set $T = t^*$ since for all $t > t^*$, the principal is optimistic enough to invest without any news. Over $T < t^*$, the solution solves the trade-off characterized in [Proposition 2](#).

REFERENCES

- Aghion, Philippe and Jean Tirole (1997), “Formal and real authority in organizations.” *Journal of Political Economy*, 105, 1–29. [574]
- Alonso, Ricardo and Niko Matouschek (2008), “Optimal delegation.” *Review of Economic Studies*, 75, 259–293. [572, 574]
- Amador, Manuel and Kyle Bagwell (2013), “The theory of optimal delegation with an application to tariff caps.” *Econometrica*, 81, 1541–1599. [572, 574]

Ambrus, Attila and Georgy Egorov (2017), “Delegation and nonmonetary incentives.” *Journal of Economic Theory*, 101–135. [574]

Armstrong, Mark and John Vickers (2010), “A model of delegated project choice.” *Econometrica*, 78, 213–244. [574]

Bergemann, Dirk and Juuso Välimäki (2019), “Dynamic mechanism design: An introduction.” *Journal of Economic Literature*, 57, 235–274. [574]

Chen, Yi (2018), “Dynamic communication with commitment.” Unpublished paper, Cornell University. [574]

Deimen, Inga and Dezső Szalay (2019), “Delegated expertise, authority, and communication.” *American Economic Review*, 109, 1349–1374. [574]

Ely, Jeffrey C. (2017), “Beeps.” *American Economic Review*, 107, 31–53. [575]

Frankel, Alexander (2016), “Discounted quotas.” *Journal of Economic Theory*, 166, 396–444. [574]

George, Stephen L. and Marc Buyse (2015), “Data fraud in clinical trials.” *Clinical Investigation*, 5, 161–173. [572]

Grenadier, Steven R., Andrey Malenko, and Nadya Malenko (2016), “Timing decisions in organizations: Communication and authority in a dynamic environment.” *American Economic Review*, 106, 2552–2581. [574]

Guo, Yingni (2016), “Dynamic delegation of experimentation.” *American Economic Review*, 106, 1969–2008. [574]

Guo, Yingni and Johannes Hörner (2020), “Dynamic allocation without money.” Working paper, Toulouse School of Economics Working papers no 1133. [574, 575]

Henry, Emeric and Marco Ottaviani (2019), “Research and the approval process: The organization of persuasion.” *American Economic Review*, 109, 911–955. [575]

Holmström, Bengt (1984), “On the theory of delegation.” In *Bayesian Models in Economic Theory* (Marcel Boyer and Richard Kihlstrom, eds.), 115–141, North-Holland, Amsterdam, Netherlands. [572, 574, 577]

Keller, Godfrey, Sven Rady, and Martin Cripps (2005), “Strategic experimentation with exponential bandits.” *Econometrica*, 73, 39–68. [576]

Lewis, Tracy R. and Marco Ottaviani (2008), “Search agency.” Unpublished paper, Northwestern University. [574]

Li, Jin, Niko Matouschek, and Michael Powell (2017), “Power dynamics in organizations.” *American Economic Journal: Microeconomics*, 9, 217–241. [574]

Lipnowski, Elliot and João Ramos (2020), “Repeated delegation.” [574]

Liptser, Robert S. and Albert N. Shiryaev (2013), *Statistics of Random Processes: General Theory*, volume 5. Springer Science & Business Media. [575]

Madsen, Erik (2018), “Expert advice and optimal project termination.” Unpublished paper, SSRN 3155328. [574]

McClellan, Andrew (2017), “Experimentation and approval mechanisms.” Unpublished paper, New York University. [575]

Melumad, Nahum and Toshiyuki Shibano (1991), “Communication in settings with no transfers.” *RAND Journal of Economics*, 22, 173–198. [572, 574]

Orlov, Dmitry, Andrzej Skrzypacz, and Pavel Zryumov (2020), “Persuading the principal to wait.” *Journal of Political Economy*, 128, 2542–2578. [575]

Pavan, Alessandro, Ilya Segal, and Juuso Toikka (2014), “Dynamic mechanism design: A Myersonian approach.” *Econometrica*, 82, 601–653. [574, 592]

Seife, Charles (2015), “Research misconduct identified by the US food and drug administration: Out of sight, out of mind, out of the peer-reviewed literature.” *JAMA Internal Medicine*, 175, 567–577. [572]

Sugaya, Takuo and Alexander Wolitzky (2020), “The revelation principle in multistage games.” *Review of Economic Studies*. [577]

Szalay, Dezsö (2005), “The economics of clear advice and extreme options.” *Review of Economic Studies*, 72, 1173–1198. [574]

Co-editor Simon Board handled this manuscript.

Manuscript received 10 March, 2020; final version accepted 9 July, 2020; available online 13 July, 2020.