

Paying with information

AYÇA KAYA

Department of Economics, University of Miami

The founder of a start-up (principal) who has a project with uncertain returns must retain and incentivize an agent using promise of future payments and information gathering. The agent's effort incrementally advances production and such advance is a prerequisite for gathering new information. The principal decides how much information to gather based on these incremental advancements. The principal faces cash constraints. The agent's outside option is large relative to his effort cost. Equilibrium features one of two outcomes: *immediate learning*, whereby the agent's compensation is low, learning is immediate and retention is possible only conditional on the project being of high quality; or *gradual learning*, whereby the agent's compensation is high, learning is gradual, the agent never quits and effort is inefficiently high.

KEYWORDS. Informational incentives, information control, agency costs.

JEL CLASSIFICATION. C72, D82, D83, D86.

1. INTRODUCTION

Consider a start-up working toward launching a new product, say a new reusable drinking straw.¹ The founder (the principal) hires an expert on industrial design (the agent), whose efforts are necessary to create a mass-produceable concept. There is little uncertainty about whether the agent will be able to complete such a design. At the same time, the reusable drinking straw market is sufficiently saturated that unless the straw has truly innovative features, it can generate only a moderate income stream. The principal does have such a novel idea: a slide apart straw for easy cleaning. There is uncertainty about whether such a straw can be implemented in a cost-effective way and whether it is marketable.

During his employment, each incremental productive contribution of the agent gives the principal an opportunity to gather information about the feasibility and marketability of the new concept, while also contributing to its eventual success. For instance, once the agent completes an initial design for the slide-apart concept, information will be revealed about what type of materials are best suited to achieve proper sealing, which the principal then can attempt to price in order to assess feasibility. The initial design can also be used to test the market's response. Thus, completing various

Ayça Kaya: a.kaya@miami.edu

See RainHydration product page on Amazon: <https://www.amazon.com/Rain-Straw-Reusable-Drinking-Cleaning/dp/B084GFJ1CK> and on Facebook: <https://www.facebook.com/rainhydration/>.

¹This is based on a product called "Rain Straw," which was launched in 2019 by RainHydration (<https://www.facebook.com/rainhydration/>).

steps of production both contributes toward the eventual output, and creates an opportunity for the principal to gather information. This paper takes this dual role of effort as a point of departure and considers the problem of a principal (she) who tries to motivate the participation and effort of an agent (he) by leveraging these two types of returns.

Production under uncertainty, which is the source of the dual roles of effort, is a common feature of innovative startups. Such startups are the main application of this paper's theory. In addition to facing uncertainty, startup founders often face cash-in-advance constraints, which compel them to pay in stocks rather than wages. Further, a quick browsing of start-up advice reveals that convincing highly skilled agents to give up high-paying employment in established ventures while lacking the liquidity to match the market salaries is a major hurdle facing founders.² Thus, the start-up owners' decisions are often shaped by the need to retain their employees rather than incentivizing their effort while employed. The principal-agent model of this paper incorporates these three features: (i) production under uncertainty and the resulting dual roles of effort; (ii) cash constraints necessitating payment in shares; and (iii) relatively large outside options of the agent. The analysis reveals that the conjunction of these features, along with the principal's inability to commit to long term information gathering strategies, can explain gradual learning as a form of agency cost, which may contribute to understanding the observed high rate of delayed attrition.³

In the formal model, the principal has a project with (binary) uncertain quality: the project is either "good" or "bad," which, once completed, will lead to a future stream of income. The expected value of this future stream depends on both the unknown quality of the project and the total amount of effort the agent exerts while working on the project. The principal and the agent hold a common prior about the quality. At the beginning of the relationship, the principal commits to a contract. Due to the principal's cash constraints and her resulting inability to offer upfront payments, the contract simply specifies the agent's share of future income streams.

Once the contract is agreed upon, a production game ensues which lasts one unit of time and, if the agent is successfully retained, culminates in the launch of a new product. The game is split into discrete (but short) time periods. In each period, the agent may work, shirk, or (irreversibly) quit. At the beginning of each period, before the agent makes his choice, the principal designs, and for the duration of the period, commits to an informative signal. The signal is modeled as a probability distribution over posterior beliefs. Effort is a prerequisite for learning: for the principal's signal to generate a realization, the agent must choose to work (and not shirk or quit) in that period. Due to the principal's lack of commitment to future signals, this stage is analyzed as a dynamic

²For instance, in a recent HBR article specifically discussing these issues, Amelia Friedman opens with this: "As a startup founder, I'm constantly struggling to recruit top talent without breaking the bank. We can't always match market salaries, but we need exceptional (read: expensive) talent in order to build from scratch. How do you recruit a developer making well into six figures, or an experienced salesperson with four kids in private school?" (<https://hbr.org/2018/07/7-compensation-strategies-for-cash-strapped-startups>).

³For instance, among all start-ups established in 2014 and failed in or before 2019, only 41.2% failed within the first year, while 10% survived 5 years before failing. (Calculations are based on the BLS data at https://www.bls.gov/bdm/us_age_naics_00_table7.txt).

game of information disclosure. I consider the perfect Bayesian equilibria of this game. I show that the game admits an (essentially) unique equilibrium, which is characterized using backwards induction.⁴

The equilibrium of the production game (almost) always features gradual learning, whereby the principal collects information slowly. In fact, at the extreme case where the agent's share of the output is large enough so that the expected marginal *monetary* return for his effort covers his marginal cost of effort as well as his forgone outside option, the principal never reveals information. When the monetary marginal return does not cover these costs, the principal supplements it with informational returns for the agent's effort. The agent values information about the underlying state because he repeatedly makes decisions about whether to quit and whether to exert effort, and his optimal choice varies with the true state. Thus, even if his expected monetary returns do not justify it, the agent may be willing to exert effort when information collection is tied to such effort. In this sense, information can be an incentive device supplementing monetary payments. Of course, once information is generated it cannot be forgotten or reused, and that is why the principal gives out the information in small increments, resulting in gradual learning.

When deciding the agent's share at the contracting stage, the principal takes into account its impact on the equilibrium of the ensuing production game. If the agent's share is relatively large, the monetary marginal return for his effort is large and, therefore, the principal is able to (and, in equilibrium, does) dispense information more slowly. For the principal, there are two benefits of extending the learning period in this manner. First, since the agent's effort is a prerequisite for information collection, slower learning means more effort by the agent. Second, an extended period of learning may help avoid project abandonment: If the agent becomes very pessimistic about the project's promise (i.e., if his belief places high weight on the state being bad), he may choose to abandon the project (quit) *unless* the project has already substantially advanced and its completion is near. If the principal can sufficiently slow down learning so that bad news is concealed until such time, project abandonment is avoided.

The main result of the paper is to show that the principal's optimal choice of contract may take one of two forms: (i) a share small enough that the principal cannot delay information gathering at all, and immediately collects full information; or (ii) a share *just* large enough that it allows the principal to extend the learning period exactly until such time when the agent would not consider quitting even when maximally pessimistic. I refer to these two outcomes as the "immediate learning" and "gradual learning" outcomes, respectively. In either case, the outcome of the production game is efficient conditional on good state (i.e., the agent exerts effort throughout). This is because the agent's effort is a prerequisite for information collection and, therefore, he works until the belief becomes degenerate. Conditional on bad state, the immediate learning outcome leads to inefficient project abandonment with probability 1 while the gradual learning outcome completely eliminates project abandonment but leads to inefficient effort by the agent during the learning period.

⁴I require that off the equilibrium path, the beliefs be updated only based on the principal's signals, and not based on other aspects of the history; see Section 2.

The intuition behind the principal's choice of contract is best understood by expressing the principal's payoff as the total surplus generated less the agent's equilibrium payoff. Once the agent's share is high enough that abandonment can be avoided, increasing it further (weakly) increases the agent's payoff and reduces the overall surplus. The latter is because, with a higher pay to the agent, the principal can and is compelled to dispense information even slower. Consequently, the principal never pays more than the share that just allows avoiding abandonment. Any share which, conditional on bad state, leads to project abandonment with probability strictly between 0 and 1 is worse than either the immediate or the gradual learning outcomes: the former if the equilibrium surplus conditional on bad state is negative, and the latter if this surplus is positive.

Some of the modeling choices are worth elaborating. The production technology is assumed to be such that each increment of the agent's effort contributes the same amount to the output, that is, the marginal product of the agent's effort is constant. Further, each increment of effort generates the same opportunities for the principal to collect information. These assumptions, which abstract away from possible inherent interaction between productive and informational returns of effort, render the analysis very tractable and make it possible to highlight how these returns to effort interact via the principal's strategic incentives to use them. Additionally, I assume that the principal can "flexibly" design the informative signals, and in particular can perfectly learn the state based on a single increment of the agent's effort. Once again, this assumption abstracts from the constraints that the principal may face in reality but allows a tractable analysis. Moreover, in the outcome of the model, gradual learning may arise even though immediate learning is feasible. Thus, the analysis highlights a set of strategic considerations—rather than exogenous constraints—leading to slow learning.

In what follows, Section 1.1 discusses the related literature, Section 2 lays out the model. Sections 3 and 4, respectively, characterize the outcomes of the production stage and the contracting stage. Section 5 discusses some variations of the model and clarifies the roles played by various assumptions of the main model.

1.1 *Related literature*

This paper combines dynamic information design with (static) optimal contracting. While the contracting problem studied in this paper is simple, the dynamic information design component incorporates aspects that are novel relative to the existing literature. However, the main contribution is the analysis of the interaction between monetary contracting and information design.

The production stage of the principal–agent model of this paper contributes to the large and growing dynamic information design literature. One of the main results identifies conditions under which gradual learning occurs. Gradual learning has been identified as equilibrium outcomes in other contexts modeled as dynamic information design problems. For instance, Ely and Szydlowski (2020), Orlov, Skrzypacz, and Zryumov (2020), and Bizzotto, Rudiger, and Vigier (2021) demonstrate that slowing down information release is a valuable tool to delay an agent's irreversible action, such as quitting. Relative to this literature, this paper introduces two novel considerations: The first is the

role of information control when incentivizing the agent on the (reversible) work/shirk margin in addition to the (irreversible) quit/participate margin, thereby incorporating the two classical constraints of principal–agent models of employment relationships. The second is the interaction between informational and monetary incentives, recognizing that in many situations a principal is likely to have some control over both. Additionally, none of these papers consider the agent's (receiver's) costly effort as a precondition for information gathering.

In spite of these differences, the analysis of the production stage of this game exhibits similarities to these papers. For instance, even though [Orlov, Skrzypacz, and Zryumov \(2020\)](#) study the impact of exogenous outside information on the principal's incentives, there are some common insights. In particular, that lack of commitment precludes delayed information release is a common theme while the cost of resulting incremental learning manifests differently in the two papers. The dynamics of the gradual learning outcome in this paper is similar to [Ely and Szydlowski \(2020\)](#) in which a principal persuades an agent to stay long enough to complete a project of unknown length. In both, the principal “leads the agent on” until project completion is sufficiently close. However, in their case these dynamics occur only under full commitment and cannot be replicated with policies that keep the agent to his autarky payoff at each instant.

Other papers in the dynamic information design literature include [Ely \(2017\)](#) and [Renault, Solan, and Vieille \(2017\)](#) which, differently from this paper, feature sequences of myopic receivers, evolving state and full commitment. [Ball \(2022\)](#) considers use of promised information to incentivize an agent's choices under full commitment and shows that the principal can use threat of cutting information to punish an agent who deviates from recommended actions. In this paper, due to the finiteness of the horizon and the principal's lack of commitment, the threat of such punishment does not play a role. [Au \(2015\)](#) and [Guo and Shmaya \(2018\)](#) consider dynamic persuasion of privately informed receivers. [Liu \(2021\)](#) considers how differences between the patience levels of the principal and the agent affect the principal's incentive to reveal information gradually under full commitment. [Smolin \(2021\)](#) studies dynamic information disclosure (with commitment) about unknown productivity to an agent in order to achieve retention, where the optimal policy leverages the agent's career concerns. [Che, Kim, and Mierendorff \(forthcoming\)](#) depart from the flexible information design assumption by introducing cost of waiting and bounds on speed of information release.

Each step of the principal's problem in the production game is a Bayesian persuasion problem ([Kamenica and Gentzkow \(2011\)](#) and [Aumann and Maschler \(1995\)](#)), with an additional participation constraint of the agent. In line with the analysis of such problems in [Le Treust and Tomala \(2019\)](#), [Doval and Skreta \(2018\)](#), the optimal signal's support may contain more realizations than the number of states.

Importantly, none of the above papers discuss the interaction of monetary rewards and information design. A paper that studies how information can be exchanged for money is [Hörner and Skrzypacz \(2016\)](#), which shows that the promise of incremental information release can be leveraged to extract larger payments from an uninformed agent and creates a motive for gradual learning. In a stylized model, [Li \(2017\)](#) augments the

classical static persuasion problem of [Kamenica and Gentzkow \(2011\)](#) by allowing imperfect transfer of surplus contingent on signal realizations. The principal's information gathering in the current paper can broadly be interpreted as monitoring the progress of the project. In an environment with no persistent uncertainty but with moral hazard, [Orlov \(2022\)](#) studies the optimal codetermination of an agent's monetary compensation and performance feedback and highlights how the agent's information affects the cost of incentive provision. In a moral hazard problem without uncertainty, [Lizzeri, Meyer, and Persico \(2002\)](#) demonstrate that a principal who controls monetary rewards is better off not providing interim performance feedback, even though with such feedback he can implement higher effort.

2. MODEL

A principal hires an agent to work on a project. The relationship lasts 1 unit of time, which is split into periods of length Δ . Assume that $K = 1/\Delta$ is an integer so that the interaction lasts K periods. Thus, time periods are indexed $\{1, \Delta, 2\Delta, \dots, 1\}$. The focus is on the limit as $\Delta \rightarrow 0$. If the agent remains employed until the end, that is, until $t = 1$, the project is completed, and results in a stream of future earnings. During each period of his employment the agent may choose to work (exert effort) or shirk. The present discounted value of the future earnings depends on an unknown state representing the project quality as well as the amount of time the agent spends exerting effort during production. The state θ is either good ($\theta = 1$) or bad ($\theta = 0$). The principal and the agent are symmetrically uninformed about θ and they have a common prior belief that assigns probability μ_0 to $\theta = 1$.

Timing The interaction starts with a take-it-or-leave-it contract offer by the principal. Once a contract is agreed upon, the “production stage” begins. In each period during production, the following sequence of events occur:

- The principal designs a “signal” informative of the underlying state;
- Agent, after observing the signal design, chooses one of three options:
 - (i) work ($e = 1$);
 - (ii) shirk: remain on the project but do not work ($e = 0$);
 - (iii) quit: abandon the project.
- If the agent chooses “work,” information is revealed according to the principal's signal. Otherwise, no information is revealed.
- At each $t < 1$, if the agent does not quit, the play moves to the next period.

All choices of the agent and the principal are publicly observed, and thus, the information remains symmetric throughout the interaction.

The value of the eventual output depends on the agent's cumulative effort as well as the unknown state of the world. When the project quality is $\theta \in \{0, 1\}$ and E is the total

duration of time that the agent spends working, the expected present discounted value of future earnings is

$$Y(\theta, E) \equiv Y_L + (Y_H - Y_L)E\theta,$$

where $Y_H > Y_L > 0$. This specification has the following implications: First, the unknown state determines the marginal product of the agent's effort. And second, the production technology is linear, that is, the marginal product of the agent's effort is independent of the so-far accumulated effort.

