

Financing Breakthroughs under Failure Risk*

Simon Mayer[†]

January 7, 2021

Abstract

In a dynamic principal-agent model, the principal, financing the project, cannot observe project failure and the agent, developing the project, can hide or fake failure. Punishments for completion delays, excessive rewards for success, and time-decreasing rewards for failure provide incentives for truthful disclosure of failure. The optimal contract does not always incentivize disclosure of failure and consists of distinct financing stages, whereby financing becomes more performance sensitive over time. Incentives are backloaded and the agent is rewarded for late but not for early failure. The optimal contract can be implemented by financing the project with a mixture of (convertible) debt and equity in multiple financing rounds.

*I especially thank Sebastian Gryglewicz and Erwan Morellec for their helpful comments. Parts of this paper were written during a research visit at the Kellogg School of Management. I am particularly grateful to Konstantin Milbradt for hosting and advising me. I also would like to thank Ron Kaniel (the editor), an anonymous referee, Andres Almazan, Philip Bond, Simon Board, Barney Hartman-Glaser, Florian Hoffmann, Mike Fishman, Sivan Frenkel, William Fuchs, Lukas Kremens, Thomas Geelen (EFA Discussant), Sebastian Pfeil, Semyon Malamud, Andrey Malenko, William Mann, Ramana Nanda, Tomasz Sadzik, Amiyatosh Purnarandam, Yingjie Qi, Zara Sharif, Kathy Yuan, Vladimir Vladimirov, Felipe Varas, and Brian Waters for useful discussions and valuable comments. I also thank the seminar participants at the Erasmus University Rotterdam, the FTG European Summer Meeting 2019, the EFA Doctoral Tutorial 2019, the Econometric Society European Winter Meeting 2019, University of Washington (Foster), National University of Singapore, Hong Kong University, City University of Hong Kong, and the Tinbergen Institute Jamboree for their comments and questions.

[†]Erasmus University Rotterdam. E-mail: mayer@ese.eur.nl.

After its founding in 2003, the US startup *Theranos* raised funds from venture capitalists and private investors, building on the promise of a novel method of blood testing. This resulted in a 10 billion dollar valuation at its peak in 2014. Between 2014 and 2018, the blood test technology developed by Theranos turned out to be inaccurate. However, instead of disclosing the technology's failure, the company issued false statements regarding the project's progress and continued to raise funds, until the pyramid of lies eventually collapsed in 2018.

The example of Theranos highlights key difficulties inherent to financing innovative projects. Such projects typically i) exhibit substantial failure risk, ii) require a high level of expertise from insiders developing the project, and iii) require capital from investors with limited expertise. To resolve agency conflicts between the insiders, developing the project, and investors, financing the project, the provision of financing must be contingent on project outcomes. However, when it is hard for investors to track project progress, insiders can hide bad outcomes, such as project failure. This paper studies dynamic contracting in light of this tension and characterizes the optimal incentive provision for truthful disclosure of bad outcomes.

In the model, a principal finances a project developed by an agent with limited liability. Project development requires funds from the principal and, absent frictions, it is efficient to finance the project until completion. The timing of completion is uncertain and for simplicity not affected by the agent. Completion results in either success or failure, whereby the agent's hidden effort during project development determines the likelihood of success. Moral hazard arises because the agent derives private benefits from shirking. Incentives are provided by paying the agent more for success than for failure. If both potential project outcomes, success and failure, are publicly observable and contractible, the principal pays the agent only for success and finances the project until completion. Because the timing of completion is uninformative about the agent's effort, the agent is not punished for completion delays and not fired before completion.

This paper studies incentive provision when it is hard for the principal to observe and verify project failure (whereas success is perfectly observable and contractible). Project failure is observed by the agent but possibly — that is, with some probability — not observed by the principal. Because it is efficient to terminate financing once the project fails, the principal would like the agent to disclose failure. However, the agent can hide or fake failure. Hiding failure averts termination and, therefore, allows the agent to continue enjoying private benefits from running the project. In addition, as failure and reports thereof are not verifiable by the principal, the agent can mis-report failure before it occurs, which is akin to faking failure or causing failure on purpose.

To incentivize disclosure of failure at the time it occurs, the contract stipulates rewards for failure. These rewards for failure must be time-decreasing, as otherwise the agent would delay disclosure and report failure at a later time than at which it occurred. Since both rewards for failure and success decrease over time, the agent is punished for completion delays beyond her influence. However, rewards for failure incentivize the agent to fake or cause (i.e., mis-report) failure. To prevent that the agent fakes or causes failure, it becomes necessary to increase the agent's stake in the project by raising rewards for success beyond what is needed to motivate effort, leading to excessive rewards for success and agency rents. That is, a tension arises between providing incentives to disclose and not to fake or cause failure.

As a result, the principal faces the following trade-off when designing the contract. On the one hand, financing a failed project is inefficient, so the principal ideally would like the agent to disclose failure and terminate financing upon failure. On the other hand, incentivizing disclosure of failure is costly, as it leads to excessive agency rents. In light of this trade-off, the optimal contract does not always incentivize disclosure of failure and consists of two distinct stages: an *unconditional financing stage* followed by a *disclosure stage*. During the unconditional financing stage, the principal does not incentivize disclosure of failure, which limits agency rents but comes at the cost that the project is potentially financed and inefficiently continued after failure. Moreover, the agent i) is not paid for failure, ii) receives low rewards for success, and iii) incurs mild punishments for delays. The unconditional financing stage ends with a soft deadline at which the principal elicits a truthful progress report from the agent on whether the project has failed so far. The principal finances the project over the next stage, i.e., the disclosure stage, if and only if the progress report reveals that the project is still profitable to pursue (and has not failed yet).

During the disclosure stage, the principal incentivizes disclosure of failure and finances the project until either completion is reported or a hard deadline is reached. Thus, financing is terminated upon failure and therefore performance-sensitive. During this stage, the agent receives high and time-decreasing rewards for failure and success and incurs harsh punishments for delays that include the threat of contract termination. The disclosure stage ends with a hard deadline at which the principal terminates financing regardless of whether the project is still profitable to pursue. In summary, within the optimal contract, financing is staged for pure incentive purposes, even though there is only a single milestone to complete the project. The provision of financing becomes more performance-sensitive over time and across stages. Likewise, the agent's incentives are backloaded, in that they become stronger over time and across stages.

The analysis of the optimal contract has implications for the design of venture capital and R&D financing contracts. The model predicts that optimal venture capital contracts involve distinct financing stages and that insiders' dollar rewards for success decrease during a given financing stage yet increase once a new financing stage begins. That is, insiders are effectively rewarded for reaching the next financing stage. In addition, the valuation of the project (i.e., startup valuation) jumps up at the beginning of a new financing stage, as new information is released and uncertainty is resolved through insiders' progress reports. Our findings also imply that within venture capital contracts, insiders' compensation and, generally, the provision of financing should become more performance sensitive over time and in each subsequent stage. Importantly, the provision of unconditional financing early on limits agency rents. As a result, optimal financing contracts for projects that are subject to severe agency conflicts involve a relatively long unconditional financing stage (and soft deadline) that is followed by a relatively short disclosure stage (and hard deadline).

Moreover, our results suggest that optimal financing of R&D projects involves several stages, whereby the continuation of financing is contingent on reported progress by insiders. Progress reports are less (more) frequent in early (later) stages of R&D financing. In addition, as progress reports are subject to moral hazard, it becomes optimal to elicit progress reports less frequently when moral hazard is severe. Compensation contracts for R&D workers (insiders) stipulate back-loaded incentive schemes, so that incentives increase once a new financing stage begins. It is optimal to reward failure that occurs at a later stage, but not failure that occurs at an early stage.

We show how to implement the optimal contract by financing the project with a mixture of debt and equity in two financing rounds. During the first financing round at the beginning of the unconditional financing stage, the principal injects funds and receives in exchange both debt and equity claims in the project, while the agent retains the remaining equity. During the second financing round at the beginning of the disclosure stage, the principal injects additional funds and receives additional equity, reducing (diluting) the agent's equity stake. During the disclosure stage (but not during the unconditional financing stage), funds allocated to the project exceed the face value of debt, so termination of financing due to failure leads to full repayment of debt and dividend payouts to equity holders generating rewards for failure for the agent. Note that during the first financing round, funds are raised by issuing both debt and equity, whereas during the second financing round, funds are raised only by issuing equity. The implementation of the optimal contract therefore rationalizes the widespread use of venture debt (see, e.g., [González-Urbe and Mann \(2020\)](#)) and suggests that venture debt and equity financing are complementary, consistent

with the findings of [Davis, Morse, and Wang \(2020\)](#). Interestingly, our results imply that startup firms (should) rely on venture debt in early financing stages, while they (should) rely more on equity financing in later stages. Alternatively, the optimal contract can be implemented using convertible equity or convertible debt, as is common in venture capital financing. Note that even though the implementation of the optimal contract is generally not unique, under any implementation the principal’s stake in the project becomes less debt-like and more equity-like over time and across stages. At the same time, the agent’s equity stake is gradually diluted.

Next, we show that agency conflicts can induce over- or under-provision of financing (i.e., over- or under-investment) relative to the net present value (NPV) criterion. Here, the under-provision of financing (under-investment) refers to a situation, wherein agency conflicts hamper financing of a project with positive net present value, and the over-provision of financing (over-investment) refers to financing of projects with negative net present value. During the contract’s unconditional financing stage, the principal may inefficiently continue and finance a failed project with negative NPV, leading to over-provision of financing. In contrast, during the contract’s disclosure stage, punishments for delays may trigger premature termination of a project with positive NPV, leading to under-provision of financing. Over-provision of financing occurs when the project development phase is short or moral hazard is mild, and under-provision of financing occurs when the project development phase is long or moral hazard is severe. In the context of venture capital financing, the model predicts venture capital over-investment in projects that generate (preliminary) results quickly and under-investment in projects that do not generate (preliminary) results quickly.

In a model extension, we study how moral hazard affects investors’ (i.e., the principal’s) endogenous project choice at inception at time zero. Interestingly, we find that under certain circumstances, riskier projects are subject to less severe agency problems and thus are preferred by investors. The reason is that for less risky projects, it is more difficult to incentivize the agent to truthfully disclose bad outcomes (i.e., failure), which exacerbates moral hazard. Thus, our model offers an explanation why venture capitalists seek to finance high risk start ups, i.e., potential *unicorns*, even if this choice is not supported by the net present value criterion.¹

Last, we analyze the role of monitoring in incentive provision. For this purpose, we assume that the principal can inspect the project’s progress at a cost. Upon inspection, the principal learns whether the project has failed so far. That is, an inspection enables the principal to detect

¹A unicorn is a privately held startup company with a valuation that exceeds one billion dollars. The most recent and largest unicorns include Airbnb, Uber and Pinterest.

whether the agent hides (or fakes) bad outcomes and allows punishment of that misbehavior with termination. As a result, the optimal contract (with monitoring) features several financing stages and periodic inspections at deterministic dates. During each financing stage, rewards for failure and success decrease over time until the project is completed or a deterministic (soft) deadline is reached. At this deadline, the principal inspects the project and grants financing for the next stage if the inspection reveals that the project is still profitable to pursue (i.e., has not failed so far). Once a new financing stage starts, stipulated pay for success and failure increase, in that the agent is rewarded for positive inspection outcomes. Optimal dynamic monitoring i) precludes premature termination of financing, ii) reduces rewards for failure and success, and iii) makes financing more performance sensitive and efficient. In the context of venture capital financing, our results imply that optimal inspections complement reports by insiders in determining project progress and whether to grant financing over the next stage.

This paper builds on the literature that studies dynamic contracting in continuous time, starting with [DeMarzo and Sannikov \(2006\)](#), [Biais, Mariotti, Plantin, and Rochet \(2007\)](#), and [Sannikov \(2008\)](#). Recent contributions include [He \(2009, 2011, 2012\)](#), [Hoffmann and Pfeil \(2010\)](#), [Biais, Mariotti, Rochet, and Villeneuve \(2010\)](#), [DeMarzo, Fishman, He, and Wang \(2012\)](#), [Hartman-Glaser, Piskorski, and Tchistyi \(2012\)](#), [Zhu \(2012\)](#), [DeMarzo and Sannikov \(2016\)](#), [Marinovic and Varas \(2019\)](#), [Gryglewicz, Mayer, and Morellec \(2019\)](#), and [Hoffmann, Inderst, and Opp \(2019\)](#). Relatedly, [DeMarzo and Fishman \(2007b,a\)](#) study dynamic financial contracting in a discrete time setting. [Piskorski and Westerfield \(2016\)](#), [Orlov \(2018\)](#), and [Malenko \(2019\)](#) analyze incentive provision with optimal dynamic contracts and monitoring. Likewise, [Halac and Prat \(2016\)](#), [Varas, Marinovic, and Skrzypacz \(2020\)](#), and [Marinovic and Szydlowski \(2019\)](#) characterize optimal monitoring in dynamic settings but do not focus on optimal contracts. None of these papers studies optimal incentive provision for disclosure of bad outcomes. The innovation in our model is that it considers the risk of project failure when the agent can hide or fake failure.

Our framework is similar to [Green and Taylor \(2016\)](#) and [Varas \(2017\)](#). [Green and Taylor \(2016\)](#) study contracting for a multistage project when intermediate progress is privately observed. They derive an optimal contract that involves a period during which financing is guaranteed whereby the agent is rewarded for good but never for bad outcomes. [Varas \(2017\)](#) studies managerial short-termism in a dynamic model and finds that the threat of termination and deferred payouts for success provide efficient incentives. Our analysis differs from these papers along several dimensions. First, in our model project development may lead to two types of adverse outcomes, failure and

completion delays, whereas in [Green and Taylor \(2016\)](#) and [Varas \(2017\)](#) only completion delays may arise. Second, our paper studies incentives to truthfully disclose, i.e., not to hide or to fake, bad outcomes. In [Green and Taylor \(2016\)](#) and [Varas \(2017\)](#) the moral hazard problem is different in that the agent is tempted to fake but never tempted to hide good outcomes.

Our paper builds on the literature on optimal incentive schemes for venture capital, such as [Bergemann and Hege \(1998, 2005\)](#), or for financing innovation in a broader sense, such as [Holmstrom \(1989\)](#), [Aghion and Tirole \(1994\)](#), and [Nanda and Rhodes-Kropf \(2017\)](#). [Hall and Lerner \(2010\)](#), [Kerr, Nanda, and Rhodes-Kropf \(2014\)](#), and [Kerr and Nanda \(2015\)](#) provide excellent reviews of the academic literature on financing entrepreneurship and innovation, and [Lerner and Nanda \(2020\)](#) discuss the role of venture capitalists in financing innovation. Conceptually, our paper is related to [Manso \(2011\)](#), who studies incentive provision in a static multi-armed bandit setting in which the agent can either exploit existing technologies or explore new technologies through experimentation. Our model is dynamic and differs from [Manso \(2011\)](#) mainly in the following aspects. First, we study incentive provision for information disclosure rather than for experimentation. Second, we demonstrate that to provide incentives it can be sometimes optimal to reward failure but it is always necessary to penalize delays; delays do not arise in [Manso \(2011\)](#). Third, our dynamic model yields the prediction that it is optimal to reward failure that occurs at a later stage, but not failure that occurs at an early stage.

Lastly, our paper is related to the literature on dynamic adverse selection — see e.g. [Daley and Green \(2012, 2016, 2019\)](#), [Morellec and Schürhoff \(2011\)](#), and [Grenadier and Malenko \(2011\)](#) — and to the literature on dynamic information disclosure (without optimal contracts) — see e.g. [Jovanovic \(1982\)](#), [Acharya, DeMarzo, and Kremer \(2011\)](#), [Guttman, Kremer, and Skrzypacz \(2014\)](#), [Marinovic and Varas \(2016\)](#), [Bertomeu, Marinovic, Terry, and Varas \(2019\)](#), and [Halac and Kremer \(2019\)](#).²

The paper is structured as follows. Section 1 presents the model and discusses the contracting problem. Section 2 solves the model and derives the optimal contract. We discuss the implications of our analysis in Sections 3 and 4. Section 5 studies incentive provision with optimal dynamic monitoring. Section 6 discusses the robustness of our model and provides several extensions. Section 7 concludes the paper. All technical developments are gathered in the Appendix.

²Papers on optimal information disclosure in static settings include [Fishman and Hagerty \(1989, 1990\)](#), [Diamond and Verrecchia \(1991\)](#) and, more recently, [Szydlowski \(2019\)](#). For a literature review on this topic, we refer to [Beyer, Cohen, Lys, and Walther \(2010\)](#).

1 Model setup

Time t is continuous and defined over $[0, \infty)$. A principal (she) finances a project that is developed by an agent (he) with limited liability and zero wealth. The principal has deep pockets, can fully commit to any long term contract, and possesses all bargaining power when signing a contract with the agent. Both the principal and the agent are risk neutral, do not discount payoffs, and have an outside option equal to zero. Per unit of time, project development requires funds $\kappa > 0$ from the principal. The principal can always terminate financing and project development and does so at an endogenous time $T_0 \in [0, \infty]$.

