

Learning from Manipulable Signals

Author(s): Mehmet Ekmekci, Leandro Gorno, Lucas Maestri, Jian Sun and Dong Wei

Source: *The American Economic Review*, DECEMBER 2022, Vol. 112, No. 12 (DECEMBER 2022), pp. 3995-4040

Published by: American Economic Association

Stable URL: <https://www.jstor.org/stable/10.2307/27181550>

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.

Your use of the JSTOR archive indicates your acceptance of the Terms & Conditions of Use, available at <https://about.jstor.org/terms>



American Economic Association is collaborating with JSTOR to digitize, preserve and extend access to *The American Economic Review*

JSTOR

Learning from Manipulable Signals[†]

By MEHMET EKMEKCI, LEANDRO GORNO,
LUCAS MAESTRI, JIAN SUN, AND DONG WEI*

We study a dynamic stopping game between a principal and an agent. The principal gradually learns about the agent's private type from a noisy performance measure that can be manipulated by the agent via a costly and hidden action. We fully characterize the unique Markov equilibrium of this game. We find that terminations/market crashes are often preceded by a spike in manipulation intensity and (expected) performance. Moreover, due to endogenous signal manipulation, too much transparency can inhibit learning and harm the principal. As the players get arbitrarily patient, the principal elicits no useful information from the observed signal. (JEL C73, D82, D83, G24, M13)

Asymmetric information is pervasive in long-term relationships; meanwhile, learning often takes place during the interactions between different parties. For instance, venture capital (VC) firms face asymmetric information in their investments: VCs only wish to invest in promising projects, but startups often have better information about the chances of success of their projects than investors (Gompers and Lerner 2004). Upon agreeing to finance a startup, the VC receives periodical performance reports (e.g., subscription growth, number of patents, media and user reviews, etc.) from the startup. These reports can provide information about the viability or profitability of the startup. However, the startup may undertake hidden actions to inflate the performance reports, tampering with their informativeness. Examples include rideshare platforms that periodically announce their numbers of users and could augment such statistics by specialized promotions, and Theranos and Bolt who have allegedly misrepresented key performance data.

We analyze a learning problem with asymmetric information and hidden actions, and investigate the equilibrium learning dynamics. In our model, a principal (VC)

*Ekmekci: Boston College (mehmet.ekmekci@bc.edu); Gorno: FGV EPGE (leandro.gorno@fgv.br); Maestri: FGV EPGE (lucas.maestri@fgv.br); Sun: Singapore Management University (jjiansun@smu.edu.sg); and Wei: University of California, Santa Cruz (dwei10@ucsc.edu). Jeffrey Ely was the coeditor for this article. This paper originated from two related but independent projects; all authors contributed equally to the current paper. We thank the four anonymous referees for their insightful suggestions which greatly improved the paper. We are grateful for the helpful comments from Alp Atakan, Costas Cavounidis, Doruk Cetemen, Brett Green, Matt Jackson, Sam Jindani, Aaron Kolb, Aditya Kuvalekar, Elliot Lipnowski, Andrey Malenko, Harry Di Pei, João Ramos, Yiman Sun, Philipp Strack, as well as seminar and conference participants at Boston College, Boston University, MIT Sloan, Stanford University, University of Bonn, University of Glasgow, University of Montreal, the 12th Econometric Society World Congress, and the 33rd Stony Brook Game Theory Conference. Leandro Gorno and Lucas Maestri acknowledge that this study was financed in part by the Coordenação de Aperfeiçoamento de Pessoal de Nível Superior—Brazil (CAPES)—Finance Code 001. Leandro Gorno and Lucas Maestri are also grateful for the financial support of the National Council for Scientific and Technological Development—Brazil (CNPq) and FAPERJ.

[†]Go to <https://doi.org/10.1257/aer.20211158> to visit the article page for additional materials and author disclosure statements.

and an agent (startup) are engaged in a relationship that takes place in continuous time. Performance reports are modeled as a public signal evolving according to a Brownian motion with a drift that depends on the agent's privately known type and action. If the agent is an *investible* type, then the drift is $\mu > 0$; if the agent is a *noninvestible* type, then the drift is 0 by default, but the agent can take a costly action to boost the drift up to μ . The signal serves only an informational role and does not affect the principal's payoffs. When a stopping opportunity stochastically arrives, the principal needs to decide whether to continue or terminate the relationship (i.e., discontinue the funding). The principal prefers to continue the relationship with the investible type and to terminate the relationship with the noninvestible type.

We study Markov equilibria where the state variable is the public belief that the agent is a noninvestible type. We call the complementary probability, i.e., the probability that the agent is an investible type, the agent's reputation. In this game, there exists a unique Markov equilibrium.

In the unique equilibrium, the principal's termination strategy has a cutoff structure—the principal terminates the relationship if and only if the agent's reputation is sufficiently bad. The agent's equilibrium strategy depends on his level of patience. If the agent is impatient, he never engages in costly performance boosting. Otherwise, if the agent is sufficiently patient, he does not engage in performance boosting when his reputation is either excellent or terrible, but will do so when his reputation is intermediate. In particular, the intensity of performance boosting is hump-shaped in the agent's reputation and peaks at the principal's termination cutoff. This leads to an interesting “*scramble-to-rescue*” effect: the agent increases his intensity of performance boosting as the relationship gets less stable.

Utilizing the equilibrium characterization, we obtain three main qualitative results. The first result concerns the relationship between the agent's reputation and the expected performance, measured by the drift of the signal assessed from an outsider's perspective. If the agent is so impatient that he never engages in performance boosting, then the expected performance is increasing in the agent's reputation (decreasing in the state variable). However, when the agent is more patient and engages in some performance boosting, the expected performance is nonmonotone in the agent's reputation. Starting from an initial good reputation, as the agent's reputation deteriorates, the expected performance first declines, reaching a local minimum, and then rises, reaching a local maximum precisely at the principal's termination cutoff, and decreases again thereafter (see Figure 3). The nonmonotonicity is driven by the scramble-to-rescue effect, as a deterioration of the agent's reputation can induce him to boost performance more intensely. This finding may help explain why some startups delivered seemingly impressive performance reports, such as extraordinary revenue flow (e.g., Theranos) and rapid expansion (e.g., WeWork, Ofo), not long before investors decided to cut funding. It is also consistent with the observation that growing market suspicion and strong (expected) performance can coexist for a period of time.

Our second qualitative result concerns the relationship between the amount of information transmission and the transparency of the performance measure. Due to random events such as demand shocks and measurement errors, performance reports are imperfect signals of the agent's type and action, and we use the signal-to-noise ratio of the signal process to capture its transparency. In reality, transparency may be

determined by the volatility of the product market and may also be affected by how much detail a startup is required to disclose. We show that, due to the agent's endogenous signal manipulation, the principal may be worse off in more transparent environments. This result suggests that VCs can sometimes fare better when the startup initially operates in a more volatile market; also, policies that require disclosing too precise information may end up hurting the investors.

Specifically, the principal's payoff is nonmonotone in the signal-to-noise ratio of the performance reports, implying that the principal can be worse off when the performance measure becomes more transparent (less noisy). We obtain this result by looking at two extreme cases. At one extreme, if the performance reports are independent of the agent's type and action (i.e., an uninformative signal), then the principal can never learn about the agent's type and will receive her "no-information" value. At the other extreme, as the signal-to-noise ratio grows without bound, we show that the principal cannot utilize any information about the agent's type either. Intuitively, in this case the agent will aggressively boost the performance, for otherwise his type would be revealed almost immediately. Such aggressive performance boosting is anticipated by the principal, and thus largely reduces the informativeness of the signal. As a result, the principal's equilibrium payoff converges to her "no-information" value. In contrast to the extreme cases, for intermediate values of the signal-to-noise ratio, the principal will learn some information about the agent's type and get a payoff strictly above her "no-information" value.

Finally, we investigate the equilibrium outcomes as players get arbitrarily patient. We find a strong manifestation of the ratchet effect in the patient limit of our model. Since the principal cannot commit to refraining from using future information against the agent, a patient agent will engage in performance boosting with almost full intensity in order to maintain his reputation. In the limit, no useful information is revealed, and the principal's lack of commitment hurts her in the most extreme way.¹

Related Literature.—Our paper is most closely related to the reputation literature and the literature on dynamic games with stopping decisions.

Most of the reputation literature—starting with Kreps and Wilson (1982) and Milgrom and Roberts (1982), later generalized by Fudenberg and Levine (1989, 1992), and most recently by Pei (2020)—investigates whether and how much a long-lived informed player can benefit from their private information in repeated games played against myopic opponents. The focus is typically on the case where the informed party is arbitrarily patient, and on bounding the informed player's equilibrium payoffs. In contrast, our analysis fully characterizes (Markov) equilibrium behavior for all discount rates and uncovers new qualitative features of the equilibrium dynamics.

Faingold and Sannikov (2011) study reputation effects in games played in continuous time with one long-lived informed player against myopic opponents, and they characterize the set of sequential equilibria. Unlike Faingold and Sannikov (2011), the uninformed player in our game is forward-looking and can terminate the

¹ This result holds if both players get arbitrarily patient at the same rate, or if the agent gets patient at a faster rate than does the principal.

game. More importantly, the termination payoffs depend on the informed player's type, creating *interdependence* of payoffs between the players (similar to Pei 2020) and thus making their characterization not applicable to our model.

There is a growing interest in dynamic games with stopping decisions. Daley and Green (2012); Kolb (2015, 2019); Dilmé (2019); Ekmekci and Maestri (2022); and Sun (2021) all study stopping games with two long-lived players, where the uninformed party receives information over time and obtains type-dependent payoffs. In Daley and Green (2012); Kolb (2015); and Dilmé (2019), the informed player makes the stopping decision while in our paper such a decision is made by the uninformed player. This makes the players' incentives in our model quite different from theirs. In Kolb (2019), the agent can only influence the information process by irreversibly changing his type, while in our model the agent can directly manipulate the signal, with his type being persistent. The qualitative results on the equilibrium dynamics in our paper do not have a counterpart in these papers. Ekmekci and Maestri (2022) study a similar setting in discrete time and focus solely on the limiting case with arbitrarily patient players. While our Theorem 4 is a continuous-time analog of their main finding, we obtain much richer equilibrium dynamics for any fixed discount rate. Finally, Sun (2021) studies dynamic censorship with Poisson news, wherein the agent can decide whether to show or hide the bad news after privately observing its realization.

Aghion and Jackson (2016) and Kuvalekar and Lipnowski (2020) also study dynamic games (in discrete and continuous time, respectively) between two long-run players with stopping decisions. Both papers look at a career-concern type of model with symmetric information between the two players, while the agent's actions affecting the signal process are *costless* to the agent and *observable* to the principal. By contrast, in our model the agent has private information about his type, and his action is *costly* and *hidden*. Moreover, in our model the agent's trade-off is between improving his reputation and saving the mimicking cost, while in their models the agent is optimizing over the speed of learning (i.e., variance, rather than drift, of the belief process).

I. Model

A. Players, Actions, and Signals

A principal (she) and an agent (he), both risk neutral, interact in continuous time $t \in [0, \infty)$. The agent can be one of two types, denoted by θ : an *investible* type ($\theta = I$), or a *noninvestible* type ($\theta = NI$). The agent's type is his private information. From the principal's viewpoint, the initial probability that the agent is a noninvestible type is $p_0 \in (0, 1)$.

The principal faces an optimal stopping problem with a small amount of friction. An exogenous stopping opportunity arrives according to a Poisson process with rate $\lambda > 0$. When said opportunity arrives, the principal can choose to *continue* or irreversibly *terminate* the relationship with the agent. We allow for random termination and denote by β_t the principal's termination probability *conditional on* the arrival of a stopping opportunity at time t . In our leading application of VC investing, we interpret termination as discontinuation of funding to a startup.

There is a public signal $\{X_t\}_{t \geq 0}$ that evolves over time. If the agent is an investible type, then the public signal evolves according to the process

$$dX_t = \mu dt + \sigma dB_t,$$

where $\{B_t\}_{t \geq 0}$ is a standard Brownian motion. Without loss, we assume that $\mu > 0$ and $\sigma > 0$, and we define the *signal-to-noise ratio* ψ of the process as $\psi \equiv \mu/\sigma$. If the agent is a noninvestible type, he chooses an $\alpha_t \in [0, 1]$, unobservable to the principal, at any time t when the game has not stopped yet. In this case, his choice controls the drift of the public signal process:

$$dX_t = \mu\alpha_t dt + \sigma dB_t.$$

The model assumes that the investible type does not have any action choice, and the evolution of the public signal is exogenous conditional on this type (always having a drift of μ).² Meanwhile, the noninvestible type chooses a *mimicking intensity*, which can be interpreted as the probability with which the noninvestible type acts the same as the investible type. In the context of VC investing, one can interpret the public signal as performance reports from the startup and the mimicking action taken by the noninvestible type as *performance boosting*.

B. Payoffs

Once the game is stopped, the agent receives his outside option which we normalize to zero. If the game is not yet stopped, the noninvestible agent receives a flow payoff that depends on his action, $[u + (1 - \alpha_t)c]$, where $u > 0$ and $c > 0$.³ That is, if the noninvestible agent does (not) mimic the investible type then his flow payoff in the relationship is u (resp., $u + c$); thus, c is the flow cost of mimicking. If the relationship is terminated at a (stochastic) time $\mathbb{T}$, the agent's payoff is

$$r_1 \int_0^{\mathbb{T}} e^{-r_1 t} [u + (1 - \alpha_t)c] dt,$$

where r_1 is the agent's discount rate.

The principal's flow payoff does not depend on the agent's action or the public signal, and we normalize her flow payoff to zero. However, the principal receives a lump-sum payoff of $w_{NI} > 0$ if the game stops against a noninvestible type, and $w_I < 0$ if the game stops against an investible type. That is, relative to continuing the relationship, the principal prefers stopping against a noninvestible type but

²The assumption that the investible type does not have an action choice makes him a "commitment" type in the reputation literature. Our analysis is robust to certain forms of strategic behavior. For instance, if we allow the investible type to costlessly choose a drift, the unique equilibrium characterized in Theorem 1 remains an equilibrium in this modified game.

³The flow payoff of the investible type in the relationship is always some positive constant, say, $u + c$.

dislikes terminating an investible type.⁴ If the relationship is terminated at a (stochastic) time $\mathbb{T}$, the principal's payoff is

$$e^{-r_2 \mathbb{T}} w_{\theta},$$

where r_2 is the principal's discount rate.

C. Discussion of Model Assumptions

The essential ingredients of our model are the following:

- (i) The agent has private information about his type and wants to stay in the relationship for as long as possible.
- (ii) The principal prefers to terminate the relationship with the noninvestible type and to continue with the investible type, and she gradually learns about the agent's type via a noisy signal.
- (iii) The noninvestible type can privately manipulate the noisy signal at a cost to mimic the investible type's performance.
- (iv) The signal only serves an informational role and does not affect the principal's payoff.

These ingredients capture the key incentives in our leading examples of VC-startup relationships. It is well-documented that businesses may adopt devious tactics ranging from aggressive marketing strategies to outright fraud, in order to maintain their reputation and/or secure more funding (Banerjee 2022). In practice, startups often have better information about the prospects of their projects than VCs, meanwhile VCs only want to invest in those promising ones. With no commitment to future financing, whether and when to abandon a project is one of the most challenging problems VCs decide on (Bergemann and Hege 1998, 2003).⁵ The public signal in our model corresponds to such observable measures as user growth, sales statistics or scientific progress reports, providing imperfect information about the viability of the startup. The manipulative action can represent, for example, (i) personalized discounts for users (e.g., rideshare and food delivery platforms), (ii) misrepresentation of data (e.g., Theranos and Bolt),⁶ and (iii) cost-ineffective expansion strategies

⁴The payoff normalization allows us to capture the principal's incentive in the simplest possible way. We refer the reader to our working paper (Ekmekci et al. 2021) for a microfoundation of the principal's payoffs with a flow investment cost, a type-dependent reward (upon initial public offering [IPO] or project failure), and a constant outside option.

⁵See Section VIB for a discussion of the case with commitment.

⁶The Securities and Exchange Commission uncovered through investigation that Theranos largely misrepresented its product efficacy and financial performance, and charged its founder and CEO with fraud (see "Theranos, CEO Holmes, and Former President Balwani Charged with Massive Fraud," *SEC Press Releases*, March 14, 2018). Bolt, a startup focusing on payment processing technology, was accused by former employees of overstating its technological capability and misrepresenting the number of merchants using its service (see Erin Woo and Maureen Farrell, "Bolt Built \$11 Billion Payment Business on Inflated Metrics and Eager Investors," *New York Times*, May 10, 2022).

(e.g., WeWork, Ofo).⁷ Clearly these manipulative actions are costly and hard to detect, further complicating the VC's learning problem. Our assumption that the public signal does not directly affect the principal's payoff is also justified in these contexts. This is because VC investments are long term by nature where an investor's payoff is largely driven by the viability (type) of the startup rather than its short-term performance in the initial financing period (Gompers and Lerner 2004), though the initial performance can be informative about the startup's type.

Finally, we assume that the principal can only terminate the relationship when a stopping opportunity arrives. This is mainly a technical assumption, which we discuss in more detail in Section VIA. Economically, these intermittent stopping opportunities can capture the frictions in a VC's decision-making process. For example, the discontinuation of funding may be decided only through board meetings which are called upon only once in a while; moreover, a VC that wants to liquidate its stake in a startup may have to wait some time until a buyer shows up.

II. Equilibrium

A. Equilibrium Concept

Let p_t be the principal's public belief at time t about the probability that the agent is noninvestible. This belief is common knowledge between the players because the signal process X_t is publicly observed.