Agent's effort also makes information collection possible. As the agent works, more components of the final outcome is completed. The principal then can collect information, for instance, by testing the completed components for performance or for end-user opinions. The principal's promise of information collection at time t specifies a probability distribution F_t over the posteriors at $t + \Delta$. A realization from F is drawn if and only if the agent exerts effort at time t .⁵ I place no restrictions on F apart from the standard "Bayesian plausibility" requirement: at each t , $\mathbb{E}_{F_t}[\mu_{t+\Delta}] = \mu_t$, where μ_t is the probability that the common belief assigns to good state at time t .⁶ Let $\mathcal{F}(\mu_t)$ represent the set of all probability distributions over $\mu_{t+\Delta}$ with mean μ_t . Focusing on the workhorse models of flexible information design and linear production provides a tractable framework in which to study the interdependence of learning and production.

At time t , the principal can commit to the current signal F_t , but not to future sequences of signals. Inability to commit to future signals is natural in an uncertain environment, where the understanding of what will be possible to test and what will be the next step of development may evolve throughout the production process.

Payoffs The agent has a flow outside option u , which he gives up while employed. Choosing $e = 1$ leads to an additional flow cost of $c > 0$. The principal is biased toward higher effort by the agent. Specifically, the agent's effort generates a flow payoff ε for the principal, satisfying $0 < \varepsilon < c$.⁷

Contracts A feasible contract is a share $\alpha \in [0, 1]$ of the future income that the agent is entitled to. This is agreed upon before production starts and is never adjusted. These contractual restrictions are further discussed in Section 5.2.

⁵In the main part of the paper, I assume that new information cannot be collected at time t unless the agent works in the current period. In Section 5.3, I demonstrate that, even if the principal could go back to previously completed components to gather new information, in equilibrium he would not choose to do so as long as going back entails a small lump-sum cost. The amount of cost that would be sufficient to dissuade the principal from generating information based on previously finished components approaches 0 as the length Δ of a period approaches 0.

⁶Throughout, I adapt the notation $\mathbb{E}_H[Z(x)]$ to represent the expected value of the function $Z(\cdot)$ of a random variable x , where H is a probability distribution over x .

⁷When $\varepsilon > 0$, however small, the equilibrium of the production stage is (essentially) unique, while if $\varepsilon = 0$ there may be multiple equilibria. I highlight the precise role of $\varepsilon > 0$ in Section 3.2. The equilibrium uniqueness and the qualitative properties of the unique equilibrium are independent of the size of ε as long as it is positive. Thus, this specification can be viewed as a perturbation facilitating equilibrium selection.

Outcomes A (payoff-relevant) outcome of this game consists of a realized total effort $E = \int_0^1 e(s) ds$ and the realized time of quitting $\tau \in [0, 1]$. By convention, if the agent never quits, then $\tau = 1$. Fixing α , and given an outcome (E, τ) , the agent's and the principal's respective payoffs conditional on state θ are

$$u(E, \tau|\theta) = \begin{cases} \alpha Y(E, \theta) - Ec - u & \text{if } \tau = 1 \\ -Ec - \tau u & \text{otherwise} \end{cases}, \quad (1)$$

$$\pi(E, \tau|\theta) = \begin{cases} (1 - \alpha)Y(E, \theta) + \varepsilon E & \text{if } \tau = 1 \\ \varepsilon E & \text{otherwise} \end{cases}. \quad (2)$$

Production equilibrium The production stage following an agreement on a share α is analyzed as a dynamic game due to the principal's lack of commitment. A time- t history h_t consists of the effort path $\{e(s)\}_{s \in (0, t)}$, signal offers $\{F_s\}_{s \in (0, t)}$, and signal realizations.

Beliefs: Each history of length t is associated with a belief μ_t interpreted as the probability assigned to the good state at such history. On and off the path of equilibrium, beliefs are updated only based on signal realizations. In particular, off-path signal offers or choices by the agent do not affect beliefs.⁸

Strategies and payoffs: A behavior strategy of the principal specifies a signal $F_t \in \mathcal{F}(\mu_t)$ at each history h_t for each $t = \Delta k$, $k = 1, \dots, 1/\Delta - 1$. A behavior strategy of the agent maps each history into a choice of whether to quit, and if not, a choice of $e(t) \in \{0, 1\}$. Let σ^P and σ^A be arbitrary mixed behavior strategies for the principal and the agent, respectively. Each strategy profile $\sigma = (\sigma^P, \sigma^A)$ induces a joint probability distribution G^σ over (E, τ, θ) . Then the principal and the agent's payoffs from such strategy profile, respectively are $\mathbb{E}_{G^\sigma}[\pi(E, \tau|\theta)]$ and $\mathbb{E}_{G^\sigma}[u(E, \tau|\theta)]$.

Equilibrium: For fixed Δ , an equilibrium is a strategy profile (σ^P, σ^A) and a belief system satisfying the above refinement, which constitutes a perfect Bayesian equilibrium of the game.⁹ I focus on the limit of the equilibrium outcomes as $\Delta \rightarrow 0$. A “limit equilibrium” refers to any strategy profile and belief system that can be obtained as the limit of a selection of equilibria for a sequence Δ_n , where $\lim_{n \rightarrow \infty} \Delta_n = 0$.

Contracting Prior to production, the principal chooses the agent's share α to maximize her expected payoff subject to delivering the agent a nonnegative payoff, where the payoff calculations anticipate a limit equilibrium of the production stage.¹⁰

Assumptions I maintain the following restrictions on parameters:

Project abandonment is inefficient

$$Y_L > u. \quad (3)$$

⁸This is natural since neither the principal nor the agent has private information, and thus their off-path actions cannot be interpreted to convey different pieces of information. This corresponds to the “no signaling what you don't know” property of PBE discussed in [Watson \(2017\)](#).

⁹The players form beliefs about the underlying state based on the signal outcomes. At every history, on or off-path, their strategies maximize their payoff given their belief and the continuation strategies of the other.

¹⁰As established in Section 3.2 all such equilibria are payoff equivalent. Thus, this problem is well-defined.

Effort is efficient in the good state

$$Y_H - Y_L > c. \quad (4)$$

Retention is (sufficiently) costlier than incentivizing effort

$$\frac{c}{u} < \frac{Y_H - Y_L}{Y_L} < \frac{u}{c}. \quad (5)$$

Assumptions (3) and (4), along with $c > \varepsilon$, pin down the surplus-maximizing outcomes in each state: in the bad state, the pair's surplus is maximized if the agent is retained but exerts no effort. In the good state, it is maximized when the agent exerts effort at each instant. The essential component of Assumption (3) is the *existence* of states where retention is efficient but effort is not, rather than the implication that there are no states in which abandonment is efficient. Indeed, if there were “worse states” in which project abandonment was efficient, the principal would never offer high enough shares that would make it optimal for her to conceal those states in the ensuing production game.¹¹

Assumption (5) requires that the agent's outside option u is larger than his effort cost c .¹² It embeds two assumptions. The first inequality in (5) states that the ratio u/c of cost of retention to cost of effort, is larger than the ratio $Y_L/(Y_H - Y_L)$ of returns to retention versus returns to effort. This is the main driver of gradual learning as shown in Proposition 3 in Section 5.1. The second inequality places a different lower bound on u/c . Under this assumption, any information collection policy that achieves retention until completion allows the principal to extract effort while concealing the bad state. I highlight its precise role in Section 3.2.

3. PRODUCTION STAGE

This section characterizes the outcome of the production stage after the share α is agreed upon, focusing on the nontrivial case where

$$c + u \leq \alpha Y_H. \quad (6)$$

In words, (6) requires that the agent's promised compensation is large enough so that if the state is known to be high, he optimally stays and works until completion.

Section 3.1 presents some preliminary results. Section 3.2 characterizes the equilibrium outcomes for α satisfying (6). In preparation for the analysis of the contracting stage, Section 3.3 states the main result about limit equilibrium outcomes and payoffs for all values of α , including those that violate (6).

¹¹By the same token, if, in this two-state environment, abandonment was efficient in the bad state, the principal could extract full surplus by using the share to align incentives in the good state (i.e., by offering $\alpha = (c + u)/Y_H$), and immediately revealing the low state.

¹²This is a common occurrence in innovative industries where highly skilled agents forego high paying jobs in established companies to join start-ups. While c represents their opportunity cost of leisure or other self-serving activity, u represents their opportunity cost of higher earnings in alternative employment.

3.1 Preliminaries

In any equilibrium, all continuation payoffs and continuation strategies can be expressed as functions of calendar time t , accumulated effort E_t , and current belief μ_t . This is thanks to the game's finite horizon, and follows by backwards induction.

Fixing an equilibrium, let (E^*, τ^*) be the random variables representing the total effort in equilibrium and date of quitting. At each history with associated (t, E_t, μ_t) and conditional on each state, the continuation equilibrium induces a probability distribution over (E^*, τ^*) , say G_{t, E_t, μ_t}^θ . It is convenient to introduce the following notation for the continuation payoffs:

$$\begin{aligned}\Pi^*(t, E_t, \mu_t) &= \mu_t \Pi_G^*(t, E_t, \mu_t) + (1 - \mu_t) \Pi_B^*(t, E_t, \mu_t) \\ U^*(t, E_t, \mu_t) &= \mu_t U_G^*(t, E_t, \mu_t) + (1 - \mu_t) U_B^*(t, E_t, \mu_t)\end{aligned}$$

with $\Pi_\theta^*(t, E_t, \mu_t) = \mathbb{E}_{G_{t, E_t, \mu_t}^\theta} [\pi(E^*, \tau^* | \theta)]$ and $U_\theta^*(t, E_t, \mu_t) = \mathbb{E}_{G_{t, E_t, \mu_t}^\theta} [u(E^*, \tau^* | \theta)]$.

The rest of this section makes some preliminary observations and introduces notation that will be used in equilibrium characterization.

Autarky payoffs and actions The agent can always ignore the principal's signal offers and choose the action he finds optimal in the absence of new information. His payoff from doing so, which I refer to as his “autarky payoff” places a lower bound on his continuation payoff in any equilibrium and any history. Let $U^a(t, E_t, \mu_t)$ denote this autarky payoff at a history associated with (t, E_t, μ_t) . Then

$$U^a(t, E_t, \mu_t) \equiv \max\{U^{\text{work}}(t, E_t, \mu_t), U^{\text{shirk}}(t, E_t, \mu_t), 0\}, \quad (7)$$

where $U^x(t, E_t, \mu_t)$, $x \in \{\text{work}, \text{shirk}\}$ is the agent's payoff if he chooses action x in all future periods without receiving any further information, and 0 is his payoff from quitting. As formally established in Lemma 4 in the [Appendix](#), unless the autarky action is to quit, it remains the same over time in the absence of new information. This justifies the calculation of payoffs $U^x(t, E_t, \mu_t)$ along paths where the agent's action is constant. I refer to the agent's action(s) that attain the maximum in (7) as his “autarky action.”

Getting “vested” As the product launch approaches, the option of remaining on the project until completion while shirking becomes attractive. When t is large enough so that $(1 - t)u < \alpha Y_L$, this option dominates quitting, and the agent is willing to stay even if he finds out that the state is bad, in which case his expected payment for completing the project is just αY_L . Specifically, letting

$$\underline{t}(\alpha) = \max\left\{0, 1 - \frac{\alpha Y_L}{u}\right\}, \quad (8)$$

for dates $t \geq \underline{t}(\alpha)$ the agent's autarky action can never be quit. Accordingly, I say that the agent is “vested” when $t \geq \underline{t}(\alpha)$.

“Work until vesting” outcome As is formally established in Section 3.2, on the path of any equilibrium, the agent always works unless the state is revealed to be bad, and the principal delays such revelation until the agent is vested whenever this is feasible. Such delay is feasible, if the belief is sufficiently high so that the agent’s payoff from working until $\underline{t}(\alpha)$ and then learning the true state at that time is no less than his autarky payoff. Namely, work until vesting is feasible if $\mu_t \geq \mu_t^*(E_t)$ where the belief cutoff $\mu_t^*(E_t)$ is defined by

$$U^a(t, E_t, \mu_t^*(E_t)) = \underbrace{\alpha Y_L - u(1-t) - c(1-\underline{t}(\alpha))}_{\text{bad state payoff if work until vesting}} + \underbrace{\mu_t^*(E_t) [\alpha(Y_H - Y_L) - c]}_{\text{additional good state payoff}} (1-\underline{t}(\alpha)).$$

The right-hand side calculates the agent’s payoff when his current belief is $\mu_t^*(E_t)$ and when he anticipates that in the good state he will work throughout, while in the bad state he will work exactly until $\underline{t}(\alpha)$ and then shirk. The defining requirement is that this payoff is equal to the agent’s autarky payoff.

Feasibility of effort If, in a given period, the agent stays and exerts effort, he incurs a cost $\Delta(c+u) > 0$. When the belief is low, the expected monetary returns to the agent’s effort would typically not cover this cost. In such cases, the principal can use information as an additional incentive tool. Nevertheless, when the belief is sufficiently low, even the promise of a fully informative signal may not be sufficient. In such cases, inducing effort is not feasible. Fixing Δ , t , and E_t , I let $\underline{\mu}_t^\Delta(E_t)$ be the lowest belief at which it is feasible to induce effort. Importantly, this cutoff approaches 0 as the period length Δ vanishes. This is intuitive since the cost the agent must incur to receive information is $\Delta(c+u)$, which approaches 0 as $\Delta \rightarrow 0$. The formal definition of the cutoff $\underline{\mu}_t^\Delta(E_t)$ as well as the proof of convergence is deferred to the [Appendix](#).

3.2 Equilibrium characterization

In this section, I characterize the equilibrium outcomes of the production stage. The characterization is via backwards induction.

Since $t = 1 - \Delta$ is the last decision point of the agent, the principal’s information revelation decision at that date is irrelevant and the agent chooses his autarky action. Consider $t = 1 - 2\Delta$. The principal’s problem has two steps: choosing (i) the optimal informative signal that induces effort; and choosing (ii) whether to induce effort.

Figure 1 illustrates the agent’s and the principal’s payoffs as functions of the posterior belief $\mu_{1-\Delta}$ conditional on the agent exerting effort at $t = 1 - 2\Delta$. The agent’s payoff is simply his autarky payoff, and thus exhibits a convex kink at the belief where his autarky action switches from shirk to work, namely $\mu_{1-\Delta} = c/\alpha(Y_H - Y_L)$. Note that at $t = 1 - \Delta$ this cutoff is both the belief threshold above which the autarky action is “work” and the threshold $\underline{\mu}_{1-\Delta}^\Delta(E_{1-\Delta})$ below which effort is no longer feasible. The principal’s payoff exhibits a discrete jump at this same belief due to the switch in the agent’s choice. Importantly, the concavification of the principal’s payoff exhibits a concave kink at this belief cutoff.¹³

¹³This is thanks to the intrinsic value $\varepsilon > 0$ that she attaches to effort. If ε were zero, the concavification of her payoff would be linear, and thus she would have multiple optimal signals, including those that would

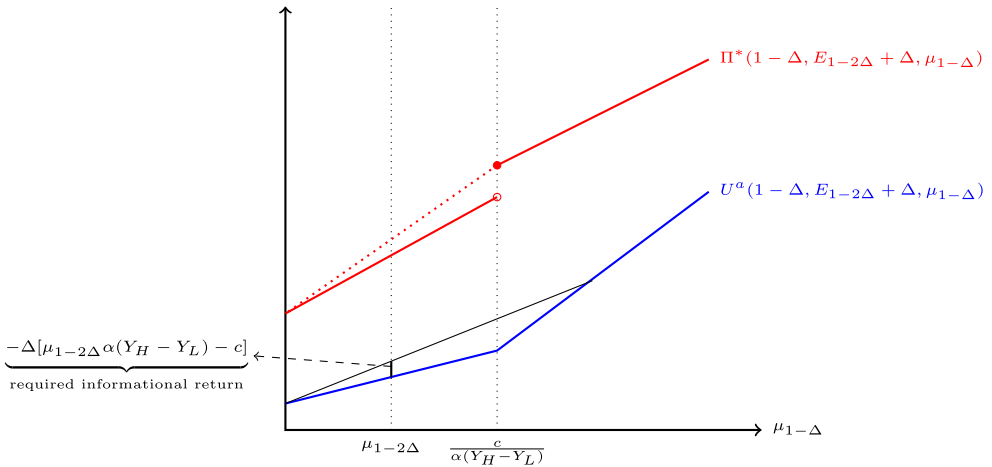


FIGURE 1. Continuation payoffs of the principal and the agent as functions of the belief at $t = 1 - \Delta$.