The project completion time τ is uncertain and arrives according to a Poisson process N_t , in that $\tau = \inf\{t \geq 0 : dN_t = 1\}$. The intensity of N_t equals $\Lambda > 0$ for $t \leq T_0$ and zero for $t > T_0$, meaning that the project cannot be completed anymore after financing is terminated. This specification implies that the expected time to completion equals $\frac{1}{\Lambda}$ provided the project receives sufficient financing. Project completion results in one of two possible outcomes, called *success* and *failure*. Upon completion at time τ , with probability $a_\tau p$, the project is successful and yields terminal payoff $\mu > \frac{2\kappa}{\Lambda p}$ to the principal. Otherwise, with probability $1 - a_\tau p$, the project fails and yields no terminal payoff. Here, $a_\tau \in \{0, 1\}$ is the agent's effort (just before time τ), and $p \in (0, 1)$ is exogenous.³ Notably, without frictions and costless effort, the project has net present value (NPV) $p\mu - \frac{\kappa}{\Lambda} > 0$ and it is efficient to finance the project until time τ .

Over a short period of time $[t, t + dt]$, the project succeeds with probability $a_t p \Lambda dt$, fails with probability $(1 - a_t p) \Lambda dt$, and does not complete with probability $1 - \Lambda dt$. Similar to [Board and Meyer-ter Vehn \(2013\)](#) and [Hoffmann and Pfeil \(2019\)](#), the agent chooses effort $a_t \in \{0, 1\}$ before the random event $dN_t \in \{0, 1\}$ realizes over $[t, t + dt]$. Figure 1 illustrates the heuristic timing over $[t, t + dt]$.

Project development is subject to the following frictions. First, project failure is hard for the principal to observe or verify. When the project fails, failure is *publicly* observed (i.e., observed by both principal and agent) with probability $\pi \in [0, 1]$. Publicly observed failure is contractible. Otherwise, with probability $1 - \pi$, failure is *privately* observed by the agent but not observed by the principal, who also cannot verify failure. For instance, $\pi \in (0, 1)$ captures the fact that the insiders developing the project can hide certain bad outcomes from outside investors, while it is difficult or

³The terminal payoff can be interpreted broadly and, for instance, can represent immediate monetary payoffs, future expected cash flows, or the principal's private value of a technological achievement.

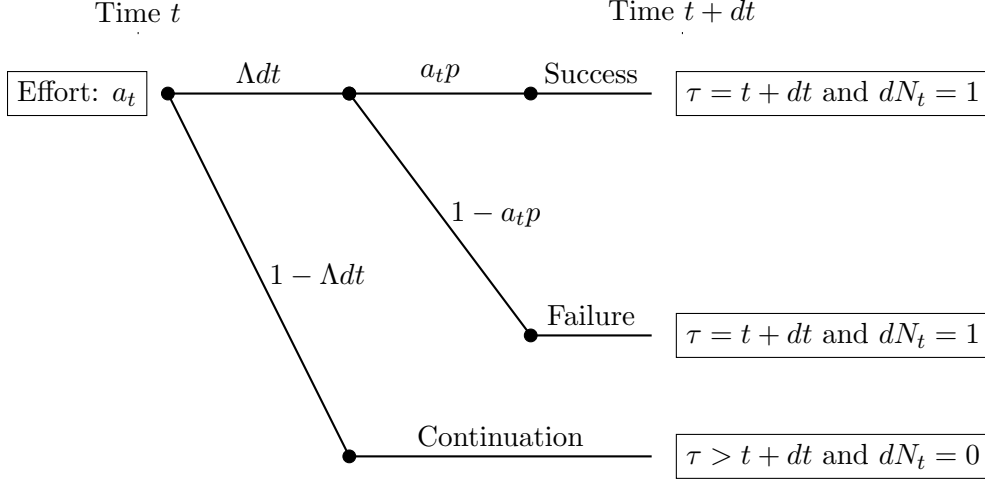


Figure 1: Heuristic Timing over $[t, t + dt)$. The branches of the tree contain the probabilities of the respective random event over $[t, t + dt)$.

even impossible to hide all types of bad outcomes.⁴ The agent may report (i.e., disclose) failure, but the principal cannot verify the reported failure. As a result, the agent can hide or fake failure. However, because it is efficient to terminate financing after failure at time τ , the principal would ideally like the agent to disclose failure truthfully. Unlike failure, success is publicly observable, verifiable, and contractible. This assumption reflects that the agent is tempted to hide bad rather than good outcomes and that it is harder to fake good rather than bad outcomes.

Second, effort a_t is not observed by the principal and the agent derives private benefits from shirking $(1 - a_t)\phi$ for $t \leq T_0$, giving rise to moral hazard. After financing is terminated, private benefits are zero. Even though not explicitly modeled here, private benefits (from shirking) may arise from the inefficient diversion of funds, as e.g. in [DeMarzo and Sannikov \(2006\)](#), and, therefore, pertain as long as the principal allocates funds to the project (i.e., only for $t \leq T_0$). Crucially, effort after time τ is redundant so that the agent optimally chooses $a_t = 0$ for $t \in (\tau, T_0]$. This leads to private benefits ϕ from operating the project after time τ , which only pertain if the agent hides project failure and averts termination after time τ . We assume $0 \leq \phi \leq \kappa$, implying that shirking is inefficient.

Last, certain modeling assumptions — e.g., observable success and exogenous completion — and the robustness of our findings are discussed in [Section 6.2](#).

⁴In the context of venture capital financing, $\pi > 0$ may capture that the principal, representing the venture capitalist, has some ability to detect bad outcomes in project development, e.g., due to expertise or oversight arising from the VC's representation on the board of the startup firm (compare [Lerner \(1995\)](#)).

Contracting problem

At time zero, the principal offers the agent a contract $\mathcal{C} = (c, T)$ specifying cumulative payments (i.e., wages) c and a deadline T . At the deadline, project development and financing are terminated. For notational convenience, payments c do not include project development costs κ , which are paid by the principal up to time T_0 . Facing the contract $\mathcal{C}$, the agent chooses effort a and time τ^A when he reports completion. This implies that the principal finances the project until either completion is reported or the deadline is reached. That is: $T_0 = T \wedge \tau^A$. We focus on contracts $\mathcal{C}$ that induce full effort $a_t = 1$ at all times $t \in [0, T \wedge \tau]$. To solve the model, we conjecture that the (optimal) deadline T is deterministic. Proposition 7 in Appendix B.4.5 verifies that the (optimal) deadline is indeed deterministic, so random termination does not improve the principal's payoff.⁵

We characterize the agent's strategy using the time at which he reports completion (rather than using the time at which he reports failure). Because success is publicly observable, this notation implies for the case of success $\tau^A = \tau$. Likewise, we adopt the notation $\tau^A = \tau$ when failure is publicly observed. Conversely, $\tau^A \neq \tau$ means that the agent mis-reports failure. In particular, the agent can mis-report failure in two ways. First, he can *hide failure* (i.e., bad outcomes), in which case $\tau^A > \tau$. Second, because reported failure is not verifiable, the agent can *fake failure* and report failure even if it has not occurred, in which case $\tau^A < \tau$. Alternatively, one can interpret "faking bad outcomes" also as "generating bad outcomes (on purpose)."⁶

As the principal and the agent do not discount, it is optimal to pay the agent only at time τ^A , when reported success or failure is informative about the agent's effort. Because the completion timing as such is not informative about effort, payments before time τ^A do not help to incentivize effort. In fact, they even render it attractive to hide failure and so generate dis-incentives. Thus

$$dc_t = \alpha_t \mathbf{1}_{\{\text{Success at time } t\}} + \beta_t \mathbf{1}_{\{\text{Failure report at time } t\}} + \gamma_t \mathbf{1}_{\{\text{Failure observed at time } t\}} \quad (1)$$

where α_t is the payment to the agent if the project succeeds at time t (in which case $\tau = \tau^A = t$). The agent's payment for failure depends on whether the principal observes failure. Specifically, β_t is the agent's pay when he reports failure at time t , while γ_t is the agent's pay when failure is publicly observed at time t . As will become clear later, payments for *privately* observed and reported failure $\beta_t > 0$ are necessary to incentivize disclosure of failure. In contrast, payments for

⁵I would like to thank an anonymous referee for encouraging me to add Proposition 7 and its formal proof.

⁶This is akin to assuming that the agent can destroy payoff, which is a frequent assumption in the financial contracting literature (see e.g. Innes (1990)).

publicly observed failure are not necessary to incentivize disclosure of failure and, in addition, do not motivate effort. Hence, we conjecture and verify that it is optimal not to pay the agent for observed failure, in that $\gamma_t = 0$.

The agent maximizes his payoff, stemming from wage payments and private benefits:

$$W_0 := \max_{a, \tau^A} \mathbb{E}^A \left[\int_0^{T \wedge \tau^A} (dc_t + \phi(1 - a_t)dt) \right]. \quad (2)$$

Here, superscript A indicates that the expectation is taken to be conditional on the agent's information, which may differ from the principal's information. The principal chooses contract $\mathcal{C}$ to maximize expected project payoff net of the cost of financing project development and compensating the agent:

$$F_0 := \max_{\mathcal{C}} \mathbb{E} \left[\int_0^{T \wedge \tau^A} (\mu dS_t - \kappa dt - dc_t) \right]. \quad (3)$$

where the expectation is conditional on the principal's information. Here, $dS_t = 1$ indicates project success at time t . Because the agent is protected by limited liability, it follows that $dc_t \geq 0$, i.e., $\alpha_t, \beta_t, \gamma_t \geq 0$, for all $t \geq 0$.

2 Model solution

We solve the model in several steps. First, to provide a starting point for our analysis, we analyze the benchmark, in which project failure is publicly observable (i.e., $\pi = 1$) and contractible, yet the moral hazard problem with respect to hidden effort remains. Second, we focus on *full disclosure contracts* that always incentivize the agent to disclose failure truthfully. Third, we argue why full disclosure contracts are in general not optimal and show how they can be improved through the provision of *unconditional financing*.

2.1 Second-best benchmark

We start by analyzing the second-best benchmark, in which project failure is observable and contractible in that $\pi = 1$. With observable failure, the optimal contract takes a simple form: there is no deadline (i.e., $T = \infty$) and the agent is only paid for success. That is, $\alpha^{SB} = \frac{\phi}{\Lambda p} > \beta^{SB} = \gamma^{SB} = 0$.

The intuition behind this result is as follows. First, note that the agent controls the project's propensity to succeed at time τ but not the completion timing τ . Therefore, the completion timing τ as such is not informative about the agent's effort, implying that the agent's compensation should

not be contingent on τ . Thus, the agent is not punished for completion delays and not fired before project completion.

Second, to motivate effort, it is necessary to pay the agent more for success than for failure, leading to the incentive constraint

$$\alpha_t - \gamma_t \geq \frac{\phi}{\Lambda p}. \quad (4)$$

To derive (4), suppose the agent shirks over a short period of time $[t, t + dt)$. Then, the agent derives private benefits ϕdt and the project completes with probability Λdt , resulting in failure and pay γ_t . In contrast, if the agent exerts effort $a_t = 1$ over $[t, t + dt)$, the agent does not derive any private benefits and the project completes with probability Λdt , resulting in failure and pay γ_t with probability $1 - p$ and in success and pay α_t with probability p . Thus, exerting effort over $[t, t + dt)$ is optimal if and only if

$$\underbrace{((1 - p)\gamma_t + p\alpha_t)\Lambda dt}_{\text{Payoff with } a_t=1} \geq \underbrace{\phi dt + \gamma_t \Lambda dt}_{\text{Payoff with } a_t=0},$$

which simplifies to (4).

Note that because of $\pi = 1$, the choice of β becomes redundant. Incentives are captured by (net) rewards for success $\alpha_t - \gamma_t$. Payments for observed failure motivate shirking and hence are optimally set to $\gamma_t = 0$. In addition, (4) is tight to minimize agency costs. Incentives $\alpha_t - \gamma_t$ must be stronger, if moral hazard is more severe. Moral hazard is severe when private benefits from shirking ϕ or the expected duration of the project development phase $1/\Lambda$ are large.

Proposition 1. *Suppose that failure is publicly observable and contractible, in that $\pi = 1$. Then, the optimal contract $\mathcal{C}^{SB}$ is stationary with $T^{SB} = \infty$ and $\alpha^{SB} = \frac{\phi}{\Lambda p} > \beta^{SB} = \gamma^{SB} = 0$. The principal's payoff equals $F^{SB} = p\mu - \frac{\phi + \kappa}{\Lambda}$. The agent's payoff equals $W^{SB} = \frac{\phi}{\Lambda}$.*

With this clean benchmark in hand, we now analyze the problem with $\pi \in [0, 1)$.

2.2 Full disclosure contracts

Because it is efficient to terminate financing at time τ , it is natural to study *full disclosure contracts* that incentivize truthful disclosure of failure. That is, a full disclosure contract induces $\tau^A = \tau$.

2.2.1 Truth telling incentives

We start by characterizing the agent's incentives to disclose failure truthfully. Note that $\beta_t(1 - \pi)$ is the agent's expected payment upon failure at time t , which is proportional to β_t . We thus refer

to β_t simply as “rewards for failure” rather than “rewards for reported failure.” Define the agent’s continuation payoff under truthful reporting (i.e., $\tau^A = \tau$) and full effort as

$$W_t = \mathbb{E}_t \left[\int_t^{T \wedge \tau} dc_s \right] = \int_t^T e^{-\Lambda(s-t)} \Lambda((1-p)(1-\pi)\beta_s + p\alpha_s) ds, \quad (5)$$

where the (second) equality uses integration by parts, and recall that $\gamma_s = 0$. Hence:

$$\dot{W}_t = \Lambda W_t - \Lambda((1-p)(1-\pi)\beta_t + p\alpha_t), \quad (6)$$

where “dot” denotes the time derivative (i.e., $\dot{W}_t = \frac{dW_t}{dt}$). Hence, $\dot{W}_t < 0$ means that the agent is *punished for delays* (beyond his influence).

First, note that the agent can always fake bad outcomes and mis-report failure before it actually occurs. This leads to a loss of the continuation value under truth telling W_t and a reward for failure β_t and, therefore, is sub-optimal if and only if

$$W_t \geq \beta_t. \quad (7)$$

That is, to provide incentives not to fake bad outcomes, it is necessary to grant the agent a sufficiently large stake W_t in the project. Note that the optimal contract always incentivizes the agent not to fake failure. If the agent finds it optimal to fake failure at some time t with $\beta_t > 0$, the agent reports failure and hence precipitates termination at time t , while being paid $\beta_t > 0$. Then, the principal could improve her payoff by stipulating termination with zero severance pay for the agent at time t .

Second, let us analyze the agent’s incentives to hide failure. To obtain some intuition, suppose the project fails at time t and failure is privately observed by the agent. Reporting failure at time t yields pay β_t . In contrast, reporting failure at time $t + dt$ not only yields pay β_{t+dt} but also benefits from operating the project over $[t, t + dt]$, ϕdt . Hence, the agent is better off disclosing failure truthfully at time t if and only if

$$\beta_t \geq \beta_{t+dt} + \phi dt.$$

Letting $dt \rightarrow 0$ yields

$$\dot{\beta}_t \leq -\phi. \quad (8)$$

Because failure can potentially occur at any time $t \in [0, T]$, a full disclosure contract must satisfy (8) for all $t \in [0, T]$. Thus, we can integrate (8) to obtain

$$\beta_t \geq \beta_s + (s - t)\phi \quad \text{for all } s \in [t, T]. \quad (9)$$

Indeed, it follows that the agent prefers to report failure at time t rather than at any time $s > t$.

Importantly, the agent's limited liability requires $\beta_t \geq 0$. Once $\beta_t = 0$, the contract cannot provide truth telling incentives anymore by decreasing β_t . Under these circumstances, the only way to ensure truth telling is contract termination, in that $T = \inf\{t \geq 0 : \beta_t = 0\}$. Contract termination also implies that the agent never profits from hiding failure completely. Hiding failure (that occurs at time t) completely allows the agent to derive private benefits up to time T , i.e., $\phi(T - t)$, which by (9) is smaller than the reward for failure β_t .

In the following, it is convenient to keep track of the agent's (off-equilibrium) continuation payoff w_t , in case he has privately observed failure at some time t' with $t' \leq t$ but not reported it yet:

$$w_t := \max_{\tau^A \in [t, T]} [\phi(\tau^A - t) + \beta_{\tau^A}]. \quad (10)$$

Note that a full disclosure contract implies for any $t \in [0, T]$ that w_t is maximized for $\tau^A = t$, leading to $\beta_t = w_t$ and $\dot{w}_t = \dot{\beta}_t$ for all $t \in [0, T]$. With the agent's expected pay for failure $(1 - \pi)w_t$, we obtain the incentive condition w.r.t. effort a_t :

$$\alpha_t \geq (1 - \pi)w_t + \frac{\phi}{\Lambda p}. \quad (11)$$

The derivation of condition (11) is analogous to the derivation of condition (4) upon replacing γ_t with $(1 - \pi)w_t$.

2.2.2 Solving for the optimal full disclosure contract

Using integration by parts and $\tau^A = \tau$, we can rewrite the principal's continuation payoff for any $t \in [0, T \wedge \tau^A]$ as

$$F_t = \mathbb{E} \left[\int_t^{T \wedge \tau^A} (\mu dS_s - \kappa ds - dc_s) \right] = \int_t^T e^{-\Lambda(s-t)} \left(\Lambda p \mu - \kappa - \Lambda((1-p)(1-\pi)w_s + p\alpha_s) \right) ds. \quad (12)$$

The optimal full disclosure contract maximizes (12) subject to the incentive constraint w.r.t. effort (11) and the incentive constraints w.r.t. truthful information disclosure (7) and (8) for all $t \in [0, T]$. The incentive condition (8) constrains the control of the level of β (or equivalently w) and thus β (or equivalently w) enters the principal's dynamic optimization problem as state variable, while the change in β or w (i.e., $\dot{\beta}$ or $\dot{w}$) is a control variable.