We look at Markov equilibrium, that is, strategies based on policy functions, $(a(p), b(p))$, which depend only on the public belief. Each player's policy function is a measurable function from $(0, 1)$ to $[0, 1]$ and prescribes the player's (possibly mixed) action given the belief p_t at each time t . As in Bolton and Harris (1999), the belief follows the law of motion,

$$(1) \quad dp_t = -\psi(1 - a(p_t))\gamma(p_t) \underbrace{\left[\frac{dX_t}{\sigma} - \psi(p_t a(p_t) + 1 - p_t)dt \right]}_{=: d\tilde{B}_t},$$

where ψ is the signal-to-noise ratio parameter, $\gamma(\cdot)$ is a function defined by $\gamma(p) := p(1 - p)$, and $\tilde{B}_t$ defined in the bracket is the innovation process associated with the filtration of the principal.

In our equilibrium concept, we require that the players optimally respond to each other's strategy and the principal's belief updating process. In addition, we impose minimum conditions that enable the utilization of continuous-time techniques. In particular, we focus on the equilibrium in which the agent's policy function $a(p)$ is well-behaved so that (i) the belief process admits a unique strong solution on and off the equilibrium path, and (ii) the agent's induced value function is regular enough

⁷ WeWork's rapid expansion increased its revenue by \$1 billion (or 106 percent) between 2017 and 2018, but this turned out to come at a much higher cost, resulting in \$1 billion decline in economic earnings (see David Trainer, "WeWork Is the Most Ridiculous IPO of 2019," *Forbes*, August 27, 2019). Ofo, a defunct bike-sharing startup headquartered in Beijing, adopted an expansion strategy between 2016 and 2018 that resulted in immense cash flow pressure and led to the appropriation of user deposits which were supposedly refundable at any time (see Yue Wang, "How China's Bike-Sharing Startup Ofo Went from Tech Darling to Near Bankruptcy," *Forbes*, December 20, 2018).

for standard stochastic calculus tools to apply. The reader is referred to Appendix A for the formal definition of equilibrium.

It is worth noting that in a Markov equilibrium, $a(p)$ serves two roles: it describes both the agent's behavior and the principal's *anticipation* about the agent's behavior. Indeed, because the agent's action is unobservable to the principal, the principal's belief updating rule is based on what *she anticipates* the agent is doing, which in equilibrium must coincide with what the agent is actually doing.⁸

As usual, one could change a player's policy function on a zero-measure set without affecting the outcome or the payoffs from the game. Therefore, in what follows, we do not distinguish Markov equilibria whose policy functions agree almost everywhere.

B. Equilibrium Characterization

We first introduce some terminology to define properties for the players' policy functions.

DEFINITION 1: A policy function b for the principal has a **cutoff structure** if there exists $\tilde{p} \in [0, 1]$ such that $b(p) = 0$ for $p < \tilde{p}$, and $b(p) = 1$ for $p > \tilde{p}$. We refer to $\tilde{p}$ as the *cutoff belief* of b .

DEFINITION 2: A policy function a for the (noninvestible) agent is **fully separating** if $a(p) = 0$ for all $p \in (0, 1)$.

DEFINITION 3: A policy function a for the (noninvestible) agent is **hump-shaped** if a is continuous and there are cutoffs $0 < p_L < p^* < p_R < 1$ such that $a(p) = 0$ for $p \leq p_L$, strictly increasing on (p_L, p^*) , strictly decreasing on (p^*, p_R) , and $a(p) = 0$ for $p \geq p_R$.

THEOREM 1: There always exists a unique Markov equilibrium (a, b) . In this equilibrium, b has a cutoff structure with some cutoff belief $p^* \in (0, 1)$. Moreover, there exists $r^* > 0$ such that:

- (i) If $r_1 \geq r^*$, then a is fully separating.
- (ii) If $r_1 < r^*$, then a is hump-shaped and is maximized at p^* .

Theorem 1 characterizes the structure of the unique Markov equilibrium. First, the principal uses a cutoff strategy. In this environment, the principal decides when to optimally stop her experimentation. If she stops the game with belief p , her

⁸ If the noninvestible agent takes action α_t and the principal anticipates him to play the equilibrium action $a(p_t)$, then the belief process from the noninvestible agent's perspective follows:

$$dp_t = \psi^2(1 - a(p_t)) [1 - \alpha_t - p_t(1 - a(p_t))] \gamma(p_t) dt - \psi(1 - a(p_t)) \gamma(p_t) dB_t.$$

Note that the law of motion of the belief process from the agent's perspective is different than that from the principal's perspective in (1), even after setting $\alpha_t = a(p_t)$ in equilibrium. This is because the agent knows his type and can use this information to better predict dX_t .

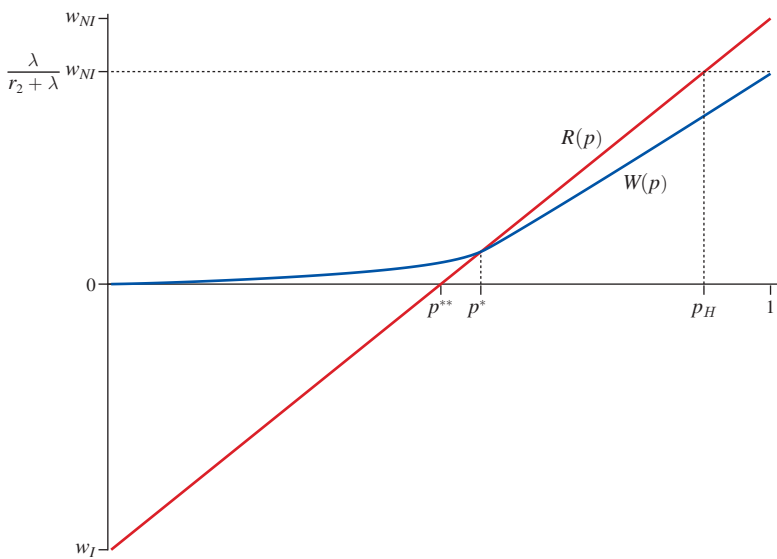


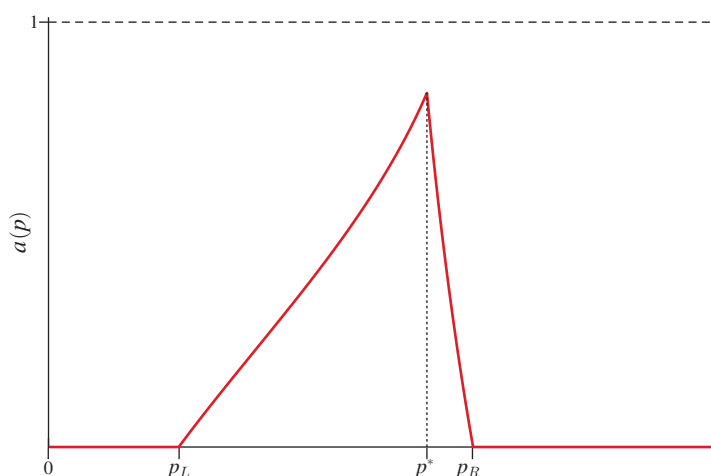
FIGURE 1. PRINCIPAL’S EQUILIBRIUM CUTOFF AND VALUE FUNCTION

Note: This figure is plotted under the following parameter values: $r_1 = 0.5$, $r_2 = 0.5$, $\lambda = 2$, $\psi = 1.5$, $u = 1$, $c = 1$, $w_{NI} = 1$, $w_I = -1$.

expected payoff is $R(p) := pw_{NI} + (1 - p)w_I$. When a stopping opportunity arrives, the principal should terminate the relationship if the continuation value $W(p)$ from the relationship is less than the termination payoff $R(p)$. This occurs when the agent’s reputation is sufficiently bad (i.e., when p is large enough), as shown in Figure 1.

Figure 1 also illustrates that the equilibrium cutoff p^* lies between $p^{**} := R^{-1}(0)$ and $p_H := R^{-1}(\lambda w_{NI}/(r_2 + \lambda))$. This happens because p^{**} is the principal’s myopic cutoff if she were completely impatient or if there were no future information, and thus $p^* > p^{**}$ reflects the option value of waiting for more information to a patient principal. Moreover, the principal has absolutely no reason to continue the game if her belief is greater than p_H , as the continuation value of the game is at most $\lambda w_{NI}/(r_2 + \lambda)$.

Turning to the agent, the noninvestible type’s behavior depends on his discount rate. If he is impatient (i.e., with a high discount rate), then he never mimics the investible type, because he always finds the saving of the mimicking cost to outweigh the benefit of having a better reputation. A richer dynamic arises if the agent is patient (i.e., with a low discount rate). In this case, his behavior can be described by three reputation phases: good, medium, and bad, as depicted in Figure 2. Both in the good and the bad reputation phases, the noninvestible type does not mimic the investible type at all, but for different reasons: when his reputation is good ($p < p_L$), the relationship is highly stable, so the noninvestible agent gains little from further improving his reputation by costly mimicking; when his reputation is bad ($p > p_R$), termination is so imminent that the noninvestible agent gives up building reputation. In the intermediate phase, however, the agent’s short-run temptation and long-run benefits are more balanced, so $a(p) \in (0, 1)$. Specifically, as his

FIGURE 2. AGENT'S EQUILIBRIUM POLICY FUNCTION WHEN $r_1 < r^*$

Note: This figure is plotted under the following parameter values: $r_1 = 0.5$, $r_2 = 0.5$, $\lambda = 2$, $\psi = 1.5$, $u = 1$, $c = 1$, $w_{NI} = 1$, $w_I = -1$.

reputation worsens, the noninvestible type first starts to mimic more often in order to slow down the principal's learning. We interpret this as a “*scramble-to-rescue*” effect: the agent increases his mimicking intensity (before p^*) as the relationship gets less stable. After a certain point, he gradually gives up as the relationship becomes doomed. His mimicking intensity is highest at belief p^* when the principal's action switches from continuing the relationship to termination.⁹

That the agent's mimicking intensity increases as the belief gets closer to the termination cutoff p^* can be understood from an equilibrium perspective. In this game, by paying the mimicking cost, the noninvestible type can boost his performance, improve his reputation, and delay termination. Thus, in general, the agent's incentive to manipulate depends on (i) how responsive the belief is to signal realizations, and (ii) how sensitive the principal's decision is to belief changes. As the belief gets closer to the termination cutoff, the principal's decision becomes more sensitive to the agent's reputation. Through (ii), this would give the agent a strict incentive to mimic had the principal's anticipation about the agent's mimicking intensity stayed constant. In equilibrium, as the belief gets closer to p^* , the agent's mimicking intensity needs to increase, with the principal adjusting her anticipation accordingly. The principal's adjustment reduces the responsiveness of her belief to signal realizations and, through (i), lowers the marginal benefit of manipulation. The equilibrium mimicking intensity is uniquely pinned down by equating the marginal cost and the marginal benefit of manipulation.

The “scramble-to-rescue” effect we find offers an important empirical prediction: that a startup's manipulative behavior can be rather severe before the collapse or

⁹Even though the principal's optimal action switches at p^* , the agent's mimicking intensity does not immediately drop to zero right after the belief passes p^* . This is because the relationship can only be terminated when a stopping opportunity arrives, leaving some hope for the agent to rebuild his reputation and avoid termination.

termination occurs. While the manipulative action is unobservable in nature (and by assumption), investigations of and admissions from troubled companies reveal some evidence consistent with this prediction. For example, prior to the unraveling of scandals, Theranos exaggerated its annual revenues in 2014 by a thousand times; using a metric that ignored certain expensive costs, WeWork reported its \$900 million loss for the first half of 2019 as a \$142 million profit, right before their scheduled IPO which was later called off; Ofo started to appropriate user deposits for its expensive daily operations and expansion, not long before investors stopped its funding in 2019.¹⁰

III. Nonmonotonicity of Expected Performance

We have seen in Theorem 1 that the noninvestible type will engage in performance boosting whenever he is sufficiently patient, in which case $a(\cdot)$ peaks at the termination cutoff p^* . We refer to it as the “scramble-to-rescue” effect: the agent increases his mimicking intensity (before p^*) as the relationship gets less stable.

What are the implications of this effect for the observables? An outsider (the principal or a modeler) does not see the agent’s type or action but can observe his performance, such as sales growth and progress reports. Moreover, before observing the agent’s performance, an outsider can form an expectation about it. Specifically, the expected performance at time t is given by

$$EP_t := \frac{E[dX_t]}{dt} = \mu \left(\underbrace{1 - p_t}_{\text{investible}} + \underbrace{p_t a(p_t)}_{\text{noninvestible}} \right).$$

As implied by equation (1), the agent’s reputation gets better (i.e., $dp_t < 0$) if and only if the observed performance beats the expectation. Hence, the expected performance serves as a *standard* for evaluating the observed performance of the agent.

Let us first examine how the agent’s expected performance varies with the public belief. Holding constant $a < 1$, the expected performance decreases with p . We call this the *belief effect*: if the agent is more likely to be noninvestible, then the expected performance tends to be lower. This is the entire story if the equilibrium is fully separating, in which case $a(p) = 0$ everywhere and $EP(p) = \mu(1 - p)$.

However, if the equilibrium is not fully separating, the noninvestible type’s mimicking intensity a is no longer constant: it is increasing below p^* due to the scramble-to-rescue effect. Hence, whether the expected performance increases or decreases with the public belief depends on which of the two effects is stronger. The following theorem characterizes the evolution of the expected performance when the stopping opportunity arrives sufficiently fast.

¹⁰See Chloe Aiello, “Theranos President Exaggerated the Company’s Revenue by 1,000 Times to Investors, Says SEC,” *CNBC*, March 14, 2018; Amanda Iacone, “SEC Flags Cash-Flow Measure That Made WeWork Look Profitable,” *Bloomberg Tax*, December 3, 2019; Jun Zhang and Siyi Li, “Across the Ruins: The Bicycle Sharing Drama Has Not Ended” [in Chinese], *Tencent News*, April 5, 2021.

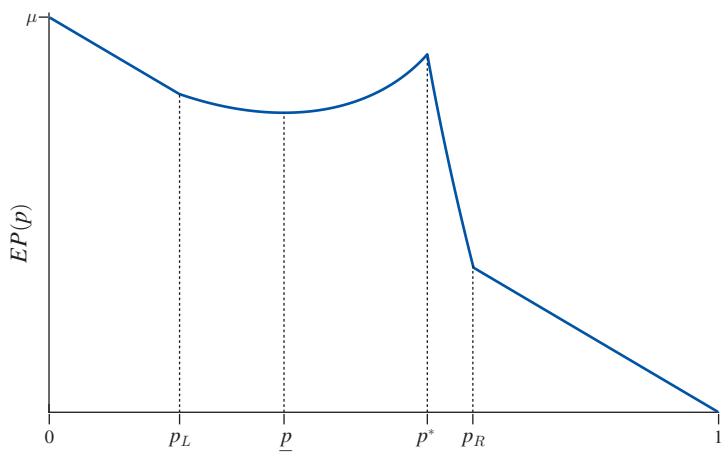


FIGURE 3. AGENT’S EXPECTED PERFORMANCE WHEN $\lambda > \bar{\lambda}$ AND $r_1 < r^*$

Note: This figure is plotted under the following parameter values: $r_1 = 0.5, r_2 = 0.5, \lambda = 2, \psi = 1.5, u = 1, c = 1, w_{NI} = 1, w_I = -1$.

THEOREM 2: *There exists $\bar{\lambda}$ such that for all $\lambda > \bar{\lambda}$, $EP(p)$ is nonmonotone whenever the equilibrium is not fully separating (i.e., whenever $r_1 < r^*$). In particular, $EP(p)$ is*

- (i) *strictly decreasing for $p \in [0, \underline{p})$, where $\underline{p}$ is in $[p_L, p^*)$;*
- (ii) *strictly increasing for $p \in (\underline{p}, p^*)$;*
- (iii) *strictly decreasing for $p \in (p^*, 1]$.*

Theorem 2 shows that if the principal’s decision frictions are small ($\lambda > \bar{\lambda}$), the scramble-to-rescue effect will dominate the belief effect when the public belief is less than but close to the termination cutoff p^* . As a result, whenever the agent is relatively patient ($r_1 < r^*$), his expected performance reaches a local maximum at p^* , as illustrated in Figure 3. This result about the spike in expected performance prior to termination also carries over to the time domain. Indeed, one can show that the stochastic process EP_t has a positive drift whenever $p_t \in (\underline{p}, p^*)$.¹¹

An interesting prediction from Theorem 2 is that, when a startup’s reputation is close to the termination cutoff (i.e., when $p \in (\underline{p}, p^*)$), observing an unsatisfactory performance in the current period will lead to a higher market expectation in the next period, making it even harder to meet the expectation and improve the reputation.

This prediction about the expected performance also speaks to some seemingly paradoxical phenomena regarding the observed performance. In particular, in various

¹¹ This follows from the facts that (i) $EP(\cdot)$ is convex on $(\underline{p}, p^*)$, and (ii) the belief is a martingale.

famous cases of corporate failure, such as WeWork, Ofo and Theranos, there was a period of strong or even increasing performance before the company suddenly fell apart.¹² Our theory offers an explanation from a reputation perspective for why such performance spikes can occur before sudden collapses or terminations.

Specifically, for a startup with a mediocre reputation, a subpar performance in one period will result in a higher evaluation standard in the next period. As a result, even if the firm performs better in the next period, it can still fall short of the higher expectation, which worsens its reputation and increases the evaluation standard even further. Therefore, despite improving its performance over time, the firm may be trapped in a spiral of substandard performance and harder-to-meet standards. Such increasing trend in the observed (and expected) performance and worsening trend in reputation can continue until the investor's belief moves beyond a tipping point, and the termination ensues.

IV. Environments with Low Volatility/High Transparency

The noise in the performance measure may come from various random events such as temporary demand shocks, measurement errors, etc., making X_t only an imperfect signal of the agent's type. We say that a performance measure is more transparent if it is less affected by the noise component, and we can use the signal-to-noise ratio of the process to capture its *transparency*. In reality, transparency may be determined by the intrinsic volatility of the product market; it may also be affected by how much detail about the market or the project a startup is required to disclose in a performance report. One may expect that, other things equal, disclosing more details about the objective market conditions can help an investor better understand the numbers in the report.