When the current belief $\mu_{1-2\Delta}$ is as shown in the figure, the monetary expected return $\Delta\mu_{1-2\Delta}\alpha(Y_H - Y_L)$ to effort does not cover the cost of effort $c\Delta$. Thus, the principal must promise a signal that delivers the agent a payoff above his expected autarky payoff in order to cover this deficit. An optimal signal that the principal may choose is illustrated in the figure. There are two posteriors in the support of this signal: 0 and a posterior strictly above the cutoff $c/\alpha(Y_H - Y_L)$. Importantly, this signal leaves the agent no rents: renders him indifferent between working and shirking. This important property results from the fact that the principal's payoff is concave and the agent's is convex.

Next, I argue that the principal always chooses to induce effort. What drives this are first, that the principal's and the agent's incentives are aligned in the good state: they both prefer that the agent works; second, that the incentives are diametrically opposed in the bad state: an increase in expected effort helps the principal and hurts the agent; and last, the agent receives his autarky payoff at $t = 1 - 2\Delta$ regardless of whether he works or not. Then, if the principal does *not* induce effort at $1 - 2\Delta$, relative to when she does, the agent is worse off conditional on good state, and thus, because of the third observation, must be better off conditional on bad state. Both of these make the principal worse off, because of the first and the second observations. Thus, the principal prefers to induce effort when feasible.

In summary, the following are the crucial aspects of the continuation equilibrium at $t = 1 - 2\Delta$: (i) the agent works when feasible, (ii) the agent receives his autarky payoff, (iii) the posterior belief is either zero, or is above $\underline{\mu}_{1-\Delta}^\Delta(E_{1-\Delta})$, so that effort remains feasible. (See Lemma 7 for the formal arguments.) It turns out that these three properties hold at any history, as formally stated in Lemma 1.

LEMMA 1. *In any equilibrium, at any history, if effort is feasible so that $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$,*

deliver the agent more than his autarky payoff. Note that for such multiplicity to be ruled out, it is sufficient that $\varepsilon > 0$, and there is no positive lower bound on it.

[C1] the agent's continuation payoff is his autarky payoff, that is, $U^*(t, E_t, \mu_t) = U^a(t, E_t, \mu_t)$

[C2] the agent works, that is, $e(t) = 1$

[C3] on the equilibrium path, effort remains feasible unless the bad state is revealed, that is, at any $t' > t$, $\mu_{t'} \geq \underline{\mu}_{t'}^\Delta(E_t + t' - t)$ or $\mu_{t'} = 0$.

The proof of Lemma 1 is described in Section 3.2.2. Before going into the proof, Section 3.2.1 characterizes equilibrium outcomes using Lemma 1.

3.2.1 Equilibrium outcomes Thanks to [C1], at any history, the agent agrees to work if the principal's signal offer delivers him at least his autarky payoff. This allows to characterize the principal's equilibrium strategies and payoff as the solution of an optimization problem, where the agent's incentive constraint can be formulated independently of future play. Namely, at any history with associated (t, E_t, μ_t) , the principal's continuation payoff $\Pi^*(t, E_t, \mu_t)$ must satisfy the following Bellman equation:¹⁴

$$\Pi^*(t, E_t, \mu_t) = \max_{F \in \mathcal{F}(\mu_t)} \mathbb{E}_F[\Pi^*(t + \Delta, E_t + \Delta, \mu_{t+\Delta})] \quad (\text{PP})$$

subject to

$$\mathbb{E}_F[U^a(t + \Delta, E_t + \Delta, \mu_{t+\Delta})] - \Delta(c + u) = U^a(t, E_t, \mu_t) \quad (9)$$

$$\mu_{t+\Delta} \in \{0\} \cup [\underline{\mu}_{t+\Delta}^\Delta(E_t + \Delta), 1] \quad (10)$$

This “principal's problem” imposes the assumed conditions [C1], [C2], and [C3] as constraints to the principal's payoff maximization problem: Constraint (9) incorporates conditions [C1] and [C2] while (10) corresponds to [C3].

An auxiliary problem (PP) is a finite horizon recursive problem each step of which is a constrained information design problem. Solving it using backwards induction, though conceptually straightforward, quickly becomes intractable. Instead, I use Lemma 1 to define and solve a more tractable auxiliary problem and later show how the solution of the latter can be used to characterize the solution of (PP).

To set up the auxiliary problem, let T_B be the random time at which the belief reaches 0 during the production stage, with the convention that $T_B = 1$ if the bad state is never revealed. By Lemma 1, all equilibrium payoffs can be expressed as a function of this random variable. Indeed, by [C3] starting from any belief where inducing effort is feasible, unless the state is revealed to be low, effort remains feasible, and by [C2], effort is induced as long as it is feasible. Thus, in any continuation equilibrium following histories covered by Lemma 1, the agent works when $t < T_B$. This in particular implies that conditional on good state agent necessarily works throughout, as the belief can never reach 0. Further, since at T_B the belief reaches 0, the agent chooses his autarky action thereafter: he quits if $T_B < \underline{t}(\alpha)$ and otherwise, shirks throughout.

¹⁴It is also necessary to specify what the payoffs will be at continuation histories when effort provision is not feasible. These terminal conditions are specified in the Appendix in equation (23).

Consider the problem of choosing a probability distribution over T_B to maximize the principal's expected payoff, subject to delivering the agent at least his autarky payoff. Since the payoffs conditional on good state are fixed across equilibria, this problem can be formulated to maximize the principal's payoff conditional on the bad state. Namely,

$$\max_H \mathbb{E}_H [\tilde{\pi}_B(T_B|t, E_t, \mu_t)] \quad (\text{PP-aux})$$

subject to

$$\mu_t U_G^*(t, E_t, \mu_t) + (1 - \mu_t) \mathbb{E}_H [\tilde{u}_B(T_B|t, E_t, \mu_t)] \geq U^a(t, E_t, \mu_t) \quad (11)$$

$$\text{supp } H \subset \{t + \Delta, \dots, 1\}, \quad (12)$$

where $\tilde{u}_B(T_B|t, E_t, \mu_t)$ and $\tilde{\pi}_B(T_B|t, E_t, \mu_t)$ are the payoffs in the bad state *as functions of* T_B starting from some (t, E_t, μ_t) , and are given by

$$\begin{aligned} \tilde{u}_B(T_B|t, E_t, \mu_t) &\equiv \begin{cases} \alpha Y_L - (1 - T_B)u - (T_B - t)(c + u) & \text{if } T_B \geq \underline{t}(\alpha) \\ -(T_B - t)(c + u) & \text{if } T_B < \underline{t}(\alpha) \end{cases} \\ \tilde{\pi}_B(T_B|t, E_t, \mu_t) &\equiv \begin{cases} (1 - \alpha)Y_L + \varepsilon(T_B - t + E_t) & \text{if } T_B \geq \underline{t}(\alpha) \\ \varepsilon(T_B - t + E_t) & \text{if } T_B < \underline{t}(\alpha) \end{cases} \end{aligned}$$

The principal's payoff is strictly increasing in T_B , while the agent's is strictly decreasing, since later revelation means more effort by the agent in the bad state.

Naturally, the principal would like to delay the revelation of the bad state as much as possible in order to extract more effort from the agent. When the expected value of T_B (hence the expected amount of effort in bad state) is sufficiently large so that (11) binds, increasing it further can only be achieved at the cost of increasing the probability of quitting (i.e., of revelation before $\underline{t}(\alpha)$). But note that both increased effort and increased probability of quitting reduce the surplus in the bad state. Since the agent's payoff in the bad state is fixed by (11), the principal would never increase the expected effort at the cost of increased probability of quitting. This has the following implications: when the initial belief is low so that a degenerate distribution with support $\{\underline{t}(\alpha)\}$ cannot satisfy (11) (i.e., when $\mu_t < \mu_t^*(E_t)$ so that work until vesting is not feasible), then the optimal probability distribution reveals the bad state either immediately or exactly at $\underline{t}(\alpha)$ in a way that (11) binds. Thus, the worker either immediately quits or works until vesting. If the belief is high ($\mu_t \geq \mu_t^*(E_t)$), the principal delays the revelation until after vesting with probability 1, and the expected time to revelation is chosen to bind (11). The following lemma formally states the solution of (PP-aux).

LEMMA 2. *Any H that solves the auxiliary problem satisfies the following:*

- If $\mu_t < \mu_t^*(E_t)$, then

$$T_B = \begin{cases} t + \Delta & \text{with probability } \gamma_\Delta^*(t, E_t, \mu_t) \\ \underline{t}(\alpha) & \text{with probability } 1 - \gamma_\Delta^*(t, E_t, \mu_t) \end{cases},$$

where $\gamma_{\Delta}^*(t, E_t, \mu_t)$ is defined by

$$U^a(t, E_t, \mu_t) = \mu_t U_G^*(t, E_t, \mu_t) + (1 - \mu_t)(1 - \gamma_{\Delta}^*(t, E_t, \mu_t))\tilde{u}_B(\underline{t}(\alpha)|t + \Delta, E_t + \Delta, \mu_t^*(E_t)).$$

- If $\mu_t \geq \mu_t^*(E_t)$, then $T_B \in [\underline{t}(\alpha), 1]$ with probability 1, and $\mathbb{E}_H[T_B] = T_B^*(t, E_t, \mu_t)$, where $T_B^*(t, E_t, \mu_t)$ is defined by

$$U^a(t, E_t, \mu_t) = \mu_t U_G^*(t, E_t, \mu_t) + (1 - \mu_t)\tilde{u}_B(T_B^*(t, E_t, \mu_t)|t).$$

Further, (11) binds.

The first item of Lemma 2 considers the case where H with support $\{\underline{t}(\alpha)\}$ cannot deliver the agent's autarky payoff, or equivalently $\mu_t < \mu_t^*(E_t)$. It states that in this case, the solution of (PP-aux) immediately reveals the bad state with probability $\gamma_{\Delta}^*(t, E_t, \mu_t)$, which holds the agent to his autarky payoff. The second item states the two characterizing properties of H in the opposite case: that its support does not intersect the interval $[t + \Delta, \underline{t}(\alpha))$, so that agent never quits, and that the expected value of T_B is $T_B^*(t, E_t, \mu_t)$, which holds the agent to his autarky payoff.

Equivalence of (PP) and (PP-aux) Relative to (PP), which also naturally induces a probability distribution over T_B , (PP-aux) relaxes the requirement that this probability distribution be implemented via a sequence of signals, which are optimal at each history. Thus, the value of (PP-aux) is no less than $\Pi_B^*(t, E_t, \mu_t)$. Lemma 8 in the Appendix shows that the two values are equal. The proof uses the solution of (PP-aux) to generate an upper bound for the value of (PP), and verifies that it can be attained by constructing optimal policies.¹⁵ Here, with the help of Figure 2, I describe these signals, which solve (PP). The formal construction follows from iterative application of the signals constructed in the proof of Lemma 8.

The left panel of Figure 2 corresponds to a history where $\mu_t < \mu_t^*(E_t)$ while the right panel corresponds to a history where $\mu_t > \mu_t^*(E_t)$. In both cases, $\mu_{t+\Delta}^*(E_{t+\Delta})$ is in the support of the signal at t , and unless this is the realized posterior, no further information is released. If the realized belief is $\mu_{t+\Delta}^*(E_{t+\Delta})$, in the subsequent periods, the principal reveals the high state at a rate just sufficient to keep the agent indifferent between quitting now and following the “work until vesting” path. Feasibility of such a sequence of signals follows by Lemma 5 in the Appendix, which shows that the “work until vesting” path is decreasing. Finally, note that conditional on jumping onto this path, in either case, $T_B = \underline{t}(\alpha)$ with probability 1.

When $\mu_t < \mu_t^*(E_t)$ (left panel), the support of the principal's signal at t includes 0 and 1 in addition to $\mu_{t+\Delta}^*(E_{t+\Delta})$. The probability of each posterior is pinned down by the two requirements: (i) Bayes plausibility and (ii) delivering the agent his autarky payoff. Since conditional on bad state the posterior cannot be 1, this signal either reveals T_B immediately, or at $\underline{t}(\alpha)$, as in the solution of (PP-aux). By requirement (ii), the probability of

¹⁵This equivalence does not imply that any solution H of the auxiliary problem can be generated by a feasible policy for the principal's problem; rather, that there exists at least one such H .

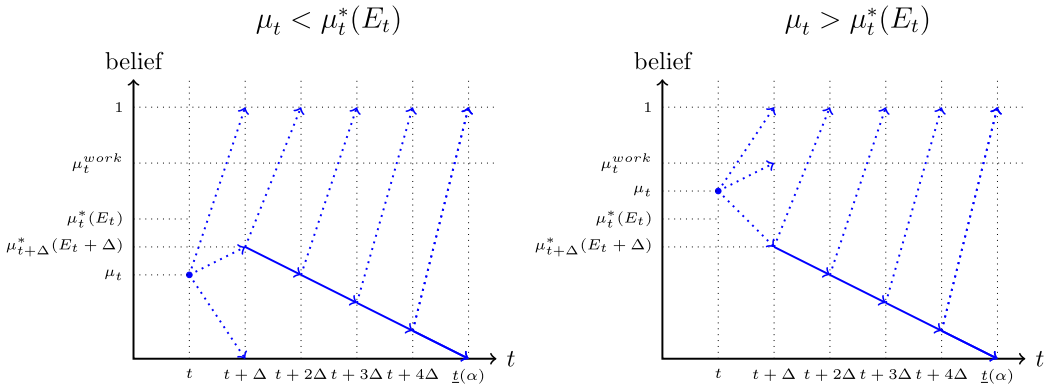


FIGURE 2. Sample equilibrium belief paths for when quit cannot be avoided (left panel) and when it can (right panel).

immediate revelation is $\gamma_{\Delta}^*(t, E_t, \mu_t)$ as described in Lemma 2, so that the agent receives exactly his autarky payoff.

When $\mu_t > \mu_t^*(E_t)$ (right panel), the support of the principal's signal at t includes μ_t^{work} (i.e., the cutoff belief above which the agent's autarky action is to work) and 1 in addition to $\mu_{t+\Delta}^*(E_t)$. Once again, the probabilities of each posterior are pinned down by requirements (i) and (ii) above. In this case, conditional on bad state, the posterior has two possible realizations: $\mu_{t+\Delta}^*(E_t + \Delta)$ or $\mu_{t+\Delta}^{\text{work}}(E_t + \Delta)$. The principal reveals no further information if the realization is $\mu_{t+\Delta}^{\text{work}}$, since the agent's autarky action is already to work. Thus, the induced probability distribution over T_B has support $\{\underline{t}(\alpha), 1\}$. By requirement (ii), its expectation is $T_B^*(t, E_t, \mu_t)$ as defined in Lemma 2, exactly delivering the agent's autarky payoff.

3.2.2 Overview of the proof of Lemma 1 The proof of Lemma 1 is by induction. The first step of induction shows that at $t = 1 - 2\Delta$, [C1], [C2], [C3] necessarily hold. This step is discussed in detail at the beginning of Section 3.2 and formalized in Lemma 7 in the Appendix. Next, the induction hypothesis fixes an arbitrary $\tilde{t} \leq 1 - 2\Delta$ and assumes that [C1], [C2], and [C3] hold for all $t' \geq \tilde{t}$. Naturally, under the induction hypothesis, for each posterior belief $\mu_{\tilde{t}}$, the principal's payoff is given by $\Pi^*(\tilde{t}, E_{\tilde{t}}, \mu_{\tilde{t}})$, which solves (PP) while the agent's payoff is his autarky payoff.