To minimize agency rents, the incentive condition (7) is tight, in that $W_t = w_t = \beta_t$ for all $t \in [0, T]$ and hence the principal's optimization features a single state variable w . This implies that one can express the principal's payoff as function of w , $F(w)$. Differentiating (12) with respect to time t and using $\frac{dF_t}{dt} = \frac{dF(w_t)}{dw_t} \frac{dw_t}{dt} = F'(w_t)\dot{w}_t$, we obtain that the principal's value function solves the following HJB equation on the endogenous state space $[0, w_0]$:

$$\Lambda F(w) = \max_{\dot{w}, \alpha} \left\{ \Lambda p \mu - \kappa - \Lambda((1-p)(1-\pi)w + p\alpha) + F'(w)\dot{w} \right\} \text{ s.t. (7), (8), and (11). } \quad (13)$$

Some observations are in order. First, because it is always possible to reduce the agent's pay for reported failure from w to some level $\hat{w} < w$, it must be that $F(\hat{w}) \leq F(w)$, implying that $F'(w) \geq 0$. Then, the maximization w.r.t. $\dot{w}$ yields $\dot{w} = -\phi < 0$ so that condition (8) is binding. Second, because of the agent's limited liability, the contract is terminated once $w = \beta = 0$, yielding $f(0) = W(0) = 0$. Third, the value function is strictly concave, reflecting that termination is inefficient. The concavity of the value function also implies that randomized termination of a full disclosure contract is not optimal (for details and a more general statement see Proposition 7).

Importantly, rewards for success $\{\alpha_t\}$ determine the size of the agent's stake in the project W_t and hence his incentives to fake bad outcomes (see (7)). To minimize agency costs, (7) is tight and $W_t = w_t = \beta_t$ for all $t \in [0, T]$, leading to⁷

$$\alpha_t = (1-\pi)w_t + \frac{\phi}{\Lambda p} + \frac{\pi}{p}w_t = \left(1-\pi + \frac{\pi}{p}\right)w_t + \frac{\phi}{\Lambda p}. \quad (14)$$

Thus, the incentive condition for effort (11) is slack, when $\pi > 0$.⁸ The reason is that the agent requires rewards $\beta > 0$ to disclose failure. Because rewards for failure provide the agent with incentives to fake failure, it becomes necessary to increase his stake in the project by raising

⁷For a derivation, note that $W_t = w_t = \beta_t$ for all $t \in [0, T]$ implies that $\dot{W}_t = \dot{w}_t$ for all $t \in [0, T]$. Hence, by (6) with $w_t = \beta_t$: $\dot{W}_t - \dot{w}_t = \Lambda(W_t - (1-\pi)w_t - p(\alpha_t - (1-\pi)w_t)) + \phi = \phi - \Lambda p(\alpha_t - (1-\pi)w_t) + \Lambda \pi w_t = 0$. Thus, $\alpha_t - (1-\pi)w_t = \phi/(\Lambda p) + \pi w_t/p$.

⁸Expected rewards for failure $\beta_t(1-\pi)$ increase in $1-\pi$ and higher (expected) rewards for failure require higher rewards for success to motivate effort. When $\pi = 0$, incentives for effort require particularly high rewards for success that are at the same time sufficient to incentivize the agent not to fake failure.

rewards for success beyond what is needed to motivate effort. Also observe that both rewards for success and failure decrease over time, inducing punishment for delays $\dot{W}_t = -\phi < 0$.

Plugging (14) and $\dot{w} = -\phi$ back into the HJB equation (13), one obtains a linear first order ODE that admits the closed form solution

$$F(w) = \left(\mu p - \frac{\kappa}{\Lambda}\right) \left(1 - \exp\left(\frac{-w\Lambda}{\phi}\right)\right) - w. \quad (15)$$

The starting value w_0 is chosen to maximize payoff at time zero, $F(w_0)$. Thus, w_0 solves the first-order optimality condition $F'(w_0) = 0$ so that

$$w_0 = \phi T = \frac{\phi}{\Lambda} \ln\left(\frac{\Lambda\mu p - \kappa}{\phi}\right). \quad (16)$$

We summarize our findings in the following proposition.

Proposition 2. *Under the optimal full disclosure contract $\mathcal{C}$, at time t with $w_t = w$, the principal's value $F(w)$ solves (13). The contract $\mathcal{C}$ stipulates $\dot{w}_t = -\phi$ and termination at time $T = \inf\{t \geq 0 : w_t = 0\}$. Payments satisfy (14), $\beta_t = w_t = W_t$, and $\gamma_t = 0$. The value w_0 solves $F'(w_0) = 0$.*

The optimal full disclosure contract differs from the second-best contract mainly in three aspects: i) rewards for failure, ii) excessive rewards for success beyond what is needed to motivate effort, and iii) punishments for delays, including the threat of contract termination. Rewards for failure and punishments for delays incentivize the agent not to *hide* bad outcomes, whereas sufficiently high rewards for success incentivize the agent not to *fake* bad outcomes.

The optimal full disclosure contract is not unique. The reason is that rewards for observed failure boost the agent's stake in the project and hence generate incentives not to fake bad outcomes. As a result, there exists an optimal full disclosure contract that pays the agent for observed failure and pays the agent less for success. This contract maintains incentive compatibility but stipulates weaker incentives to exert effort. We focus without loss of generality on the optimal (full disclosure) contract that maximizes incentives. Notably, all optimal (full disclosure) contracts share the same key characteristics and stipulate punishments for delays, rewards for failure, and excess pay for success (i.e., $\alpha_t > w_t(1 - \pi) + \frac{\phi}{\Lambda p}$). We present the generalization of Proposition 2 in the Appendix in Proposition 5.

Notably, the value function $F(w)$ does not depend on π . The reason is that within a full disclosure contract, agency costs are independent of π . An increase in π makes it easier to incentivize

the agent not to hide failure, leading to lower rewards for failure $\beta_t(1 - \pi)$. This, however, reduces agency rents W_t , which generates incentives to fake failure and so requires higher rewards for success α_t . That is, a tension arises between incentivizing the agent not to hide and not to fake bad outcomes. In light of this tension, a full disclosure contract is not optimal and, therefore, the optimal contract does not always incentivize disclosure of failure.

2.3 The optimal contract

Suppose that the contract stipulates no rewards for failure over $[t_0, t_1)$, in that $\beta_t = \gamma_t = 0$ for $t \in [t_0, t_1)$. Consequently, the agent never reports and never fakes failure over $[t_0, t_1)$. As the contract does not incentivize disclosure of failure over $[t_0, t_1)$, it is also not terminated over $[t_0, t_1)$ due to disclosure of failure. In other words, the principal provides *unconditional financing* over $[t_0, t_1)$. Crucially, the provision of unconditional financing generates incentives not to fake failure (i.e., relaxes incentive constraint (7)) and hence limits agency rents but comes at the expense that the project may be continued and financed after failure, which is inefficient. Although a contract could stipulate unconditional financing over several distinct time intervals, we focus in the following discussion on the latest interval in time denoted by $[t_0, t_1)$. Thereafter, we verify that under the optimal contract there is only one (connected) time interval with unconditional financing.

2.3.1 Incentives

Consider that $\beta_t = 0$ for $t \in [t_0, t_1)$ and that the contract incentivizes truthful information disclosure from time t_1 onwards up to a deadline $T > t_1$. If the project fails at some time $t \in [t_0, t_1)$ and failure is privately observed by the agent, the agent does not report failure up to time t_1 , when he receives pay $\beta_{t_1} > 0$. Thus, the agent's continuation payoff after failure at time t is:

$$w_t := (t_1 - t)\phi + \beta_{t_1} \quad \text{for } t \in [t_0, t_1],$$

so that $\dot{w}_t = -\phi < 0$. To incentivize effort, the agent's payoff after success must sufficiently exceed his payoff after failure in that (11) holds, that is, $\alpha_t \geq (1 - \pi)w_t + \frac{\phi}{\Lambda p}$.

2.3.2 Unconditional financing

Next, we heuristically determine when it is optimal to provide unconditional financing. First, consider $t_0 = 0$. Over the time interval $[0, t_1)$, the agent merely requires incentives for effort and it

is optimal to provide minimal incentives, in that

$$\alpha_t = (1 - \pi)w_t + \frac{\phi}{\Lambda p} \quad \text{for all } t \in [0, t_1]. \quad (17)$$

That is, the provision of unconditional financing facilitates the stipulation of low rewards for success. Low rewards for success in turn reduce both punishments for delays and the agent's stake in the project, thereby reducing incentives and agency costs relative to a full disclosure contract. Formally, $W_t < w_t$ and $0 > \dot{W}_t > \dot{w}_t = -\phi$ for $t < t_1$. Overall, unconditional financing limits agency rents W_t but may lead to inefficient financing of a failed project.⁹

Second, we argue that unconditional financing starting from $t_0 > 0$ is sub-optimal. If $t_0 > 0$, the contract incentivizes truthful disclosure of failure just before time t_0 and just after time t_1 , which by (7) requires that $W_{t_i} \geq w_{t_i}$ for $i = 0, 1$. That is, the principal cannot reduce the agent's stake W_{t_0} (i.e., agency costs) by providing unconditional financing after time t_0 , while incentivizing truth telling before time t_0 . More intuitively, unconditional financing after time t_0 dilutes truth telling incentives before time t_0 . To avoid this inefficient dilution of incentives, it must be that $t_0 = 0$. Formally, at any time t with $W_t \geq w_t$, the optimal continuation contract is a full disclosure contract, as characterized in Proposition 2. It readily follows that there is maximally one connected time interval, during which the optimal contract stipulates unconditional financing.

As a result, the optimal contract stipulates unconditional financing over some time period $[0, t_1]$. After time t_1 , the optimal contract incentivizes truthful disclosure of failure and hence becomes a full disclosure (continuation) contract with deadline $T \geq t_1$, yielding payoff $F(w_{t_1})$ to the principal and payoff $w_L := w_{t_1} = W_{t_1}$ to the agent. Also note that $w_t = \beta_t = W_t$ for all $t \geq t_1$.

2.3.3 Solving for the optimal contract

The principal incentivizes disclosure of failure at time t_1 and forms a belief, q_t , of whether the project has failed so far, for times $t < t_1$. One can derive that:¹⁰

$$q_t = q(w_t) = 1 - e^{-\Lambda(1-p)(1-\pi)t} = 1 - e^{-\frac{\Lambda(1-p)(1-\pi)(w_0 - w_t)}{\phi}}. \quad (18)$$

⁹Note that $W_t = \int_t^T e^{-\Lambda(s-t)} \Lambda((1-p)(1-\pi)w_s + p\alpha_s) ds$. Hence: $\dot{W}_t = \Lambda W_t - \Lambda(w_t(1-\pi) + p(\alpha_t - w_t(1-\pi)))$ and $\dot{W}_t - \dot{w}_t = \Lambda W_t - \Lambda(w_t(1-\pi) + p(\alpha_t - w_t(1-\pi))) - \dot{w}_t = \Lambda(W_t - w_t) + \Lambda\pi w_t$, where we plugged in $\alpha_t = w_t(1-\pi) + \phi/(\Lambda p)$ and $\dot{w}_t = -\phi$ for $t < t_1$. Integrating this ODE for $t < t_1$ subject to $W_{t_1} = w_{t_1}$ yields $W_t - w_t = -\int_t^{t_1} e^{-\Lambda(s-t)} \Lambda\pi w_s ds < 0$ and hence also $0 > \dot{W}_t > \dot{w}_t = -\phi$.

¹⁰For a derivation, note that Bayes' rule implies $q_{t+dt} = q_t + (1 - q_t)(1 - p)(1 - \pi)\Lambda dt$, which simplifies in the limit $dt \rightarrow 0$ to $\dot{q}_t = (1 - q_t)(1 - p)(1 - \pi)\Lambda$. This linear first order ODE is solved subject to $q_0 = 0$, yielding solution (18).

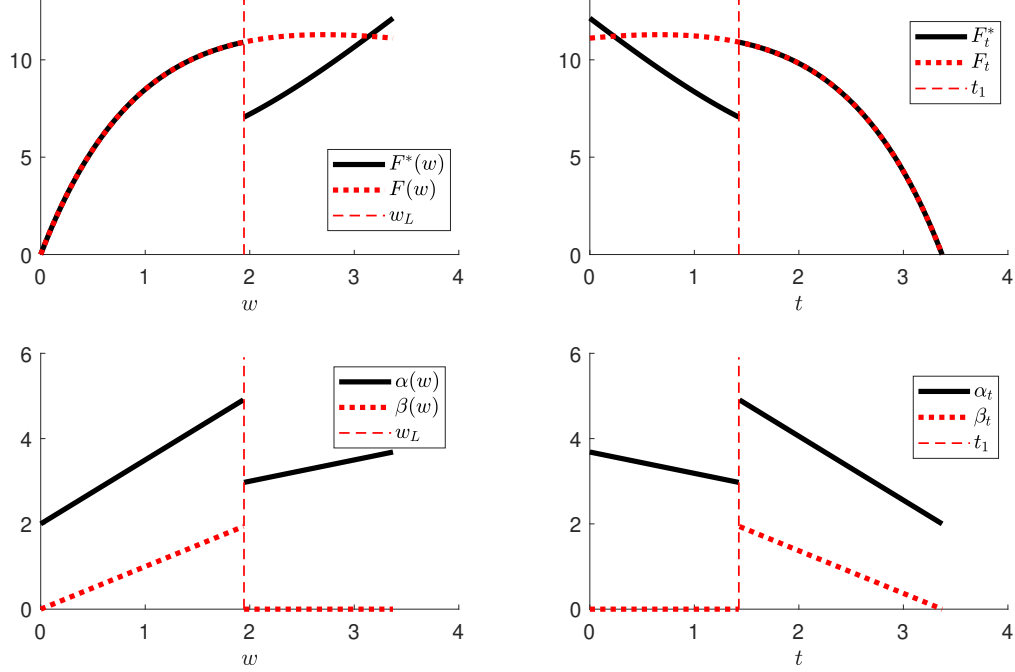


Figure 2: Numerical Example in both the “ w -space” and “time space”. The parameters are $\mu = 50$, $\kappa = 10$, $\Lambda = 1$, $p = 0.5$, $\pi = 0.5$, and $\phi = 1$.

Intuitively, the principal faces in addition to the moral hazard problem also an *adverse selection* problem when providing unconditional financing over $[0, t_1]$. In particular, the principal may inefficiently extend financing to a failed project (which is the case with probability q_t).

Given the (previously derived) optimal controls $\beta_t = 0$ and $\alpha_t = w_t(1 - \pi) + \frac{\phi}{\Lambda p}$ for $t \in [0, t_1]$ and the fact that w_t co-moves with time t via $\dot{w}_t = -\phi$, one can express the agent’s continuation value over the time interval $[0, t_1]$ (i.e., for $w > w_{t_1} = w_L$) as function of w , $W(w)$.¹¹ As the principal’s belief $q(w)$, the agent’s continuation value $W(w)$, and rewards for success (failure) $\alpha(w)$ ($\beta(w)$) are along the optimal path functions of w , the state variable w summarizes all contract relevant information. Thus, it is possible to also express the principal’s value function over the time interval $[0, t_1]$ (i.e., for $w > w_{t_1} = w_L$) as function of w , $f(w)$.

During the unconditional financing stage, over a short period of time $[t, t + dt)$, the principal incurs funding costs κdt , and the following two observable outcomes that trigger termination are possible: i) (with probability $(1 - q(w))\Lambda p dt$) the project has not failed so far and succeeds over

¹¹Conjecture that W is a function of w , in that $W_t = W(w_t)$. Recall that $\dot{W}_t = \Lambda W_t - \Lambda(p\alpha_t + (1 - p)(1 - \pi)\beta_t)$. Using $\beta_t = \beta(w_t) = 0$ and $\alpha_t = \alpha(w_t) = w_t(1 - \pi) + \frac{\phi}{\Lambda p}$, it follows that $\dot{W}_t = \Lambda(W_t - w_t) + \Lambda\pi w_t - \phi$. Using $\dot{W}_t = \frac{dW_t}{dw_t} \frac{dw_t}{dt} = W'(w_t)\dot{w}_t$, one obtains that the agent’s continuation value solves $W'(w)\dot{w} = \Lambda(W(w) - w) + \Lambda\pi w - \phi$ for $w > w_{t_1} = w_L$ subject to $W(w_L) = w_L$. This confirms that W can be expressed as function of w .