In this section, we investigate the question, do improvements in transparency always benefit the principal? We show that, under some parameters, improving transparency can inhibit learning and hurt the principal, due to the agent's endogenous response through signal manipulation.

Recall that the signal-to-noise ratio parameter is defined by $\psi = \mu/\sigma$. In what follows, we will fix all parameters other than ψ , and analyze the principal's payoff as ψ increases. Hence, we make explicit the dependence of any variable or function on ψ .

As a benchmark, suppose that the principal never receives any information about the agent's type. Recall that the principal's termination payoff at belief p is $R(p) = pw_{NI} + (1-p)w_I$, and her myopic cutoff is $p^{**} = R^{-1}(0)$. In this case, the principal would continue the relationship if $p < p^{**}$, and she would terminate

¹²For example, WeWork once expanded to over 86 cities in 32 countries, despite growing suspicion about its business model and profitability (see Rebecca Aydin, "WeWork Isn't Even Close to Being Profitable," *Business Insider*, July 3, 2019). However, in September 2019, the company delayed its IPO, followed by a 90 percent slash in valuation and enormous layoffs. Similarly, Ofo experienced a period of rapid growth in 2016–2018, expanding its services to more than 250 cities around the world, before all investors discontinued their funding in 2019 (see Alison Griswold, "Bike-Sharing Company Ofo Is Dramatically Scaling Back in North America," *Quartz*, July 19, 2018). In the case of Theranos, before the scandals unraveled, the company falsely reported \$100 million of annual revenues in 2014 and even obtained US Food and Drug Administration approval of one of its blood tests in 2015 (see Lauren Friedman, "Controversial Multibillion-Dollar Health Startup Theranos Just Got a Huge Seal of Approval from the US Government," *Business Insider*, July 2, 2015).

the relationship at the first stopping opportunity if $p > p^{**}$. This leads to the following “no-information value function” for the principal:

$$\underline{W}(p) := \frac{\lambda}{r_2 + \lambda} \max\{0, R(p)\}.$$

Note that when ψ is near zero, the signal process is close to pure noise regardless of the agent’s action. So we have the following observation.

OBSERVATION 1: $\lim_{\psi \rightarrow 0} W(p_0; \psi) = \underline{W}(p_0)$ and $\lim_{\psi \rightarrow 0} p^*(\psi) = p^{**}$.

Now suppose that the agent’s type is exogenously and immediately revealed to the principal. In this case, after information revelation the principal will always obtain her highest possible payoff (that is, $\max\{0, \lambda w_\theta / (r_2 + \lambda)\}$), which leads to her “full-information value function”:

$$\bar{W}(p) := \frac{\lambda}{r_2 + \lambda} p w_{NI}.$$

Note that the principal’s equilibrium payoff is always strictly between $\underline{W}(p_0)$ and $\bar{W}(p_0)$. This is because some learning will take place in equilibrium (as the noninvestible type never fully mimics), but the agent’s type is not immediately revealed (as $\psi < \infty$).

OBSERVATION 2: For all $\psi \in (0, \infty)$ and $p_0 \in (0, 1)$, $W(p_0; \psi) \in (\underline{W}(p_0), \bar{W}(p_0))$.

The next theorem characterizes the principal’s limiting behavior and payoff as ψ goes to infinity.

THEOREM 3: There exists $\tilde{\lambda}$ such that whenever $\lambda > \tilde{\lambda}$, as the signal-to-noise ratio $\psi \rightarrow \infty$, the principal’s equilibrium cutoff $p^*(\psi)$ converges to the myopic cutoff p^{**} , and her value function $W(\cdot; \psi)$ converges uniformly to the no-information value function $\underline{W}(\cdot)$.

Theorem 3 states that, whenever the stopping frictions are small, as transparency grows without bound (i.e., as the volatility of the performance measure shrinks to zero), the principal’s equilibrium behavior and value function converge in a way as if no information would ever arrive (see Figure 4).

This result may be counterintuitive at the first sight, but it can be understood from the agent’s incentives and equilibrium response. Specifically, as the signal-to-noise ratio ψ grows, *other things equal*, the principal can learn the agent’s type more quickly. In turn, it gives the noninvestible type a stronger incentive to mimic the investible type. This is reflected in the higher levels of the agent’s mimicking intensity across the belief space, but most importantly, around the termination cutoff, leading to a stronger scramble-to-rescue effect. In particular, as ψ increases without bound, the equilibrium mimicking intensity at the termination cutoff converges to one, which largely slows down the principal’s learning.

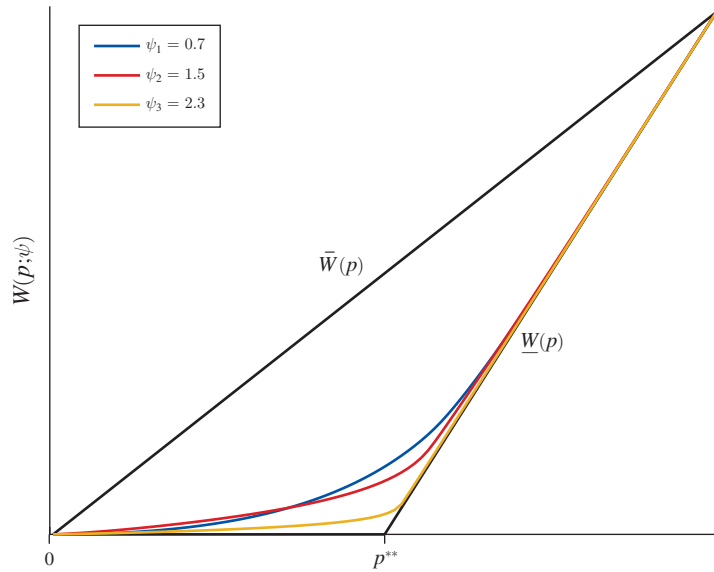


FIGURE 4. CONVERGENCE OF THE PRINCIPAL'S EQUILIBRIUM VALUE FUNCTION

Note: This figure is plotted under the following parameter values: $r_1 = 0.5$, $r_2 = 0.5$, $\lambda = 2$, $u = 1$, $c = 1$, $w_{NI} = 1$, $w_I = -1$.

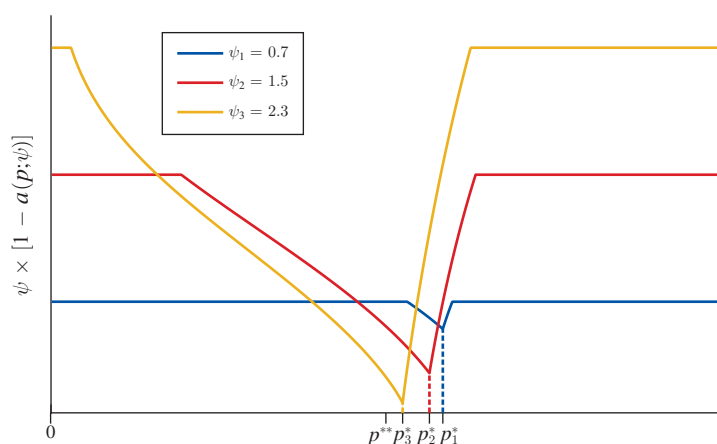
Note, however, that the increase in mimicking intensity does not fully explain the result because the signal-to-noise also grows substantially. Indeed, from equation (1) we see that the diffusion coefficient of the principal's belief process at time t , which measures the variance of her belief and thus her speed of learning, is proportional to $\psi(1 - a_t)$. For any fixed p , we can write it as

$$\underbrace{\psi}_{\text{direct effect}} \left[1 - \underbrace{a(p; \psi)}_{\text{equilibrium effect}} \right].$$

As ψ increases, the *direct effect* through the multiplier accelerates information revelation while the *equilibrium effect* through the agent's response slows down learning. We show that, not only does $a(p^*; \psi)$ converge to one, the speed of this convergence is so fast that the variance of the belief process vanishes there (as shown in Figure 5), making p^* (almost) an absorbing state. This creates a *learning barrier* at the termination cutoff, because the belief can hardly move across p^* . In the limit, even though learning does not stop when the principal's belief is above or below the termination cutoff, that her belief never crosses the cutoff means that the principal's action never changes with the information she learns. Hence, payoff-wise it is as if no information is ever revealed.

The next corollary follows immediately from Theorem 3.

COROLLARY 1: For any $\lambda > \tilde{\lambda}$ and $p_0 \in (0, 1)$, there exist ψ_1, ψ_2 such that $\psi_1 > \psi_2$ and $W(p_0; \psi_1) < W(p_0; \psi_2)$.

FIGURE 5. EMERGENCE OF A LEARNING BARRIER AS $\psi \rightarrow \infty$

Notes: This figure is plotted under the following parameter values: $r_1 = 0.5$, $r_2 = 0.5$, $\lambda = 2$, $u = 1$, $c = 1$, $w_{NI} = 1$, $w_I = -1$. On the horizontal axis, $p_1^* = p^*(\psi_1)$, $p_2^* = p^*(\psi_2)$, and $p_3^* = p^*(\psi_3)$.

Corollary 1 illustrates that a more transparent performance measure (higher ψ) can sometimes reduce the principal's payoff. This result indicates that VCs can sometimes fare better when the startup operates in a more volatile environment. In addition, policies that require the startup to disclose too precise information may end up hurting the investors. This provides a rationale for the views against excessive due diligence in VC investing.¹³ Intuitively, when the signal-to-noise ratio is very high, the mimicking incentives of the noninvestible type can be so strong that the public signal provides little useful information to the principal about the agent's type. In that case, introducing more noise to the public signal can lessen such perverse incentives of the noninvestible type, which leads to a more informative equilibrium signal process, benefiting the principal.

The insight that opaqueness can enhance information transmission by discouraging manipulation also applies to other contexts. For example, credit bureaus such as TransUnion never publish their exact credit scoring formulae in order to prevent consumers from “gaming the system” (Weston 2016). For similar reasons, central banks do not (fully) reveal their stress test models to individual banks (Leitner and Williams forthcoming).

V. Information at the Patient Limit

We now investigate the equilibrium outcomes as players get arbitrarily patient. First, consider an extreme case where r_1 is constant and r_2 goes to zero (i.e., the principal gets arbitrarily more patient than the agent). In this case, the agent's mimicking intensity (which is independent of r_2) stays bounded away from one everywhere, implying that the public signal always reveals some information about the

¹³ See, for example, Connie Loizos, “The Truth about Due Diligence,” *TechCrunch*, May 19, 2016; Fred Wilson, “You Can Do Too Much Due Diligence,” *Business Insider*, May 13, 2013.

agent's type. As the principal gets more patient, her marginal cost of waiting for new information becomes lower. Consequently, a patient principal will terminate the relationship only when p is very high. Indeed, the termination cutoff converges to one, and the principal's payoff converges to the full-information value $\bar{W}$.

Next, consider the other extreme case where r_2 is constant and r_1 goes to zero (i.e., the agent gets arbitrarily more patient than the principal). As the agent gets more patient, he cares more about staying in the relationship for long and less about the instantaneous mimicking cost. Thus, the noninvestible type has stronger incentives to mimic the investible type. Similar to what happens in Section IV, the equilibrium mimicking intensity approaches one at and below the termination cutoff, leading to an extreme form of scramble-to-rescue effect in the limit. This again creates a learning barrier at the termination cutoff, so that the principal's action does not change with the information she learns. Hence, from the principal's perspective, it is as if no information ever arrives, so the termination cutoff converges to p^{**} and the principal's payoff converges to the no-information value $\underline{W}$.

What happens in between the two extreme cases? Take a sequence $\{r_{1,n}, r_{2,n}\}_n$ of discount rates such that $r_{i,n} \rightarrow 0$ for both $i = 1, 2$ and $\lim_n (r_{2,n}/r_{1,n}) = \chi \in (0, \infty)$. Consider a sequence of games along which all other parameters are fixed, and let $\{W_n, V_n\}_n$ be the corresponding sequence of value functions for the principal and the agent, respectively. The theorem below displays their limits.

We say a function $V^*: (0, 1) \rightarrow \mathbb{R}$ is the agent's full-mimicking value function if it satisfies

$$V^*(p) := \begin{cases} u, & \text{if } p < p^{**}; \\ 0, & \text{if } p > p^{**}. \end{cases}$$

Such function delivers the value of a patient agent who fully mimics all the time, when the principal optimally responds by terminating the relationship whenever $p > p^{**}$.¹⁴

THEOREM 4: *Suppose that $\lim_n r_{i,n} = 0$ for both $i = 1, 2$ and $\lim_n \frac{r_{2,n}}{r_{1,n}} = \chi \in (0, \infty)$.*

- (i) *For the principal, $W_n(\cdot)$ converges uniformly to her no-information value function.*
- (ii) *For the agent, $V_n(\cdot)$ converges pointwise to a full-mimicking value function.*

Theorem 4 shows that if both players get arbitrarily patient at comparable rates, then the principal gets a payoff as if she does not receive any information, and the agent gets a payoff as if he commits to full mimicking all the time.

The key driving force of this result is again a learning barrier created by the scramble-to-rescue effect. Indeed, the agent getting arbitrarily patient induces him to approach full mimicking at any belief below the termination cutoff, creating a learning barrier at the cutoff. That is, from the agent's perspective, the belief never

¹⁴The definition of a full-mimicking value function leaves the value at p^{**} indefinite, leading to a class of functions satisfying this definition. As a consequence, the second part of Theorem 4 applies to all points in $(0, 1)$ but p^{**} .

moves across the termination cutoff. As a result, if the current belief is below the cutoff, the agent can stay in the relationship forever while paying the mimicking cost, leading to a payoff of u ; if the current belief is above the cutoff, he will be terminated soon with a payoff of zero. For the principal, because she gets increasingly patient at a rate comparable to the agent, from the principal's perspective the termination cutoff also becomes a learning barrier.¹⁵ Such a learning barrier makes her action unresponsive to the information she learns, so that payoff-wise it is as if no information is ever revealed.

We view this result as a strong manifestation of the ratchet effect in the patient limit of our model. Since the principal cannot commit to not using future information against the agent, the noninvestible type will engage in performance boosting with almost full intensity in order to maintain his reputation. In the end, no useful information is ever revealed, and the principal's lack of commitment hurts her in the most extreme way. In our applications, this result suggests that the use of other instruments, such as some form of commitment (e.g., setting a deadline and/or a grace period), additional screening devices (e.g., performance-based investment levels and/or salaries), or huge fines that increase the expected cost of performance boosting, may be necessary to help the principal get more information.

VI. Discussion and Conclusion

A. The Poisson Arrival of Stopping Opportunities

In our model, the principal can terminate the relationship only when a stopping opportunity arrives. We now explain the technical advantages of this modeling approach, and the robustness of our results as the frictions vanish (i.e., $\lambda \rightarrow \infty$).

Frictions as a Technical Assumption.—It is natural to consider a frictionless environment in which the principal can terminate the relationship at any time. In that case, in any Markov equilibrium, the principal will still use a cutoff strategy, terminating the relationship if and only if her belief is greater than some threshold p_∞^* . Two technical issues then arise. First, one has to deal with off-equilibrium histories and beliefs. In particular, any belief greater than p_∞^* is off the equilibrium path because the relationship should have been terminated the moment the belief moves beyond p_∞^* . Second, a plethora of equilibria emerge. To see why, one might expect that the agent's best reply when $p > p_\infty^*$ is not mimicking at all, because the relationship is about to be terminated in the next instant and paying the mimicking cost is futile. However, this intuition is not quite correct in continuous time because if the agent expects the relationship to be terminated right away, his action choice has no payoff consequences, implying that he is *indifferent* between all choices. The agent's indifference can also feed back to the principal's optimal choice of the termination cutoff. As a result, there is a continuum of Markov equilibria, with the

¹⁵ If the principal gets patient much faster than does the agent, then a learning barrier for the agent does not necessarily translate to a learning barrier for the principal, as the principal may put much higher values on future movements of the belief.

principal using different termination cutoffs and the agent choosing different levels of mimicking intensity when the relationship is to be terminated immediately.

By introducing small frictions to the principal's stopping decision, all histories are on path because the relationship can last for an arbitrarily long time due to the absence of a stopping opportunity. Moreover, it allows us to obtain equilibrium uniqueness (within the class of Markov strategy profiles) without additional refinements, as the agent's sequential rationality is now a binding condition at all beliefs. In addition, as we will see, taking the limit as the frictions vanish provides an equilibrium refinement for the frictionless environment.

Robustness of Results as Frictions Vanish.—Next, we explain what happens to our main results when the arrival rate λ of the stopping opportunity gets arbitrarily large.

For equilibrium strategies, the principal always follows a cutoff rule, and as $\lambda \rightarrow \infty$, the equilibrium threshold converges to some p_∞^* greater than the myopic cutoff p^{**} . The limit of the agent's policy function is depicted in Figure 6.

Specifically, when $p < p_\infty^*$, the agent does not mimic at all if his reputation is sufficiently good (i.e., if $p < p_{L,\infty}$), and as his reputation deteriorates he starts to manipulate the signal more intensely due to the scramble-to-rescue effect. When $p > p_\infty^*$, the agent stops mimicking completely, expecting the relationship to end in the next instant. As we just discussed, there are multiple equilibria in a frictionless environment because of the agent's indifference, but the most intuitive equilibrium is arguably the one where the agent does not mimic whenever termination is immediate. Hence, the limiting policy functions as $\lambda \rightarrow \infty$ provides a reasonable equilibrium selection for the frictionless environment.¹⁶

For Theorem 2, given that the scramble-to-rescue effect continues to be present and dominant as $\lambda \rightarrow \infty$, our qualitative characterization of the expected performance remains the same: it is nonmonotone with respect to the agent's reputation and has a peak at termination cutoff.