Figure 3 represents the principal's payoff as a function of the posterior $\mu_{\tilde{t}}$ and is a dynamic analogue of the one-period payoff illustrated in Figure 1. The right panel corresponds to the case where $\tilde{t} > \underline{t}(\alpha)$, so that the agent's payoff exhibits only one convex kink where his autarky action switches to work from shirk. The left panel corresponds to the case where $\tilde{t} < \underline{t}(\alpha)$ and the agent's payoff here exhibits an additional convex kink when his autarky action switches from quit to shirk. In each case, the principal's payoff exhibits a concave kink corresponding to each of the kinks in the agents payoff. In the left panel, the principal's payoff exhibits an additional kink when the “work until vesting” becomes feasible. Finally, in each panel, the principal's payoff exhibits a right-continuous upward jump at the belief $\underline{\mu}_{\tilde{t}}^{\Delta}(E_{\tilde{t}})$ where effort becomes feasible.

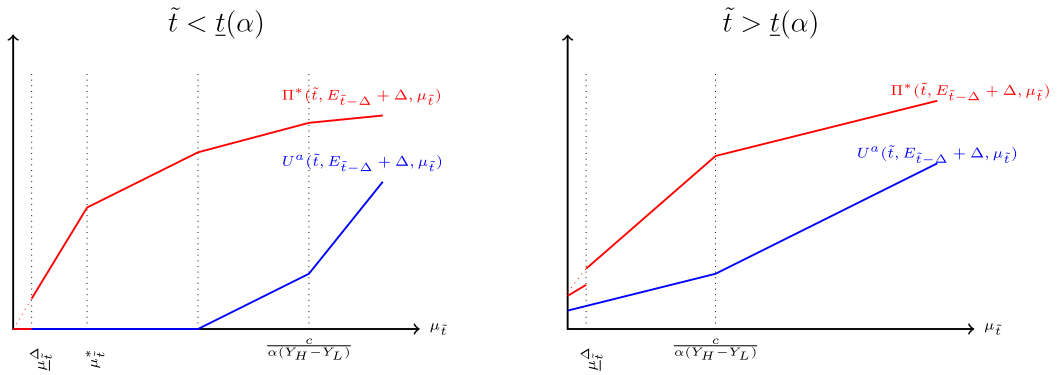


FIGURE 3. Equilibrium continuation payoffs of the principal and the agent for when the agent is not vested (left panel) and for when he is vested (right panel).

The principal's decision problem at $\tilde{t} - \Delta$ is subject to considerations that are analogous to those at $t = 1 - 2\Delta$: if the principal induces effort, due to the upward jump in her payoff at $\mu_{\tilde{t}}^{\Delta}(E_{\tilde{t}})$, she would never induce a posterior between 0 and $\mu_{\tilde{t}}^{\Delta}(E_{\tilde{t}})$, so that effort remains feasible ([C3]). That the agent receives his autarky payoff ([C1]) follows from the concavity of the principal's payoff and the convexity of the agent's. Then it remains to argue that the principal always wants to induce effort ([C2]). The arguments are again similar to those for $t = 1 - 2\Delta$: forgoing effort in the current period reduces the agent's payoff conditional on bad state, and thus must be compensated for with increased payoff in the bad state. Both of these adjustments reduce the principal's payoff.

Remark (the role of Assumption (5)) Naturally, any informative signal must make the agent more pessimistic with some probability. Before the agent is vested inducing effort may require that with positive probability, the posterior is less than $\mu_{t+\Delta}^*(E_{t+\Delta})$, leading to positive probability of project abandonment. If at the same time the agent's current autarky action is to shirk, the principal may prefer to forego effort to avoid taking this chance. This possibility is ruled out precisely by Assumption (5) via the second part of Lemma 5: whenever the agent's autarky action is to shirk, it is feasible to follow the “work until vesting” path, thus inducing effort does not require posteriors less than $\mu_{t+\Delta}^*(E_{t+\Delta})$.

3.3 Equilibrium outcomes as $\Delta \rightarrow 0$

In preparation for the next section, which characterizes the principal's optimal choice of α , the next proposition summarizes the equilibrium outcomes at the initial history $(t, E_t, \mu_t) = (0, 0, \mu_0)$ as $\Delta \rightarrow 0$. The proposition also characterizes equilibrium outcomes for shares α outside of the range (6). This is based on the characterization in Lemma 2 and obtained by considering the limits of outcomes stated in that lemma as $\Delta \rightarrow 0$. Define $\mu^*(\alpha)$ to be the lowest belief at which “work until vesting” is feasible at the initial history, that is, $\mu^*(\alpha) \equiv \mu_0^*(0)$.

PROPOSITION 1. *Fix the initial belief μ_0 and the agent's share α . Then in all limit equilibria,*

- *If α satisfies (6),*
 - *Conditional on good state ($\theta = 1$) the agent never quits, works throughout.*
 - *Conditional on bad state ($\theta = 0$)*
 - * *if $\mu_0 \geq \mu^*(\alpha)$, the agent never quits and works an expected duration $T_B^*(0, 0, \mu_0) \geq \underline{t}(\alpha)$ where $T_B^*(t, E_t, \mu_t)$ is characterized in Lemma 2.*
 - * *if $\mu_0 < \mu^*(\alpha)$, with probability $\gamma^*(\mu_0, \alpha)$ the agent quits at $t = 0$. With the remaining probability, the agent never quits and works until $\underline{t}(\alpha)$, where*

$$\gamma^*(\mu_0, \alpha) = \lim_{\Delta \rightarrow 0} \gamma_{\Delta}^*(0, 0, \mu_0) = 1 - \frac{\mu^*(\alpha)}{1 - \mu^*(\alpha)} \frac{1 - \mu_0}{\mu_0},$$
where $\gamma_{\Delta}^(t, E_t, \mu_t)$ is characterized in Lemma 2.*
- *If α does not satisfy (6), then regardless of the true state, the agent quits immediately and never works.*

4. OPTIMAL SHARING CONTRACT

This section characterizes the principal's optimal choice of α given the initial belief, anticipating that the outcome of the ensuing production game will be as described in Proposition 1. To indicate the dependence of ex ante payoffs on the contract α as well as on the initial belief μ_0 , define $\Pi^-(\alpha, \mu_0)$ and $U^-(\alpha, \mu_0)$ to be the equilibrium payoffs of the principal and the agent, respectively, when α is the agent's share. Then the principal's optimal choice of α solves

$$\max_{\alpha} \Pi^-(\alpha, \mu_0) \quad \text{subject to} \quad U^-(\alpha, \mu_0) \geq 0. \quad (13)$$

The main result of this paper is to demonstrate that the solution of (13) features one of two possible choices, leading to either an “immediate learning” outcome or a “gradual learning” outcome, which I describe below.

Immediate learning This outcome features a low monetary reward for the agent. In particular, the share $\underline{\alpha}$ of the agent satisfies

$$\underline{\alpha} Y_H = c + u. \quad (14)$$

In the best case scenario, that is, conditional on good state, this share leaves the agent a payoff of 0. Thus, to preclude quitting, the principal is compelled to offer a fully revealing signal immediately at $t = 0$. Consequently, conditional on good state, the agent is incentivized to exert effort throughout the production stage while conditional on bad state, he quits immediately.

Gradual learning In this outcome, the agent's share is larger, the project is never abandoned, and the principal reveals information gradually. Define $\bar{\alpha}(\mu_0)$ by

$$\mu_0[\bar{\alpha}(\mu_0)Y_H - u - c] + (1 - \mu_0)[\bar{\alpha}(\mu_0)Y_L - u - \underline{t}(\bar{\alpha}(\mu_0))c] = 0. \quad (15)$$

Thus, $\mu^*(\bar{\alpha}(\mu_0)) = \mu_0$. By the analysis of Section 3.2, this implies that the principal slowly releases information so that the state is fully revealed exactly when the agent is vested. The next lemma establishes useful properties of $\bar{\alpha}(\mu_0)$.

LEMMA 3. *The share $\bar{\alpha}(\mu_0)$ is decreasing in μ_0 and satisfies $\frac{u}{Y_L} \geq \bar{\alpha}(\mu_0) \geq \frac{c+u}{Y_H}$.*

Since $\underline{\alpha} < \bar{\alpha}(\mu_0) < u/Y_L$, in both the gradual and the immediate learning outcomes, the agent's autarky action at $t = 0$ is to quit, and thus his ex ante payoff is 0. Conditional on good state, the outcome under gradual learning and immediate learning are identical: the agent works throughout the production stage. Conditional on bad state, gradual learning eliminates inefficient project abandonment but leads to inefficient effort. Consequently, the principal's preference between these two outcomes depends on the relative severity of the two types of distortion in the bad state. Theorem 1 formally states that these two outcomes are the only possibilities and specifies the exact conditions under which each may occur.

THEOREM 1. *The unique equilibrium outcome is “immediate learning” if*

$$Y_L - u < \underline{t}(\bar{\alpha}(\mu_0))(c - \varepsilon), \quad (16)$$

and is “gradual learning” if the opposite strict inequality holds.

The left-hand side of (16) is the loss of surplus, conditional on the bad state, resulting from immediate project abandonment. The right-hand side is the corresponding loss due to the agent working when $t \in [0, \underline{t}(\bar{\alpha}(\mu_0))]$. Thus, the comparison of the two outcomes is immediate. That immediate and gradual learning are equilibrium outcomes is established in Section 3.2.

In order to rule out other possibilities as solutions of (13), first recall that all equilibria lead to the same outcome conditional on high state. Thus, it suffices to consider the impact of varying α on the principal's payoff in the bad state. First, by Proposition 1 whenever $\alpha < \underline{\alpha}$, the agent immediately quits. Thus, the principal's payoff is 0, implying that such α cannot be optimal.

When $\alpha = \bar{\alpha}(\mu_0)$, the agent is just optimistic enough that he is willing to exert effort until he is vested before learning the true state. Since $\bar{\alpha}(\mu_0)$ achieves retention with probability 1, increasing the agent's share any further only affects the amount of effort the agent exerts. In particular, it allows and compels the principal to extract more inefficient effort in the bad state. Further, it (weakly) increases the agent's equilibrium payoff, and thus unambiguously hurts the principal. Therefore, the principal's offer must be in the interval $[\underline{\alpha}, \bar{\alpha}(\mu_0)]$.

Consider $\alpha \in (\underline{\alpha}, \bar{\alpha}(\mu_0))$. For such α , $\mu_0 < \mu_0^*(0)$, and thus at the initial history, the agent's autarky action is to quit, implying that the principal captures all surplus. The

analysis of the previous section reveals that for this range of α , the equilibrium features a positive probability of project abandonment at time 0, and conditional on continuing, agent working until he is vested. Thus, the surplus conditional on bad state takes the following form:

$$\underbrace{\left[\frac{\mu_0}{1 - \mu_0} \frac{\alpha Y_H - c - u}{u + \underline{t}(\alpha)c - \alpha Y_L} \right]}_{\text{probability of retention: } 1 - \gamma^*(\mu_0, \alpha)} \times \underbrace{\left[Y_L - u - \underline{t}(\alpha)(c - \varepsilon) \right]}_{\text{bad state surplus conditional on retention}}.$$

Since the date $\underline{t}(\alpha)$ at which the agent gets vested *decreases* in α , this surplus is negative whenever (16) holds, implying that $\underline{\alpha}$ delivers a higher payoff than any α in this range. When the opposite inequality holds, the surplus is positive. Further, both the probability of continuing and the surplus upon continuing is increasing in α . Therefore, $\bar{\alpha}(\mu_0)$ performs better than any α in this range.

The principal's preference between gradual versus immediate learning depends on various features of the environment, which leads to some immediate comparative statics. For instance, smaller outside option u as well as higher Y_L and lower μ_0 favors gradual learning. Formally, we have the following.

PROPOSITION 2 (Comparative Statics). *Consider a set of parameters for which the unique equilibrium outcome is gradual learning. This continues to be the case if u or μ_0 decreases, or if Y_L increases.*

The comparative statics with respect to u and Y_L are intuitive: if u is small or Y_L is large, the efficiency loss due to project abandonment is more significant *and* the agent gets vested earlier so that efficiency loss in the gradual learning outcome is smaller. The impact of μ_0 is via the fact that $t^*(\mu_0)$ is increasing in μ_0 . This is because with larger μ_0 the principal is able to recruit the agent with a smaller share α , which implies that the agent gets vested later.

5. DISCUSSION

The baseline model makes several assumptions that render the model tractable or allows focus on the most interesting cases. This section discusses the roles of some of these assumptions.

5.1 Motivating effort and immediate learning

In this section, I show that Assumption (5) is a key driver of the gradual learning outcome because in its absence, gradual learning cannot occur in equilibrium.

PROPOSITION 3. *Assume that $(Y_H - Y_L)/Y_L < c/u$. There exists a cutoff $\tilde{\mu}$ such that if $\mu_0 > \tilde{\mu}$, learning is immediate, total surplus is maximized and the agent receives a positive payoff. If $\mu_0 < \tilde{\mu}$, the unique outcome involves the agent shirking throughout, no learning takes place, and the agent's payoff is 0.*

To understand Proposition 3, note that incentivizing effort requires $\alpha(Y_H - Y_L) \geq c$, which by $(Y_H - Y_L)/Y_L < c/u$, also implies $\alpha Y_L \geq u$. Thus, retention is automatically achieved if effort is incentivized. Thus, the principal chooses between offering $c/(Y_H - Y_L)$ and u/Y_L .¹⁶ With the former offer, she immediately reveals full information, which achieves surplus maximization but leaves the agent rents equal to $cY_L/(Y_H - Y_L) - u$. With the latter offer, regardless of the state and belief, the agent's optimal action is to shirk. Thus, the timing of learning does not affect the outcome, which is efficient conditional on bad state but inefficient conditional on good state. When μ_0 is high, principal prefers to induce effort since the expected returns to effort is large, even when this means she must leave rents to the agent.

5.2 Contractual constraints and full surplus extraction

If the principal can extract full surplus, she would not delay information release. Surplus maximization requires immediate learning. Thus, to extract full surplus, the principal must give the agent ex post incentives to stay and shirk in the bad state, and stay and work in the good state, without leaving him any rents. With a sharing contract, this is not possible: retention in the bad state requires $\alpha \geq u/Y_L$. This leaves the agent rents at least equal to $u(Y_H - Y_L)/Y_L - c > 0$ in the good state.

Two alternative modifications of the model would allow the principal to extract full surplus and eliminate his need to slow down learning. First, if the principal can offer a noncontingent payment equal to $u - Y_L c/(Y_H - Y_L) > 0$ in addition to a share $c/(Y_H - Y_L)$, she could extract full surplus. A natural condition that precludes this is cash constraints typical of start-ups.¹⁷ Second, if the principal can costlessly adjust the agent's share offer based on the information learned, she can trivially extract full surplus, by first offering $(c + u)/Y_H$, and then adjusting it to u/Y_L if the state is revealed to be low. The rigidity of contracts as assumed in the main model is justified by costs of renegotiation, perhaps due to the need to receive approval from funders or upper management or due to the need to pause while renegotiating.

5.3 Costly information collection without the agent's effort

The main model assumes that for the signal designed by the principal to generate a realization, it is necessary that the agent exerts effort in the current period. This assumption captures the idea that in order to generate new information, new components of the project on which tests can be run must be completed. However, perhaps unrealistically, it rules out the principal's ability to generate new information on the previously

¹⁶The proof of Proposition 3 rules out other possible choices of α .

¹⁷It is worth noting, however, that such contracts may be possible to implement via more sophisticated stocks such as put options. For instance, the principal may offer a share $\underline{\alpha} = (c + u)/Y_H$ with a promise to buy back the agent's shares at a total price equal to u after product launch. Such contracts may not be available if the returns conditional on θ still have a large variance and are expected to be realized far into the future so that with positive probability the principal lacks the funds to buy back the shares even after launch.

completed components of production. Reassuringly, the model can be modified to allow for information collection without the agent's effort. The main result applies to the case when the period length is small, and states that if generating information without the agent's effort entails even a slight cost, the principal chooses never to do that. The formal result is stated in Proposition 4. Its proof is deferred to Appendix D.2.