$[t, t + dt)$ yielding payoff $\mu p - \alpha(w)$, or ii) (with probability $(1 - q(w))\Lambda(1 - p)(1 - \pi)dt$) the project has not failed so far but fails over $[t, t + dt)$ and failure is observed which yields payoff zero. This leads to the HJB equation:

$$(1 - q(w))\Lambda(p + (1 - p)(1 - \pi))f(w) = (1 - q(w))\Lambda p(\mu - \alpha(w)) - \kappa + f'(w)\dot{w}, \quad (19)$$

with $\dot{w} = -\phi < 0$. At the end of the unconditional financing stage at time t_1 when $w_t = w_L$, the principal asks the agent whether the project has already failed. With probability $q(w_L)$, the project has failed and the principal must compensate the agent for failure $\beta_{t_1} = w_L$. With probability $1 - q(w_L)$, the project has not failed yet and the principal realizes the continuation payoff $F(w_L)$, leading to the value matching condition

$$f(w_L) = (1 - q(w_L))F(w_L) - q(w_L)w_L. \quad (20)$$

In addition, optimal w_L is pinned down by the smooth pasting condition:

$$f'(w_L) = \frac{\partial}{\partial w_L} ((1 - q(w_L))F(w_L) - q(w_L)w_L). \quad (21)$$

Taking stock, on the state interval $[0, w_L]$ (i.e., the time interval $[t_1, T]$), the value function is characterized by (15). On the state interval $(w_L, w_0]$ (i.e., the time interval $[0, t_1]$), the value function is characterized by function $f(w)$, solving (19). That is:

$$F^*(w) = \begin{cases} F(w) & \text{for } w \in [0, w_L], \\ f(w) & \text{for } w \in (w_L, w_0]. \end{cases} \quad (22)$$

The optimization of the initial payoff $F(w_0)$ with respect to w_0 determines the optimal deadline $T = \inf\{t \geq 0 : w_t = 0\}$. Lemma 4 in Appendix A presents a closed-form expression for $f(w)$ and Figure 2 provides a numerical example of the optimal contract. The upper two panels display the value function under the optimal contract (solid black line) and the value function under the optimal full disclosure contract (dotted red line) both in dependence of w and time, t . Note that the value function exhibits a jump at time t_1 (i.e., at $w = w_L$), when the agent makes a progress report and uncertainty is resolved. The lower two panels display the agent's rewards for success and failure both in dependence of w and time, t . Rewards for success (and failure) decrease during

a given financing stage but increase once a new financing stage begins at time t_1 .

Proposition 3. *The optimal contract does not incentivize disclosure of failure over some time period $[0, t_1)$ and becomes a full disclosure contract, characterized in Proposition 2, after time t_1 .*

1. *With $w_{t_1} = w_L$, the optimal time t_1 is characterized by (20) and (21), while the value function F^* is characterized by (22). In addition, $t_1 \rightarrow 0$ as $\pi \rightarrow 0$ and $t_1 \rightarrow T$ as $\pi \rightarrow 1$.*
2. *The contract is terminated at time $T = \inf\{t \geq 0 : w_t = 0\}$ and w_0 (and equivalently T) maximizes the principal's initial payoff $F(w_0)$.*
3. *$\alpha_t = w_t(1 - \pi) + \phi/(\Lambda p) \geq \beta_t = \gamma_t = 0$, $\dot{w}_t = -\phi < \dot{W}_t < 0$, and $W_t < w_t$ for all $t \in [0, t_1)$.*
4. *$\alpha_t = w_t(1 - \pi + \pi/p) + \phi/(\Lambda p)$, $w_t = \beta_t = W_t$, and $\gamma_t = 0$ for all $t \in [t_1, T)$.*

The optimal contract is not unique, as the optimal full disclosure contract is not unique. Proposition 3 describes the optimal contract that maximizes incentives, and its generalization is presented in Proposition 6 in the Appendix. In addition, because the principal does not ask the agent to disclose failure over $[0, t_1)$, the exact value of β_t for any $t \in [0, t_1)$ is not payoff relevant.¹² It is therefore without loss of generality to set $\beta_t = 0$ for all $t \in [0, t_1)$.

In summary, the optimal contract involves two stages: an *unconditional financing stage* $[0, t_1)$ and a *disclosure stage* $[t_1, T]$. During the unconditional financing stage, the contract does not incentivize disclosure of failure and financing is guaranteed. During that stage, the agent i) is not paid for failure, ii) receives relatively low rewards for success, and iii) incurs mild punishments for delays. The unconditional financing stage ends with the soft deadline t_1 at which the principal incentivizes a truthful progress report of whether the project has failed yet. The principal continues financing for the next stage if and only if the progress report reveals that the project is still profitable to pursue (i.e., has not failed yet).

After time t_1 , during the disclosure stage, the principal incentivizes truthful disclosure of failure and finances the project until either completion is reported or the hard deadline T is reached. Note that at the hard deadline T financing is terminated regardless of whether the project has failed yet. During the disclosure stage, the contract stipulates high but time-decreasing rewards for success and failure, inducing harsh punishments for delays. That is, the agent's incentives and performance pay — captured by (net) rewards for success $\alpha_t - w_t(1 - \pi)$ and punishments for delays $-\dot{W}_t$ — are

¹²That is, one could stipulate another value for β_t at times $t < t_1$ as long as $\beta_t \leq \min\{W_t, w_t\}$. When $\beta_t \leq \min\{W_t, w_t\}$, the agent never prefers to disclose failure or to fake failure for $t < t_1$, as is desired during the unconditional financing stage $[0, t_1)$.

stronger during the disclosure stage and hence increase following completion delays (i.e., following poor performance).

In summary, within the optimal contract, financing is staged for pure incentive purposes, even though there is only a single milestone to reach to complete the project. The provision of financing becomes more performance-sensitive over time and across stages. Likewise, the agent’s incentives are backloaded, in that they become stronger over time and across stages.

3 Analysis

3.1 Applications

Venture capital financing. In this context, the principal represents the venture capitalist and the agent represents the entrepreneur or founder (insider) of the startup financed by the venture capitalist. Consistent with the empirical findings of [Kaplan and Strömberg \(2003, 2004\)](#), the model implies that optimal venture capital contracts involve distinct financing stages and that insiders’ dollar rewards for success decrease during a given financing stage yet increase whenever a new financing stage begins.¹³ That is, insiders are effectively rewarded for reaching a new financing stage. Any financing stage concludes with a deadline, but depending on the financing stage a different type of deadline is optimal. At a soft deadline, insiders make a progress report and receive financing for the next stage if and only if the progress report reveals that the project is still profitable to pursue. At a hard deadline, financing is terminated regardless of whether the project is still profitable to pursue. Soft deadlines are optimal early on (in the unconditional financing stage) and hard deadlines are optimal in later stages.

Moreover, the model predicts an upward jump in startup valuation (i.e., project valuation) at the beginning of a new financing stage, as new information is released and uncertainty is resolved through insiders’ progress reports. Figure 2 illustrates that the principal’s value function F_t^* , representing the valuation of the project (startup firm), jumps at time t_1 either up, if the agent discloses no failure up to time t_1 , or jumps down to zero, if the agent discloses prior failure.

We also show that within venture capital contracts, insiders’ compensation and the provision of financing should become more performance sensitive over time and in each subsequent stage. The provision of unconditional (i.e., not performance sensitive) financing at the early stages of project development limits agency rents. Finally, Section 3.2 below demonstrates that the optimal contract

¹³A review of the academic literature on venture capital financing can also be found in [Gompers and Lerner \(2001\)](#).

can be implemented by financing the project with debt and equity in multiple financing rounds. This implementation has implications for the use of venture debt ([González-Urbe and Mann \(2020\)](#) and [Davis et al. \(2020\)](#)). Alternatively, the optimal contract can be implemented using convertible equity and convertible debt, as is common in venture capital financing. Note that even though the implementation of the optimal contract is generally not unique, under any implementation the principal’s stake in the project becomes less debt-like and more equity-like over time and across stages. At the same time, the agent’s equity stake is gradually diluted.

R&D financing. Our model also has implications for optimal R&D financing within more general types of firms as well as for the design of R&D projects and compensation contracts of R&D workers (insiders). In particular, optimal financing of R&D projects involves several stages, whereby insiders make occasional progress report and the continuation of financing is contingent on the outcomes of these reports. Crucially, it is optimal for the financier to elicit less frequent progress reports at the early stages of project development and more frequent progress reports at later stages.

Over time, the provision of financing becomes more sensitive to reported progress and performance. Specifically, optimal financing of R&D projects involves an initial soft deadline before which no progress report is made and funding is not terminated after poor outcomes. Thereafter, financing is subject to a hard deadline at which financing is terminated with certainty. Also note that as self-reported progress is subject to moral hazard, it is optimal to elicit progress reports less frequently when moral hazard is severe. The compensation of R&D workers should involve tolerance for failure, high rewards for success, and penalties for delays in project completion. R&D workers’ incentives should be backloaded (across stages), which is achieved by increasing rewards for success and failure upon entering a new stage. Importantly, R&D workers should be rewarded for failure at later stages, but not for failure at early stages.

Executive compensation contracts. Alternatively, one can interpret the agent as the manager (i.e., CEO) of a more general type of firm while the principal represents the firm’s investors or outside shareholders. Golden parachutes, golden handshakes, and severance pay are instruments that induce tolerance for failure in executive compensation contracts ([Edmans and Gabaix \(2016\)](#)). Our model highlights that such instruments incentivize executives to fake or generate bad outcomes. As a consequence, golden parachutes and severance pay should not be in place in early stages of the contract (i.e., CEO tenure) and should decrease in size over time and in contract duration. In

addition, our model implies that incentive pay (e.g., through stock based compensation) and the threat of termination should be stronger in the presence of golden parachutes. Broadly interpreted, stock based compensation and golden parachutes are complements for incentive provision.

3.2 Implementation

We demonstrate how to implement the optimal contract by financing the project with a mixture of debt and equity. For the implementation, we consider a slightly different version of the optimal contract that differs from the contract in Proposition 3 only in the values of (α_t, γ_t) during the disclosure stage (but this contract features the same deadlines (t_1, T) and the same unconditional financing stage). In this contract, $\beta_t = \gamma_t = w_t$ during the disclosure stage and $\alpha_t = w_t + \phi/(\Delta p)$. Proposition 6 in the Appendix demonstrates that this alternative contract is optimal and yields the same payoffs for principal and agent as the contract from Proposition 3.¹⁴ Appendix D.1 presents the details of our implementation. We give some intuition and the main findings below.

In our implementation, there are two financing rounds, at time $t = 0$ at the beginning of the unconditional financing stage and at time $t = t_1$ at the beginning of the disclosure stage. During the first financing round at $t = 0$, the principal injects funds to facilitate project development over the next stage until time t_1 . In exchange, the principal receives both debt and equity claims in the project, while the agent retains the remaining equity. That is, the project is financed by issuing debt and equity. During the unconditional financing stage, equity pays dividends only upon project success. Debt including its interest is paid back during the unconditional financing stage only if the project succeeds at time $t \in [0, t_1)$. On the other hand, if the project fails during the unconditional financing stage, debt defaults.

Provided the project does not fail before time t_1 , there is a second financing round that is held at the beginning of the disclosure stage at time t_1 . During the second financing round, the principal contributes additional funds to continue project development. In exchange, the principal receives additional equity in the project (from the agent), reducing (i.e., diluting) the agent's equity stake. That is, funds are raised by issuing equity. Existing debt (including interest) is not paid back during the second financing round. Instead, debt including interest is paid back during the disclosure stage upon project completion or at the deadline T (from the remaining funds that were not used in project development). Any funds that are left after paying back debt are distributed as dividends

¹⁴The only difference between the alternative contract and the contract from Proposition 3 is that to incentivize the agent not to fake failure during the disclosure stage, the principal boosts the agent's rents by stipulating rewards for publicly observed failure $\gamma_t = \beta_t$ rather than by stipulating excessive rewards for success.

to equity holders. During the disclosure stage (but not during the unconditional financing stage), funds allocated to the project exceed the face value of debt (plus accrued interest), so termination of financing due to failure leads to full repayment of debt and dividend payouts to equity holders generating rewards for failure for the agent. Note that in the above implementation, the project is financed with debt and equity. Remarkably, during the first financing round, funds are raised by issuing both debt and equity. In contrast, during the second financing round, funds are raised by only issuing (selling) equity. Interestingly, in our implementation, debt resembles a credit line with time-increasing interest rate that the principal grants to the project.

In the context of venture capital financing, our findings suggest that financing with both equity and debt is optimal, which rationalizes the use of venture debt (González-Uribe and Mann (2020), and Davis et al. (2020)) and credit lines granted by venture capitalists.¹⁵ That is, the use of venture debt and equity financing are complementary, consistent with Davis et al. (2020) who find that venture debt is often a complement to equity financing, with over 40% of all financing rounds including some amount of debt. The implementation of the optimal contract also predicts that the use of venture debt is more prevalent in early financing stages. Specifically, startup firms (should) rely on venture debt mostly in early financing stages, while they (should) rely more on equity financing in later stages.

Finally, as discussed in Appendix D.1, we emphasize that the implementation of the optimal contract is generally not unique. Therefore, the optimal contract can be implemented in an alternative way, also using convertible debt and convertible equity, as is common in venture capital financing. Although the implementation of the optimal contract is generally not unique, it holds that under any implementation of the optimal contract, the principal’s stake in the project becomes less debt-like and more equity-like over time and across stages. At the same time, the agent’s equity stake is gradually diluted.

3.3 Project characteristics and financing contracts

We study how project characteristics shape the design of optimal financing contracts. Figure 3 plots the financing deadline T , the reward for failure at the beginning of the disclosure stage β_{t_1} , and the relative length of the unconditional financing stage t_1/T against $1/\Lambda$ (upper three panels) and ϕ (lower three panels). Observe that β_{t_1} is the highest reward for failure attainable and proxies

¹⁵Relatedly, Robb and Robinson (2014) and Hochberg, Serrano, and Ziedonis (2018) document that young firms and startup firms heavily rely on debt financing.

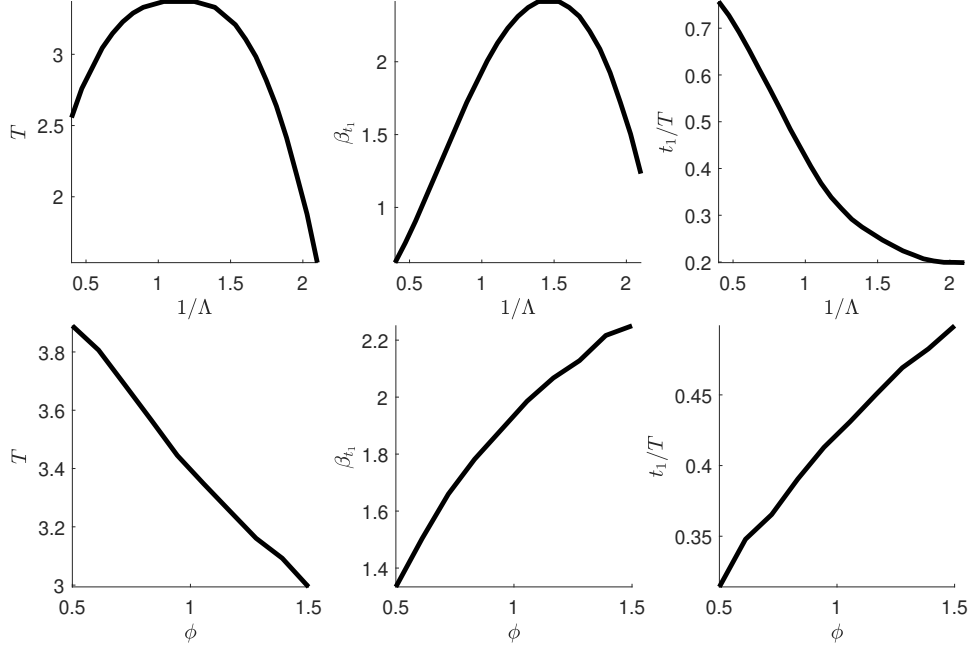


Figure 3: Comparative statics w.r.t. length of project development phase $1/\Lambda$ and the agency parameter ϕ . The baseline parameters are $\mu = 50$, $\kappa = 10$, $p = 0.5$, $\pi = 0.5$, and $\phi = 1$.

overall/average rewards for failure within the optimal contract.¹⁶

Notably, the financing deadline T is hump-shaped in the (average) duration of the project development phase $1/\Lambda$, as shown in Corollary 1 and illustrated in Figure 3. The intuition is as follows. Projects with a long development phase naturally require financing over a longer horizon. However, projects with a long development phase are at the same time subject to more severe moral hazard, which makes financing such projects less profitable.

Figure 3 also illustrates that projects with a sufficiently short development phase are more suitable (or likely) to receive unconditional financing early on. The reason is that projects with a short development phase are financed with a short deadline, which curbs the temptation to hide bad outcomes and hence facilitates the provision of unconditional financing. As a result, the model implies that for short term projects the provision of financing is less sensitive to reported progress yet subject to a strict deadline.

In addition, rewards for failure tend to be higher for moderate values of $1/\Lambda$. Hence, optimal financing contracts for projects with a sufficiently short or long development phase stipulate low rewards for failure yet exhibit tolerance towards failure through the provision of unconditional financing. These findings imply for executive compensation contracts that the use of golden

¹⁶The reason for this is that $\beta_t = \beta_{t_1} - \phi(t_1 - t)$ for $t > t_1$.

parachutes should be contingent on the horizon of corporate policies (i.e., the project horizon). In particular, golden parachutes are less valuable for incentive provision when the horizon of corporate policies is either sufficiently long or sufficiently short.

Last, the financing deadline T decreases in ϕ , while rewards for failure β_{t_1} and the relative length of the unconditional financing stage t_1/T increase in ϕ . As a result, optimal financing contracts for projects that are subject to severe agency conflicts involve a relatively long unconditional financing stage (and soft deadline) that is followed by a relatively short disclosure stage (and hard deadline). This finding also implies that the provision financing is less performance sensitive, when moral hazard is more severe. The intuition is that because the provision of unconditional (not performance-sensitive) financing limits agency rents, it is especially valuable when agency conflicts are severe and ϕ is high. That is, because self-reported progress is subject to moral hazard, it becomes optimal not to ask the agent to disclose failure when moral hazard is severe.

We conclude this section with the following corollary that provides analytical results regarding the effects discussed above.