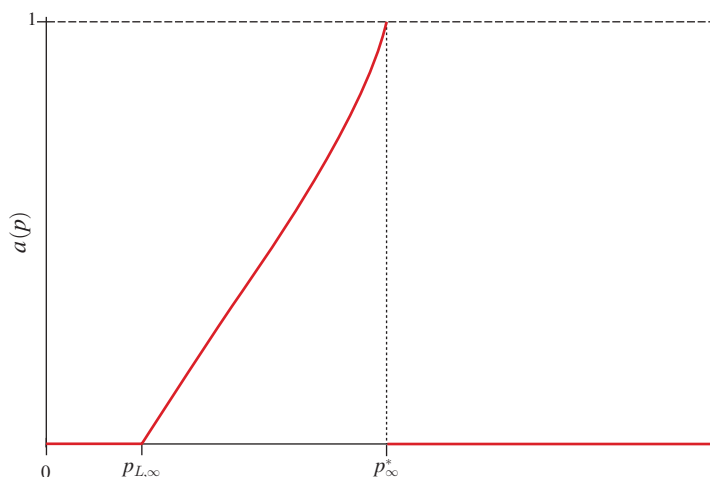
For Theorem 3, one can show that $\lim_{\lambda \rightarrow \infty, \psi \rightarrow \infty} |W(\cdot; \lambda, \psi) - \underline{W}(\cdot; \lambda)| = 0$. That is, the principal's equilibrium value function continues to converge to her no-information value function (as $\psi \rightarrow \infty$), even when λ increases to infinity at the same time. Moreover, because the convergence is uniform for any fixed λ , the order of taking the limits does not matter.

For Theorem 4, when the players get arbitrarily patient, the arrival rate of the stopping opportunity becomes irrelevant because any future events are almost immediate for payoff purposes. Consequently, the result is unaffected by taking $\lambda \rightarrow \infty$.

B. Commitment

We have focused on the case in which the principal cannot commit to a strategy. But what would happen if the principal had some commitment power and, for example, were able to commit to terminating the relationship at the first opportunity

¹⁶Kuvalekar and Lipnowski (2020) introduce an alternative equilibrium refinement that will also select the most intuitive Markov equilibrium in the frictionless environment (i.e., the agent does not mimic whenever $p > p_\infty^*$).

FIGURE 6. AGENT'S LIMITING EQUILIBRIUM POLICY FUNCTION AS $\lambda \rightarrow \infty$

Note: This figure is plotted under the following parameter values: $r_1 = 0.5$, $r_2 = 0.5$, $\psi = 1.5$, $u = 1$, $c = 1$, $w_{NI} = 1$, $w_I = -1$.

whenever p_t lies above a certain cutoff? Assume that the agent's initial reputation does not warrant immediate termination ($p_0 < p^{**}$). Provided frictions are sufficiently low (i.e., λ is high), one can show that the best commitment cutoff lies above the equilibrium cutoff. The intuition is simple. Consider a marginal increase from the equilibrium cutoff. If the principal commits to this higher cutoff at the beginning of the game, the agent expects a longer delay until being terminated. Because of that, the agent has weaker incentives to engage in performance boosting. This in turn speeds up information revelation, and in the absence of other first-order effects, the principal's payoff is improved. Interestingly, the principal's commitment also benefits the agent, as a higher termination cutoff implies a longer duration of the relationship. This Pareto improvement can be achieved because the principal's commitment helps to reduce performance boosting, which is a wasteful action intended to forestall termination by inhibiting learning.¹⁷

Of course, fixing a cutoff represents only a restricted form of commitment. What would happen if the principal were able to fully commit to a general mechanism? In this case, we can describe some qualitative features of the optimal mechanism in a frictionless environment ($\lambda = \infty$). In the optimal mechanism, the principal commits to a type-dependent contract aimed at eliminating wasteful signaling by the noninvestible type. On the one hand, if the agent reports being noninvestible, there is no reason to make him engage in performance boosting. The optimal contract entails a deterministic deadline $T > 0$. Faced with this deadline, the non-

¹⁷Note, however, that the benefits from committing to a higher cutoff gradually vanish as the players become more patient. Similar to the finding in Theorem 4, as they get arbitrarily patient, the agent's manipulation incentives become so strong that the principal is unable to learn any useful information from the signal. Therefore, regardless of the termination cutoff, in the patient limit the principal always gets her no-information value and the agent gets his full-mimicking value.

investible agent never chooses to boost his performance as it has no effect on the termination date. On the other hand, the contract intended for the investible type is designed to preclude the noninvestible type from pretending to be investible. In this contract, the principal calibrates the sensitivity of the contract with respect to the observable performance in such a way that if the noninvestible type were to take this contract then he would (weakly) prefer to shirk at every point in time. The deadline evolves stochastically according to a score based on the observable performance.

C. Concluding Remarks

Several ways to extend our analysis are worth pursuing. While our main focus is the adverse selection problem, in some settings moral hazard is a prominent issue. Thus, it would be interesting to allow the agent's action to directly influence the principal's payoff. Another possibility is to expand the choice set of the principal by allowing her to elevate the status of the relationship, such as promoting the agent or upgrading the terms of financing. We leave these interesting directions for future research.

APPENDIX A. EQUILIBRIUM DEFINITION

In this section, we provide a formal definition of our equilibrium concept. We freely use basic concepts from the theory of stochastic processes. We refer to Karatzas and Shreve (1998) for the definitions. We assume that the Brownian motion B driving the observable signal X is defined on a filtered probability space $(\Omega, \mathcal{F}, \{\mathcal{F}_t\}_{t \geq 0}, \mathbb{P})$ satisfying the usual conditions.

A *strategy* for the principal (respectively agent) is a measurable stochastic process taking values on $[0, 1]$ and adapted to the filtration generated by X (respectively B). A *policy function* is a measurable function $(0, 1) \rightarrow [0, 1]$.

DEFINITION A1: A *Markov equilibrium* is a pair of policy functions (a, b) that satisfies the following:

- (i) **Well-defined Beliefs on and off Path:** The belief process (A1) from the agent's perspective admits a unique strong solution for any choice of the agent's strategy α .

$$(A1) \quad dp_t = \psi^2(1 - a(p_t)) [1 - \alpha_t - p_t(1 - a(p_t))] \gamma(p_t) dt - \psi(1 - a(p_t)) \gamma(p_t) dB_t.$$

- (ii) **Principal Optimality:** Policy b maximizes the principal's expected payoff, attaining

$$W(p_0) := \sup_{\beta} E \left\{ \int_0^{\infty} e^{-r_2 t} e^{\int_0^t -\lambda \beta_{\tau} d\tau} \lambda \beta_t [p_t w_{NI} + (1 - p_t) w_I] dt \right\},$$

where beliefs follow equation (1) and β is chosen among all of the principal's strategies.

(iii) **Agent Optimality:** Policy a maximizes the agent's expected payoff, attaining

$$V(p_0) := \sup_{\alpha} E \left[r_1 \int_0^{\infty} e^{-r_1 t} e^{\int_0^t -\lambda b(p_{\tau}) d\tau} [u + (1 - \alpha_t)c] dt \right],$$

where beliefs follow equation (A1) and α is chosen among all of the agent's strategies.

(iv) **Regularity:** The agent's value function $V(\cdot)$ is C^1 everywhere and C^2 everywhere except perhaps at a finite number of points.

Condition (i) ensures a well-defined belief process even if the agent were to deviate. Optimality conditions (ii) and (iii) are part of any equilibrium definition. Condition (iv) enables standard stochastic calculus tools.

APPENDIX B. PROOFS

Remark.—Whenever there is no confusion we simply refer to the “noninvestible-type agent” as the agent. We denote by Φ and ϕ the cumulative distribution function and probability density function of the standard normal distribution, respectively.

A. Equilibrium Characterization: Toward a Proof of Theorem 1

Consider a Markov equilibrium, (a, b) , and the underlying probability space $(\Omega, \mathcal{F}, \mathbb{P})$. For each $p \in (0, 1)$, we define $\Psi(p) := \{\omega \in \Omega : \exists t \leq \mathbb{T} \text{ such that } p_t(\omega) = p\}$, where $\mathbb{T}$ is the random stopping time of the relationship induced by (a, b) . The belief span, $SP(a)$, is the set of all p such that $\mathbb{P}(\Psi(p)) > 0$.

LEMMA A1: If (a, b) is a Markov equilibrium, then $SP(a) = (0, 1)$.

PROOF:

See online Appendix.

Principal's Equilibrium Behavior.—Recall that $R(p) := pw_{NI} + (1 - p)w_I$ is the principal's expected payoff if the relationship is terminated at belief p . Define $p^{**} := R^{-1}(0) > 0$ and $p_H := R^{-1}[\lambda w_{NI}/(r_2 + \lambda)] < 1$.

LEMMA A2: If (a, b) is a Markov equilibrium, then b has a cutoff structure with a cutoff belief $p^* \in [p^{**}, p_H]$.

PROOF:

For any $p \in SP(a) = (0, 1)$, define $F(p) := W(p) - R(p)$. At any time t such that the stopping opportunity arrives, given her belief $p_t = p$, if the principal terminates the relationship, her expected payoff is $R(p)$; if the principal continues the relationship, her continuation value is $W(p)$. Thus, principal optimality requires that $b(p) = 1$ if $F(p) < 0$ and that $b(p) = 0$ if $F(p) > 0$.

We first establish two useful properties of F .

PROPERTY 1: If $F(\tilde{p}) > 0$ at some $\tilde{p} \in (0, 1)$, then $F(p) > 0$ for all $p < \tilde{p}$.

To see this, suppose that $F(\tilde{p}) > 0$ at some $\tilde{p} \in (0, 1)$. Let (p_a, p_b) be the largest interval containing $\tilde{p}$ such that $F(p) > 0$ for all $p \in (p_a, p_b)$. We want to show that $p_a = 0$. Suppose, toward a contradiction, that $p_a > 0$. Since W is continuous, we have $F(p_a) = 0$, i.e., $W(p_a) = R(p_a)$. Moreover, principal optimality requires that $b(p) = 0$ for all $p \in (p_a, p_b)$. We consider two cases, and will reach a contradiction in each of these cases.

CASE 1: $p_b < 1$.

In this case, continuity of W also implies that $F(p_b) = 0$, i.e., $W(p_b) = R(p_b)$. Consider now a history that leads to the belief $\tilde{p}$ and the continuation play starting from this history. Let $\mathbb{T}^\dagger$ be the first time that the posterior belief reaches $(p_a, p_b)^c$ (setting $\mathbb{T}^\dagger = \infty$ if this event is not reached in finite time). Let φ represent the probability measure (from the principal's perspective) induced by the distribution of $p_{\mathbb{T}^\dagger}$. Then,

$$\begin{aligned} R(\tilde{p}) &< W(\tilde{p}) \\ &= W(p_a)E[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p_a]\varphi(p_{\mathbb{T}^\dagger} = p_a) \\ &\quad + W(p_b)E[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p_b]\varphi(p_{\mathbb{T}^\dagger} = p_b) \\ &= R(p_a)E[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p_a]\varphi(p_{\mathbb{T}^\dagger} = p_a) \\ &\quad + R(p_b)E[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p_b]\varphi(p_{\mathbb{T}^\dagger} = p_b) \\ &< R(p_a)\varphi(p_{\mathbb{T}^\dagger} = p_a) + R(p_b)\varphi(p_{\mathbb{T}^\dagger} = p_b) \\ &= R(\tilde{p}). \end{aligned}$$

The second equality trivially holds if $E[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p] = 0$, and otherwise if $\mathbb{T}^\dagger < \infty$, then it holds because $W(p_{\mathbb{T}^\dagger}) = R(p_{\mathbb{T}^\dagger})$. The second inequality uses the facts that $0 \leq W(p_a) = R(p_a)$ implies $R(p) > 0$ for all $p > p_a$, and that $\mathbb{T}^\dagger > 0$ almost surely. The final equality holds because p_t is a bounded martingale (so that $p_a\varphi(p_{\mathbb{T}^\dagger} = p_a) + p_b\varphi(p_{\mathbb{T}^\dagger} = p_b) = \tilde{p}$) and $R(\cdot)$ is an affine function. But then, we have an obvious contradiction.

CASE 2: $p_b = 1$.

In this case, we have $b(p) = 0$ for all $p \in (p_a, 1)$. Consider again a history that leads to the belief $\tilde{p}$ and the continuation play starting from this history. Let $\mathbb{T}^\dagger$ be the first time that the posterior reaches p_a (setting $\mathbb{T}^\dagger = \infty$ if this event is not

reached in finite time). Let φ represent the probability measure (from the principal’s perspective) induced by the distribution of $p_{\mathbb{T}^\dagger}$. Then,

$$R(\tilde{p}) < W(\tilde{p}) = W(p_a)E\left[e^{-r\mathbb{T}^\dagger}\right] < R(p_a),$$

which contradicts R being increasing.

PROPERTY 2: Let $p^* := \sup\{p : F(p) > 0\}$.¹⁸ Then, $F(p) < 0$ for all $p > p^*$.

By definition of p^* , we know that $F(p) \leq 0$ for all $p \geq p^*$, i.e., $W(p) \leq R(p)$ for all $p \geq p^*$, so it is weakly optimal for the principal to terminate the relationship whenever $p \in (p^*, 1)$. Suppose, toward a contradiction, that $F(\tilde{p}) = 0$ for some $\tilde{p} > p^*$. Consider a history that leads to the belief $\tilde{p}$ and the continuation play starting from this history. Let $\mathbb{T}^\dagger$ be the first time that the stopping opportunity arrives or the posterior reaches p^* (setting $\mathbb{T}^\dagger = \infty$ if this event is not reached in finite time). Let φ represent the probability measure (from the principal’s perspective) induced by the distribution of $p_{\mathbb{T}^\dagger}$. Then,

$$R(\tilde{p}) = W(\tilde{p}) = \int_{p^*}^1 R(p)E\left[e^{-r_2\mathbb{T}^\dagger} | p_{\mathbb{T}^\dagger} = p\right]\varphi(dp) < \int_{p^*}^1 R(p)\varphi(dp) = R(\tilde{p}).$$

The first equality follows from the contradiction assumption that $F(\tilde{p}) = 0$, the inequality holds because $0 \leq W(p^*) \leq R(p^*)$ implies that $R(p) > 0$ for all $p > p^*$, and the final equality holds because p_t is a bounded martingale (so that $\int_{p^*}^1 p\varphi(dp) = \tilde{p}$) and $R(\cdot)$ is an affine function. But then, we have an obvious contradiction, establishing Property 2.

These two properties of F immediately deliver our result. Specifically, let $p^* := \sup\{p : F(p) > 0\}$. Then, Property 1 implies that $F(p) > 0$ (and thus $b(p) = 0$) for all $p \in (0, p^*)$, and Property 2 implies that $F(p) < 0$ (and thus $b(p) = 1$) for all $p \in (p^*, 1)$. Finally, we argue that $p^* \in [p^{**}, p_H]$. Since $R(p) < 0$ for all $p \in (0, p^{**})$, principal optimality requires that $b(p) = 0$ for all such p , and so $p^* \geq p^{**}$. Moreover, because $W(\cdot)$ is the principal’s value function conditional on *not* receiving a stopping opportunity, it is bounded above by $\lambda w_{NI}/(r_2 + \lambda)$. Thus, $R(p) > W(p)$ for all $p \in (p_H, 1)$, and principal optimality requires that $b(p) = 1$ for all $p \in (p_H, 1)$, so $p^* \leq p_H$. ■

Agent’s Equilibrium Behavior.—

DEFINITION A2: A policy function $a(\cdot)$ of the agent is a **pseudo-best reply** to a policy function $b(\cdot)$ of the principal if $a(\cdot)$ satisfies conditions (i), (iii), and (iv) in Definition A1.

We call such $a(\cdot)$ pseudo-best reply (instead of “best reply”) because condition (iii) in Definition A1 implicitly assumes that the principal’s anticipation about the agent’s behavior is correct. In words, $a(\cdot)$ is a pseudo-best reply to $b(\cdot)$, if $a(\cdot)$ is

¹⁸By convention, if $\{p : F(p) > 0\} = \emptyset$, we set $\sup\{p : F(p) > 0\} = 0$.

optimal for the agent given that the principal uses policy $b(\cdot)$ and that the principal anticipates the agent to use $a(\cdot)$.

LEMMA A3: Suppose $b(\cdot)$ is a cutoff policy function of the principal with cutoff belief p^* . Then, $b(\cdot)$ has a unique continuous pseudo-best reply $a(\cdot)$. Moreover, there exists $r^* > 0$ (independent of p^*) such that $a(\cdot)$ is fully separating if $r_1 \geq r^*$, and it is hump-shaped if $r_1 < r^*$. Finally, any other pseudo-best reply to $b(\cdot)$ agrees with $a(\cdot)$ almost everywhere.

We prove Lemma A3 through a sequence of claims. Define a new state variable $Z_t := \log(p_t/(1-p_t))$, which is a strictly increasing transformation of p_t . Note that Z_t is defined on $(-\infty, \infty)$. Define $v(z) := V(e^z/(1+e^z))$ and $z^* := \log(p^*/(1-p^*))$. Because we work with Z_t most of the time in this Appendix, we will write $a(z)$ to mean $a(e^z/(1+e^z))$ whenever there is no confusion.

If the principal anticipates the agent to use policy $a(\cdot)$ but the agent actually uses $\tilde{a}(\cdot)$, the law of motion of Z_t from the agent's perspective (after applying Itô's lemma to (A1)) is

$$(A2) \quad dZ_t = \psi^2(1-a_t) \left[1 - \tilde{a}_t - \frac{1}{2}(1-a_t) \right] dt - \psi(1-a_t) dB_t.$$

The Hamilton-Jacobi-Bellman equation for the agent is

$$(A3) \quad \begin{aligned} [r_1 + b(z)\lambda]v(z) &= \max_{\tilde{a} \in [0,1]} r_1 [u + (1-\tilde{a})c] \\ &\quad + \psi^2 [1 - a(z)] \left\{ 1 - \tilde{a} - \frac{1}{2} [1 - a(z)] \right\} v'(z) \\ &\quad + \frac{1}{2} \psi^2 [1 - a(z)]^2 v''(z). \end{aligned}$$

Let $b(\cdot)$ be a cutoff policy function of the principal with cutoff belief p^* . The following claims establish some necessary properties of any *continuous* pseudo-best reply $a(\cdot)$.