PROPOSITION 4. *Suppose that, at the beginning of each period, by paying a cost κ , the principal can flexibly design an informative signal, which generates a realization at the end of the period, without the agent's effort. There exists $\bar{\kappa}(\Delta)$ with $\lim_{\Delta \rightarrow 0} \bar{\kappa}(\Delta) = 0$ such that for any period length Δ , whenever $\kappa \geq \bar{\kappa}(\Delta)$, no equilibrium features information gathering without the agent's current effort.*

To understand this result, note that at the very end of the interaction, the principal's gains from such information generation is limited. Further, if in the future no such information gathering is expected to take place, then the gains from one-time generation of information without effort is once again limited. The latter follows because (i) by the equilibrium characterization in Section 3.2, the principal prefers to induce effort rather than allow the agent shirk (or quit) by not generating information; and (ii) as $\Delta \rightarrow 0$, the principal's continuation payoff becomes concave, so that he prefers not to generate information.

APPENDIX A: PRELIMINARIES

In this section, I formalize the concepts introduced in Section 3.1 and present supporting analysis.

Autarky actions and payoffs The agent's autarky actions introduced in Section 3.1 are formally given by

$$U^{\text{work}}(t, E_t, \mu_t) = \mu_t \alpha (Y_H - Y_L) E_t + \alpha Y_L - (1 - t)u + [\mu_t \alpha (Y_H - Y_L) - c](1 - t)$$

$$U^{\text{shirk}}(t, E_t, \mu_t) = \mu_t \alpha (Y_H - Y_L) E_t + \alpha Y_L - (1 - t)u$$

The next lemma characterizes histories where each action is the agent's autarky action and also shows that unless the autarky action is to quit, it remains the same over time in the absence of new information. This observation justifies the calculation of payoffs $U^x(t, E_t, \mu_t)$ along paths where the agent's action is constant.

LEMMA 4. *There exists cutoffs $\mu_t^{\text{work}}(E_t) \geq \mu_t^{\text{shirk}}(E_t)$ such that the agent's autarky action is to work if $\mu_t \geq \mu_t^{\text{work}}(E_t)$, shirk if $\mu_t \in [\mu_t^{\text{shirk}}(E_t), \mu_t^{\text{work}}(E_t))$, and quit otherwise. Further:*

- (i) *If agent's autarky action is either work or shirk (i.e., not quit), it remains the same over time unless new information is revealed.*
- (ii) *If $t \geq \underline{t}(\alpha)$, then $\mu_t^{\text{shirk}}(E_t) \equiv 0$.*

PROOF OF LEMMA 4. Define the cutoff beliefs $\mu_t^{w/s}(E_t)$, $\mu_t^{s/q}(E_t)$, $\mu_t^{w/q}(E_t)$ as the beliefs that leave the agent indifferent between a pair of his autarky actions work (w), shirk (s), or quit (q). Thus,

$$U^{\text{work}}(t, E_t, \mu_t^{w/s}(E_t)) = U^{\text{shirk}}(t, E_t, \mu_t^{w/s}(E_t))$$

$$U^{\text{shirk}}(t, E_t, \mu_t^{s/q}(E_t)) = 0$$

$$U^{\text{work}}(t, E_t, \mu_t^{w/q}(E_t)) = 0$$

Since $U^{\text{shirk}}(t, E_t, \mu_t)$ is increasing in μ_t , the agent's payoff from shirking throughout is nonnegative if and only if $\mu_t \geq \mu_t^{s/q}(E_t)$. Further, U^{work} and $U^{\text{work}} - U^{\text{shirk}}$ are also increasing in μ_t . Thus, work is the autarky action of the agent if and only if $\mu_t \geq \max\{\mu_t^{w/q}, \mu_t^{w/s}\}$ and quit is the autarky action of the agent if and only if $\mu_t \leq \min\{\mu_t^{w/q}, \mu_t^{s/q}\}$. Accordingly, let

$$\mu_t^{\text{work}}(E_t) \equiv \max\{\mu_t^{w/q}(E_t), \mu_t^{w/s}(E_t)\} \quad \text{and} \quad \mu_t^{\text{quit}}(E_t) \equiv \min\{\mu_t^{w/q}(E_t), \mu_t^{s/q}(E_t)\}. \quad (17)$$

That $\mu_t^{\text{shirk}}(E_t) \leq \mu_t^{\text{work}}(E_t)$ follows by definition. Next, consider the two items:

- (i) Fix t, E_t, μ_t . If $\mu_t(Y_H - Y_L) - c \leq 0$, the claim follows because $U^{\text{work}}(t, E_t, \mu_t)$, $U^{\text{shirk}}(t, E_t, \mu_t)$, and $U^{\text{work}}(t, E_t, \mu_t) - U^{\text{shirk}}(t, E_t, \mu_t)$ are nondecreasing in t when E_t and μ_t remain constant. If $\mu_t(Y_H - Y_L) - c > 0$, the claim follows because $U^{\text{shirk}}(t, E_t, \mu_t)$ is increasing in t when E_t and μ_t remain constant and $U^{\text{work}}(t, E_t, \mu_t) - U^{\text{shirk}}(t, E_t, \mu_t) > 0$.

- (ii) Follows because for $t > \underline{t}(\alpha)$, $U^{\text{shirk}}(t, E_t, 0) > 0$. □

“Work until vesting” outcome Here, I replicate the definition of the cutoff belief $\mu_t^*(E_t)$ defined in Section 3.1:

$$\begin{aligned} U^a(t, E_t, \mu_t^*(E_t)) &= \alpha Y_L - u(1-t) - c(1-\underline{t}(\alpha)) \\ &\quad + \mu_t^*(E_t)[\alpha(Y_H - Y_L) - c](1-\underline{t}(\alpha)). \end{aligned} \quad (18)$$

The next lemma records two useful properties of this cutoff.

LEMMA 5. For any $t < \underline{t}(\alpha)$ and $E_t \leq t$, necessarily $\mu_t^*(E_t) < \mu_t^{\text{quit}}(E_t)$. Further, for $t > t'$, $\mu_t^*(t - t' + E_t)$ is decreasing in t .

PROOF OF LEMMA 5. I first show that $\mu_t^*(E_t) < \mu_t^{\text{quit}}(E_t)$ when $t < \underline{t}(\alpha)$. The claim trivially holds when $\mu_t^{\text{quit}}(E_t) = \mu_t^{w/s}(E_t)$, because by definition when $\mu_t = \mu_t^{w/s}(E_t)$, the agent is willing to exert effort without any further information. When $t < \underline{t}(\alpha)$ and $\mu_t^{\text{quit}}(E_t) = \mu_t^{s/q}(E_t)$, item 4 is satisfied if and only if for any μ and (t, E_t) , that $\mu[\alpha(Y_H - Y_L)E_t - (1-t)u] + (1-\mu)[\alpha Y_L - (1-t)u] \geq 0$ implies $\mu[\alpha(Y_H - Y_L)(1-t + E_t) - (1-t)(c+u)] + (1-\mu)[\alpha Y_L - (1-t)u - (\underline{t}(\alpha) - t)c] \geq 0$. Thus, item 4 is equivalent to

$$\frac{u(1-t) - \alpha Y_L}{\alpha E_t(Y_H - Y_L)} \geq \frac{u(1-t) - \alpha Y_L + (\underline{t}(\alpha) - t)c}{\alpha E_t(Y_H - Y_L) + (1-t)\alpha(Y_H - Y_L) - (1-\underline{t}(\alpha))c}.$$

This is equivalent to

$$\frac{u(1-t) - \alpha Y_L}{\alpha E_t(Y_H - Y_L)} \geq \frac{(\underline{t}(\alpha) - t)c}{(1-t)\alpha(Y_H - Y_L) - (1 - \underline{t}(\alpha))c}.$$

Using $\underline{t}(\alpha) = 1 - \alpha Y_L/u$ and $t < \underline{t}(w)$, the latter inequality is reexpressed as

$$\frac{c}{(1-t)\alpha(Y_H - Y_L) - (1 - \underline{t}(\alpha))c} \leq \frac{u}{\alpha E_t(Y_H - Y_L)}.$$

The left-hand side is increasing in t , while the right-hand side is decreasing in E_t . Thus, the inequality holds for all $t < \underline{t}(w)$ and $E_t \leq t$ if and only if

$$\frac{c}{(1 - \underline{t}(\alpha))(\alpha(Y_H - Y_L) - c)} \leq \frac{u}{\alpha \underline{t}(\alpha)(Y_H - Y_L)}.$$

Once again, using $\underline{t}(\alpha) = 1 - \alpha Y_L/u$ this is reexpressed as

$$\frac{c}{u} \leq \frac{1 - \underline{t}(\alpha)}{\underline{t}(\alpha)} \frac{\alpha(Y_H - Y_L) - c}{\alpha(Y_H - Y_L)} = \frac{Y_L}{u - \alpha Y_L} \frac{\alpha(Y_H - Y_L) - c}{Y_H - Y_L},$$

which is equivalent to

$$\frac{c}{u} \frac{Y_H - Y_L}{Y_L} \leq \frac{\alpha(Y_H - Y_L) - c}{u - \alpha Y_L}.$$

By Assumption (5), the left-hand side is less than 1. Thus, a sufficient condition is $\alpha Y_H \geq u + c$, which is satisfied for any $\alpha \geq \underline{\alpha}$, establishing the claim.

Next, that $\mu_t^*(t - t' + Et')$ is decreasing follows by noting that by the previous claim, $U^a(t, t - t' + E_{t'}, \mu_t^*(t - t' + E_{t'})) = 0$. Thus, (18) can be rearranged to solve for $\mu_t^*(t - t' + Et')$ and that it is decreasing follows by inspection. $\square$

Feasibility of effort The belief cutoff $\underline{\mu}_t^\Delta(E_t)$ is defined by

$$\begin{aligned} & \underline{\mu}_t^\Delta(E_t)U^a(t + \Delta, E_t + \Delta, 1) + (1 - \underline{\mu}_t^\Delta(E_t))U^a(t + \Delta, E_t + \Delta, 0) - \Delta(c + u) \\ & = U^a(t, E_t, \underline{\mu}_t^\Delta). \end{aligned} \tag{19}$$

The left-hand side of the equality is the agent's expected payoff if he exerts effort at time t and in return, perfectly learns the true state. Thus, only when $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$ effort in return for full information delivers the agent no less than his autarky payoff.

The next lemma formally establishes that as $\Delta \rightarrow 0$ so that the agent's cost of choosing to stay and work for the next period vanishes, this cutoff approaches 0.

LEMMA 6. *For any E_t and t , $\underline{\mu}_t^\Delta(E_t) \rightarrow 0$ as $\Delta \rightarrow 0$.*

PROOF OF LEMMA 6. Follows by inspecting (19) and noting that $U^a(t, E_t, \mu_t)$ is a (weakly) convex function of μ_t . $\square$

APPENDIX B: OMITTED PROOFS FOR SECTION 3

B.1 Proof of Lemma 1

Throughout this section, I consider Δ small enough so that $1 - 2\Delta > \underline{t}(\alpha)$.

B.1.1 First step of the induction argument

LEMMA 7. *Conditions [C1], [C2], and [C3] are satisfied at $t = 1 - 2\Delta$.*

PROOF. First, I characterize continuation payoffs at time $1 - \Delta$ for any $E_{1-\Delta}$ and $\mu_{1-\Delta}$. Since $t = 1 - \Delta$ is the latest decision time, the agent has no use for further information. Thus, the principal's offer of information collection at $t = 1 - \Delta$ is irrelevant, and the agent takes his autarky action, which is to shirk if $\mu_{1-\Delta} < c/\alpha(Y_H - Y_L)$ and work, otherwise. Thus, the agent's equilibrium payoff is $U^a(1 - \Delta, E_{1-\Delta}, \mu_{1-\Delta})$, which is piecewise linear and continuous in $\mu_{1-\Delta}$. The principal's payoff at $1 - \Delta$ is

$$\Pi^*(1 - \Delta, E_{1-\Delta}, \mu_{1-\Delta}) = \begin{cases} (1 - \alpha)[\mu_{1-\Delta}E_{1-\Delta}(Y_H - Y_L) + Y_L] + \varepsilon E_{1-\Delta} & \text{if } \mu_{1-\Delta} < \frac{c}{\alpha(Y_H - Y_L)} \\ (1 - \alpha)[\mu_{1-\Delta}(\Delta + E_{1-\Delta})(Y_H - Y_L) + Y_L] + \varepsilon(\Delta + E_{1-\Delta}) & \text{otherwise} \end{cases},$$

which is also piecewise linear but is discontinuous at $c/\alpha(Y_H - Y_L)$ with a right-continuous jump up. Also, note that $\underline{\mu}^\Delta(E_{1-\Delta}) = c/\alpha(Y_H - Y_L)$ for any $E_{1-\Delta}$.

Consider a signal offer F at time $1 - 2\Delta$. The agent chooses $e = 1$ if and only if

$$\mathbb{E}_F[U^a(1 - \Delta, E_{1-2\Delta} + \Delta, \mu_{1-\Delta})] - \Delta(u + c) \geq U^a(1 - \Delta, E_{1-2\Delta}, \mu_{1-2\Delta}). \quad (20)$$

The left-hand side of the inequality is the agent's payoff from working, and the right-hand side is his payoff for shirking at $1 - 2\Delta$. First, consider the following problem:

$$\max_{F \in \mathcal{F}(\mu_{1-2\Delta})} \mathbb{E}_F[\Pi(1 - \Delta, E_{1-2\Delta} + \Delta, \mu_{1-\Delta})] \quad \text{subject to} \quad (20). \quad (21)$$

This problem maximizes the principal payoff by choosing a signal offer and subject to inducing effort. The problem is feasible since $\mu_{1-2\Delta} \geq \underline{\mu}_{1-2\Delta}^\Delta(E_{1-2\Delta})$.

I first characterize the value of (21) and show that its solution satisfies [C1] and [C3]. Then I show that this value strictly exceeds the principal's payoff from not inducing effort.

Solution of (21) First, $U^a(1 - \Delta, E_{1-\Delta}, \mu_{1-\Delta})$ is linear and continuous in $\mu_{1-\Delta}$ over $[0, \underline{\mu}_{1-\Delta}^\Delta(E_{1+\Delta})]$ while $\Pi^*(1 - \Delta, E_{1-\Delta}, \mu_{1-\Delta})$ exhibits a right-continuous jump up at $\underline{\mu}_{1-\Delta}^\Delta(E_{1+\Delta})$. Thus, any $F_{1-2\Delta}$ that puts positive probability on $(0, \underline{\mu}_{1-\Delta}^\Delta(E_{1+\Delta}))$ can be improved upon by one that distributes the weight of this interval to its end points while keeping the conditional expectation constant. This has no impact on the payoff of the agent, thus (20) continues to be satisfied, and it strictly increases the principal's payoff.

This establishes that *if* the principal's optimal choice would be to induce effort at $1 - 2\Delta$, [C3] would be satisfied.

An implication of the above, coupled with the observation that $\underline{\mu}^\Delta(E_{1-\Delta}) = c/\alpha(Y_H - Y_L)$ is that conditional on $\theta = 1$, the agent works for sure at $t = 1 - \Delta$. Consider the following auxiliary problem:

$$\max_{\tau \in [0,1]} (1 - \alpha)Y_L + (1 + \tau)\varepsilon \quad (22)$$

subject to

$$\begin{aligned} & \mu_{1-2\Delta}U_G^*(1 - 2\Delta, E_{1-2\Delta}, \mu_{1-2\Delta}) + (1 - \mu_{2\Delta})[\alpha Y_L - (1 + \tau)\Delta c - 2\Delta u] \\ & \geq U^a(1 - \Delta, E_{1-2\Delta}, \mu_{1-2\Delta}). \end{aligned}$$

In (22), the choice variable τ is the probability with which the agent works conditional on $\theta = 0$ at $t = 1 - \Delta$. The objective is the principal's payoff conditional on $\theta = 0$ while the constraint requires that the agent's expected payoff, if he chooses work, is no less than his autarky payoff.