Corollary 1. *Let $\pi > 0$ be sufficiently small. Then, the following holds:*

1. *T and β_{t_1} increase in $1/\Lambda$ for $1/\Lambda$ sufficiently small and decrease in $1/\Lambda$ for $1/\Lambda$ sufficiently large.*
2. *T decreases in ϕ and β_{t_1} decreases in ϕ for ϕ sufficiently large.*

3.4 Over- and under-provision of financing

In our model, moral hazard can lead to under-provision of financing (i.e., under-investment), when contract termination before time τ precludes financing of a project with positive NPV. Moral hazard can also lead to over-provision of financing (i.e., over-investment), which refers to financing a project with negative NPV. Over-provision of financing may arise when the principal inefficiently extends financing to a failed project (with negative NPV) during the unconditional financing stage.

To assess financing efficiency, we calculate the ex-ante probability of under-investment

$$\mathcal{P}^U := \mathbb{P}(T < \tau) = e^{-\Lambda T},$$

i.e., the likelihood that the project is terminated before completion. Likewise, we calculate the

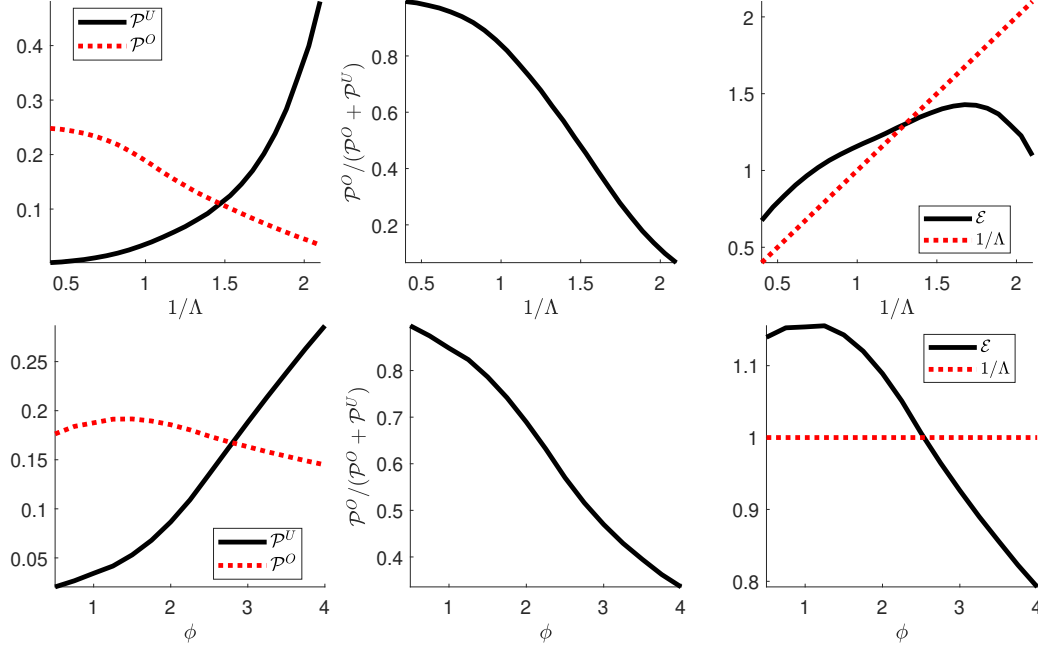


Figure 4: Over- vs. under-investment. The baseline parameters are $\mu = 50$, $\kappa = 10$, $p = \pi = 0.5$, and $\phi = 1$.

ex-ante probability of over-investment

$$\mathcal{P}^O = \mathbb{P}(\text{Unobserved Failure during } [0, t_1]) = (1 - e^{-\Lambda t_1})(1 - p)(1 - \pi).$$

Another measure of financing efficiency is the expected length of the financing period $\mathcal{E} := \mathbb{E}[T \wedge \tau^A]$, which captures the principal's investment horizon. This quantity equals $1/\Lambda$ in the second-best case and without frictions. Hence, $\mathcal{E} < 1/\Lambda$ indicates under-provision of financing and $\mathcal{E} > 1/\Lambda$ indicates over-provision of financing. Appendix D.4 shows how to calculate $\mathcal{E}$. Note that we have optimal financing (i.e., $\mathcal{P}^O = \mathcal{P}^U = \mathcal{E} - 1/\Lambda = 0$) in the second-best environment (when $\pi = 1$) but under-investment (i.e., $0 = \mathcal{P}^O < \mathcal{P}^U$ and $\mathcal{E} < 1/\Lambda$) under a full disclosure contract. Conversely, the optimal contract with distinct financing stages or, broadly interpreted, stage financing as such is more likely to cause over-investment (i.e., $\mathcal{P}^O > 0$ and $\mathcal{E} > 1/\Lambda$).

Figure 4 displays the ex-ante probabilities of under- and over-investment $\mathcal{P}^U$ and $\mathcal{P}^O$, the probability of over-investment conditional on under- or over-investment $\mathcal{P}^O/(\mathcal{P}^O + \mathcal{P}^U)$, and the average length of the financing period $\mathcal{E}$ in dependence of both $1/\Lambda$ (upper three panels) and ϕ (lower three panels). Figure 4 shows that, on average, there is under-provision of financing for projects with a long development phase, whereas there is over-provision of financing for projects

with a short development phase. That is, investors tend to terminate financing for long-term (short-term) projects inefficiently early (late). In other words, moral hazard increases the principal’s investment horizon (i.e., financing deadline) for short term projects but decreases the investment horizon for long-term projects, relative to the benchmark without frictions. Likewise, Figure 4 also illustrates that mild agency conflicts induce over-provision of financing (i.e., over-investment) but severe agency conflicts induce under-provision of financing (i.e., under-investment).

Note that over-investment in our model is an extreme case in that it corresponds to financing a project that does not produce a payoff at all. The reason is that we have normalized payoffs after failure to zero. Section 4 considers a model variant in which the project may generate payoff, when it receives financing despite failure (which might be inefficient). Under these circumstances, over-investment implies that the project generates more payoffs than is efficient according to the NPV criterion.

Finally, in the context of venture capital financing, the model predicts venture capital over-investment in projects that generate (preliminary) results relatively quickly and allow to capitalize on these results — such as pharmaceutical or biotech projects — and under-investment in projects that do not generate (preliminary) results quickly — such as clean energy projects — which is consistent with the findings of [Nanda, Younge, and Fleming \(2014\)](#).¹⁷

4 Financing unicorns

We relax the assumption that a failed project does not generate payoff. Consider that a failed project may also generate (terminal) payoff μ . However, unlike a successful project, which generates payoff μ immediately at time τ , a failed project generates payoff μ with delay at time $\tau^\lambda > \tau$. Terminal payoff μ is observable and contractible. As in the baseline model, failure at time τ is observed by the agent but observed by the principal only with probability $\pi \in [0, 1]$. Thus, the principal may not be able to distinguish between whether it is immediate success at time τ or “failure” and later “success” at time τ^λ that has led to the terminal payoff. In addition, the principal cannot verify failure or reports thereof. This setting is relevant when the terminal payoff represents a certain goal or milestone in project development, whereby failure captures interim bad outcomes that are hard for the principal to observe or verify. Thus, this setting can also be interpreted as a multi-stage project.

¹⁷Pharmaceutical or biotech projects can generate intermittent results through experiments and/or have to adhere to certain milestones set by the FDA. For a discussion, see, e.g., [Kerr and Nanda \(2015\)](#).

The time τ^λ arrives at exogenous rate $\lambda > 0$. Notably, over $[\tau, \tau^v)$ the principal still incurs financing costs κ so that delay is costly and the payoff of a failed project is lower than the payoff of a successful project.¹⁸ On average, a failed project takes $1/\lambda$ units of time to generate payoff and, therefore, possesses at time τ *after completion* value (i.e., NPV):

$$\mu^f := \mu - \frac{\kappa}{\lambda}$$

while the value of a successful project equals μ . The NPV of the project *before completion* is then

$$NPV^\lambda := p\mu + (1-p)\max\{\mu^f, 0\} - \frac{\kappa}{\lambda},$$

where financing is optimally terminated after failure, if $\mu^f < 0$.

To characterize the agent's incentives to disclose failure, suppose the project has failed at time t . Reporting failure immediately yields payoff β_t . Delaying disclosure for dt units of time not only yields benefits from running the project ϕdt but also entails the chance, i.e., λdt , that the failed project generates payoff μ , leading to a reward α_t . Hence, the agent is better off disclosing failure at time t if and only if

$$\beta_t \geq (\beta_{t+dt} + \phi dt)(1 - \lambda dt) + \alpha_t \lambda dt,$$

which becomes for $dt \rightarrow 0$:

$$\dot{\beta}_t \leq -(\phi + \lambda(\alpha_t - \beta_t)) \tag{23}$$

Notably, the incentive condition w.r.t. effort (11) implies $\alpha_t > \beta_t$ and hence an increase in λ tightens the incentive condition (23). Intuitively, the prospect of a future breakthrough despite (interim) failure provides the agent with incentives to hide failure and to continue project development under the motto “Fake it till you make it.” The remainder of the solution is similar to the solution of the baseline model and is thus relegated to Appendix C.

A measure of project risk is given by the variance of potential project outcomes, i.e., the difference in project value after success and failure $\mu - \mu^f = \kappa/\lambda$. This measure of project risk clearly decreases in λ (and does not depend on μ). An implication is that riskier projects are easier to incentivize and hence easier for the principal to finance. Formally:

Corollary 2. *Suppose that $\mu^f > 0$ and λ is sufficiently small so that full effort $a_t = 1$ is optimal for*

¹⁸The cost of delay may alternatively arise due to displacement risk from other technologies or due to investors' time preference.

all $t \geq 0$. A mean preserving spread, which increases μ but decreases λ while holding the project's net present value NPV^λ fixed, increases the principal's payoff.

The underlying reason is that higher risk $1/\lambda$ relaxes the incentive constraint (23). When the project is sufficiently risky, the agent is less inclined to continue operating the project in the hope of reaching a breakthrough in the future, which generates incentives to disclose failure.

As a result, our model offers an explanation why venture capitalists seek to finance high risk start ups, i.e., potential unicorns, even if this choice is not necessarily supported by the NPV criterion. Interestingly, Gornall and Strebulaev (2020) and Gompers, Gornall, Kaplan, and Strebulaev (2020) find evidence that unicorn startups are frequently over-valued. This is broadly consistent with the notion that venture capital investors seek investments with high potential but high risk, thereby boosting the valuation for such startup firms (possibly above the fundamental value).

In other words, our model predicts that risk-taking and risky investments reduce agency costs and therefore might be optimal for venture capitalists. This implication is broadly consistent with the findings of Nanda and Rhodes-Kropf (2013) who document venture capitalists' risk-taking in "hot markets."

5 Optimal contracts with monitoring

In this section, we introduce the possibility that the principal can inspect (i.e., monitor) project progress at a cost $K > 0$. Upon an inspection at time t , the principal learns whether the project has failed up to time t . In contrast, an inspection does not generate information about the agent's effort up to and including time t .¹⁹ Also note that monitoring takes the form of a costly state verification in the spirit of Townsend (1979).

As in Piskorski and Westerfield (2016) and Varas et al. (2020), we assume that the principal can fully commit (at time zero) to an inspection policy. In addition, we assume that the cost K is not prohibitively high which is formalized in Assumption 1, presented in Appendix B. In the following, $dM_t = 1$ indicates that the principal inspects the project at time t , whereby $dM_t \in \{0, 1\}$.

¹⁹Similar to Piskorski and Westerfield (2016), it is possible to accommodate the possibility that an inspection at time t allows the principal to also observe effort a_t .

5.1 Monitoring and incentives

We start by discussing how inspections generate incentives to disclose failure truthfully, i.e., truth telling incentives. For this sake, let us consider that the contract incentivizes disclosure of failure. When the principal inspects the firm and learns that the agent is hiding failure, she can punish the agent. The threat of punishment generates truth telling incentives. Because of the agent's limited liability, these truth telling incentives are maximized, when the principal terminates financing and fires the agent upon detecting misbehavior.

Suppose that the project has failed at time t and that failure is privately observed by the agent. The agent can hide failure maximally up to the next inspection date $\tau_t^M = \inf\{s \geq t : dM_s = 1\}$ after time t . Notably, we conjecture and verify that the principal inspects the project periodically at deterministic dates. The reason is that the principal's value function is concave, which renders randomization sub-optimal. Anticipating the next inspection at the (deterministic) date τ_t^M , the agent prefers to disclose failure truthfully at time t if and only if $\beta_t \geq \beta_s + (s-t)\phi$ for all $s \in (t, \tau_t^M)$. Taking the limit $s \rightarrow t$ yields $\dot{\beta}_t \leq -\phi$, that is (8).

As failure can potentially occur at any time t , the incentive condition (8) must hold for all times t at which there is no inspection (i.e., $dM_t = 0$). That is, the principal provides incentives to disclose failure both through inspections at deterministic dates and through punishments for delays at all other dates. Also note that contract terms *after* time τ_t^M do not affect the agent's incentives to disclose failure *before* time τ_t^M .

As in the baseline version of the model, the incentive condition w.r.t. effort a_t is given by $\alpha_t \geq w_t(1-\pi) + \frac{\phi}{\Lambda p}$ (that is, (11)), whereby $w_t = \beta_t$ in optimum. Last, the principal can incentivize the agent not to fake failure by verifying reported failure (i.e., inspecting the project upon disclosure of failure). However, it turns out that within the optimal contract this is not necessary.²⁰

5.2 Solving for optimal contract with monitoring

Let us consider a full disclosure contract with deadline T , and conjecture that the agent is never paid for observed failure (i.e., $\gamma_t = 0$). The contract incentivizes truthful disclosure of failure for all times $t \in [0, T \wedge \tau]$ so that $\beta_t = w_t$ and β_t satisfies (8) whenever there is no inspection. Condition (8) constrains the choice of the level of w (i.e., β) and thus w (i.e., β) enters the principal's dynamic

²⁰If the principal checks reported failure with probability $\hat{\Theta}_t \in [0, 1]$, the agent is detected faking failure with probability $\hat{\Theta}_t$, in which case her payoff is zero. As a result, the incentive constraint (7) is modified to $W_t \geq (1-\hat{\Theta}_t)\beta_t$. Within the optimal contract, $W_t > \beta_t$ for $t > 0$ so that the principal optimally chooses $\hat{\Theta}_t = 0$.

optimization as state variable, while the change in w (i.e., β) is a control variable. Along the optimal path the agent's continuation value W can be expressed as function of w (that is, $W_t = W(w_t)$), so that w summarizes all contract relevant information and hence is the only state variable for the principal's dynamic optimization.

As a result, the principal's payoff at any time t with $w_t = w$ equals $F(w)$, where $F(\cdot)$ is the principal's value function. The starting value w_0 is chosen without constraints to maximize the principal's payoff at time zero $F(w_0)$, so that $F'(w_0) = 0$. Because rewards for failure after an inspection do not affect truth telling incentives before an inspection, the principal can choose w at (i.e., just after) an inspection without constraints to maximize her continuation payoff. Therefore, the principal optimally sets $w = w_0$ and realizes continuation payoff $F(w_0)$, whenever an inspection indicates that the project is still profitable to pursue (i.e., has not failed so far).

Once $w = 0$, the principal cannot provide truth telling incentives anymore by decreasing rewards for failure β and punishing the agent for delays. To maintain sufficient incentives, the principal must therefore either terminate financing or inspect the project. Because in optimum the agent does not hide failure, the inspection yields positive outcomes and the rewards for failure are set to $w = w_0$, leading to

$$F(0) = \max\{F(w_0) - K, 0\}. \quad (24)$$

We assume that monitoring is optimal, i.e., $F(w_0) > K$. Hence, the principal inspects the project at all times t with $w_t = 0$.

The principal's value function $F(w)$ solves the HJB equation

$$\Lambda F(w) = \max_{\dot{w}, \alpha} \left\{ \Lambda p \mu - \kappa - \Lambda((1-p)(1-\pi)w + p\alpha) + F'(w)\dot{w} \right\} \text{ s.t. (7), (8), and (11).} \quad (25)$$

The value function satisfies the boundary conditions (24) and $F'(w_0) = 0$. Since $F(w)$ increases for $w \in [0, w_0]$, it follows that $\dot{w} = -\phi$, so that condition (8) is tight. The HJB equations (13) and (25) only differ in their boundary conditions and hence can be derived using similar arguments.

Notably, there is no financing deadline, and project development is not terminated before completion. The lack of early termination increases the value of the agent's stake in the project so that $W_t > w_t$ for all t . This in turn generates incentives not to fake bad outcomes and hence makes it possible to reduce the agent's rewards for success, subject to motivating effort. As a result, $\alpha_t = w_t(1-\pi) + \frac{\phi}{\Lambda p}$ and, unlike in the baseline model without monitoring, the incentive condition w.r.t. effort (11) is tight. That is, monitoring reduces the cost of truth telling incentives. Therefore,

the optimal contract with monitoring always incentivizes disclosure of failure and, in particular, does not feature an unconditional financing stage.

Lemmata 2 and 3 in Appendix A derive

$$\begin{aligned}
w_0 &= \left(\frac{2K\phi}{\Lambda(1-\pi)} \right)^{1/2} - \chi \quad \text{with} \quad 0 < \chi \in o\left(K^{\frac{3}{2}}\right), \\
F(w) &= \mu p - \frac{\kappa + \phi}{\Lambda} - w(1-\pi) - (K + w_0(1-\pi)) \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right), \\
W(w) &= \frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + \frac{\pi\phi}{\Lambda} + (1-\pi)w.
\end{aligned} \tag{26}$$

The following Proposition summarizes the findings of this section.