CLAIM A1: $a(z) = 1 - \frac{r_1 c}{\max\{r_1 c, -\psi^2 v'(z)\}}$, for all $z \in (-\infty, \infty)$.

PROOF:

By (the proof of) Lemma A1, any pseudo-best reply to a cutoff strategy must induce a full belief span. Because the right-hand side (RHS) of (A3) is affine in the choice variable, agent optimality requires that, for almost every $z \in \mathbb{R}$,

$$a(z) \begin{cases} = 0, & \text{if } r_1 c + \psi^2 [1 - a(z)] v' > 0; \\ \in [0, 1], & \text{if } r_1 c + \psi^2 [1 - a(z)] v' = 0; \\ = 1, & \text{if } r_1 c + \psi^2 [1 - a(z)] v' < 0. \end{cases}$$

This implies that, for almost every $z \in \mathbb{R}$, we have

$$(A4) \quad a(z) = 1 - \frac{r_1 c}{\max\{r_1 c, -\psi^2 v'(z)\}}.$$

Since both sides of (A4) are continuous in z by assumption, (A4) must hold for every $z \in \mathbb{R}$. ■

CLAIM A2: Fix any $z_1 < z_2 \leq z^*$ and suppose that $a(z) = 0$ for all $z \in (z_1, z_2)$. Then,

$$(A5) \quad v(z) = u + c + A_1 e^{\xi_L z} + A_2 e^{\xi'_L z}, \quad \forall z \in (z_1, z_2)$$

for some $A_1, A_2 \in \mathbb{R}$, where $\xi_L > 0 > \xi'_L$ are the two roots of the characteristic equation $\xi^2 + \xi = 2r_1/\psi^2$.

PROOF:

Since $b(z) = 0$ and $a(z) = 0$ for all $z \in (z_1, z_2)$, equation (A3) becomes

$$r_1 v(z) = r_1(u + c) + \frac{1}{2}\psi^2[v'(z) + v''(z)].$$

It is easy to verify that its general solution is given by (A5). ■

CLAIM A3: Fix any $z^* \leq z_1 < z_2$ and suppose that $a(z) = 0$ for all $z \in (z_1, z_2)$. Then,

$$(A6) \quad v(z) = \frac{r_1}{r_1 + \lambda}(u + c) + B_1 e^{\xi_R z} + B_2 e^{\xi'_R z}, \quad \forall z \in (z_1, z_2)$$

for some $B_1, B_2 \in \mathbb{R}$, where $\xi_R < 0 < \xi'_R$ are the two roots of the characteristic equation $\xi^2 + \xi = 2(r_1 + \lambda)/\psi^2$.

PROOF:

Since $b(z) = 1$ and $a(z) = 0$ for all $z \in (z_1, z_2)$, equation (A3) becomes

$$(r_1 + \lambda)v(z) = r_1(u + c) + \frac{1}{2}\psi^2[v'(z) + v''(z)].$$

It is easy to verify that its general solution is given by (A6). ■

CLAIM A4: Fix any $z_1 < z_2 \leq z^*$ and suppose that $a(z) \in (0, 1)$ for all $z \in (z_1, z_2)$. Then,

$$(A7) \quad v(z) = u + \sqrt{\kappa_L} \Phi^{-1}(C_1 e^z + C_2), \quad \forall z \in (z_1, z_2)$$

and

$$(A8) \quad a(z) = 1 + \frac{\sqrt{2r_1}}{\psi} \frac{\phi(\Phi^{-1}(C_1 e^z + C_2))}{C_1 e^z}$$

for some $C_1 < 0$ and $C_2 \in \mathbb{R}$, where $\kappa_L := r_1 c^2 / (2\psi^2)$.

Moreover, $a(z)$ is strictly increasing, or strictly decreasing, or first strictly decreasing and then strictly increasing on (z_1, z_2) .

PROOF:

Fix any $z_1 < z_2 \leq z^*$ such that $a(z) \in (0, 1)$ for all $z \in (z_1, z_2)$. Claim A1 implies that

$$(A9) \quad a(z) = 1 + \frac{r_1 c}{\psi^2 v'(z)}, \quad \forall z \in (z_1, z_2).$$

Substituting (A9) into (A3) and setting $b(z) = 0$, we have

$$(A10) \quad v(z) = u + \kappa_L \frac{v''(z) - v'(z)}{v'(z)^2}.$$

It is easy to verify that its general solution is given by (A7), and that the resulting $a(\cdot)$ implied by (A9) is given by (A8). Moreover, since $a(z) \in (0, 1)$, we must have $v'(z) < 0$, i.e., $C_1 < 0$.

To analyze the monotonicity of $a(\cdot)$ on (z_1, z_2) , we first establish the following equality which links $a'(z)$ to $a(z)$:

$$(A11) \quad 1 - a(z) - a'(z) = 2 \left[\frac{v(z) - u}{c} \right].$$

By (A9), $1 - a(z) = -r_1 c / [\psi^2 v'(z)]$. Differentiating this expression, we obtain $a'(z) = v''(z) [1 - a(z)] / v'(z)$. Recall, from the agent's HJB (A10) in this case, that $v''(z)/v'(z) = 1 - 2 \left[\frac{v(z) - u}{c} \right] / [1 - a(z)]$, where equation (A9) and $\kappa_L = r_1 c^2 / (2\psi^2)$ are used. Equation (A11) then follows immediately.

Equation (A11) implies the following:

- (i) If $a'(\tilde{z}) < 0$ for some $\tilde{z} \in (z_1, z_2)$, then $a'(z) < 0$ for all $z \in (z_1, \tilde{z})$.
- (ii) There does not exist an interval $I \subseteq (z_1, z_2)$ such that $a'(z) = 0$ for all $z \in I$.
- (iii) If $a(\tilde{z}) \geq 0$ for some $\tilde{z} \in (z_1, z_2)$, then $a(\cdot)$ is strictly increasing on $(\tilde{z}, z_2)$.

To see (i), suppose that $a'(\tilde{z}) < 0$ for some $\tilde{z} \in (z_1, z_2)$. Since $a(\cdot)$ given by (A8) is a smooth function on (z_1, z_2) , we can define $\underline{z} = \inf \{ z \in [z_1, \tilde{z}) : a'(\cdot)|_{(z, \tilde{z}]} < 0 \}$. Result (i) is proved if $\underline{z} = z_1$. Suppose (for a contradiction) that $\underline{z} > z_1$. Continuity of a' implies that $a'(\underline{z}) = 0$. Moreover, since $a'(z) < 0$ for all $z \in (\underline{z}, \tilde{z}]$, we have $a(\underline{z}) > a(\tilde{z})$. Then,

$$2 \left[\frac{v(\underline{z}) - u}{c} \right] = 1 - a(\underline{z}) - a'(\underline{z}) < 1 - a(\tilde{z}) - a'(\tilde{z}) = 2 \left[\frac{v(\tilde{z}) - u}{c} \right],$$

where the equalities follow from (A11) and the strict inequality follows from $a(\underline{z}) > a(\tilde{z})$ and $a'(\underline{z}) = 0 > a'(\tilde{z})$. But this is a contradiction to $v(\underline{z}) > v(\tilde{z})$ because $C_1 < 0$ and v is strictly decreasing on (z_1, z_2) .

To see (ii), suppose (for a contradiction) that there exists an interval $I \subseteq (z_1, z_2)$ such that $a'(z) = 0$ for all $z \in I$. Then, the left-hand side (LHS) of (A11) is constant on I while the RHS is strictly decreasing, a contradiction.

To see (iii), suppose that $a'(\tilde{z}) \geq 0$ for some $\tilde{z} \in (z_1, z_2)$. Then, we must have $a'(z) \geq 0$ for all $z \in (\tilde{z}, z_2)$, for otherwise we would reach a contradiction to (i). Further, take any z_3, z_4 such that $\tilde{z} \leq z_3 < z_4 \leq z_2$. Since $a'(z) \geq 0$, we know that $a(z_3) \leq a(z_4)$. But this inequality must be strict, for otherwise $a(z_3) = a(z_4) = a(z)$ for all $z \in (z_3, z_4)$ contradicting result (ii). So, $a(\cdot)$ is strictly increasing on $(\tilde{z}, z_2)$.

Finally, to establish the monotonicity of $a(\cdot)$, suppose first that $a'(z) \geq 0$ for all $z \in (z_1, z_2)$. The same argument for (iii) above proves that a must be strictly increasing on (z_1, z_2) . Suppose now that $a'(\tilde{z}) < 0$ for some $\tilde{z} \in (z_1, z_2)$. Let $Z \subseteq (z_1, z_2)$ be the largest interval containing $\tilde{z}$ such that $a'(z) < 0$ for all $z \in Z$. Result (i) immediately implies that $\inf Z = z_1$. If $\sup Z = z_2$, then $a(\cdot)$ is strictly decreasing on (z_1, z_2) . If $\sup Z < z_2$, continuity of a' implies that $a'(\sup Z) = 0$. Then, result (iii) implies that $a(\cdot)$ is strictly increasing on $(\sup Z, z_2)$. In summary, $a(\cdot)$ is either strictly increasing, or strictly decreasing, or first strictly decreasing and then strictly increasing on (z_1, z_2) . ■

CLAIM A5: Fix any $z^* \leq z_1 < z_2$ and suppose that $a(z) \in (0, 1)$ for all $z \in (z_1, z_2)$. Then,

$$(A12) \quad v(z) = \frac{r_1}{r_1 + \lambda} u + \sqrt{\kappa_R} \Phi^{-1}(D_1 e^z + D_2), \quad \forall z \in (z_1, z_2)$$

and

$$(A13) \quad a(z) = 1 + \frac{\sqrt{2(r_1 + \lambda)}}{\psi} \frac{\phi(\Phi^{-1}(D_1 e^z + D_2))}{D_1 e^z}$$

for some $D_1 < 0$ and $D_2 \in \mathbb{R}$, where $\kappa_R := r_1^2 c^2 / [2(r_1 + \lambda)\psi^2]$.

Moreover, $a(z)$ is strictly increasing, or strictly decreasing, or first strictly decreasing and then strictly increasing on (z_1, z_2) .

PROOF:

The proof is analogous to that of Claim A4, and is thus omitted.

CLAIM A6: If $a(\tilde{z}) > 0$, then there are z_L, z_R such that $z_L < \tilde{z} < z_R$ and $a(z_L) = a(z_R) = 0$.

PROOF:

Let $a(\tilde{z}) > 0$ and suppose, seeking a contradiction, that $a(z) > 0$ for all $z > \tilde{z}$. Then, for all such z , we would have (by Claim A1) $v'(z) = -r_1 c / [\psi^2(1 - a(z))] \leq -r_1 c / \psi^2$. Taking the limit for arbitrarily large z , we obtain $\lim_{z \rightarrow +\infty} v(z) = -\infty$, a contradiction as v is always nonnegative. Similarly, if $a(z) > 0$ for all $z < \tilde{z}$, then $\lim_{z \rightarrow -\infty} v(z) = +\infty$, which contradicts that v is bounded above by $u + c$. ■

CLAIM A7: For any z_1, z_2 such that either $z_1 < z_2 \leq z^*$ or $z^* \leq z_1 < z_2$, $a(z_1) = a(z_2) = 0$ implies $a(z) = 0$ for all $z \in [z_1, z_2]$.

PROOF:

Fix any $z_1 < z_2 \leq z^*$ such that $a(z_1) = a(z_2) = 0$. Suppose (for a contradiction) that there exists $\tilde{z} \in (z_1, z_2)$ such that $a(\tilde{z}) \in (0, 1)$. Let Z be the largest interval containing $\tilde{z}$ such that $a(z) \in (0, 1)$ for all $z \in Z$. Obviously, $z_1 \leq \inf Z < \sup Z \leq z_2$, and $a(\inf Z) = a(\sup Z) = 0$ because $a(\cdot)$ is continuous. By Claim A4, $a(\cdot)$ is strictly increasing, or strictly decreasing, or first strictly decreasing and then strictly increasing on Z . Since $a(\inf Z) = 0$, $a(\cdot)$ can only be strictly increasing on Z , but this contradicts the continuity of $a(\cdot)$ at $\sup Z$. An analogous argument which invokes Claim A5 establishes same result for any $z^* \leq z_1 < z_2$ such that $a(z_1) = a(z_2) = 0$. ■

COROLLARY A1: One of the following must hold for any continuous pseudo-best reply $a(\cdot)$:

- $a(z) = 0$ for all $z \in \mathbb{R}$;
- $a(z)$ is hump-shaped (and maximized at z^*).

PROOF:

Suppose that $a(\cdot)$ is not always equal to zero. Then there exists $\tilde{z}$ such that $a(\tilde{z}) > 0$. We now show that $a(\cdot)$ must be hump-shaped and maximized at z^* . Let Z be the largest interval containing $\tilde{z}$ such that $a(z) > 0$ for all $z \in Z$. Let $z_L = \inf Z$ and $z_R = \sup Z$. By Claim A6, $-\infty < z_L < z_R < \infty$. By continuity of a , $a(z_L) = a(z_R) = 0$. Then we must have $z^* \in (z_L, z_R)$, for otherwise we would reach a contradiction to Claim A7. Moreover, for any $z \in (z_L, z_R)^c$, we must have $a(z) = 0$, for otherwise we can construct another Z' which also contains z^* such that $Z' - Z \neq \emptyset$, contradicting the maximality of Z . Since $a(z_L) = 0$, Claim A4 implies that $a(\cdot)$ must be strictly increasing on (z_L, z^*) . Similarly, since $a(z_R) = 0$, Claim A5 implies that a must be strictly decreasing on (z^*, z_R) . ■

CLAIM A8: $a(\cdot)$ is hump-shaped if and only if $r_1 < r^*$, where r^* is the unique solution to

$$(A14) \quad r^* \left(\sqrt{1 + 8r^*/\psi^2} + \sqrt{1 + 8(r^* + \lambda)/\psi^2} \right) + \lambda \left(\sqrt{1 + 8r^*/\psi^2} + 1 \right) \\ = 4\lambda \left(\frac{u}{c} + 1 \right).$$

PROOF:

“Only If” Part.—Suppose first that $a(\cdot)$ is hump-shaped. We will show that this implies $r_1 < r^*$. By Definition 3, there exist z_L, z_R such that $-\infty < z_L < z^* < z_R < \infty$ such that $a(z) = 0$ on $(-\infty, z_L] \cup [z_R, \infty)$ and $a(z) \in (0, 1)$ on (z_L, z_R) . Let $v^* := v(z^*)$, $v_L := v(z_L)$, and $v_R := v(z_R)$.

We now calculate the undetermined coefficients in Claims A2 to A5 ($A_1, A_2, B_1, B_2, C_1, C_2, D_1, D_2$ and z_R, z_L, v_R, v_L) in various parts of the agent’s policy and value, as

functions of v^* and model parameters. First, consider $z < z^*$. As $z \rightarrow -\infty$, the agent's value function is given in Claim A2 by (A5). Because $v(\cdot)$ is bounded, we must have

$$(A15) \quad A_2 = 0.$$

Note that Claim A1 implies that $r_1 c = -\psi^2 v'(z_L)$, that is, $v'(z_L) = -r_1 c / \psi^2$. By Claims A2 and A4, the value function $v(\cdot)$ must satisfy the following value-matching and smooth-pasting conditions at z_L and z^* :

$$(\text{value matching at } z_L) \quad u + c + A_1 e^{\xi_L z_L} = v_L,$$

$$(\text{value matching at } z_L) \quad u + \sqrt{\kappa_L} \Phi^{-1}(C_1 e^{z_L} + C_2) = v_L,$$

$$(\text{value matching at } z^*) \quad u + \sqrt{\kappa_L} \Phi^{-1}(C_1 e^{z^*} + C_2) = v^*,$$

$$(\text{smooth pasting at } z_L) \quad A_1 \xi_L e^{\xi_L z_L} = -\frac{r_1 c}{\psi^2},$$

$$(\text{smooth pasting at } z_L) \quad \frac{\sqrt{\kappa_L} C_1 e^{z_L}}{\phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right)} = -\frac{r_1 c}{\psi^2}.$$

These five conditions can uniquely pin down the undetermined vector $(v_L, z_L, A_1, C_1, C_2)$:

$$(A16) \quad v_L = u + c - \frac{r_1 c}{\xi_L \psi^2},$$

$$(A17) \quad e^{z_L} = e^{z^*} \left[\frac{\frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right)}{\frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right) + \Phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right) - \Phi\left(\frac{v^* - u}{\sqrt{\kappa_L}}\right)} \right],$$

$$(A18) \quad A_1 = -\frac{r_1 c}{\xi_L \psi^2} e^{-\xi_L z_L},$$

$$(A19) \quad C_1 = e^{-z^*} \left[\Phi\left(\frac{v^* - u}{\sqrt{\kappa_L}}\right) - \Phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right) - \frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right) \right],$$

$$(A20) \quad C_2 = \Phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right) + \frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{v_L - u}{\sqrt{\kappa_L}}\right).$$

Now, consider $z > z^*$. As $z \rightarrow \infty$, the agent's value function is given in Claim A3 by (A6). Because $v(\cdot)$ is bounded, we must have

$$(A21) \quad B_2 = 0.$$

Note that Claim A1 implies that $r_1 c = -\psi^2 v'(z_R)$, that is, $v'(z_R) = -r_1 c / \psi^2$. By Claims A3 and A5, the value function $v(\cdot)$ must satisfy the following value-matching and smooth-pasting conditions at z_R and z^* :

$$(\text{value matching at } z_R) \quad (u + c) \frac{r_1}{r_1 + \lambda} + B_1 e^{\xi_R z_R} = v_R,$$

$$(\text{value matching at } z_R) \quad \frac{r_1}{r_1 + \lambda} u + \sqrt{\kappa_R} \Phi^{-1}(D_1 e^{z_R} + D_2) = v_R,$$

$$(\text{value matching at } z^*) \quad \frac{r_1}{r_1 + \lambda} u + \sqrt{\kappa_R} \Phi^{-1}(D_1 e^{z^*} + D_2) = v^*,$$

$$(\text{smooth pasting at } z_R) \quad B_1 \xi_R e^{\xi_R z_R} = -\frac{r_1 c}{\psi^2},$$

$$(\text{smooth pasting at } z_R) \quad \frac{\sqrt{\kappa_R} D_1 e^{z_R}}{\phi\left(\frac{v(z_R) - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right)} = -\frac{r_1 c}{\psi^2}.$$