Let τ^* be a solution of (22). Define μ^* by $\mu^*/(1 - \mu^*) = \mu_{1-2\Delta}/(1 - \mu_{1-2\Delta})\tau^*$. Then, naturally, F with support $\{0, \mu^*\}$ satisfying Bayesian plausibility solves (21).¹⁸ Also, since the principal's payoff increases in τ , while the agent's payoff decreases, it follows that at any solution of (22), the agent receives his autarky payoff. This establishes that *if* the principal's optimal choice is to induce effort at $1 - 2\Delta$, then [C1] is satisfied.

Second, I show that the principal's optimal choice is to induce effort at $1 - 2\Delta$. This is trivial if the agent's autarky action at $(\mu_{1-2\Delta}, E_{1-2\Delta})$ is work. Next, assume that the agent's autarky action is to shirk. To see that the principal strictly prefers to induce effort, observe that if the agent shirks at $1 - 2\Delta$, then he also shirks at $1 - \Delta$. This is because $\mu_{1-2\Delta} < c/\alpha(Y_H - Y_L) = \mu_{1-\Delta}^{\text{work}}(E_{1-\Delta}) = \mu_{1-2\Delta}^{\text{work}}(E_{1-2\Delta})$. Further, if the agent shirks, the belief is not updated, that is, $\mu_{1-\Delta} = \mu_{1-2\Delta}$. Thus, if the agent shirks at $1 - 2\Delta$, total effort he exerts conditional on either state is 0, while in the solution to (21), he exerts total effort of 2Δ conditional on $\theta = 1$ and $(1 + \tau^*)\Delta$ conditional on $\theta = 0$. Note that the principal strictly prefers higher effort in each state, establishing [C2]. $\square$

B.1.2 Induction hypothesis and its implications

Induction hypothesis Fix t , E_t , μ_t with $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$. Assume that in each continuation equilibrium, at any $t' \geq t$, the following conditions hold:

[C1] The agent's payoff is $U^a(t', E_{t'}, \mu_{t'})$.

[C2] The agent works at t' .

[C3] $\mu_{t'+\Delta} \in \{0\} \cup [\underline{\mu}_{t'+\Delta}^\Delta(E_{t'} + \Delta), 1]$.

¹⁸Suppose not. Take any solution F' of (21) and let $\tau' = \text{Prob}[\mu_{1-\Delta} \geq c/(\alpha(Y_H - Y_L)) | F']$. Then τ' is feasible and delivers a strictly higher value for (22).

Characterization of continuation payoffs under the induction hypothesis Next, I characterize the principal's equilibrium payoff under the induction hypothesis. This is necessary for the induction step which will show that [C1], [C2], and [C3] also hold at $t - \Delta$.

Principal's problem Fix $t \geq t'$, E_t , $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$. I first characterize the value $\Pi^*(\cdot, \cdot, \cdot)$ of the problem (PP) subject to (9) and (10), and subject also to the following terminal conditions: For $t = 1 - \Delta$ or $\mu_t < \underline{\mu}_t^\Delta(E_t)$,

$$\Pi^*(t, E_{1-\Delta}, \mu_t) = \begin{cases} \varepsilon E_t & \text{if } \mu_t < \underline{\mu}_t^\Delta(E_t) \text{ and } t < \underline{t}(\alpha) \\ [Y_L + \mu_t(Y_H - Y_L)E_t](1 - \alpha) + \varepsilon E_t & \text{if } \mu_{1-\Delta} < \underline{\mu}_t^\Delta \text{ and } 1 - \Delta > t \geq \underline{t}(\alpha) \\ [Y_L + \mu_t(Y_H - Y_L)(E_t + \Delta)](1 - \alpha) + \varepsilon(E_t + \Delta) & \text{if } \mu_{1-\Delta} \geq \underline{\mu}_t^\Delta \text{ and } 1 - \Delta = t \end{cases} \quad (23)$$

The terminal condition (23) uses the following facts: (i) if $\mu_t < \underline{\mu}_t^\Delta(E_t)$ and $t < \underline{t}(\alpha)$, the agent quits, because for such t , $\underline{\mu}_t^\Delta(E_t) \leq \mu_t^*(E_t) < \mu_t^{\text{quit}}(E_t)$; (ii) if $\mu_t < \underline{\mu}_t^\Delta(E_t)$ and $t \geq \underline{t}(\alpha)$, no further learning takes place, and the agent chooses his autarky action; (iii) $\underline{\mu}_{1-\Delta}^\Delta(E_{1-\Delta}) = \mu_{1-\Delta}^{\text{work}}(E_{1-\Delta})$.

An auxiliary problem Next, I prove Lemma 2, which characterizes the solution to the auxiliary problem (PP-aux) subject to (11) and (12). For brevity, throughout the Appendix, I write $\tilde{\pi}_B(T_B|t)$ and $\tilde{u}_B(T_B|t)$ instead of $\tilde{\pi}_B(T_B|t, E_t, \mu_t)$ and $\tilde{u}_B(T_B|t, E_t, \mu_t)$, suppressing the dependence on E_t, μ_t whenever this causes no confusion.

PROOF OF LEMMA 2. Fix Δ and (t, E_t, μ_t) . First, consider $\mu_t \geq \mu_t^{\text{work}}(E_t)$. Then the optimal solution trivially is the degenerate H placing all weight on $T_B = 1$. Note that in this case, (11) binds. Next, consider $t \geq \underline{t}(\alpha) - \Delta$. Then, by (12), the support of H is contained in $[\underline{t}(\alpha), 1]$. Since the payoffs are linear over this interval, the principal's payoff is increasing and the agent's payoff is decreasing in T_B ; any H with expectation T_B^* with support contained in $\{\underline{t}(\alpha), \underline{t}(\alpha) + \Delta, \dots, 1\}$ solves the problem. Note that (11) binds.

Next, assume that $t < \underline{t}(\alpha) - \Delta$ and $\mu_t < \mu_t^{\text{work}}(E_t)$. This is possible only when $\alpha < 1$, because otherwise $\underline{t}(\alpha) = 0$. Since the agent's payoff is linear and continuous over $[t + \Delta, \underline{t}(\alpha)]$ and the principal's has a right-continuous upward jump at $\underline{t}(\alpha)$, any H placing positive weight on the interior of interval $[t + \Delta, \underline{t}(\alpha)]$ can be improved upon by a mean-preserving spread of it, which is supported over the interval's end points. Further, since Π_B, U_B are both linear over $[\underline{t}(\alpha), 1]$, all payoff relevant aspects of H can be summarized by a pair $(\gamma, \bar{T})$, where γ is the probability of $T_B = t + \Delta$ and $\bar{T} = \mathbb{E}_G[T_B | T_B \geq \underline{t}(\alpha)]$. Thus, the problem (PP-aux) can be reexpressed as

$$\max_{\bar{T} \geq \underline{t}(\alpha), \gamma} (1 - \gamma)\tilde{\pi}_B(\bar{T}|t + \Delta, E_{t+\Delta}, \mu_{t+\Delta}) + \varepsilon \Delta \quad (24)$$

subject to

$$(1 - \gamma)\tilde{u}_B(\bar{T}|t + \Delta) - \Delta(c + u) \geq \underline{U}_B(t, E_t, \mu_t), \quad (25)$$

where $\underline{U}_B(t, E_t, \mu_t)$ is defined by

$$U^a(t, E_t, \mu_t) = \mu_t \tilde{u}_G(t, E_t) + (1 - \mu_t) \underline{U}_B(t, E_t, \mu_t).$$

Note that in any solution of (24), necessarily $\gamma < 1$: since $\mu_t > \underline{\mu}_t^\Delta(E_t)$, $\gamma = 1$ can not satisfy (25). This is because $\gamma = 1$ corresponds to a fully informative signal, and would necessarily imply that left-hand side of (25) is larger. Then the first-order conditions are

$$(1 - \gamma) \left[\underbrace{\frac{\partial \tilde{\pi}_B(\bar{T}|t + \Delta)}{\partial \bar{T}}}_{\varepsilon} + \lambda \underbrace{\frac{\partial \tilde{U}_B(\bar{T}|t + \Delta)}{\partial \bar{T}}}_{-c} \right] \begin{cases} \leq 0 & \text{if } \bar{T} = \underline{t}(\alpha) \\ = 0 & \text{if } \bar{T} \in (\underline{t}(\alpha), 1) \\ \geq 0 & \text{if } \bar{T} = 1 \end{cases} \quad (\text{FOC}_{\bar{T}})$$

and

$$-\tilde{\pi}_B(\bar{T}|t + \Delta) - \lambda U_B(\bar{T}|t + \Delta) \begin{cases} \leq 0 & \text{if } \gamma = 0 \\ = 0 & \text{if } \gamma \in (0, 1) \end{cases}, \quad (\text{FOC}_\gamma)$$

where $\lambda \geq 0$ is a Lagrange multiplier. First, note that $\Pi_B(T_B|t)$ is strictly increasing in T_B . Thus, if $\lambda = 0$, that is, (11) does not bind, then the optimal H would assign probability 1 to $T_B = 1$. But such H violates (11) since $\mu_t < \mu_t^{\text{work}}(E_t)$. Thus, (11) binds, and $\lambda > 0$.

If $\gamma > 0$, by (FOC $_\gamma$), $\lambda = -\tilde{\pi}_B(\bar{T}|t + \Delta)/\tilde{u}_B(\bar{T}|t + \Delta) > 0$. Then I claim that the left-hand side of (FOC $_{\bar{T}}$) is negative. This is equivalent to

$$\varepsilon + c \frac{\tilde{\pi}_B(\bar{T}|t)}{\tilde{u}_B(\bar{T}|t)} < 0 \quad \Leftrightarrow \quad \varepsilon \tilde{u}_B(\bar{T}|t) + c \tilde{\pi}_B(\bar{T}|t) > 0,$$

where the inequality reverses because $\tilde{u}_B(\bar{T}|t) < 0$. This can be reexpressed as

$$\varepsilon [\alpha Y_L - (1 - t - \Delta)u - (\bar{T} - t - \Delta)c] + c [(1 - \alpha)Y_L + \varepsilon(\bar{T} - t - \Delta)] > 0,$$

which is equivalent to

$$\varepsilon [\alpha Y_L - (1 - t - \Delta)u] + c [(1 - \alpha)Y_L] > 0.$$

Since $c > \varepsilon$, a sufficient condition is $Y_L > (1 - t)u$, which is satisfied.

Thus, for any γ , $\bar{T}$ that satisfy first-order conditions, $\gamma > 0$ if and only if $\bar{T} = \underline{t}(\alpha)$. If $\mu_t \geq \mu_t^*(E_t)$, for any $\gamma \geq 0$ and $\bar{T} = \underline{t}(\alpha)$ the left-hand side of (25) exceeds the right. Thus, necessarily $\gamma = 0$, and the supp $H \subset [\underline{t}(\alpha), 1]$, as claimed. If $\mu_t < \mu_t^*(E_t)$, $\gamma = 0$ is not feasible. Thus, $\bar{T} = \underline{t}(\alpha)$, and $\gamma = \gamma_\Delta^*(t, E_t, \mu_t)$ by (25), as claimed.

Finally, the expression for $\Pi_t^*(t, E_t, \mu_t)$ is obtained by evaluating $\Pi_B(t, E_t)$ at the solution of (PP-aux) subject to (11) and (12). $\square$

The solution of the principal's problem

LEMMA 8. Fix t, E_t and $\mu_t > \underline{\mu}_t^\Delta(E_t)$. Let $\Pi_B^{\text{AUX}}(t, E_t, \mu_t)$ be the value of (PP-aux). Then

$$\Pi^*(t, E_t, \mu_t) = \mu_t \Pi_G^*(t, E_t, \mu_t) + (1 - \mu_t) \Pi_B^{\text{AUX}}(t, E_t, \mu_t) \quad (26)$$

PROOF OF LEMMA 8. I show that the value defined in (26) satisfies the finite horizon Bellman equation defined in (PP). First, consider $t = 1 - 2\Delta$. The proof of Lemma 7 solves the principal's problem at this date and demonstrates that, under its solution, conditional on good state the agent works throughout, and conditional on bad state, he works an expected amount $(1 + \tau^*)\Delta$, where τ^* is chosen such that the constraint of (22), which is equivalent to (11), binds. Therefore, the principal's payoff at this date for $\mu_t > \underline{\mu}_t(E_t)$ is given by the value function in (26). Fix $t < 1 - 2\Delta$ and assume that for all $t' \geq t + \Delta$, the principal's continuation payoff, as a function of $\mu_{t'}$ is given by (26). Now I show that this must also be true at t .

For $\mu_t < \mu_t^*(E_t)$, consider the unique $F_1 \in \mathcal{F}(\mu_t)$, which has support $\{0, \mu_{t+\Delta}^*(E_t + \Delta), 1\}$ and satisfies (11) with equality. For $\mu_t \geq \mu_t^*(E_t)$, consider the unique $F_2 \in \mathcal{F}(\mu_t)$, which has support $\{\mu_{t+\Delta}^*(E_t + \Delta), c/(\alpha(Y_H - Y_L)), 1\}$ and satisfies (11) with equality. F_1 and F_2 satisfies (11) by construction. Given that from $t + \Delta$ on, the principal's payoff is given by (26), by the choice of such F , the principal attains the value (26) at time t . Thus, the principal's payoff is no less than the value in (26).

Suppose the value of the problem (PP) strictly exceeds $\Pi^*(t, E_t, \mu_t)$ given in (26). Choose an optimal policy for (PP). Let H^* be the probability distribution induced on T_B by this policy. By the constraints (11) and (10), H^* is feasible for (PP-aux) and attains a strictly higher value than given in (26), a contradiction. $\square$

B.1.3 The induction step

LEMMA 9 (Induction Step). Fix t, E_t, μ_t with $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$. Assume that the induction hypothesis holds. Then, [C1], [C2], and [C3] also hold at $t - \Delta$.

PROOF. I first solve the problem of maximizing the principal's payoff *conditional on* inducing effort and show that this solution satisfies [C1] and [C3]. Later, I show that this value is strictly larger than the payoff the principal could obtain if she did not induce effort, establishing [C2].

For this purpose, first note that, under the induction hypothesis, when $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$, the principal's continuation payoff is given by (26). When $\mu_t < \underline{\mu}_t^\Delta(E_t)$, the agent cannot be induced to work in the continuation game, and thus the principal's continuation payoff starting at t is

$$\Pi^*(t, E_t, \mu_t) = \begin{cases} \varepsilon E_t & \text{if } t < \underline{t}(\alpha) \\ [Y_L + \mu_t(Y_H - Y_L)E_t](1 - \alpha) + \varepsilon E_t & \text{if } t \geq \underline{t}(\alpha) \end{cases}.$$

Consider the problem

$$\max_{F \in \mathcal{F}(\mu_{t-\Delta})} \mathbb{E}_F[\Pi^*(t, E_t, \mu_t)] \quad (27)$$

subject to

$$\mathbb{E}_F[U^a(t, E_{t-\Delta} + \Delta, \mu_t)] - \Delta(c + u) \geq U^a(t - \Delta, E_{t-\Delta}, \mu_{t-\Delta}) \quad (28)$$

Since the agent expects to receive his autarky payoff starting at t , he prefers to work rather than take his autarky action at $t - \Delta$ as long as (28) is satisfied.