Proposition 4. *The optimal contract features no deadline (i.e., $T = \infty$), always incentivizes truthful disclosure of failure, and inspections occur at all times t with $w_t = 0$. For all $t \geq 0$, $\alpha_t = w_t(1-\pi) + \frac{\phi}{\Lambda p}$, $\beta_t = w_t$, $dw_t = -\phi dt + w_0 dM_t$. The principal's value function $F(w)$ solves (25) subject to (24) and $F'(w_0) = 0$.*

In summary, the optimal contract features several identical financing stages and periodic inspections at deterministic dates. During each financing stage, the contract stipulates time-decreasing rewards for success and failure, effectively punishing the agent for delays. If the project is not completed by the end of the financing stage, the principal inspects the project and grants financing for the next stage if and only if an inspection reveals that the project is still profitable to pursue. Whenever a new financing stage begins, stipulated pay for success and failure increase so that the agent is effectively rewarded for positive inspection outcomes.²¹

5.3 Results and implications

The form of the optimal contract with monitoring has interesting implications. First, recall that the principal inspects the project whenever $w = 0$, in which case w is set to $w = w_0$ after the inspection. As w decreases with $\dot{w}_t = -\phi$ at all other dates, the time between two inspection dates equals $\frac{w_0}{\phi}$ (provided the project is not completed in the meantime). For K sufficiently small, $w_0 \simeq \left(\frac{2K\phi}{\Lambda(1-\pi)} \right)^{1/2}$ increases in $1/\Lambda$, implying that monitoring is less frequent for projects with a

²¹If monitoring produced false-negative results (i.e., wrongly indicates failure) with some probability $\pi' > 0$, the agent's continuation value W_t would indeed jump up after a positive inspection outcome. With perfect monitoring technology (i.e., $\pi' = 0$), continuation value W_t is smooth. Thus, the agent is rewarded by an increase in pay for success and failure (α_t, β_t) but not by an increase in continuation value.

long development phase. Conversely, incentives by means of punishments for delays and rewards for failure are stronger for such projects. The reason is that long term projects not only have a long development phase but also are less likely to produce bad outcomes early, which renders it optimal to monitor later and less frequently. In addition, the frequency of monitoring — captured by the inverse of $\frac{w_0}{\phi}$ — increases in the severity of moral hazard ϕ and $1 - \pi$.²² That is, more severe moral hazard ϕ requires both more frequent monitoring and harsher punishments for delays in incentive provision.

Second, punishments for delays — such as contract termination — and monitoring are substitutes for the provision of truth telling incentives. In addition, monitoring leads to lower rewards for success and failure. Overall, monitoring substitutes for performance pay and rewards for failure in incentive provision. In the context of executive compensation, the model predicts that monitoring by shareholders and the board of directors or shareholder activism is negatively related to the use of golden parachutes.

Third, (the possibility of) monitoring induces more performance sensitive and hence more efficient financing. In particular, financing is terminated at the time of project completion and therefore not subject to under- or over-investment. As Figure 3 highlights that short-term (long-term) projects are prone to over-provision (under-provision) of financing, the model implies that monitoring curbs the provision of financing (i.e., the principal’s investment horizon) for short-term projects while boosting the provision of financing (i.e., the principal’s investment horizon) for long-term projects. Likewise, as Figure 3 also shows that projects with mild (severe) moral hazard are prone to over-provision (under-provision) of financing, monitoring curbs (boosts) the provision of financing when moral hazard is mild (severe).

6 Further results and robustness

6.1 Moral hazard versus adverse selection

In our model, *moral hazard* arises because effort is hidden and costly. The severity of moral hazard is captured by the agent’s private benefits from shirking ϕ . In addition, failure is hard for the principal to observe or verify. Imperfectly observable failure induces another agency problem that can be interpreted as an *adverse selection* problem and its severity is captured by $1 - \pi$. The reason is that the principal provides truth telling incentives to the agent under the assumption that the

²²Note that for $w_0 \simeq \left(\frac{2K\phi}{\Lambda(1-\pi)}\right)^{1/2}$, it follows that $\frac{w_0}{\phi} \simeq \left(\frac{2K}{\Lambda(1-\pi)\phi}\right)^{1/2}$ decreases in ϕ .

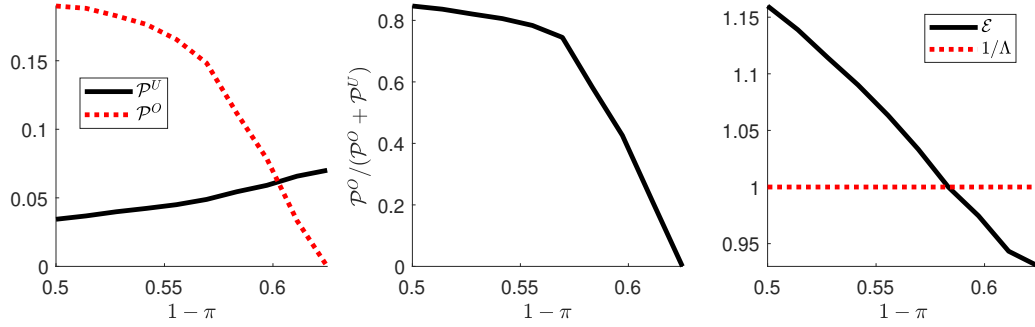


Figure 5: Over- vs. under-investment. The baseline parameters are $\mu = 50$, $\kappa = 10$, $p = 0.5$, and $\phi = 1$.

agent (already) has privately observed project failure.

Proposition 3 shows that $t_1 \rightarrow 0$ as $\pi \rightarrow 0$ and $t_1 \rightarrow T$ as $\pi \rightarrow 1$ and Figure 3 shows that t_1/T increases in ϕ . That is, the provision of unconditional financing is valuable when moral hazard is severe (i.e., when ϕ is large) but adverse selection concerns are mild (i.e., when $1 - \pi$ is low). In other words, the provision of unconditional financing is suitable for dealing with moral hazard but less suitable for dealing with adverse selection. As a result, the provision of financing is more (less) performance sensitive when adverse selection (moral hazard) is more severe.

The lower three panels of Figure 4 plot the (scaled) probabilities of over- and under-investment and the average length of the financing period against ϕ , capturing the severity of moral hazard, while Figure 5 plots the same quantities against $1 - \pi$, capturing the severity of adverse selection. Figures 4 and 5 demonstrate that mild adverse selection concerns (moral hazard problems) induce over-provision of financing (i.e., over-investment) whereas severe adverse selection concerns (moral hazard problems) induce under-provision of financing (i.e., under-investment).

6.2 Robustness

Our model entails a number of assumptions that are mainly designed to enhance simplicity and to facilitate a clear analysis of the main forces in a tractable model. Below, we discuss these assumptions and the robustness of the results.

Exogenous project completion. Existing dynamic contracting papers — such as Mason and Välimäki (2015), Green and Taylor (2016), and Varas (2017) — study how to motivate the agent to complete a project. In these papers, completion always corresponds to “success” and the only bad outcomes that can arise are completion delays. Notably, the agent is never tempted to hide

project completion/success, but — in [Green and Taylor \(2016\)](#) and [Varas \(2017\)](#) — the agent is tempted to fake success (i.e., good outcomes). The innovation in our paper is that it considers the risk of project failure when the agent can hide or fake project failure (i.e., bad outcomes).

In our model, the assumption of an exogenous completion rate is made for simplicity and theoretical clarity and is not consequential. It offers the advantage that we obtain a clean second-best benchmark when failure is observable and contractible. Departing from this benchmark, we are able to clearly identify how imperfect observability of failure affects incentive provision. [Appendix D.2](#) shows that we obtain the same results employing an alternative framework, in which the agent controls project completion while the project is subject to failure risk. Likewise, one could extend our baseline model to a model of multi-tasking in which the agent controls both project completion and the project’s propensity to succeed or fail.

Unobservable success. Throughout this paper, we have assumed that success is perfectly observable and contractible. This assumption intuitively reflects that the agent is tempted to conceal bad outcomes rather than good outcomes. Even though not modelled explicitly, disclosure of good outcomes could yield private benefits to the agent, e.g., related to the agent’s reputation or career concerns, whereas disclosure of bad outcomes could yield private dis-utility.

We show that our findings do not change substantially when success is imperfectly observable but verifiable. Consider that both failure and success are (publicly) observed by the principal only with probability π and privately observed by the agent otherwise. To obtain a non-trivial solution, at least one of the two possible outcomes, success and failure, must be verifiable. We assume that success is verifiable, as it is likely to be more difficult to fake good outcomes rather than bad outcomes.

To characterize the agent’s incentives to disclose success, note that the agent can always delay disclosure for a unit of time and derive private benefits ϕdt . By the same arguments leading to [\(8\)](#), we obtain that the agent prefers to disclose success truthfully if and only if $\dot{\alpha}_t \leq -\phi$. This condition is obviously not satisfied in the optimal contract from [Proposition 3](#) because α_t exhibits a jump at time t_1 when the unconditional financing stage ends. However, during the disclosure stage and within an (optimal) full disclosure contract, it follows that $\dot{\alpha}_t \leq -\phi$. This implies that the disclosure stage of the optimal contract is unaffected by whether success is observable.^{[23](#)}

[Appendix D.3](#) extends this intuition and demonstrates that the optimal contract does not change

²³Recall that — by [Propositions 2 and 3](#) — for $t \geq t_1$: $\alpha_t = (1 - \pi + \pi/p)w_t + \phi/(\Lambda p)$, so that $\dot{\alpha}_t \leq \dot{w}_t = -\phi$.

much when success is imperfectly observable: it features i) an unconditional financing stage $[0, t_1)$ during which the agent discloses neither success nor failure (yet is rewarded if success is observed), and ii) a disclosure stage that looks similar to the disclosure stage of Proposition 3.

7 Conclusion

We study a dynamic contracting model in which a principal hires an agent to develop an innovative project. Crucially, project failure is hard for the principal to observe or verify and the agent can hide or fake failure, leading to a tension in incentive provision. The optimal contract consists of two distinct stages: i) an unconditional financing stage and ii) a disclosure stage. During the unconditional financing stage, financing is guaranteed and the contract does not incentivize disclosure of failure. During the disclosure stage, the contract incentivizes disclosure of failure and the principal finances the project until either a deadline is reached or the project is completed. Then, the agent receives time decreasing rewards for success and failure and incurs harsh punishments for delays. In the optimal contract, the provision of financing becomes more performance sensitive following completion delays, i.e., following poor performance. Our results also imply that moral hazard may lead to over- or under-provision of financing relative to the net present value criterion. Last, we characterize optimal dynamic monitoring and study the role of monitoring in incentive provision. The paper generates a set of implications for venture capital financing and the design of executive compensation contracts.

We employ a tractable framework that features a single project development stage and only two possible outcomes upon project completion. While this allows us to derive the optimal contract analytically, it would be interesting to analyze more involved settings. A natural extension is to study incentive provision for information disclosure in multistage settings, as in [Green and Taylor \(2016\)](#), or when project outcomes are drawn from a continuum rather than from a binary set. However, we note that the analysis is technically challenging if outcomes that can take values from a continuum are unobservable to one player (i.e., the principal) in a dynamic game, as is illustrated in a different context by [Duffie and Lando \(2001\)](#). Another avenue for future research is to study the agent's incentives to hide and fake positive rather than negative outcomes.

Appendix

A Closed form expressions

Lemma 1. *The ODE (13) subject to $F(0) = F'(w_0) = 0$ has the closed form solution (15).*

Proof. We just verify that the proposed function indeed is the desired solution to the ODE (13) subject to the stipulated boundary conditions $F(0) = F'(w_0) = 0$.

Take (15):

$$F(w) = \left(\mu p - \frac{\kappa}{\Lambda}\right) \left(1 - \exp\left(\frac{-w\Lambda}{\phi}\right)\right) - w$$

and differentiate to obtain

$$F'(w) = \left(\mu p - \frac{\kappa}{\Lambda}\right) \frac{\Lambda}{\phi} \exp\left(\frac{-w\Lambda}{\phi}\right) - 1.$$

Define for any function $F(w)$ the operator

$$\mathcal{D}F(w) = \Lambda F(w) + F'(w)\phi - (\Lambda p\mu - \kappa - \Lambda w - \phi),$$

and note that $\mathcal{D}F(w) = 0$ if and only if $F(w)$ solves (13) (under the optimal controls). We use above expressions for $F(w)$ and $F'(w)$ and obtain

$$\begin{aligned} \mathcal{D}F(w) = & \Lambda \left[\left(\mu p - \frac{\kappa}{\Lambda}\right) \left(1 - \exp\left(\frac{-w\Lambda}{\phi}\right)\right) - w \right] \\ & - (\Lambda p\mu - \kappa - \Lambda w - \phi) - \phi \left[\left(\mu p - \frac{\kappa}{\Lambda}\right) \frac{\Lambda}{\phi} \exp\left(\frac{-w\Lambda}{\phi}\right) - 1 \right] = 0. \end{aligned}$$

Next, we verify that $F(0) = 0$ and $F'(w_0) = 0$ with w_0 from (16). It is immediate to see that $F(0) = 0$. Next recall that

$$w_0 = \frac{\phi}{\Lambda} \ln \left(\frac{\Lambda \mu p - \kappa}{\phi} \right),$$

so that

$$F'(w_0) = \left(\mu p - \frac{\kappa}{\Lambda}\right) \frac{\Lambda}{\phi} \frac{\phi}{\Lambda \mu p - \kappa} - 1 = 0.$$

The proof is complete by virtue of the Picard-Lindelöf theorem, ensuring uniqueness of the solution. $\square$

Lemma 2. *The ODE (25) subject to $F'(w_0) = F(w_0) - K - w_0 = 0$ has the closed form solution as stipulated in (26). In addition, $w_0 < \bar{w} = \left(\frac{2K\phi}{\Lambda}\right)^{1/2}$.*

Proof. We just verify that the proposed function indeed is the desired solution to the ODE (25) subject to the stipulated boundary conditions $F(w_0) - F(0) - K = F'(w_0) = 0$.

Take

$$F(w) = \mu p - \frac{\kappa + \phi}{\Lambda} - w(1 - \pi) - (K + w_0(1 - \pi)) \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right).$$

and differentiate w.r.t. w to obtain

$$F'(w) = -(1 - \pi) + \frac{(K + w_0(1 - \pi))\Lambda}{\phi} \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right)$$

Use $\dot{w} = -\phi$ and $\alpha(w) = w(1 - \pi) + \phi/(\Lambda p)$ and define for any function $F(w)$ and w^* the operator

$$\begin{aligned} \mathcal{D}F(w) &= \Lambda F(w) - F'(w)\dot{w} - (\Lambda p\mu - \kappa - \Lambda(w(1 - \pi) + p(\alpha - w(1 - \pi)))) \\ &= \Lambda F(w) + F'(w)\phi - (\Lambda p\mu - \kappa - \Lambda w(1 - \pi) - \phi), \end{aligned}$$

and note that $\mathcal{D}F(w) = 0$ if and only if F solves (25) under the optimal controls. We use the above expressions for $F(w)$ and $F'(w)$ and obtain

$$\begin{aligned} \mathcal{D}F(w) &= \Lambda \left[\mu p - \frac{\kappa + \phi}{\Lambda} - w(1 - \pi) - (K + w_0(1 - \pi)) \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right) \right] \\ &\quad - (\Lambda p\mu - \kappa - \Lambda w(1 - \pi) - \phi) - \phi \left[(1 - \pi) - \frac{(K + w_0(1 - \pi))\Lambda}{\phi} \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right) \right] \end{aligned}$$

Simple algebra yields indeed $\mathcal{D}F(w) = 0$.

The proof is complete after w_0 from (26) and verifying that

$$\begin{aligned} F(w_0) - F(0) - K &= 0 \\ F'(w_0) &= 0. \end{aligned}$$

holds. This can be done by straightforward algebra. The first condition can be verified calculating

$$F(w_0) - F(0) = -w_0(1 - \pi) - (K + w_0(1 - \pi)) \frac{e^{-\frac{\Lambda w_0}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + (K + w_0(1 - \pi)) \frac{1}{1 - e^{-\frac{\Lambda w_0}{\phi}}} = K,$$

as desired.

For the second condition take the first order condition $F'(w_0) = 0$. Note that

$$\begin{aligned} F'(w_0) &\propto \Lambda(K + w_0(1 - \pi))e^{-\frac{\Lambda w_0}{\phi}} - \left(1 - e^{-\frac{\Lambda w_0}{\phi}} \right) \phi(1 - \pi) \\ &\propto \Lambda(K + w_0(1 - \pi)) - \left(e^{\frac{\Lambda w_0}{\phi}} - 1 \right) \phi(1 - \pi) \end{aligned}$$

A unique solution to $F'(w_0) = 0$ exists because $F'(0) > 0$, $\lim_{w_0 \rightarrow \infty} F'(w_0) < 0$ and $F''(w) < 0$.