These five conditions can uniquely pin down the undetermined vector $(v_R, z_R, B_1, D_1, D_2)$:

$$(A22) \quad v_R = \frac{r_1}{r_1 + \lambda} (u + c) - \frac{r_1 c}{\xi_R \psi^2},$$

$$(A23) \quad e^{z_R} = e^{z^*} \left[\frac{\frac{\sqrt{2(r_1 + \lambda)}}{\psi} \phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right)}{\frac{\sqrt{2(r_1 + \lambda)}}{\psi} \phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) + \Phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) - \Phi\left(\frac{v^* - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right)} \right],$$

$$(A24) \quad B_1 = -\frac{r_1 c}{\xi_R \psi^2} e^{-\xi_R z_R},$$

$$(A25) \quad D_1 = e^{-z^*} \left[\Phi\left(\frac{v^* - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) - \Phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) - \frac{\sqrt{2(r_1 + \lambda)}}{\psi} \phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) \right],$$

$$(A26) \quad D_2 = \Phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right) + \frac{\sqrt{2(r_1 + \lambda)}}{\psi} \phi\left(\frac{v_R - \frac{r_1}{r_1 + \lambda} u}{\sqrt{\kappa_R}}\right).$$

Given v^* and model parameters, equations (A5) through (A8), (A12), and (A13) with coefficients given by (A15) through (A26) fully determine the agent's policy function $a(\cdot)$ and his value function $v(\cdot)$ on $\mathbb{R}$. Since $a(z) \in (0, 1)$ on (z_L, z_R) , Claim A1 implies that $v'(z) < 0$ on (z_L, z_R) , which in turn implies that $v_L > v_R$. Note that v_L and v_R are given by (A16) and (A22), both of which are independent

of v^* . In particular, $v_L > v_R$ amounts to $u + c - r_1 c / (\xi_L \psi^2) > r_1(u + c) / (r_1 + \lambda) - r_1 c / (\xi_R \psi^2)$. Straightforward calculation shows that this is equivalent to

$$r_1 \left(\sqrt{1 + 8r_1/\psi^2} + \sqrt{1 + 8(r_1 + \lambda)/\psi^2} \right) + \lambda \left(\sqrt{1 + 8r_1/\psi^2} + 1 \right) < 4\lambda \left(\frac{u}{c} + 1 \right),$$

that is, $r_1 < r^*$ (recall that r^* is defined in (A14)).

“If” Part.—Suppose now that $a(\cdot)$ is not hump-shaped. We will show that this implies $r_1 \geq r^*$. By Corollary A1, we know that $a(z) = 0$ for all $z \in \mathbb{R}$. Then, by Claims A2 and A3, $v(\cdot)$ is given by (A5) for $z < z^*$ and by (A6) for $z > z^*$. We now pin down the undetermined coefficients (A_1, A_2, B_1, B_2) as functions of model parameters. Since $v(\cdot)$ is bounded as $z \rightarrow -\infty$ or $z \rightarrow +\infty$, we must have

$$(A27) \quad A_2 = 0,$$

$$(A28) \quad B_2 = 0.$$

Also, the value function $v(\cdot)$ must satisfy the following value-matching and smooth-pasting conditions at z^* :

$$(\text{value matching at } z^*) \quad u + c + A_1 e^{\xi_L z^*} = \frac{r_1}{r_1 + \lambda} (u + c) + B_1 e^{\xi_R z^*},$$

$$(\text{smooth pasting at } z^*) \quad A_1 \xi_L e^{\xi_L z^*} = B_1 \xi_R e^{\xi_R z^*}.$$

These two conditions uniquely pin down (A_1, B_1) as

$$(A29) \quad A_1 = \frac{\xi_R}{\xi_L - \xi_R} \frac{\lambda}{r_1 + \lambda} (u + c) e^{-\xi_L z^*},$$

$$(A30) \quad B_1 = \frac{\xi_L}{\xi_L - \xi_R} \frac{\lambda}{r_1 + \lambda} (u + c) e^{-\xi_R z^*}.$$

Since $a(z) = 0$ for all $z \in \mathbb{R}$, by Claim A1 we must have $r_1 c \geq -\psi^2 v'(z)$ for all $z \in \mathbb{R}$. In particular, this should hold at z^* : $r_1 c \geq -\frac{\xi_R \xi_L}{\xi_L - \xi_R} \frac{\lambda}{r_1 + \lambda} (u + c) \psi^2$. Straightforward calculation shows that this is equivalent to

$$r_1 \left(\sqrt{1 + 8r_1/\psi^2} + \sqrt{1 + 8(r_1 + \lambda)/\psi^2} \right) + \lambda \left(\sqrt{1 + 8r_1/\psi^2} + 1 \right) \geq 4\lambda \left(\frac{u}{c} + 1 \right),$$

that is, $r_1 \geq r^*$ (recall that r^* is defined in (A14)). ■

CLAIM A9: Suppose that $r_1 < r^*$, and let v_L and v_R be given by (A16) and (A22), respectively. Define $a_-^*, a_+^* : [v_R, v_L] \rightarrow \mathbb{R}$ by

$$(A31) \quad a_-^*(x) := 1 - \frac{\frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{x-u}{\sqrt{\kappa_L}}\right)}{\frac{\sqrt{2r_1}}{\psi} \phi\left(\frac{v_L-u}{\sqrt{\kappa_L}}\right) + \Phi\left(\frac{v_L-u}{\sqrt{\kappa_L}}\right) - \Phi\left(\frac{x-u}{\sqrt{\kappa_L}}\right)},$$

$$(A32) a_+^*(x) := 1 - \frac{\frac{\sqrt{2(r_1+\lambda)}}{\psi} \phi\left(\frac{x-\frac{r_1}{r_1+\lambda}u}{\sqrt{\kappa_R}}\right)}{\frac{\sqrt{2(r_1+\lambda)}}{\psi} \phi\left(\frac{v_R-\frac{r_1}{r_1+\lambda}u}{\sqrt{\kappa_R}}\right) + \Phi\left(\frac{v_R-\frac{r_1}{r_1+\lambda}u}{\sqrt{\kappa_R}}\right) - \Phi\left(\frac{x-\frac{r_1}{r_1+\lambda}u}{\sqrt{\kappa_R}}\right)}.$$

Then, $a_-^*(\cdot)$ is strictly decreasing on $[v_R, v_L]$ with $a_-^*(v_R) \in (0, 1)$ and $a_-^*(v_L) = 0$; $a_+^*(\cdot)$ is strictly increasing on $[v_R, v_L]$ with $a_+^*(v_R) = 0$ and $a_+^*(v_L) \in (0, 1)$.

PROOF:

The proof of this technical claim can be found in our working paper (Ekmekci et al. 2021).

PROOF OF LEMMA A3:

Suppose first that $r_1 \geq r^*$. By Corollary A1 and Claim A8, any continuous pseudo-best reply $a(\cdot)$ must be such that $a(z) = 0$ for all $z \in \mathbb{R}$. Obviously, such a function is unique. To verify that $a(z) = 0$ for all $z \in \mathbb{R}$ satisfies agent optimality, note that we have shown in the proof of Claim A8 that, together with the $v(\cdot)$ given by (A5) and (A6) and the coefficients given by (A27) through (A30), it satisfies the agent's HJB equation (A3).¹⁹ Since $v(\cdot)$ is bounded, a standard verification argument shows that $a(z) = 0$ for all $z \in \mathbb{R}$ indeed satisfies agent optimality (see, for instance, Pham 2009, Theorem 3.5.3). In addition, regularity is satisfied because the functions given by (A5) and (A6) are smooth, and value-matching and smooth-pasting conditions are imposed at z^* .

Suppose now that $r_1 < r^*$. By Claim A8, any continuous pseudo-best reply $a(\cdot)$ is hump-shaped. We first show that such a function, if it exists, must be unique. By Claims A2 to A5 and the proof of the “only if” part of Claim A8, such a policy function and the associated value function must satisfy (A5) through (A8), (A12) and (A13) with coefficients given by (A15) through (A26), given the value $v^* := v(z^*)$.

Recall that the functions a_-^* and a_+^* are defined in (A31) and (A32), respectively. By Claim A4 and equations (A19) and (A20), it is easy to verify that $\lim_{z \uparrow z^*} a(z; v^*) = a_-^*(v^*)$. Similarly, by Claim A5 and equations (A25) and (A26), it is easy to verify that $\lim_{z \downarrow z^*} a(z; v^*) = a_+^*(v^*)$. Since $a(\cdot)$ is continuous at z^* , v^* must satisfy

$$(A33) \quad a_-^*(v^*) = a_+^*(v^*).$$

¹⁹Everywhere except at z^* where v'' does not exist.

By Claim A9, there is a unique $v^* \in (v_R, v_L)$ satisfying (A33), rendering the unique (candidate) policy function.

Take such unique v^* , and let $a(\cdot)$ and $v(\cdot)$ be defined by (A5) through (A8), (A12), and (A13) with coefficients given by (A15) through (A26). Exactly the same verification argument as in the case of $r_1 \geq r^*$ confirms that $a(\cdot)$ indeed satisfies agent optimality. In addition, regularity is satisfied because the functions given by (A5), (A6), (A7), and (A12) are smooth, and value-matching and smooth-pasting conditions are imposed at z_L, z^* , and z_R .

Finally, for any other pseudo-best reply $\bar{a}(\cdot)$, it must satisfy condition (A4) almost everywhere on $\mathbb{R}$, leading to the same set of differential equations for the agent's value function. Consequently, agent optimality requires that $\bar{a}(\cdot)$ agrees with $a(\cdot)$ almost everywhere. ■

Proof of Theorem 1.—Lemmas A2 and A3 establish the unique *structure* of Markov equilibria. To prove Theorem 1, we still need an argument for equilibrium existence and uniqueness.

PROOF OF THEOREM 1:

Suppose first that $r_1 \geq r^*$. By Lemma A3, the agent's continuous pseudo-best reply to any cutoff termination rule satisfies that $a(p) = 0$ for all $p \in (0, 1)$. On the other hand, given this policy function of the agent under which the belief span is $(0, 1)$, the proof of Lemma A2 can be used verbatim to show that the principal has a unique best reply whose policy function b admits a cutoff $p^* \in (0, 1)$. By Lemma A3, r^* is independent of p^* and any other pseudo-best reply agrees with $a(\cdot)$ almost everywhere, so such (a, b) is the unique Markov equilibrium in this case.²⁰

Suppose now that $r_1 < r^*$. Lemmas A2 and A3 imply that any Markov equilibrium (a, b) must be such that a is hump-shaped and b has a cutoff structure. We now show that there exists a unique Markov equilibrium.

To show that a Markov equilibrium exists, note that the continuous pseudo-best reply of the noninvestible agent can be described by a function φ_1 that maps a conjecture $\tilde{p}^*$ about the cutoff p^* used by the principal into a policy function $a(\cdot)$ defined on $(0, 1)$, the probability-domain version of the policy function in z -space constructed in Claim A8's proof. The function φ_1 is continuous.²¹ Moreover, the best reply of the principal can be described by a function φ_2 that maps any policy function of the noninvestible agent into a unique cutoff p^* , and by (the proof of) Lemma A2 we have $p^* \in [p^{**}, p_H] \subset (0, 1)$.

Define the composition $\varphi := \varphi_2 \times \varphi_1$ that maps each conjecture $\tilde{p}^*$ into the associated optimal cutoff p^* . The mapping φ satisfies

$$\varphi(\tilde{p}^*) = \arg \max_{p' \in [p^{**}, p_H]} \hat{M}(p, p', \tilde{p}^*),$$

where

$$\hat{M}(p, p', \tilde{p}^*) := r_2 \int_0^{+\infty} \int_{p'}^1 e^{-r_2 t} R(q) d\Gamma(t, q | \tilde{p}^*, p),$$

²⁰Remember that we do not distinguish Markov equilibria whose policy functions agree almost everywhere.

²¹This follows from Lemma OA.1 in the online Appendix, i.e., the translation invariance of the agent's problem.

and Γ is the joint probability measure of getting the first Poisson shock in the stopping region $[p', 1)$ at time t and state q , when the prior belief at time zero is p . Notice that $\varphi(p^*)$ is the unique solution to this maximization problem and is independent of the prior due to its Markovian nature. Moreover, $\hat{M}(p, p', \tilde{p}^*)$ is jointly continuous in $(p', \tilde{p}^*)$. Since the choice space is compact and the objective is continuous in both the choice variable p' and the “parameter” $\tilde{p}^*$, the maximum theorem implies that φ is continuous.

Since φ is a continuous function mapping from $(0, 1)$ to $[p^{**}, p_H]$ such that $\liminf_{p \rightarrow 0} [\varphi(p) - p] \geq p^{**} > 0$ and $\limsup_{p \rightarrow 1} [\varphi(p) - p] \leq p_H - 1 < 0$, the intermediate value theorem implies that there must be $p^* \in (0, 1)$ such that $\varphi(p^*) = p^*$. Taking any such p^* , it is easy to verify that $(a^* := \varphi_1(p^*), b^* := \mathbf{1}_{\{p_i \geq p^*\}})$ represents a Markov equilibrium. This proves equilibrium existence.

To show that the Markov equilibrium is unique, let (a, b) be a Markov equilibrium in which the principal uses a threshold p^* , associated with a likelihood z^* . Consider the payoff of the principal when deviating to a different threshold z' when the state is z .

$$M(z, z', z^*) := p(z)E[e^{-r_2 T_{NI}(z, z', z^*)}]w_{NI} + (1 - p(z))E[e^{-r_2 T_I(z, z', z^*)}]w_I,$$

where $T_\theta(z, z', z^*)$ is the random time of occurrence of the first Poisson shock that arrives while the state lies in the stopping interval $[z', +\infty)$, provided the initial state is z and the dynamics is conditioned on type $\theta \in \{I, NI\}$. By the conditional translation invariance property proved in Lemma OA.2 of the online Appendix, $E[e^{-r_2 T_\theta(z, z', z^*)}] = E[e^{-r_2 T_\theta(0, z' - z, z^* - z)}]$ for all $z, z', z^* \in \mathbb{R}$. The first-order condition of the principal is

$$\frac{\partial M(z, z', z^*)}{\partial z'} = p(z)D_{NI}(z, z', z^*)w_{NI} + (1 - p(z))D_I(z, z', z^*)w_I = 0,$$

where we define $D_\theta(z, z', z^*) := \partial E[e^{-r_2 T_\theta(z, z', z^*)}]/\partial z'$ for each $\theta \in \{I, NI\}$, $z, z', z^* \in \mathbb{R}$. Note that this condition should hold for every $z \in \mathbb{R}$. In equilibrium, the principal's choice of z' must coincide with z^* , so the first-order condition becomes $p(z)D_{NI}(z, z^*, z^*)w_{NI} + (1 - p(z))D_I(z, z^*, z^*)w_I = 0$, which can be rewritten as

$$p(z) = \frac{D_I(z, z^*, z^*)w_I}{D_I(z, z^*, z^*)w_I - D_{NI}(z, z^*, z^*)w_{NI}}.$$

Evaluating the limit from below as $z \uparrow z^*$, we have

$$\begin{aligned} p^* = p(z^*) &= \frac{D_I(z^*, z^*, z^*)w_I}{D_I(z^*, z^*, z^*)w_I - D_{NI}(z^*, z^*, z^*)w_{NI}} \\ &= \frac{D_I(0, 0, 0)w_I}{D_I(0, 0, 0)w_I - D_{NI}(0, 0, 0)w_{NI}}, \end{aligned}$$

where the last equality follows from Lemma OA.2 in the online Appendix, i.e., the conditional translation invariance of the principal’s payoff function. Since the RHS is independent of p^* , we conclude that there can be at most one value of p^* consistent with a Markov equilibrium, establishing the uniqueness claim. ■

COROLLARY A2: *The following hold:*

- (i) $W(\cdot)$ is (weakly) increasing, nonnegative and convex on $(0, 1)$, and it satisfies $\lim_{p \rightarrow 0} W(p) = 0$ and $\lim_{p \rightarrow 1} W(p) = \lambda w_{NI}/(r_2 + \lambda)$.
- (ii) $v(\cdot)$ is strictly decreasing on $\mathbb{R}$, and it satisfies $\lim_{z \rightarrow -\infty} v(z) = u + c$ and $\lim_{z \rightarrow \infty} v(z) = r_1(u + c)/(r_1 + \lambda)$. Moreover, $v(\cdot)$ is concave on $(-\infty, z^*)$ and convex on (z^*, ∞) .

PROOF:

See online Appendix.

B. Expected Performance: Toward a Proof of Theorem 2

Given a Markov equilibrium (a, b) (where the equilibrium policy functions are defined on the z -space), let v be the agent’s value function, z^* be the principal’s termination cutoff, and recall that the agent’s expected performance is given by

$$(A34) \quad EP(z) = \psi[1 - (1 - a(z))p(z)],$$

where $p(z) = e^z/(1 + e^z)$. Our analysis in this section fixes all model parameters, except r_1 and/or λ .

LEMMA A4: *If $r_1 \geq r^*$, then $EP(\cdot)$ is strictly decreasing on $\mathbb{R}$. If $r_1 < r^*$, then either $EP(\cdot)$ is strictly decreasing on $\mathbb{R}$, or $EP(\cdot)$ is*

- *strictly decreasing for $z < z'$, where z' is some number in $[z_L, z^*)$;*
- *strictly increasing for $z \in (z', z^*)$;*
- *strictly decreasing for $z > z^*$.*

PROOF:

If $r_1 \geq r^*$, Theorem 1 tells us that $a(z) = 0$ for all $z \in \mathbb{R}$. Thus, $EP(z) = \psi[1 - p(z)]$, which is strictly decreasing on $\mathbb{R}$.