I first show that the support of any F that solves (27) subject to (28) does not intersect with $(0, \underline{\mu}_t^\Delta(E_{t-\Delta} + \Delta))$. For ease of notation, I write $\underline{\mu}$ instead of $\underline{\mu}_t^\Delta(E_{t-\Delta} + \Delta)$. Note that

$$\begin{aligned} & \lim_{\mu_t \rightarrow \underline{\mu}^-} \Pi^*(t, E_{t-\Delta} + \Delta, \mu_t) \\ &= \begin{cases} \varepsilon(E_{t-\Delta} + \Delta) & \text{if } t < \underline{t}(\alpha) \\ [Y_L + \underline{\mu}(Y_H - Y_L)(E_{t-\Delta} + \Delta)](1 - \alpha) + \varepsilon(E_{t-\Delta} + \Delta) & \text{if } t \geq \underline{t}(\alpha) \end{cases} \end{aligned}$$

while

$$\begin{aligned} & \Pi^*(t, E_{t-\Delta} + \Delta, \underline{\mu}) \\ &= \begin{cases} \varepsilon(E_{t-\Delta} + \Delta) + \underline{\mu}[(1 - \alpha)(Y_L + (Y_H - Y_L)(E_{t-\Delta} + 1 - t)) \\ \quad + \varepsilon(1 - t - E_{t-\Delta} - \Delta)] & \text{if } t < \underline{t}(\alpha) \\ \varepsilon(E_{t-\Delta} + \Delta) + (1 - \alpha)Y_L + \underline{\mu}[(Y_H - Y_L)(E_{t-\Delta} + 1 - t)) \\ \quad + \varepsilon(1 - t - E_{t-\Delta} - \Delta)] & \text{if } t \geq \underline{t}(\alpha) \end{cases}. \end{aligned}$$

This calculation takes into account the fact that if $\mu_{t-\Delta} = \underline{\mu}$, then the principal offers a fully informative signal at time $t - \Delta$. Thus, conditional on bad state the agent quits at time t if $t < \underline{t}(\alpha)$.

By inspection, $\Pi^*(t, E_{t-\Delta} + \Delta, \mu_t)$ as a function of μ_t features a right-continuous upward jump at $\underline{\mu}$. Further, the left-hand side of (28) is linear in μ_t over this interval. Thus, the support of the principal's optimal offer F cannot intersect with $(0, \underline{\mu})$. This establishes that if the principal's induces effort, [C3] would be satisfied.

Next, I argue that at any solution of (27), the constraint (28) binds. By the induction hypothesis and the previous claim, if at t , the principal offers F that solves (27) subject to (28), in the continuation path starting at t , conditional on $\theta = 1$ the worker always works. And the outcome of the continuation path naturally generates a probability distribution $\tilde{H}$ over T_B (which recall is the random date at which the low state is revealed). Note that such $\tilde{H}$ satisfies (11) and is thus feasible for (PP-aux). Thus, the principal's payoff cannot exceed the value described in (26). Further, the policy described in the proof of Lemma 8 is feasible for (27), and thus the principal's payoff is as given in (26). Finally, since $\tilde{\pi}_B(T_B|t - \Delta)$ is strictly increasing and $\tilde{u}(T_B|t - \Delta)$ is strictly decreasing in T_B , this value cannot be attained unless (28) is satisfied with equality. This establishes that in the solution of this problem, the agent's continuation payoff at time $t - \Delta$ is equal to his autarky payoff. This establishes that if the principal's optimal action were to induce effort, [C1] would be satisfied.

Given the previous two claims, it now suffices to show that the principal prefers to induce effort, that is, [C2] is satisfied. First, consider $\mu_{t-\Delta} < \mu_{t-\Delta}^{\text{shirk}}(E_{t-\Delta})$ so that the agent's autarky action is "quit." Thus, if the principal does not induce effort, the agent quits, which leaves the principal with strictly lower payoff than (26). Next, consider $\mu_{t-\Delta}^{\text{shirk}}(E_{t-\Delta}) \leq \mu_{t-\Delta} < \mu_{t-\Delta}^{\text{work}}(E_{t-\Delta})$, so that the agent's autarky action is to shirk. If the

principal does not induce effort, the belief is not updated. If $\mu_{t-\Delta} < \underline{\mu}_t^\Delta(E_{t-\Delta})$, then the agent never works thereafter conditional on either state, and thus the principal is worse off than in the solution of (27) subject to (28). Next, if $\mu_{t-\Delta} \geq \underline{\mu}_t^\Delta(E_{t-\Delta})$, then conditional on $\theta = 1$, the agent works throughout starting at t . Conditional on $\theta = 0$, the agent never quits, since by Lemma 4, his autarky action remain “shirk” at time t and by Lemma 5, $\mu_t^*(E_{t-\Delta}) < \mu_t^{\text{shirk}}(E_{t-\Delta})$. Let E_B^* be the expected amount of effort the agent exerts in the continuation path following shirk at $t - \Delta$, conditional on $\theta = 0$. Analogously, let E_B^{**} be the expected amount of effort that the agent exerts conditional on $\theta = 0$ starting at $t - \Delta$ when the principal follows the solution of (27) subject to (28). Then it is necessary that

$$\begin{aligned} U^a(t - \Delta, E_{t-\Delta}, \mu_{t-\Delta}) &= \alpha Y_L - u(1 - t) + \mu_t((Y_H - Y_L)\alpha - c)(1 - t) + (1 - \mu_t)(-c)E_B^* \\ &= \alpha Y_L - u(1 - t + \Delta) + \mu_{t-\Delta}((Y_H - Y_L)\alpha - c)(1 - t) + (1 - \mu_{t-\Delta})(-c)E_B^{**} \end{aligned}$$

since regardless of whether the principal induces effort at $t - \Delta$ or not, the agent receives his autarky payoff at that history. Since $(Y_H - Y_L)\alpha - c > 0$ and $1 - t < 1 - t + \Delta$, such equality is possible only when $E_B^{**} > E_B^*$. Thus, the agent exerts more effort conditional on either state if the principal induces effort at time $t - \Delta$, and consequently, the principal strictly prefers to induce effort at $t - \Delta$. This establishes [C2] and completes the proof. $\square$

B.2 Proof of Proposition 1

Item 1 follows immediately from Lemma 1 and the first part of item 2 immediately follows from Lemma 2. The second part of item 2 also follows from Lemma 2 by noting that $\lim_{\Delta \rightarrow 0} \gamma_\Delta^*(t, E_t, \mu_t) = \gamma^*(\alpha, \mu_0)$ since

$$1 - \lim_{\Delta \rightarrow 0} \gamma_\Delta^*(t, E_t, \mu_t) = \frac{\mu_t}{1 - \mu_t} \frac{U_G^*(t, E_t, \mu_t)}{\tilde{u}_B(\underline{t}(\alpha)|t, E_t, \mu_t)} = \frac{\mu_t}{1 - \mu_t} \frac{1 - \mu_t^*(E_t)}{\mu_t^*(E_t)},$$

where the first equality is obtained by taking the limits of both sides of the equation defining $\gamma_\Delta^*(t, E_t, \mu_t)$ in Lemma 2 and the second equality is by definition of $\mu_t^*(E_t)$.

The case for $\alpha < (u + c)/Y_H$ immediately follows by noting that in this case, even with immediate learning, the agent’s payoff conditional on each state is negative unless he quits.

APPENDIX C: OMITTED PROOFS FOR SECTION 4

PROOF OF LEMMA 3. By rearranging (15), we obtain

$$\frac{\mu_0}{1 - \mu_0} = \frac{u + c}{u} \frac{u - \bar{\alpha}(\mu_0)Y_L}{\bar{\alpha}(\mu_0)Y_H - u - c}.$$

That $\bar{\alpha}(\mu_0)$ is decreasing follows immediately by inspection. The bounds on $\bar{\alpha}(\mu_0)$ follows by this monotonicity and by noting that $\bar{\alpha}(0) = u/Y_L$ and $\bar{\alpha}(1) = (u + c)/Y_H$. $\square$

PROOF OF THEOREM 1. Fix μ_0 and let $S^*(\alpha)$ be the surplus generated in equilibrium when the agent's share is α . Then

$$S^*(\alpha) = \begin{cases} \mu_0[Y_H - c - u + \varepsilon] + (1 - \mu_0)[Y_L - u - (c - \varepsilon)T_B^*(0, 0, \mu_0)] & \text{if } \mu_0 \geq \mu^*(\alpha) \\ \mu_0[Y_H - c - u + \varepsilon] + (1 - \mu_0)(1 - \gamma^*(\alpha, \mu_0))[Y_L - u - (c - \varepsilon)\underline{t}(\alpha)] & \text{if } \mu_0 < \mu^*(\alpha) \end{cases},$$

where $T_B^*(t, E_t, \mu_t)$ is defined in Lemma 2 and $\gamma^*(\alpha, \mu_0)$ is defined in Proposition 1. Note that $T^*(t, E_t, \mu_t)$ increases if α increases. The principal's payoff as a function of α then is

$$S^*(\alpha) - U^-(\alpha, \mu_0),$$

where $U^-(\alpha, \mu_0) = \max\{0, \alpha Y_L - u, \mu_0 \alpha (Y_H - Y_L) + \alpha Y_L - c - u\}$, is the agent's equilibrium payoff, which in turn is equal to his autarky payoff at the initial history.

Consider $\alpha \in (\bar{\alpha}(\mu_0), 1]$, so that $\mu_0 > \mu^*(\alpha)$ and $T_B^*(0, 0, \mu_0) > \underline{t}(\alpha)$. Since $T_B^*(0, 0, \mu_0)$ is increasing in α , $S^*(\alpha)$ is decreasing in α over this interval. Since the agent's payoff $U^-(\alpha, \mu_0)$ is nondecreasing in α , the principal's payoff is strictly higher with $\bar{\alpha}(\mu_0)$ than with any α in this interval. Second, by Proposition 1 if $\alpha < \underline{\alpha}$, the agent immediately quits, and thus the principal's payoff is 0, while $\alpha = \underline{\alpha}$ delivers strictly positive payoffs. Thus, $\alpha < \underline{\alpha}$ cannot be optimal.

Third, consider $\alpha \in (\underline{\alpha}, \bar{\alpha}(\mu_0))$. The agent's payoff over this range is 0 by Lemma 5, thus the principal's payoff is equal to the total surplus, which for this case is

$$\mu_0[Y_H - c - u + \varepsilon] + (1 - \mu_0)(1 - \gamma^*(\alpha, \mu_0))[Y_L - u - (c - \varepsilon)\underline{t}(\alpha)].$$

Note that for this case $\gamma^*(\alpha, \mu_0)$ is positive and decreasing in α and $\gamma^*(\underline{\alpha}, \mu_0) = 0$. Further, $\underline{t}(\alpha)$ is decreasing in α . Consider two possibilities: (i) $Y_L - u - (c - \varepsilon)\underline{t}(\alpha) < 0$. In this case, the surplus conditional on the bad state is negative for such α , while it is 0 when $\alpha = \underline{\alpha}$, establishing that such α cannot be optimal. (ii) $Y_L - u - (c - \varepsilon)\underline{t}(\alpha) \geq 0$. Since $\gamma^*(\alpha, \mu_0)$ is positive and decreasing in α while $\underline{t}(\alpha)$ is decreasing in α , the surplus conditional on bad state increases in α . Thus, the principal strictly prefers $\bar{\alpha}(\mu_0)$ to any α in this range.

The only remaining candidates are $\underline{\alpha}$ and $\bar{\alpha}(\mu_0)$. It is immediate by observation that (16) is necessary and sufficient for $\underline{\alpha}$ to deliver a larger payoff to the principal. $\square$

PROOF OF PROPOSITION 2. It suffices to show that $Y_L - u - \underline{t}(\bar{\alpha}(\mu_0))(c - \varepsilon)$ is decreasing in u, c, μ_0 and increasing in Y_L . Note that $\underline{t}(\bar{\alpha}(\mu_0))$ also varies with Y_L, u, c . Using (15), one obtains

$$\underline{t}(\bar{\alpha}(\mu_0)) = \frac{1}{1 + \frac{u + c}{\mu_0 \left[\frac{Y_H}{Y_L} u - c - u \right]}}.$$

By inspection, $\underline{t}(\bar{\alpha}(\mu_0))$ increases in μ_0, Y_L , which implies that $Y_L - u - \underline{t}(\bar{\alpha}(\mu_0))(c - \varepsilon)$ is decreasing in μ_0 and increasing in Y_L , establishing the claimed comparative statics with

respect to these variables. Taking derivative of $\underline{t}(\bar{\alpha}(\mu_0))$ with respect to u establishes that it is also increasing in u , so that $Y_L - u - \underline{t}(\bar{\alpha}(\mu_0))(c - \varepsilon)$ is decreasing in u , establishing the claimed comparative statics with respect to u . $\square$

APPENDIX D: PROOFS FOR SECTION 5

D.1 *Motivating effort and immediate learning*

PROOF OF PROPOSITION 3. Assume that $(Y_H - Y_L)/Y_L < c/u$. As in the analysis under Assumption 5, I first characterize equilibrium outcomes for various values of α . For analogous reasons to when (5) holds, in any equilibrium of the production game, at any history, the agent receives his autarky payoff. Thus, he exerts effort whenever the principal's signal delivers him at least this payoff.

Consider $\alpha < u/Y_L$. Such α cannot deliver the agent nonnegative payoff regardless of state even with immediate full information, thus in this case the agent immediately quits.

Consider $u/Y_L \leq \alpha < c/(Y_H - Y_L)$. In this case, regardless of the principal's information disclosure policy, the agent never works nor quits.

Consider $c/(Y_H - Y_L) < \alpha$. In this case, the agent is immediately vested, and never quits. Consider two subcases: (i) If $\mu_0\alpha(Y_H - Y_L) \geq c$, the agent's autarky action is to work. Thus, the principal provides no information, and the agent works throughout. (ii) If $\mu_0\alpha(Y_H - Y_L) < c$, the agent's autarky action is to shirk. Since the agent is already vested, he never quits. Thus, it remains to characterize the equilibrium effort paths. Consider an auxiliary problem where the principal chooses E_G and E_B , that is, the amount of the agent's effort conditional on each state, subject to delivering the agent his autarky payoff. More specifically consider the following problem:

$$\max_{E_G, E_B} \mu_0((1 - \alpha)(Y_H - Y_L) + \varepsilon)E_G + (1 - \mu_0)\varepsilon E_B + (1 - \alpha)Y_L$$

subject to

$$\mu_0(\alpha(Y_H - Y_L) - c)E_G + (1 - \mu_0)(-E_B c) - u + \alpha Y_L \geq \alpha Y_L - u.$$

The objective function is the principal's payoff from (E_G, E_B) . The left-hand side of the constraint is the agent's payoff from (E_G, E_B) , while the right-hand side is his autarky payoff. In any solution of this problem, the constraint must bind and $E_G = 1$ must hold, since the principal's payoff is increasing in E_G, E_B , and the agent's is increasing in E_G and decreasing in E_B while $E_G = E_B = 1$ violates the constraint.