A Taylor expansion (i.e., $e^x = 1 + x + x^2/2 + o(x^3)$) and simplifications yield

$$\Lambda(K + w_0(1 - \pi)) - \left(e^{\frac{\Lambda w_0}{\phi}} - 1 \right) \phi(1 - \pi) = \Lambda K - \frac{\Lambda^2 w_0^2 (1 - \pi)}{2\phi} - \xi,$$

where $0 < \xi \in o(K^3)$. It follows that

$$w_0 = \left(\frac{2K\phi}{\Lambda(1 - \pi)} \right)^{1/2} - \chi,$$

where $\chi \in o(K^{3/2})$, implying that

$$w_0 < \bar{w} = \left(\frac{2K\phi}{\Lambda(1-\pi)} \right)^{1/2}.$$

The proof is complete by virtue of the Picard-Lindelof theorem, ensuring uniqueness of the solution. $\square$

Lemma 3. *Let $\Delta^W(w) = W(w) - w$ under the optimal contract with monitoring. The ODE*

$$(\Delta^W)'(w)\dot{w} = \Lambda\Delta^W(w) + \Lambda\pi w \quad (27)$$

with $\dot{w} = -\phi$ and $\Delta^W(0) = \Delta^W(w_0) + w_0$ has the unique solution on $[0, w_0]$.

$$\Delta^W(w) = \frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + \pi \left(\frac{\phi}{\Lambda} - w \right).$$

Hence:

$$W(w) = \frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + \pi \left(\frac{\phi}{\Lambda} - w \right) + w.$$

Proof. We just verify that the proposed function indeed is the desired solution to the ODE subject to the stipulated boundary conditions.

We differentiate the candidate solution:

$$(\Delta^W)'(w) = -\frac{\phi}{\Lambda} \frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} - \pi.$$

Define

$$\mathcal{D}\Delta^W(w) = \Lambda\Delta^W(w) - (\Delta^W)'(w)\dot{w} + \Lambda\pi w,$$

and note that $\mathcal{D}F(w) = 0$ if and only if $F(w)$ solves (27) with $\dot{w} = -\phi$. Then use $(\Delta^W)'(w)$ and $\Delta^W(w)$ to evaluate $\mathcal{D}\Delta^W(w)$:

$$\mathcal{D}\Delta^W(w) = \Lambda \left[\frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + \pi \left(\frac{\phi}{\Lambda} - w \right) \right] - \phi \left[\frac{\phi}{\Lambda} \frac{(1-\pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} - \pi \right] + \Lambda\pi w = 0.$$

In addition, note that

$$\Delta^W(0) - \Delta^W(w_0) = \pi w_0 + \frac{(1-\pi)w_0(1 - e^{-\frac{\Lambda w_0}{\phi}})}{1 - e^{-\frac{\Lambda w_0}{\phi}}} = w_0,$$

as desired. This concludes the proof by virtue of the Picard-Lindelof Theorem, ensuring uniqueness of the solution. $\square$

Lemma 4. *Fix $w_0 > w_L > 0$. The ODE (19) with $\dot{w} = -\phi$, $\alpha - w(1-\pi) = \phi/\Lambda p$, and (20) has the following closed form solution on $(w_L, w_0]$:*

$$f(w) = \left(e^{-B(w_L)}(F(w_L)(1 - q(w_L)) - w_L) \right) e^{B(w)} + e^{B(w)} \int_{w_L}^w e^{-B(x)} a(x) dx.$$

with

$$F(w) = \left(\mu p - \frac{\kappa}{\Lambda}\right) \left(1 - \exp\left(\frac{-w\Lambda}{\phi}\right)\right) - w \quad \text{and} \quad B(w) = \frac{-(p + (1 - \pi)(1 - p))}{(1 - p)(1 - \pi)} e^{-\frac{\Lambda(1-p)(w_0-w)}{\phi}}$$

$$a(w) = \frac{(1 - q(w))\Lambda p(\mu - w(1 - \pi)) - (1 - q(w))\phi - \kappa}{\phi} \quad \text{and} \quad q(w) = \left(1 - e^{-\frac{\Lambda(1-p)(1-\pi)(w_0-w)}{\phi}}\right)$$

Proof. After substituting the optimal controls $\alpha = w(1 - \pi) + \phi/(\Lambda p)$ and $\dot{w} = -\phi$, the ODE to solve becomes

$$f'(w) = \frac{1}{\phi} \left((1 - q(w))\Lambda p(\mu - \alpha(w)) - \kappa - (\Lambda(p + (1 - p)(1 - \pi))(1 - q(w))f(w) \right).$$

This is a first order linear ODE of the general form

$$f'(w) = a(w) + b(w)f(w)$$

with

$$a(w) = \frac{(1 - q(w))\Lambda p(\mu - \alpha(w)) - \kappa}{\phi} = \frac{(1 - q(w))\Lambda p(\mu - w(1 - \pi)) - (1 - q(w))\phi - \kappa}{\phi}$$

and

$$b(w) = \frac{-\Lambda(1 - q(w))(p + (1 - p)(1 - \pi))}{\phi} = \frac{-\Lambda(p + (1 - p)(1 - \pi))}{\phi} e^{-\frac{\Lambda(1-p)(1-\pi)(w_0-w)}{\phi}}$$

Note that

$$B(w) = \frac{-(p + (1 - \pi)(1 - p))}{(1 - p)(1 - \pi)} e^{-\frac{\Lambda(1-p)(w_0-w)}{\phi}}$$

is anti-derivative of $b(w)$ in that $B'(w) = b(w)$. The fundamental theorem of calculus implies that

$$A(w) = \int_{w_L}^w e^{-B(x)} a(x) dx$$

is anti-derivative of $e^{-B(w)} a(w)$.

It is well known that the first order linear differential equation of form $f'(w) = a(w) + b(w)f(w)$ admits the general solution

$$f(w) = C e^{B(w)} + e^{B(w)} \int_{w_L}^w e^{-B(x)} a(x) dx$$

with constant C . We solve for C , using the boundary condition

$$f(w_L) = F(w_L)(1 - q(w_L)) - q(w_L)w_L,$$

which yields

$$C = e^{-B(w_L)} (F(w_L)(1 - q(w_L)) - w_L),$$

as desired. □

B Omitted Proofs

To start with, note that, because $dc_t = 0$ for $t > \tau^A$, wage payments c (i.e., (28)) can be rewritten as

$$dc_t = \left(\alpha_t \mathbf{1}_{\{\text{Success at time } t\}} + \beta_t \mathbf{1}_{\{\text{Failure report at time } t\}} + \gamma_t \mathbf{1}_{\{\text{Failure observed at time } t\}} \right) \mathbf{1}_{\{t \leq \tau^A\}}. \quad (28)$$

That is, because the project yields maximally once failure or success, the values of α_t , β_t , and γ_t after time τ^A (i.e., for $t > \tau^A$) are irrelevant. This observation is convenient since we do not (always) have to explicitly distinguish between the two scenarios $\tau^A < t$ and $\tau^A > t$, when describing $\alpha_t, \beta_t, \gamma_t$.

Throughout, we define the agent's expected reward for failure at time t as

$$r_t := (1 - \pi)w_t + \pi\gamma_t, \quad (29)$$

where w is defined as in (10). To ease the exposition, we refer to “the project fails at time t but failure is not observed by the principal” as “hidden failure”. Likewise, “the project fails at time t and failure is observed by the principal” is called “public failure”.

Last, as is standard in the literature on optimal contracts, we call a contract $\mathcal{C}$ *incentive compatible* if $\mathcal{C}$ induces full effort (i.e., $a_t = 1$ for $t \leq T \wedge \tau$) and truthful disclosure of failure, whenever the principal asks the agent to disclose failure.

B.1 Agent's incentive compatibility

Lemma 5. *A contract $\mathcal{C}$ induces truthful disclosure of failure (i.e., $\tau^A = \tau$ with certainty) from time t' onwards if and only if (7), (8) hold for all $t \in [t', T]$ and $T < \infty$. It induces full effort $a_t = 1$ for all $t \in [0, T \wedge \tau]$ if and only if $\alpha_t \geq r_t + \phi/(\Delta p)$ for all $t \in [0, T \wedge \tau]$.*

Proof. Without loss of generality, consider for the proof $t' = 0$. First, consider any time $t \geq \tau$ and that the project has failed already at time τ and failure has been privately observed by the agent. Then, if the agent has not reported failure yet up to time t , his (continuation) payoff becomes

$$w_t := \max_{\tau^A \in [t, T]} [\phi(\tau^A - t) + \beta_{\tau^A}], \quad (30)$$

given a contract with deadline $T \geq t$. The above expression for w_t is maximized for $\tau^A = t$, only if $\frac{\partial w_t}{\partial \tau^A} = \dot{\beta}_{\tau^A} + \phi \leq 0$ for $\tau^A = t$, which is equivalent to (8) and a necessary condition for truthful disclosure of failure. Also note that since the project may complete at any time $t \in [0, T]$, truthful disclosure of failure requires $\frac{\partial w_t}{\partial \tau^A}|_{\tau^A=t} \leq 0$ for all $t \in [0, T]$. That is, $\dot{\beta}_t = \dot{w}_t \leq -\phi$ must hold for all $t \in [0, T]$. After integrating, we obtain $\beta_s \leq \beta_t - (s - t)\phi$ for all $s \in [t, T]$. Because the agent's limited liability requires $\beta_s \geq 0$ at any time s and because it is clear that $\beta_t < \infty$, it follows that $T < \infty$.

On the other hand, if (8) holds for any $t \in [0, T]$ with $\beta_0 < \infty$ and $T < \infty$, it follows that $\beta_t \geq \beta_s + (s - t)\phi$ for any $s \leq T$, so that w_t is maximized for $\tau^A = t$ for any $t \in [0, T]$ and the contract induces $\tau^A \leq \tau$. Hence, $w_t = \beta_t$ for any $t \in [0, T \wedge \tau]$.

Third, take now $t < \tau$, let $r_t = (1 - \pi)w_t + \pi\gamma_t$ and note that the agent's continuation payoff

reads

$$W_t := \int_t^T e^{-\int_t^s \Lambda(s-t)} \left(\Lambda((1-pa_s)r_s + pa_s\alpha_s) + \phi(1-a_s) \right) ds \quad (31)$$

$$= \int_t^T e^{-\int_t^s \Lambda(s-t)} \left(\Lambda(r_s + pa_s(\alpha_s - r_s)) + \phi(1-a_s) \right) ds, \quad (32)$$

if he chooses $\tau^A \geq \tau > t$. On the other hand, deviating and reporting $\tau^A = t < \tau$ yields payoff β_t , so that truthful disclosure of failure, i.e., $\tau^A = \tau$ with certainty, requires (7) (i.e., $W_t \geq \beta_t$) to hold for any t , with W_t defined in (31).

Fourth, note that at any time $t < \tau$, effort $\{a_s\}_{s \in [t, T \wedge \tau]}$ maximizes W_t (see (31)) if and only if it maximizes pointwise (i.e., for all $s \geq t$) the (scaled) integrand of the expression for W_t :

$$\Lambda(r_s + pa_s(\alpha_s - r_s)) + \phi(1-a_s)$$

for all $s \in [t, T]$. As a result, $a_s = 1$ for all $s \geq t$ and therefore $a_t = 1$ for all $t \in [0, T \wedge \tau]$ if and only if $\alpha_t \geq r_t + \phi/(\Lambda p)$, i.e., with $\beta_t = w_t$, if and only if (11) holds for all $t \in [0, T \wedge \tau]$. This holds for any $T \geq 0$, even for $T = \infty$. The proof is now complete. $\square$

B.2 Proof of Proposition 1

Proof. Throughout, due to the observability of failure it follows that $w_t = \beta_t$. The previous Lemma tells us that incentive compatibility requires $\alpha_t \geq \gamma_t + \phi/(\Lambda p)$, where $\gamma_t = r_t$.

For any $t < \tau$, the principal's payoff can be written as

$$F_t = \int_t^T e^{-\Lambda(s-t)} \left(\Lambda p(\mu - \alpha_s) - \Lambda(1-p)\gamma_s - \kappa \right) ds.$$

Given a deadline T , the payoff F_t is maximized if $\{\alpha_s, \gamma_s\}_{s \geq t}$ maximize pointwise the integrand, while respecting incentive compatibility and limited liability. Subject to these constraints, the integrand is maximized pointwise for $\gamma_s = 0 < \alpha_s = \phi/(\Lambda p)$, so that the principal's payoff becomes $F_t = \int_t^T e^{-\Lambda(s-t)} (\Lambda p \mu - \phi - \kappa) ds$. Because parameters satisfy $\phi \leq \kappa$ and $\mu p > 2\kappa/\Lambda$, the integrand is positive for any $s \geq t$, so that the principal's payoff is maximized by setting $T = \infty$. Hence, the principal's payoff becomes $\mu p - \frac{\phi + \kappa}{\Lambda}$, while the agent's payoff equals $W_t = \frac{\phi}{\Lambda}$. As a result, the proposed contract is optimal and incentive compatible (i.e., induces full effort). $\square$

B.3 Proof of Proposition 2

We prove a more general version of Proposition 2.

Proposition 5. *Under the optimal full disclosure contract $\mathcal{C}$, at time t with $w_t = w$, the principal's value is given by (13). The contract $\mathcal{C}$ stipulates $\dot{w}_t = -\phi$ and termination at time $T = \inf\{t \geq 0 : w_t = 0\}$. Payments satisfy*

$$\alpha_t - r_t = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_t - \gamma_t) \quad \text{with} \quad \gamma_t \in [0, w_t], \quad (33)$$

and $\beta_t = w_t = W_t$. The value w_0 solves $F'(w_0) = 0$.

The claim of Proposition 2 is attained by setting $\gamma_t = 0$.

Proof. A full disclosure contract $\mathcal{C} = (c, T)$ — by definition — induces $\beta_t = w_t$ for all $t \in [0, T]$. Denote the overall (expected) surplus by $\mathcal{S}$, which is given at time t by

$$\mathcal{S}_t = \mathcal{S}(w_t) = \int_t^T e^{-\Lambda(s-t)} (\Lambda p \mu - \kappa + \phi(1 - a_s)) ds = \int_t^T e^{-\Lambda(s-t)} (\Lambda p \mu - \kappa) ds,$$

with $T = \inf\{t \geq 0 : w_t = 0\}$ and full effort $a_s = 1$ for $s \geq t$ in optimum (as shirking is inefficient). The surplus is split between the agent and the principal, so that the principal's payoff $\hat{F}_t$ (at any time t) is given by:

$$\hat{F}_t = \mathcal{S}_t - W_t \leq \mathcal{S}_t - w_t =: F_t = F(w_t),$$

where we used the incentive compatibility condition (7), $W_t \geq w_t = \beta_t$. As a result, a contract is optimal in the class of full disclosure contracts, if it maximizes for any time $t \geq 0$ the continuation surplus $\mathcal{S}_t$ subject to the incentive constraints (11), (8), subject to the agent's limited liability, i.e., $T = \inf\{t \geq 0 : w_t = \beta_t = 0\}$, and achieves $W_t = w_t$.

Notably, for any $t \leq T$, $\mathcal{S}_t$ does not depend α_t and monotonically increases in T . Incentive compatibility requires $\dot{w}_s \leq -\phi$. Limited liability requires $w_s \geq 0$, which yields combined with $\dot{w}_s \leq -\phi$ that $T - t \leq w_t/\phi$ with equality if and only $\dot{w}_s = -\phi$ for all $s \geq t$. That is, setting $\dot{w}_s = -\phi$ for all $s \geq t$ and hence binding the constraint (8) maximizes the deadline and continuation surplus $\mathcal{S}_t$ at time t . Note that the proposed contract from Proposition 2 sets $\dot{w}_s = -\phi$ for all $s \geq t$ and therefore maximizes for any t with given value w_t the (time to) deadline T and hence continuation surplus $\mathcal{S}_t$.

Next, for any w_t , the proposed contract from Proposition 5 stipulates payments according to (33); that is:

$$\alpha_t - r_t = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_t - \gamma_t) \quad \text{and} \quad \gamma_t \in [0, w_t]$$

and, therefore, the incentive constraint w.r.t. effort (11) is met for all $t \in [0, T]$. Differentiate (31) to obtain

$$\dot{W}_t = \Lambda W_t - \Lambda(r_t + p(\alpha_t - r_t)) = -\phi,$$

whereby the second equality follows from plugging in (33) and $r_t = (1 - \pi)w_t + \pi\gamma_t$. Thus, $\dot{W}_t = \dot{w}_t = -\phi$ and, therefore, due to $W_T = w_T = 0$ it follows that $W_t = w_t$ for all $t \in [0, T]$ (so that (7) is met). The starting value w_0 is determined to maximize the principal's ex-ante value (and so must solve $F'(w_0) = 0$). As a result, the proposed contract must be optimal in the class of full disclosure contracts.

The principal's value can be written as

$$\begin{aligned} F_t &= \mathbb{E}_t \left[\int_t^{T \wedge \tau^A} (\mu dS_t - \kappa dt - dc_t) \right] = \int_t^T e^{-\Lambda(s-t)} \left(\Lambda[p(\mu - \alpha_s) - (1 - p)r_s] - \kappa \right) ds \\ &= \int_t^T e^{-\Lambda(s-t)} \left(\Lambda[p(\mu - (\alpha_s - r_s)) - r_s] - \kappa \right) ds \\ &= \int_t^T e^{-\Lambda(s-t)} \left(\Lambda(p\mu - w_s) - \phi - \kappa \right) ds, \end{aligned}$$

where the first equality uses integration by parts and the second equality plugs in $\alpha_t - r_t = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_t - \gamma_t)$. Hence, the value function F_t indeed does not depend on the exact values of α_t, γ_t as long as (14) and $\gamma_t \in [0, w_t]$ hold.