If $r_1 < r^*$, by Theorem 1 the agent’s equilibrium policy function a is hump-shaped, with cutoffs denoted by z_L and z_R such that $a(z) > 0$ if and only if $z \in (z_L, z_R)$. Obviously, $EP(\cdot)$ is strictly decreasing on $(-\infty, z_L)$ and on (z^*, ∞) , because on each of these intervals $a(\cdot)$ is weakly decreasing and $p(\cdot)$ is strictly increasing in z . We now focus on the monotonicity of $EP(\cdot)$ on (z_L, z^*) .

First, analogous to (A11), we establish the following equality which links $EP(z)$ to $EP'(z)$:

$$(A35) \quad \psi - EP(z) - \frac{EP'(z)}{p(z)} = 2\psi \left[\frac{v(z) - u}{c} \right].$$

To see this, note first that since $a(z) > 0$ on (z_L, z^*) , by equation (A34) and Claim A1 we have

$$(A36) \quad \psi - EP(z) = \psi[1 - a(z)]p(z) = -\frac{r_1 c}{\psi} \frac{p(z)}{v'(z)}.$$

Differentiating this expression, we obtain

$$\begin{aligned} -EP'(z) &= -\frac{r_1 c}{\psi} \left\{ -\frac{v''(z)p(z)}{v'(z)^2} + \frac{p(z)[1 - p(z)]}{v'(z)} \right\} \\ &= \psi[1 - a(z)]p(z) \left[-\frac{v''(z)}{v'(z)} + 1 - p(z) \right], \end{aligned}$$

where the first equality follows from $p'(z) = p(z)[1 - p(z)]$ and the second equality follows from (A36). Recall, from the agent's HJB (A10) in this case, that $v''(z)/v'(z) = 1 - 2\left[\frac{v(z) - u}{c}\right]/[1 - a(z)]$, where Claim A1 and $\kappa_L = r_1 c^2/(2\psi^2)$ are used. Thus, $-EP'(z)/p(z) = 2\psi[v(z) - u]/c - [\psi - EP(z)]$, where (A36) is used. Equation (A35) then follows immediately.

From equation (A35), we can apply the same argument as that after equation (A11) to show that $EP(\cdot)$ must be either strictly increasing, or strictly decreasing, or first strictly decreasing and then strictly increasing on (z_L, z^*) , a property that echoes what we have shown for the policy function $a(\cdot)$ in Claim A4. Since we know that $EP(\cdot)$ is strictly decreasing on $(-\infty, z_L)$ and on (z^*, ∞) , the result in the lemma follows. ■

COROLLARY A3: $EP(\cdot)$ is nonmonotone if and only if $r_1 < r^*$ and $EP'(z^*) > 0$.

LEMMA A5: $EP(\cdot)$ is nonmonotone if and only if $r_1 < r^*$ and $[1 - a(z^*)]p(z^*) > 2\left(\frac{v^* - u}{c}\right)$, where $v^* := v(z^*)$.

PROOF:

Since $a(z) > 0$ on (z_L, z^*) , by Claim A1 and equation (A7) we have

$$(A37) \quad 1 - a(z) = -\frac{r_1 c}{\psi^2 v'(z)} \propto \frac{\phi\left(\frac{v(z) - u}{\sqrt{\kappa_L}}\right)}{e^z}, \quad \forall z \in (z_L, z^*].$$

Then,

$$EP'(z^*) > 0$$

$$\text{(by (A34))} \quad \Leftrightarrow \frac{d}{dz} \{ [1 - a(z)]p(z) \} \Big|_{z=z_-^*} < 0$$

$$\text{(because } p'(z) = \frac{p(z)}{1-p(z)} \text{)} \quad \Leftrightarrow -a'(z)p(z) + [1 - a(z)]p(z)[1 - p(z)] \Big|_{z=z_-^*} < 0$$

$$\text{(because } a(z) < 1 \text{ for all } z) \quad \Leftrightarrow \frac{d \ln[1 - a(z)]}{dz} + 1 - p(z) \Big|_{z=z_-^*} < 0$$

$$\text{(by (A37))} \quad \Leftrightarrow \frac{d}{dz} \left[\ln \phi \left(\frac{v(z) - u}{\sqrt{\kappa_L}} \right) - z \right] + 1 - p(z) \Big|_{z=z_-^*} < 0$$

$$\text{(because } -v'(z^*) > 0) \quad \Leftrightarrow \frac{v^* - u}{\kappa_L} < -\frac{p(z^*)}{v'(z^*)}$$

$$\text{(by } \kappa_L = \frac{r_1 c^2}{2\psi^2} \text{ and (A37))} \quad \Leftrightarrow 2 \left(\frac{v^* - u}{c} \right) < [1 - a(z^*)]p(z^*).$$

By Corollary A3, the result follows. ■

COROLLARY A4: $EP(\cdot)$ is nonmonotone if $r_1 < r^*$ and $v^* < u$.

For a given λ , recall that $r^*(\lambda)$ is the unique solution to (A14). It is easy to see that there exists a unique $\lambda_1 > 0$ such that

$$(A38) \quad r^*(\lambda_1) = \psi^2.$$

Consequently, $r^*(\lambda) > \psi^2$ if and only if $\lambda > \lambda_1$.

The following lemma deals with the agent's discount rates that are close to $r^*(\lambda)$.

LEMMA A6: If $\lambda > \lambda_1$ and $\psi^2 < r_1 < r^*(\lambda)$, then $EP(\cdot)$ is nonmonotone.

PROOF:

Recall that, in the proof of Claim A8, we have shown that if $a(\cdot)$ is hump-shaped, then $v_L := v(z_L)$ and $v_R := v(z_R)$ are calculated in (A16) and (A22), respectively. Moreover, $v_L > v_R$ if (and only if) $r_1 < r^*(\lambda)$, and since $v'(\cdot) < 0$, we have $v^* \in (v_R, v_L)$.

From (A16), it is easy to verify that $v_L < u$ if and only if $r_1 > \psi^2$. Therefore, if $\lambda > \lambda_1$ and $\psi^2 < r_1 < r^*(\lambda)$, we have $v_R < v^* < v_L < u$. By Corollary A4, $EP(\cdot)$ is nonmonotone. ■

Recall that the functions $a_-^*(\cdot; r_1), a_+^*(\cdot; r_1, \lambda): [v_R, v_L] \rightarrow \mathbb{R}$ are defined by (A31) and (A32), respectively. Recall also, from the proof of Lemma A3, that $v^* \in (v_R, v_L)$ is the unique solution to $a_-^*(x; r_1) = a_+^*(x; r_1, \lambda)$.

CLAIM A10: *If $\lambda > \lambda_1$ and $r_1 \leq \psi^2$, then $u \in [v_R, v_L]$. Moreover, $v^* < u$ if and only if $a_-^*(u; r_1) < a_+^*(u; r_1, \lambda)$.*

PROOF:

Suppose that $\lambda > \lambda_1$ and $r_1 \leq \psi^2$. First, by definition of λ_1 in (A38), we have $r^*(\lambda) > \psi^2 \geq r_1$, so the equilibrium $a(\cdot)$ is hump-shaped and $v_L > v_R$. Substituting the expressions of ξ_L and ξ_R into (A16) and (A22), we can rewrite v_L and v_R as

$$v_L(r_1) = u + c \left(1 - \frac{\sqrt{1 + 8r_1/\psi^2} + 1}{4} \right),$$

$$v_R(r_1, \lambda) = \frac{r_1}{r_1 + \lambda} \left[u + c \left(1 + \frac{\sqrt{1 + 8(r_1 + \lambda)/\psi^2} - 1}{4} \right) \right].$$

It is easy to verify that v_L is strictly decreasing in r_1 and v_R is strictly increasing in r_1 . Thus, for $\lambda > \lambda_1$ and $r_1 \leq \psi^2$,

$$v_R(r_1, \lambda) < v_R(r^*(\lambda), \lambda) = v_L(r^*(\lambda)) < v_L(\psi^2) \leq v_L(r_1),$$

where the inequalities follow from the monotonicity of v_L and v_R in r_1 , and the equality follows from the definition of r^* . Note that $v_L(\psi^2) = u$, so we have $u \in [v_R, v_L]$.

The fact that $u \in [v_R, v_L]$ implies that $a_-^*(u; r_1)$ and $a_+^*(u; r_1, \lambda)$ are well-defined. By Claim A9, the function $g: [v_R, v_L] \rightarrow \mathbb{R}$ defined by $g(x; r_1, \lambda) := a_-^*(x; r_1) - a_+^*(x; r_1, \lambda)$ is strictly decreasing in x , and v^* is the unique zero point of $g(\cdot; r_1, \lambda)$ on $[v_R, v_L]$. Therefore, $v^* < u$ if and only if $g(u; r_1, \lambda) < 0$, if and only if $a_-^*(u; r_1) < a_+^*(u; r_1, \lambda)$. ■

The next lemma deals with the case where the agent's discount rate r_1 is small.

LEMMA A7: *There exist $\lambda_2 \geq \lambda_1$ and $\underline{r} > 0$, such that if $\lambda > \lambda_2$ and $0 < r_1 < \underline{r}$, then $EP(\cdot)$ is nonmonotone.*

PROOF:

See online Appendix.

We now turn to the last case where $r_1 \in [\underline{r}, \psi^2]$.

CLAIM A11: *There exist $\lambda'_3 \geq \lambda_1$ and $A > 1$ such that if $\lambda > \lambda'_3$, then*

$$(A39) \quad a_+^*(u; r_1, \lambda) > 1 - A \times \exp\left(-\frac{\psi^2}{c^2 r_1^2} \frac{\lambda^2}{r_1 + \lambda} u^2\right), \quad \forall r_1 \in [\underline{r}, \psi^2].$$

PROOF:

See online Appendix.

LEMMA A8: *There exists $\lambda_3 \geq \lambda_1$ such that if $\lambda > \lambda_3$ and $\underline{r} \leq r_1 \leq \psi^2$, then $EP(\cdot)$ is nonmonotone.*

PROOF:

Let $\lambda'_3 \geq \lambda_1$ and $A > 1$ be delivered by Claim A11. For each $r_1 \in [\underline{r}, \psi^2]$, let $\lambda(r_1)$ be the unique solution on $\mathbb{R}_+$ to

$$A \times \exp\left(-\frac{\psi^2}{c^2 r_1^2} \frac{\lambda^2}{r_1 + \lambda} u^2\right) = 1 - a_-(u; r_1).$$

Note that $\lambda(r_1)$ is well-defined for all $r_1 \in [\underline{r}, \psi^2]$ because the LHS is strictly decreasing in λ while the RHS is independent of λ . Hence, we have

$$(A40) \quad A \times \exp\left(-\frac{\psi^2}{c^2 r_1^2} \frac{\lambda^2}{r_1 + \lambda} u^2\right) < 1 - a_-(u; r_1), \quad \forall \lambda > \lambda(r_1).$$

Also, $\lambda(r_1)$ is continuous in r_1 by the implicit function theorem. Let $\lambda_3'' := \max_{r \in [\underline{r}, \psi^2]} \lambda(r)$ and $\lambda_3 := \max\{\lambda_3', \lambda_3''\}$. Combining (A39) and (A40), we have

$$a_+(u; r_1, \lambda) > 1 - A \times \exp\left(-\frac{\psi^2}{c^2 r_1^2} \frac{\lambda^2}{r_1 + \lambda} u^2\right) > a_-(u; r_1),$$

$$\forall \lambda > \lambda_3 \text{ and } \forall r_1 \in [\underline{r}, \psi^2].$$

Then, by Claim A10 and Corollary A4, we conclude that $EP(\cdot)$ is nonmonotone if $\lambda > \lambda_3$ and $r_1 \in [\underline{r}, \psi^2]$. ■

PROOF OF THEOREM 2:

Let λ_1 be defined in (A38), and let λ_2 and λ_3 be delivered by Lemmas A7 and A8, respectively. Define $\bar{\lambda} := \max\{\lambda_1, \lambda_2, \lambda_3\}$. Lemmas A6 through A8 imply that if $\lambda > \bar{\lambda}$ and $r_1 < r^*(\lambda)$, then $EP(\cdot)$ is nonmonotone. Lemma A4 then leads to the conclusion of the theorem. ■

C. Effects of Better Transparency: Toward a Proof of Theorem 3

Take any sequence $\{\psi_n\}_n$ such that $\lim_n \psi_n = +\infty$. For each $n \in \mathbb{N}$, take the unique Markov equilibrium (a_n, b_n) associated with the signal-to-noise ratio ψ_n . Let $V_n(\cdot)$ be the agent's value function in the equilibrium (a_n, b_n) and $W_n(\cdot)$ be the principal's value function. We will often use $z \equiv \log(p/(1-p))$ as state variable when analyzing the agent's behavior. When doing so, we denote by $v_n(z) := V_n(p(z))$ the agent's value function in the z -space. Write z_n^* for the principal's equilibrium cutoff. Write $z_{L,n}$ for the infimum belief z at which the agent plays $a_n(z) > 0$ and write $z_{R,n}$ for the supremum. Write $\mathbb{T}$ for the equilibrium stopping time that stops the play of the game. *Without labeling explicitly, we note that the distribution of $\mathbb{T}$ depends on n and the current state z .* For $i = 1, 2$ and $\theta \in \{NI, I\}$, let $E_n^\theta[e^{-r_i \mathbb{T}}]$ be the expected discount factor when the stopping action is

taken in the equilibrium (a_n, b_n) discounted at rate r_i . When the game starts at state z , let

$$E_n[e^{-r_i \mathbb{T}}] := p(z)E_n^{NI}[e^{-r_i \mathbb{T}}] + [1 - p(z)]E_n^I[e^{-r_i \mathbb{T}}].$$

ASSUMPTION A1: $\lambda > r_1(c/u)$, i.e., $u > \frac{r_1}{r_1 + \lambda}(u + c)$.

CLAIM A12: *There exists $N \in \mathbb{N}$ such that whenever $n \geq N$, $a_n(\cdot)$ is hump-shaped.*

PROOF:

From condition (A14), it is easily verified that $\lim_n r_n^* = \lambda(2u + c)/c$. Assumption A1 implies that $r_1 < \lambda u/c < \lambda(2u + c)/c$. By Theorem 1, the result follows. ■

We assume $n \geq N$ for the rest of the proofs in this section.

CLAIM A13: *Take any compact set $[p_1, p_2] \subset (0, 1)$. We have*

$$\limsup_{n \rightarrow \infty} \left\{ \max_{z \in [z(p_1), z(p_2)]} v_n(z) \right\} \leq u.$$

PROOF:

Because $p^* \in [p^{**}, p_H]$ (by Lemma A2), we can without loss assume that $p_n^* \in [p_1, p_2]$.²² Since $v_n(\cdot)$ is decreasing, it suffices to show that $\limsup v_n(z(p_1/2)) \leq u$.

By Claim A12, $a_n(\cdot)$ is hump-shaped for all n . First assume that $a_n(z_n^*) \rightarrow 1$. Take any $\epsilon > 0$, and let z_n^ϵ be the smallest z such that $a_n(z) = 1 - \epsilon$, which is well-defined for every large n such that $a_n(z_n^*) > 1 - \epsilon$. Consider the stochastic process Z_t in the equilibrium (a_n, b_n) under the noninvestible-type strategy and the initial condition $Z_0 = z(p_1/2)$. Let $\mathbb{T}_n^\dagger$ be the stopping time that stops the game at the first time that $Z_t \geq z_n^\epsilon$. From the law of motion (A2), as $\psi_n \rightarrow \infty$, we have $E_n^{NI}[e^{-r_1 \mathbb{T}_n^\dagger}] \rightarrow 1$ and hence $v_n(z(p_1/2)) \rightarrow v_n(z_n^\epsilon)$. Moreover, since $v_n(\cdot)$ is decreasing and concave to the left of z_n^ϵ (by Corollary A2), we have

$$\begin{aligned} r_1 v_n(z_n^\epsilon) &= r_1 [u + (1 - a_n(z_n^\epsilon))c] + \frac{1}{2} \psi^2 [1 - a_n(z_n^\epsilon)]^2 [v'(a_n(z_{n-}^\epsilon)) + v''(a_n(z_{n-}^\epsilon))] \\ &\leq r_1 [u + (1 - a_n(z_n^\epsilon))c], \end{aligned}$$

which implies $v_n(z_n^\epsilon) \leq u + \epsilon c$, delivering the result as ϵ is arbitrary.

Next assume that $\liminf a_n(z_n^*) < 1$. Take an $\epsilon > 0$ such that $u > \frac{r_1}{r_1 + \lambda}(u + c) + \epsilon$; such an ϵ exists because of Assumption A1. Notice that we can

²²If p_n^* is outside $[p_1, p_2]$, we then expand the closed interval to include p_n^* by defining $\hat{p}_1 = \min\{p_1, p_n^*\}$ and $\hat{p}_2 = \max\{p_2, p_n^*\}$. Notice that $\max_{z \in [z(p_1), z(p_2)]} v_n(z) \leq \max_{z \in [z(\hat{p}_1), z(\hat{p}_2)]} v_n(z)$.

find a $z^\dagger$ sufficiently large such that $v_n(z^\dagger) < r_1(u + c)/(r_1 + \lambda) + \epsilon$, $\forall n$.²³ Let $T_n^\dagger$ be the stopping time that stops the game at the first time that $Z_t = z^\dagger$. As $\psi_n \rightarrow \infty$, we have $v_n(z(p_1/2)) \rightarrow v_n(z^\dagger) < r_1(u + c)/(r_1 + \lambda) + \epsilon < u$. ■

CLAIM A14: $\lim_{n \rightarrow \infty} z_{L,n} = -\infty$.