The principal's equilibrium payoff cannot exceed the value of the above auxiliary problem, since each equilibrium induces a probability distribution over E_G, E_B that is feasible for the auxiliary problem. Next, I show that the principal can achieve this payoff with a one-time informative signal offer at time 0. This offer reveals the bad state with probability $\beta(\alpha)$ so that

$$\frac{\mu_0}{1 - \mu_0} \frac{1}{\beta(\alpha)} = \frac{c}{\alpha(Y_H - Y_L) - c},$$

guaranteeing that if the bad state is not revealed, the agent is exactly indifferent between working and shirking. The expected amount of time $\hat{t}(\alpha)$ he spends working is given by $1 - \beta(\alpha)$, and is calculated as

$$\mu_0[\alpha(Y_H - Y_L) - c] = (1 - \mu_0)c\hat{t}(\alpha).$$

Then, in any equilibrium of the production stage, total surplus as a function of α is

$$S^*(\alpha) = \begin{cases} Y_L - u & \text{if } Y_H - Y_L < c \\ Y_L - u + \mu_0(Y_H - Y_L + \varepsilon - c) - (1 - \mu_0)(c - \varepsilon)\hat{t}(\alpha) & \text{if } Y_H - Y_L \geq c \end{cases}$$

The agent's equilibrium payoff as a function of α is

$$U(\alpha) = \begin{cases} \alpha Y_L - u & \text{if } \mu_0\alpha(Y_H - Y_L) < c \\ \alpha[Y_L + \mu_0(Y_H - Y_L)] - u - c & \text{if } \mu_0\alpha(Y_H - Y_L) \geq c \end{cases}.$$

Note that $\hat{t}(\alpha)$ is increasing in α , and thus $S^*(\alpha)$ is decreasing over the relevant range. Moreover, $U(\alpha)$ is always increasing in α . Then the principal's payoff $S^*(\alpha) - U(\alpha)$ is decreasing when $\alpha \geq c/(Y_L - Y_H)$ and when $\alpha \in [u/Y_L, c/(Y_H - Y_L))$ with a right-continuous upward jump at $\alpha = c/(Y_H - Y_L)$. Thus, the principal's optimal share offer is either u/Y_L or $c/(Y_H - Y_L)$. With the former, the agent receives no rents and shirks throughout, and thus the principal's payoff is $Y_L - u$. With the latter, the principal optimally immediately offers a fully informative signal (i.e., $\hat{t} = 0$), the agent shirks throughout in the bad state and works throughout in the good state, implying that the total surplus is maximized. In this case, the agent's payoff is $cY_L/(Y_H - Y_L) - u > 0$. Thus, the principal's payoff is

$$Y_L + \mu_0(Y_H - Y_L - c + \varepsilon) - c \frac{Y_L}{Y_H - Y_L}.$$

Then the principal prefers the former if and only if

$$\mu_0(Y_H - Y_L - c + \varepsilon) < c \frac{Y_L}{Y_H - Y_L} - u.$$

Choosing $\tilde{\mu}$ to satisfy the latter expression with equality establishes the claim. $\square$

D.2 Information collection without effort

The proof of Proposition 4 relies on the concavity of $\Pi^*(t, E_t, \mu_t)$ over $[\underline{\mu}_t^\Delta(E_t), 1]$, which is established in Lemma 10.

LEMMA 10. *For any t, E_t , the equilibrium payoff $\Pi^*(t, E_t, \mu_t)$ of the principal is a piecewise linear and concave function of μ_t over the interval $[\underline{\mu}_t^\Delta(E_t), 1]$.*

PROOF. First, note that $\Pi^*(t, E_t, \mu_t) = \mu_t \Pi_G^*(t, E_t, \mu_t) + (1 - \mu_t) \Pi_B^*(t, E_t, \mu_t)$, where for $\mu_t \geq \underline{\mu}_t^\Delta(E_t)$,

$$\Pi_G^*(t, E_t, \mu_t) = (1 - \alpha)[Y_L + (Y_H - Y_L)(1 - t + E_t)] + \varepsilon(1 - t + E_t),$$

and

$$\Pi_B^*(t, E_t, \mu_t) = \begin{cases} (1 - \alpha)Y_L + \varepsilon(1 - t - \Delta) + \varepsilon(E_t + \Delta) & \text{if } \mu_t \geq \mu_t^{\text{work}}(E_t) \\ (1 - \alpha)Y_L + \varepsilon(T_B^*(t, E_t, \mu_t) - t - \Delta) + \varepsilon(E_t + \Delta) & \text{if } \mu_t^*(E_t) \leq \mu_t < \mu_t^{\text{work}}(E_t) \\ (1 - \gamma_\Delta^*(t, E_t, \mu_t))\{(1 - \alpha)Y_L + \varepsilon[\underline{t}(\alpha) - t - \Delta]\} + \varepsilon(E_t + \Delta) & \text{if } \mu_t^*(E_t) > \mu_t \geq \underline{\mu}_t^\Delta(E_t) \end{cases},$$

where $T_B^*(t, E_t, \mu_t)$ and $\gamma_\Delta^*(t, E_t, \mu_t)$ are defined in Lemma 2. To establish concavity, first note that $\mu_t \Pi_G^*(t, E_t, \mu_t) + (1 - \mu_t)\varepsilon[E_t + \Delta]$ is linear in μ_t as $\Pi_G^*(t, E_t, \mu_t)$ is independent of μ_t . Then it suffices to show that $(1 - \mu_t)[\Pi_B^*(t, E_t, \mu_t) - \varepsilon(E_t + \Delta)]$ is piecewise linear, continuous, and concave. For this purpose, fix t, E_t, μ_t , and using the definition of $T_B^*(t, E_t, \mu_t)$, express $(1 - \mu_t)(T_B^*(t, E_t, \mu_t) - t - \Delta)$ as

$$\begin{aligned} & (1 - \mu_t)(T_B^*(t, E_t, \mu_t) - t - \Delta) \\ &= \frac{(1 - \mu_t)(\alpha Y_L - (1 - t)u) + \mu_t U_G(E_t, t) - U^a(t, E_t, \mu_t)}{c} - (1 - \mu_t)\Delta. \end{aligned}$$

Then, for any $\mu_t \geq \mu_t^*(E_t)$,

$$\begin{aligned} & (1 - \mu_t)[\Pi_B^*(t, E_t, \mu_t) - \varepsilon(E_t + \Delta)] \\ &= (1 - \mu_t)(1 - \alpha)Y_L \\ &+ \varepsilon \left(\underbrace{\frac{(1 - \mu_t)(\alpha Y_L - (1 - t)u) + \mu_t U_G(E_t, t) - U^a(t, E_t, \mu_t)}{c} - (1 - \mu_t)\Delta}_{=(1 - \mu_t)(T_B^*(t, E_t, \mu_t) - t - \Delta)} \right). \end{aligned} \quad (29)$$

Since $U_G(E_t, t)$ is independent of μ_t and $U^a(t, E_t, \mu_t)$ is continuous, piecewise linear, and convex, the result follows for $\mu_t \geq \mu_t^*(E_t)$.

Consider $\underline{\mu}_t^\Delta(E_t) \leq \mu_t < \mu_t^*(E_t)$. By definition, when $\mu_t = \mu_t^*(E_t)$, $\gamma_\Delta^*(t, \mu_t, E_t) = 0$ and $T_B^*(t, E_t, \mu_t) = \underline{t}(\alpha)$, establishing continuity at $\mu_t^*(E_t)$. Express $(1 - \gamma_\Delta^*)(1 - \mu_t)$ as

$$(1 - \mu_t)(1 - \gamma_\Delta^*) = \frac{\Delta(c + u) - \mu_t U_G(t + \Delta, E_t + \Delta)}{U_B(\underline{t}(\alpha)|t + \Delta)},$$

where I drop the arguments of $\gamma_\Delta^*(\cdot, \cdot, \cdot)$ for ease of notation. This establishes linearity. Next,

$$\begin{aligned} & (1 - \mu_t)[\Pi_B^*(t, E_t, \mu_t) - \varepsilon(E_t + \Delta)] \\ &= \frac{\Delta(c + u) - \mu_t U_G(t + \Delta, E_t + \Delta)}{U_B(\underline{t}(\alpha)|t + \Delta)} \underbrace{\{(1 - \alpha)Y_L + \varepsilon[\underline{t}(\alpha) - t - \Delta]\}}_{\Pi_B(\underline{t}(\alpha)|t + \Delta)}. \end{aligned}$$

Thus, concavity is equivalent to

$$\begin{aligned} & -\frac{U_G(t+\Delta, E_t+\Delta)}{U_B(\underline{t}(\alpha)|t+\Delta)}\Pi_B(\underline{t}(\alpha)|t+\Delta) \\ & \geq -(1-\alpha)Y_L + \varepsilon\left(\frac{-(\alpha Y_L - (1-t-\Delta)u) + U_G(E_t+\Delta, t+\Delta)}{c}\right), \end{aligned}$$

where the right-hand side is the slope calculated from (29). This is reexpressed as

$$\begin{aligned} & -\frac{U_G(t+\Delta, E_t+\Delta)}{U_B(\underline{t}(\alpha)|t+\Delta)}\Pi_B(\underline{t}(\alpha)|t+\Delta) \\ & \geq -\Pi_B(\underline{t}(\alpha)|t+\Delta) + \frac{\varepsilon}{c}[U_G(t+\Delta, E_t+\Delta) - U_B(t+\Delta, E_t+\Delta)], \end{aligned}$$

which simplifies to

$$c[(1-\alpha)Y_L + \varepsilon(\underline{t}(\alpha) - t - \Delta)] + \varepsilon[\alpha Y_L - u(1-t-\Delta) - c(\underline{t}(\alpha) - t - \Delta)] \geq 0,$$

or equivalently $c[(1-\alpha)Y_L] + \varepsilon[\alpha Y_L - u(1-t-\Delta)] \geq 0$, which is strictly satisfied since $c > \varepsilon$ and $Y_L > u$. $\square$

PROOF OF PROPOSITION 4. Fix Δ, t, E_t, μ_t . Define

$$\kappa(\Delta, t, E_t, \mu_t) = \mathbf{cav}\Pi^*(t+\Delta, E_t, \mu_t) - \Pi^*(t, E_t, \mu_t),$$

where $\mathbf{cav}\Pi^*(t+\Delta, E_t, \mu_t)$ is the concavification of $\Pi^*(t+\Delta, E_t, \mu_t)$ as a function of μ_t . Then define $\bar{\kappa}(\Delta)$ by $\bar{\kappa}(\Delta) = \sup_{t, E_t, \mu_t} \kappa(\Delta, t, E_t, \mu_t)$.

First, fix Δ . I show that if $\kappa > \kappa(\Delta)$, then at any history, the principal never reveals information without the agent's effort. The proof is by induction.

Consider $t = 1 - 2\Delta$. Since the next period $t = 1 - \Delta$ is the last decision point for the agent, the principal would not have any incentive to reveal information at $t = 1 - \Delta$, and thus, the payoffs of both the principal and the agent are as in the proof of Lemma 7. Suppose at $t = 1 - 2\Delta$ the principal does not induce effort but gathers information. Since the agent does not work at $t = 1 - 2\Delta$, it must be that $E_{1-\Delta} = E_{1-2\Delta}$. Thus, the principal's maximum payoff from information design without the agent's effort is $\mathbf{cav}\Pi^*(1-\Delta, E_{1-2\Delta}, \mu_{1-2\Delta})$. Instead, if the principal chooses not to provide information without the agent's effort, his payoff is $\Pi^*(1-\Delta, E_{1-2\Delta}, \mu_{1-2\Delta})$. Thus, the principal's maximum gain from revealing information without the agent's effort is

$$\mathbf{cav}\Pi^*(1-\Delta, E_{1-2\Delta}, \mu_{1-2\Delta}) - \Pi^*(1-\Delta, E_{1-2\Delta}, \mu_{1-2\Delta}) \leq \kappa(\Delta) < \kappa,$$

where the inequality follows by the construction of $\kappa(\Delta)$ and choice of κ . Thus, at $1 - 2\Delta$ the principal chooses not to gather information without the agent's effort.

Fix t and assume that at any history with $t' \geq t + \Delta$, the principal chooses not to gather information without the agent's effort. Then, if at t the principal chooses to gather information without the agent's effort, his maximum payoff is given by $\mathbf{cav}\Pi^*(t+\Delta, E_t, \mu_t)$ while his maximum payoff if he chooses not to do so is $\Pi^*(t, E_t, \mu_t)$. Thus, his

gain from revealing information without the agent's effort is smaller than the cost κ , and thus such choice cannot be optimal. This completes the induction step.

Next, I argue that $\lim_{\Delta \rightarrow 0} \kappa(\Delta) = 0$. For this purpose, note that

$$\lim_{\Delta \rightarrow 0} \Pi^*(t + \Delta, E_t, \mu_t) - \Pi^*(t, E_t, \mu_t) = 0. \quad (30)$$

Further, since $\Pi^*(t + \Delta, E_t, \mu_t)$ is concave over $[\underline{\mu}_t^\Delta(E_t)]$ and $\lim_{\Delta \rightarrow 0} \underline{\mu}_t^\Delta(E_t) = 0$, we have that

$$\lim_{\Delta \rightarrow 0} \text{cav} \Pi^*(t + \Delta, E_t, \mu_t) - \Pi^*(t + \Delta, E_t, \mu_t) = 0. \quad (31)$$

Combining (30) and (31) establishes the claim. $\square$

REFERENCES

- Au, Pak Hung (2015), "Dynamic information disclosure." *The RAND Journal of Economics*, 46, 791–823. ISSN 07416261. <http://www.jstor.org/stable/43895617>. [673]
- Aumann, Robert J. and Michael Maschler (1995), *Repeated Games With Incomplete Information*. MIT Press. [673]
- Ball, Ian (2022), "Dynamic information provision: Rewarding the past and guiding the future." Working Paper. [673]
- Bizzotto, Jacopo, Jesper Rudiger, and Adrien Vigier (2021), "Dynamic persuasion with outside information." *American Economic Journal: Microeconomics*, 184. [672]
- Che, Yeon-Koo, Kyungmin Kim, and Konrad Mierendorff (forthcoming), "Keeping the listener engaged: A dynamic model of Bayesian persuasion." *Journal of Political Economy*. [673]
- Doval, Laura and Vasiliki Skreta (2018), "Constrained information design: Toolkit." Working Paper. [673]
- Ely, Jeffrey C. (2017), "Beeps." *American Economic Review*, 107, 31–53, doi: 10.1257/aer.20150218. [673]
- Ely, Jeffrey C. and Martin Szydlowski (2020), "Moving the goalposts." *Journal of Political Economy*, 128, 468–506, doi: 10.1086/704387. [672, 673]
- Guo, Yingni and Eran Shmaya (2018), "The value of multistage persuasion." Working Paper. [673]
- Hörner, Johannes and Andrzej Skrzypacz (2016), "Selling information." *Journal of Political Economy*, 124, 1515–1562, doi: 10.1086/688874. [673]
- Kamenica, Emir and Matthew Gentzkow (2011), "Bayesian persuasion." *American Economic Review*, 101, 2590–2615, doi: 10.1257/aer.101.6.2590. [673, 674]
- Le Treust, Mael and Tristan Tomala (2019), "Persuasion with limited communication capacity." *Journal of Economic Theory*, 184. [673]

- Li, Cheng (2017), “A model of Bayesian persuasion with transfers.” *Economics Letters*, 161, 93–95. [673]
- Liu, Chang (2021), “Motivating effort with information about future rewards.” Working Paper. [673]
- Lizzeri, Alessandro, Margaret Meyer, and Nicola Persico (2002), “The incentive effects of interim performance evaluations.” Penn CARESS Working Papers, Penn Economics Department. <https://EconPapers.repec.org/RePEc:cla:penntw:592e9328faf6e775bf331e1c08707dd2>. [674]
- Orlov, Dmitry (2022), “Frequent monitoring in dynamic contracts.” *Journal of Economic Theory*, 206, 105550, ISSN 0022-0531, doi: 10.1016/j.jet.2022.105550. [674]
- Orlov, Dmitry, Andrzej Skrzypacz, and Pavel Zryumov (2020), “Persuading the principal to wait.” *Journal of Political Economy*, 128, 2542–2578, doi: 10.1086/706687. [672, 673]
- Renault, Jérôme, Eilon Solan, and Nicolas Vieille (2017), “Optimal dynamic information provision.” *Games and Economic Behavior*, 104, 329–349. [673]
- Smolin, Alex (2021), “Dynamic evaluation design.” *American Economic Journal: Microeconomics*, 13, 300–331, doi: 10.1257/mic.20170405. [673]
- Watson, Joel (2017), “A general, practicable definition of perfect Bayesian equilibrium.” Working Paper. [676]

Co-editor Marina Halac handled this manuscript.

Manuscript received 19 August, 2020; final version accepted 16 June, 2022; available online 23 June, 2022.