PROOF:

Recall that $v_n(z_{L,n}) = v_{L,n}$ whose expression is given by (A16), i.e., $v_{L,n} = u + c \left(1 - \frac{\sqrt{1 + 8r_1/\psi_n^2 + 1}}{4} \right)$. It is easy to see that $\lim_n v_n(z_{L,n}) = u + c/2$. Assume toward a contradiction, taking a subsequence if necessary, that $\lim_n z_{L,n} = z' > -\infty$. Then for any $\epsilon > 0$, we have $z_{L,n} \in [z - \epsilon, z + \epsilon]$ and $v_n(z_{L,n}) \geq u + c/4$ when n is sufficiently large. Take any $\epsilon < (0, c/4)$. By Claim A13 and the monotonicity of $v_n(\cdot)$, there exists n^* such that for every $n > n^*$ and for every $z \in [z - \epsilon, z + \epsilon]$, we have $v_n(z) < u + \epsilon < u + c/4$, a contradiction to $v_{L,n} \rightarrow u + c/2$ and $z_{L,n} \rightarrow z'$. ■

CLAIM A15: For any $\kappa > 0$, we have $\lim_{n \rightarrow \infty} a_n(z_n^* - \kappa) = 1$.

PROOF:

Assume toward a contradiction that, taking a subsequence if necessary, $\lim a_n(z_n^* - \kappa) = 1 - 2\epsilon$ for some $\epsilon > 0$. Take any large $M > 0$ and notice that Claim A14 tells us that $z_{L,n} < z_n^* - \kappa - M$ for large n . Since $a_n(\cdot)$ is increasing on $[z_n^* - \kappa - M, z_n^* - \kappa]$ and $\lim a_n(z_n^* - \kappa) = 1 - 2\epsilon$, we know that $a_n(z) < 1 - \epsilon$ (infinitely often) for all $z \in [z_n^* - \kappa - M, z_n^* - \kappa]$. Recall from condition (A11) that for $z \in [z_n^* - \kappa - M, z_n^* - \kappa]$, $a'_n(z) = 1 - a_n(z) - 2[v_n(z) - u]/c$. In light of Claim A13, this implies that for n sufficiently large, we have $a'_n(z) > \epsilon/2$ for all $z \in [z_n^* - \kappa - M, z_n^* - \kappa]$. But then, we can take M large enough such that $a_n(z_n^* - \kappa - M) < 0$, a contradiction. ■

CLAIM A16: For any $\kappa > 0$, we have $\lim_{n \rightarrow \infty} v_n(z_n^* - \kappa) = u$.

PROOF:

From Claim A13 and Lemma A2, we know $\limsup v_n(z_n^* - \kappa) \leq u$. Assume toward a contradiction, taking a subsequence if necessary, that $\lim_{n \rightarrow \infty} v_n(z_n^* - \kappa) = u - 2\epsilon$ for some $\epsilon > 0$. This implies that $v_n(z_n^* - \kappa) < u - \epsilon$ for n sufficiently large. Note also that Claim A14 tells us that $z_n^* - \kappa > z_{L,n}$ for n sufficiently large. From condition (A11), $a'_n(z) = 1 - a_n(z) - 2[v_n(z) - u]/c$ for all $z \in [z_n^* - \kappa, z_n^* - \kappa/2]$. Then by Claim A15 and the monotonicity of $a_n(\cdot)$, we know that for n sufficiently large, $a'_n(z) > \epsilon/c$ for all $z \in [z_n^* - \kappa, z_n^* - \kappa/2]$. But then, we have $a_n(z_n^* - \kappa) < 1 - \kappa\epsilon/(2c)$, a contradiction to Claim A15. ■

²³To see this, note first that p_n^* is bounded above by $p_H < 1$ (Lemma A2). Next observe that the posterior is a submartingale according to the strategy of the noninvestible type. This implies that for every $\eta > 0$ we can find $p_\eta < 1$ such that, conditional on the noninvestible-type strategy, $p \in (p_\eta, 1)$ implies that the posterior goes below p_H with probability less than η . This immediately implies the existence of said $z^\dagger$.

CLAIM A17: Fix a prior $p_0 \in (0, 1)$ and some $\bar{p} \in (p_0, 1)$. For each $\psi > 0$, consider an adapted Markov function $\alpha_\psi(\cdot)$ and a belief process defined by substituting $\alpha_\psi(\cdot)$ into (A1). Take $\epsilon > 0$ and let $\bar{\mathbb{T}}$ be the random time that stops the play in the first time that $p \geq \bar{p}$. Then we have

$$\limsup_{\psi \uparrow \infty} E^{NI} \left[r_1 \int_0^{\bar{\mathbb{T}}} e^{-r_1 t} \mathbf{1}_{\{\alpha_\psi(p_t) \leq 1-\epsilon\}} dt \right] = 0.$$

PROOF:

See online Appendix.

LEMMA A9: Fix any $\kappa > 0$ and, for each $n \in \mathbb{N}$, assume that the game starts at the prior $z_n^* - \kappa$. Let $\mathbb{T}_n^*$ be the stopping time that stops the play in the first time that a posterior reaches $[z_n^*, \infty)$. We have

$$\limsup_{n \rightarrow \infty} E_n^{NI} [e^{-r_1 \mathbb{T}_n^*}] = 0.$$

PROOF:

Let $\hat{z}_n := z_n^* - \kappa$ and $\underline{z}_n := z_n^* - 2\kappa$. Let $\mathbb{T}_n(\underline{z}_n)$ be the stopping time that stops the play in the first time that the posterior reaches $\underline{z}_n$ and $\mathbb{T}_n(z_n^*)$ be the stopping time that stops the play in the first time that the posterior reaches z_n^* . Observe that for any n , $\mathbb{P}_n^{NI}[\mathbb{T}_n(\underline{z}_n) < \infty] + \mathbb{P}_n^{NI}[\mathbb{T}_n(z_n^*) < \infty] = 1$.

By Claim A15, we know that for any $\epsilon > 0$, there exists $n_1 \in \mathbb{N}$ such that $n > n_1$ implies that $v_n(\hat{z}_n^*)$ is bounded above by

$$\begin{aligned} \text{(A41)} \quad & \mathbb{P}_n^{NI}[\mathbb{T}_n(\underline{z}_n) < \infty] E_n^{NI} \left[\int_0^{\mathbb{T}_n(\underline{z}_n)} u e^{-r_1 t} dt + e^{-r_1 \mathbb{T}_n(\underline{z}_n)} v_n(\underline{z}_n) \mid \mathbb{T}_n(\underline{z}_n) < \infty \right] \\ & + \mathbb{P}_n^{NI}[\mathbb{T}_n(z_n^*) < \infty] E_n^{NI} \left[\int_0^{\mathbb{T}_n(z_n^*)} u e^{-r_1 t} dt + e^{-r_1 \mathbb{T}_n(z_n^*)} v_n(z_n^*) \mid \mathbb{T}_n(z_n^*) < \infty \right] \\ & + \epsilon. \end{aligned}$$

Next we obtain an upper bound for $v_n(z_n^*)$. For that, we let $\mathbb{T}_\lambda$ be the random time of the next Poisson shock. Note that

$$\begin{aligned} v_n(z_n^*) &= \mathbb{P}_n^{NI}[z_{\mathbb{T}_\lambda} > z_n^*] E_n^{\theta_s} \left\{ \int_0^{\mathbb{T}_\lambda} e^{-r_1 t} [u + (1 - a_n(z_t))c] dt \mid z_{\mathbb{T}_\lambda} > z_n^* \right\} \\ &+ \mathbb{P}_n^{NI}[z_{\mathbb{T}_\lambda} \leq z_n^*] E_n^{\theta_s} \left\{ \int_0^{\mathbb{T}_\lambda} e^{-r_1 t} [u + (1 - a_n(z_t))c] dt \right. \\ &\quad \left. + e^{-r_1 \mathbb{T}_\lambda} v_n(z_{\mathbb{T}_\lambda}) \mid z_{\mathbb{T}_\lambda} \leq z_n^* \right\}. \end{aligned}$$

Now we use the following facts to bound the expected value above:

(i) Because $a_n(z_t) \geq 0$, we have

$$\begin{aligned} E_n^{NI} \left[\int_0^{\mathbb{T}_\lambda} e^{-r_1 t} [u + (1 - a_n(z_t))c] dt \mid z_{\mathbb{T}_\lambda} > z_n^* \right] \\ \leq E_n^{NI} \left[\int_0^{\mathbb{T}_\lambda} e^{-r_1 t} (u + c) dt \mid z_{\mathbb{T}_\lambda} > z_n^* \right]. \end{aligned}$$

(ii) From Claim A13, $\limsup E_n^{NI} [\mathbf{1}_{\{z_{\mathbb{T}_\lambda} \leq z_n^*\}} v_n(z_{\mathbb{T}_\lambda})] \leq \limsup E_n^{NI} [\mathbf{1}_{\{z_{\mathbb{T}_\lambda} \leq z_n^*\}} u]$.

(iii) For every $\epsilon > 0$, from Claim A17,

$$\limsup \mathbb{P}_n^{NI} \left[\{z_{\mathbb{T}_\lambda} \leq z_n^*\} \cap \left\{ \int_0^{\mathbb{T}_\lambda} e^{-r_1 t} (1 - a(z_t)) dt > \epsilon \right\} \right] = 0.$$

Conditions (i), (ii) and (iii) above imply that for every $\epsilon > 0$, we can find $n_2 \in \mathbb{N}$ with $n_2 > n_1$ such that $n > n_2$ implies

$$\begin{aligned} v_n(z_n^*) &\leq \mathbb{P}_n^{NI} [z_{\mathbb{T}_\lambda} > z_n^*] E_n^{NI} \left[\int_0^{\mathbb{T}_\lambda} e^{-r_1 t} (u + c) dt \mid z_{\mathbb{T}_\lambda} > z_n^* \right] \\ &\quad + \mathbb{P}_n^{\theta_S} [z_{\mathbb{T}_\lambda} \leq z_n^*] u + \epsilon. \end{aligned}$$

Since Poisson shocks are independent of the Brownian motion,

$$E_n^{NI} \left[\int_0^{\mathbb{T}_\lambda} e^{-r_1 t} (u + c) dt \mid z_{\mathbb{T}_\lambda} > z_n^* \right] = \left(\frac{r}{r + \lambda} \right) (u + c) < u.$$

Moreover, notice that the z_t (and p_t) are submartingales conditional on $\theta = NI$, so $\mathbb{P}_n^{NI} [z_{\mathbb{T}_\lambda} > z_n^*] \geq 1/2$. Therefore, the last two observations imply $v_n(z_n^*) \leq r_1(u + c)/[2(r_1 + \lambda)] + u/2 + \epsilon$. Because $r_1(u + c)/(r_1 + \lambda) < u$ (i.e., Assumption A1) and ϵ is arbitrary, we conclude that $\limsup v_n(z_n^*) < u$.

Going back to the upper bound (A41) for $v_n(\hat{z}_n^*)$, we have shown that $\limsup v_n(z_n^*) < u$, and from Claim A16, we know that $\lim v_n(\hat{z}_n^*) = \lim v_n(z_n^*) = u$. So for (A41) to be a valid upper bound of $v_n(\hat{z}_n^*)$ for arbitrary ϵ , we must have $\limsup_{n \rightarrow \infty} E_n^{NI} [e^{-r_1 \mathbb{T}_n(z_n^*)}] = 0$, i.e., $\limsup_{n \rightarrow \infty} E_n^{NI} [e^{-r_1 \mathbb{T}_n^*}] = 0$, as desired. ■

LEMMA A10: $\lim_{n \rightarrow \infty} z_n^* = z^{**}$, where z^{**} is the myopic cutoff satisfying $R(p(z^{**})) = 0$.

PROOF:

Recall that we always have $z_n^* \geq z^{**}$. Hence assume (toward a contradiction) that we can find some $\epsilon > 0$ such that, taking a subsequence if necessary, $z_n^* > z^{**} + \epsilon$ for every n . For every n , consider the game starts at $z_n^* - \epsilon/2 > z^{**} + \epsilon/2$. Let $\mathbb{T}_n^*$ be the stopping time that stops the play in the first time that a posterior reaches $[z_n^*, \infty)$. By Lemma A9, we have $E_n^{NI} [e^{-r_1 \mathbb{T}_n^*}] \rightarrow 0$, implying that $\limsup W_n(p(z_n^* - \epsilon/2)) \leq 0$. But since $z_n^* - \epsilon/2 > z^{**} + \epsilon/2$, the principal can get a strictly positive payoff by terminating the relationship when the next stopping

opportunity arrives. So the principal has a profitable deviation at $z_n^* - \epsilon/2$ when n is sufficiently large, a contradiction. ■

PROOF OF THEOREM 3:

In light of Corollary A2, we extend each W_n continuously from $(0, 1)$ to $[0, 1]$ by setting $W_n(0) = 0$ and $W_n(1) = \lambda w_{NI}/(r_2 + \lambda)$.

We first show that $W_n(\cdot)$ converges to $\underline{W}(p^{**}) = 0$. By Lemmas A9 and A10, it is easy to see that $\lim_{n \rightarrow \infty} W_n(p) = 0$ for all $p < p^{**}$. Suppose toward a contradiction, taking a subsequence if necessary, that $\lim_{n \rightarrow \infty} W_n(p^{**}) = \delta > 0$. Let $\epsilon = \delta(1 - p^{**})/(2w_{NI})$. For n large enough, we have $W_n(p^{**} - \epsilon) < \delta/2$. Since $W_n(\cdot)$ is convex (by Corollary A2), we have

$$\frac{W_n(1) - W_n(p^{**})}{1 - p^{**}} \geq \frac{W_n(p^{**}) - W_n(p^{**} - \epsilon)}{\epsilon} \geq \frac{w_{NI}}{1 - p^{**}},$$

which implies $W_n(1) > w_{NI}$, a contradiction.

Next, for each n , because $W_n(0) = \underline{W}(0)$, $W_n(1) = \underline{W}(1)$, and $W_n(\cdot)$ is increasing and convex, we have $0 \leq W_n'(\cdot) \leq \lambda/(r_1 + \lambda)$. But then, we always have $p^{**} = \arg \max_{p \in [0,1]} |W_n(p) - \underline{W}(p)|$. Hence, uniform convergence of W_n follows immediately from its pointwise convergence at p^{**} . ■

REFERENCES

- Aghion, Philippe, and Matthew O. Jackson. 2016. "Inducing Leaders to Take Risky Decisions: Dismissal, Tenure, and Term Limits." *American Economic Journal: Microeconomics* 8 (3): 1–38.
- Banerjee, Robin. 2022. *Corporate Frauds: Business Crimes Now Bigger, Broader, Bolder*. New Delhi, India: SAGE Publications.
- Bergemann, Dirk, and Ulrich Hege. 1998. "Venture Capital Financing, Moral Hazard, and Learning." *Journal of Banking and Finance* 22 (6): 703–35.
- Bergemann, Dirk, and Ulrich Hege. 2003. "The Value of Benchmarking." In *Venture Capital Contracting and the Valuation of High Tech Firms*, edited by Joseph McCahery and Luc Renneboog, 83–107. Oxford, UK: Oxford University Press.
- Bolton, Patrick, and Christopher Harris. 1999. "Strategic Experimentation." *Econometrica* 67 (2): 349–74.
- Daley, Brendan, and Brett Green. 2012. "Waiting for News in the Market for Lemons." *Econometrica* 80 (4): 1433–1504.
- Dilmé, Francisc. 2019. "Dynamic Quality Signaling with Hidden Actions." *Games and Economic Behavior* 113: 116–36.
- Ekmekci, Mehmet, Leandro Gorno, Lucas Maestri, Jian Sun, and Dong Wei. 2021. "Learning from Manipulable Signals." Unpublished. <https://arxiv.org/abs/2007.08762>.
- Ekmekci, Mehmet, and Lucas Maestri. 2022. "Wait or Act Now? Learning Dynamics in Stopping Games." *Journal of Economic Theory*: 105541.
- Faingold, Eduardo, and Yuliy Sannikov. 2011. "Reputation in Continuous-Time Games." *Econometrica* 79 (3): 773–876.
- Fudenberg, Drew, and David K. Levine. 1989. "Reputation and Equilibrium Selection in Games with a Patient Player." *Econometrica* 57 (4): 759–78.
- Fudenberg, Drew, and David K. Levine. 1992. "Maintaining a Reputation When Strategies Are Imperfectly Observed." *Review of Economic Studies* 59 (3): 561–79.
- Gompers, Paul, and Josh Lerner. 2004. *The Venture Capital Cycle*. Cambridge, MA: MIT Press.
- Karatzas, Ioannis, and Steven Shreve. 1998. *Brownian Motion and Stochastic Calculus*. New York: Springer.
- Kolb, Aaron M. 2015. "Optimal Entry Timing." *Journal of Economic Theory* 157: 973–1000.
- Kolb, Aaron M. 2019. "Strategic Real Options." *Journal of Economic Theory* 183: 344–83.

- Kreps, David, and Robert Wilson.** 1982. "Reputation and Imperfect Information." *Journal of Economic Theory* 27 (2): 253–79.
- Kuvalekar, Aditya, and Elliot Lipnowski.** 2020. "Job Insecurity." *American Economic Journal: Microeconomics* 12 (2): 188–229.
- Leitner, Yaron, and Basil Williams.** Forthcoming. "Model Secrecy and Stress Tests." *Journal of Finance*.
- Milgrom, Paul, and John Roberts.** 1982. "Predation, Reputation, and Entry Deterrence." *Journal of Economic Theory* 27 (2): 280–312.
- Pei, Harry.** 2020. "Reputation Effects under Interdependent Values." *Econometrica* 88 (5): 2175–2202.
- Pham, Huyền.** 2009. *Continuous-time Stochastic Control and Optimization with Financial Applications*. Berlin, Germany: Springer.
- Sun, Yiman.** 2021. "A Dynamic Model of Censorship." SSRN 4078301.
- Weston, Liz.** 2016. *Your Credit Score: How to Improve the 3-Digit Number That Shapes Your Financial Future*. Old Tappan, NJ: FT Press.