Differentiating for $t \in [0, T)$ yields

$$\dot{F}_t = \Lambda F_t - \Lambda p(\mu - w_t) + \phi + \kappa.$$

One arrives at the ODE (13) (under the optimal controls) using

$$\dot{F}_t = \frac{dF_t}{dt} = \frac{dF_t}{dw_t} \frac{dw_t}{dt} = F'(w_t) \dot{w}_t. \quad (34)$$

The closed-form solution is provided in Lemma 1 in Appendix A. Clearly, F is concave, so that the first order condition $F'(w_0) = 0$ is sufficient in determining optimal w_0 . The proof is now complete. $\square$

B.4 Proof of Proposition 3

We prove a more general version of Proposition 3.

Proposition 6. *The optimal contract does not incentive disclosure of failure over some time period $[0, t_1)$ and becomes a full disclosure contract with deadline T , as characterized in Proposition 2, after time t_1 .*

1. *With $w_{t_1} = w_L$, the optimal time t_1 is characterized by (20) and (21), while the value function F^* is characterized by (22). In addition, $t_1 \rightarrow 0$ as $\pi \rightarrow 0$ and $t_1 \rightarrow T$ as $\pi \rightarrow 1$.*
2. *The contract is terminated at time $T = \inf\{t \geq 0 : w_t = 0\}$ and w_0 (and equivalently T) maximizes the principal's initial payoff $F(w_0)$.*
3. *$\alpha_t = w_t(1 - \pi) + \phi/(\Lambda p) \geq \beta_t = \gamma_t = 0$, $\dot{w}_t = -\phi < \dot{W}_t < 0$, and $W_t < w_t$ for all $t \in [0, t_1)$.*
4. *$w_t = \beta_t = W_t$ and $\alpha_t - r_t = \phi/(\Lambda p) + \pi/p(w_t - \gamma_t)$, and $\gamma_t \in [0, w_t]$ for all $t \in [t_1, T)$.*

The claim of Proposition 3 is attained by setting $\gamma_t = 0$.

The proof of Proposition 6 involves several steps, that are — for a better overview — separately presented in the following Lemmata in this Section. Importantly, in Steps I through IV of the proof, we conjecture that the optimal contract termination time (i.e., financing deadline) T is deterministic. Then, Proposition 7 in Step V of the proof verifies that the optimal contract termination time (i.e., financing deadline) T is indeed deterministic in that random termination does not improve the principal's payoff.

We give a brief overview on how we proceed. Step I demonstrates that whenever $W_t \geq w_t$, the optimal contract incentivizes disclosure of failure from time t onward (until termination at time T). Step II shows that a full disclosure contract is not optimal when $\pi > 0$. Step III characterizes the agent's incentives, and her pay for success and failure. In addition, Step III shows that when $\pi > 0$, then the optimal contract involves exactly two stages: i) an unconditional financing stage, in which the agent is not incentivized to disclose failure, followed by a ii) disclosure stage, in which the agent is incentivized to disclose failure. Step IV characterizes the principal's value function. Step V verifies that the optimal financing deadline T is deterministic, so that random termination does not improve the principal's payoff.

B.4.1 Step I

Lemma 6. *At any time $t < \tau$ with $W_t \geq w_t > 0$, the optimal continuation contract is a full disclosure contract with $T - t = w_t/\phi$, $\beta_s = w_s = W_s$, $\alpha_s - r_s = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_s - \gamma_s)$, and $\gamma_s \in [0, w_s]$ for $T \wedge \tau > s \geq t$. In addition, there does not exist a time t with $W_t > w_t$.*

Proof. Take any time $t < T \wedge \tau$ with $W_t \geq w_t > 0$. Given a deadline T , the continuation surplus is clearly maximized only if i) the agent does not shirk (i.e., $a_s = 1$ for $s \in [t, T \wedge \tau]$) and ii) financing is terminated once the project completes, in that $\tau = \tau^A$ and $T_0 = T \wedge \tau$. This is because shirking is inefficient, i.e., $\phi < \kappa$.

Note that $\dot{w}_s \leq -\phi$ and limited liability, $w_s \geq 0$, imply that $T - t \leq w_t/\phi$, whereby the continuation surplus increases in $T - t$. As a result, the maximum continuation surplus attainable at time t is given by

$$\mathcal{S}_t = \mathcal{S}(w_t) = \int_t^T e^{-\Lambda(s-t)}(\Lambda p\mu - \kappa)ds \quad \text{with} \quad T - t = \frac{w_t}{\phi}.$$

The continuation surplus at time t is split between the agent and the principal so that

$$F_t + W_t \leq \mathcal{S}_t \implies F_t \leq \mathcal{S}_t - W_t. \quad (35)$$

A full disclosure continuation contract from t onwards achieves the continuation surplus $\mathcal{S}_t$, since it i) precludes shirking, ii) induces $\tau^A = \tau$ and $T_0 = T \wedge \tau$, and iii) optimally sets $\dot{w}_s = -\phi$ for all $s \geq t$ and hence maximizes the time to deadline $T - t$. That is, (35) holds in equality, in that $F_t = \mathcal{S}_t - W_t$.

Note that the principal's payoff F_t only depends on $\{\alpha_s, \beta_s, \gamma_s\}_{s \geq t}$ via

$$W_t = \int_t^T e^{-\Lambda(s-t)} \Lambda(p\alpha_s + (1-p)r_s)ds.$$

Because $W_t \geq w_t$, one way to deliver continuation payoff to the agent is by setting $\alpha_s - r_s \geq \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_s - \gamma_s)$, $\beta_s = w_s$, and $\gamma_s \in [0, w_s]$ for $s \in [t, T]$, where the inequality holds in equality for all times s with $W_s = w_s$. This is because Proposition 2 shows that within a full disclosure contract, one way to deliver continuation payoff $W_t = w_t$ to the agent is to set $\alpha_s - r_s = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_s - \gamma_s)$, $\beta_s = w_s$, and $\gamma_s \in [0, w_s]$ for $s \in [t, T]$. Hence, higher pay for success is needed to deliver a higher value $W_t \geq w_t$ to the agent within a full disclosure contract.

This choice of payments satisfies incentive compatibility w.r.t. effort (4). In addition, this choice of payments implies that W_s cannot drift below w_s , so that $W_s \geq w_s = \beta_s$ for all $s \in [t, T]$ and the incentive condition (7) is satisfied for all $s \in [t, T]$. Moreover, due to $\beta_s = w_s$ for all $s \in [t, T]$, it follows that $\dot{w}_s = \dot{\beta}_s = -\phi$, so that (8) is met for all $s \in [t, T]$. Thus, the requirement to deliver value W_t to the agent implies all incentive constraints that are relevant to incentivize truthful disclosure of failure.

As a result, a full disclosure continuation contract from t onwards maximizes $\mathcal{S}_t$ and — because it achieves (35) to hold in equality — the principal's continuation payoff F_t subject to all relevant incentive constraints, limited liability, and the requirement to deliver payoff W_t to the agent.

Let $t = \inf\{s \geq 0 : W_s = w_s\}$. Because the optimal full disclosure continuation contract from t onwards induces $W_s = w_s$ (when $W_t = w_t$) for $s \geq t$ (see Proposition 5 and its proof) and because a full disclosure continuation contract becomes optimal at time t once $W_t = w_t$, it cannot be that $W_s > w_t$ at any time s within the optimal contract. It follows that a full disclosure continuation contract becomes optimal the first time t with $W_t = w_t$, in which case $\alpha_s - r_s = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_s - \gamma_s)$, $\gamma_s \in [0, w_s]$, $\dot{w}_s = \dot{W}_s = -\phi$ and $W_s = w_s = \beta_s$ for $s \geq t$. $\square$

B.4.2 Step II

Lemma 7. *Let $\pi \in [0, 1)$. A full disclosure contract is not optimal if and only if $\pi > 0$. Put differently, if and only if $\pi > 0$, there exists time $t_1 > 0$ such that the optimal contract does not incentivize disclosure of failure over $[0, t_1)$.*

Proof. Let $\pi > 0$ and suppose to the contrary that a (optimal) full disclosure contract $\mathcal{C}^0$ with payoff F^0 is optimal. By Proposition 2 and its proof, this full disclosure contract (optimally) sets $w_t = W_t = \beta_t$, $\dot{\beta}_t = -\phi$, $T = \inf\{t \geq 0 : \beta_t = 0\}$ and $\alpha_t = \beta_t(1 - \pi + \pi/p) + \phi/(\Lambda p)$ and $\gamma_t = 0$ for all $t \in [0, T \wedge \tau]$. This particular choice of $\{\alpha_t, \gamma_t\}_{t \geq 0}$ simplifies the notation.

Take some $T > \Delta > 0$ and define a contract $\tilde{\mathcal{C}}$ as follows. $\tilde{\mathcal{C}}$ sets $\alpha_t = w_t(1 - \pi) + \phi/(\Lambda p)$, $\beta_t = \gamma_t = 0$ for $t < \Delta$ and, if the project is not completed by time Δ , it switches at time Δ to the (optimal) full disclosure contract in state $w_\Delta = w_0 - \Delta\phi$, with deadline $T = w_\Delta/\phi$, $\beta_t = w_t$, $\dot{w}_t = -\phi$, $\gamma_t = 0$ and $\alpha_t = w_t(1 - \pi + \pi/p) + \phi/(\Lambda p)$ for all $t \in [\Delta, T \wedge \tau]$. Note that the agent optimally reports failure at time Δ , if the project fails before time Δ and failure is privately observed, implying that $w_t = (\Delta - t)\phi + \beta_\Delta$ for $t \leq \Delta$. We refer to “the project fails at time t but failure is not observed by the principal” as “hidden failure”. Likewise, “the project fails at time t and failure is observed by the principal” is called “public failure”.

We calculate for any $t \in [0, \Delta]$ the likelihood of hidden failure up to time t , given by (18):

$$q_t = \left(1 - e^{-\Lambda(1-p)(1-\pi)t}\right) = q_0 + \dot{q}_0 t + o(t^2) = \Lambda(1-p)(1-\pi)t + o(\Delta^2),$$

where we used a Taylor expansion around $t = 0$ and that $o(t^2) = o(\Delta^2)$. Hence, the likelihood of hidden failure over $[0, \Delta]$ equals $\Lambda(1-p)(1-\pi)\Delta + o(\Delta^2)$, the likelihood of public failure over $[0, \Delta]$ equals $\Lambda(1-p)\pi\Delta + o(\Delta^2)$, and the likelihood that the project succeeds over $[0, \Delta]$ equals $\Lambda p\Delta + o(\Delta^2)$. In addition, for all $t \in [0, \Delta]$

$$\alpha_t = \alpha_s + \dot{\alpha}_s(t - s) + o(\Delta^2) = \alpha_s + o(\Delta) \quad \text{for any } s \in [0, \Delta],$$

as $|s - t| = o(\Delta)$.

It follows that

$$\bar{\alpha}_\Delta := \mathbb{E}[\alpha_\tau | \text{Success over } [0, \Delta]] = \alpha_s + o(\Delta) \quad \text{for any } s \in [0, \Delta],$$

which is the expected compensation in case of success over $[0, \Delta]$. Likewise, one calculates that the expected financing costs over $[0, \Delta]$ equal

$$\bar{\kappa}_\Delta = \kappa(\Delta - \mathbb{P}(\text{Success or public failure over } [0, \Delta])\mathbb{E}[\Delta - \tau | \text{Success or public failure over } [0, \Delta]]).$$

Hence, $\bar{\kappa}_\Delta = \kappa\Delta - o(\Delta^2)$. Also note that

$$\bar{\beta}_\Delta := \mathbb{E}[\beta_\tau | \text{Hidden failure over } [0, \Delta]] = \beta_s + o(\Delta) \quad \text{for any } s \in [0, \Delta].$$

Here, $\bar{\beta}_\Delta$ is the expected compensation for hidden failure under the full disclosure contract.

Take arbitrary $s \in [0, \Delta)$. The contract $\tilde{\mathcal{C}}$ implements the same deadline T as $\mathcal{C}^0$ and yields

payoff at time zero

$$\begin{aligned}
F^1 &= (1 - e^{-\Lambda\Delta}) (p(\mu - \bar{\alpha}_\Delta) - (1 - p)(1 - \pi)w_\Delta) - \bar{\kappa}_\Delta + e^{-\Lambda\Delta}F(w_\Delta) + o(\Delta^2) \\
&= \Lambda\Delta(p(\mu - \alpha_s) - (1 - p)(1 - \pi)w_\Delta) - \kappa\Delta + (1 - \Lambda\Delta)F(w_\Delta) + o(\Delta^2) \\
&= \Lambda\Delta\left(p\mu - w_\Delta(1 - \pi) - \frac{\phi}{\Lambda}\right) - \kappa\Delta + (1 - \Lambda\Delta)F(w_\Delta) + o(\Delta^2),
\end{aligned}$$

where it was used in the last inequality that $\alpha_s = \phi/(\Lambda p) + w_\Delta(1 - \pi) + o(\Delta)$ for $s \in [0, \Delta)$ under the contract $\tilde{\mathcal{C}}$.

Next, we can write the payoff under the full disclosure contract as

$$\begin{aligned}
F^0 &= (1 - e^{-\Lambda\Delta}) (p(\mu - \bar{\alpha}_\Delta) - (1 - p)(1 - \pi)\bar{\beta}_\Delta) - \bar{\kappa}_\Delta + e^{-\Lambda\Delta}F(w_\Delta) + o(\Delta^2) \\
&= \Lambda\Delta(p(\mu - \alpha_s) - (1 - p)(1 - \pi)w_s) - \kappa\Delta + (1 - \Lambda\Delta)F(w_\Delta) + o(\Delta^2) \\
&= \Lambda\Delta\left(p\mu - w_\Delta - \frac{\phi}{\Lambda}\right) - \kappa\Delta + (1 - \Lambda\Delta)F(w_\Delta) + o(\Delta^2),
\end{aligned}$$

where it was used in the last inequality that $\alpha_s = \phi/(\Lambda p) + w_\Delta(1 - \pi + \pi/p) + o(\Delta)$ and $\beta_s = w_s$ for $s \in [0, \Delta)$ under the full disclosure contract.

Hence:

$$F^1 - F^0 = \Lambda\Delta w_\Delta \pi + o(\Delta^2),$$

which exceeds zero for Δ sufficiently small, contradicting the optimality of $\mathcal{C}^0$. Note that $w_\Delta > 0$ as $T > \Delta$.

On the other hand, if $\pi = 0$, it follows, due to the incentive condition $\alpha_t \geq w_t + \phi/(\Lambda p)$, that

$$\dot{W}_t = \Lambda(W_t - w_t - p(\alpha_t - w_t)) \geq \Lambda(W_t - w_t) - \phi \geq \dot{w}_t = -\phi$$

and, therefore, due to $W_T = w_T = 0$ that $W_t \geq w_t$. Hence, a full disclosure contract is optimal by virtue of Lemma 6. $\square$

B.4.3 Step III

Lemma 8. Define $t_1 = \inf\{t \geq 0 : W_t \geq w_t\}$. Then, within the optimal contract, $t_1 > 0$ if $\pi > 0$, and $\alpha_t = w_t(1 - \pi) + \phi/(\Lambda p)$, $\beta_t = \gamma_t = 0$ for $t < t_1$. And, $\alpha_t - r_t = \frac{\phi}{\Lambda p} + \frac{\pi}{p}(w_t - \gamma_t)$, $\gamma_t \in [0, w_t]$, and $\beta_t = w_t$ for $t \geq t_1$. In addition, $W_t < w_t$ and $0 > \dot{W}_t > \dot{w}_t$ for $t < t_1$ and $W_t = w_t$ and $\dot{W}_t = \dot{w}_t$ for $t \geq t_1$.

Proof. Note that $W_{t_1} = w_{t_1}$ by continuity. Lemma 6 implies that the continuation contract from time t_1 up to the deadline T is a full disclosure contract, which — by Proposition 5 — yields the second claim of the Lemma. In addition, recall that by Lemma 7 a full disclosure contract is not optimal if and only if $\pi > 0$. As a full disclosure contract is optimal from time t_1 onward (see Lemma 6), it follows that $t_1 > 0$ if and only if $\pi > 0$.

To prove the first claim, fix a time t_1 after which the contract implements a full disclosure contract, while the contract does not incentivize disclosure of failure over $[0, t_1)$ and accordingly sets $\beta_t = 0$ for $t < t_1$. Define

$$\bar{\kappa}_{t_1} = \kappa(t_1 - \mathbb{P}(\text{Success or public failure over } [0, t_1])\mathbb{E}[t_1 - \tau | \text{Success or public failure over } [0, t_1)]).$$

It is clear that — in optimum — the contract implements full effort $a_t = 1$ for all $t \in [0, T \wedge \tau]$. The surplus at time zero generated by such a contract equals

$$\mathcal{S}_0 = e^{-\Lambda t_1}(F(w_{t_1}) + w_{t_1}) - \bar{\kappa}_{t_1} + (1 - e^{-\Lambda t_1})\left(p\mu + \phi(1 - p)(1 - \pi)\mathbb{E}(t_1 - \tau | \tau < t_1)\right),$$

which does not depend on $\{\alpha_t, \gamma_t\}_{t \leq t_1}$. Also note that the continuation surplus at time t_1 is split between the principal and the agent and — due to $w_{t_1} = W_{t_1}$ — equals $F(w_{t_1}) + w_{t_1}$. The principal's payoff equals:

$$F_0^* = \mathcal{S}_0 - W_0,$$

and, therefore, is maximized by the choice of $\{\alpha_t, \gamma_t\}_{t \leq t_1}$, that minimizes W_0 subject to $W_{t_1} = w_{t_1}$ and incentive compatibility $\alpha_t \geq \phi/(\Lambda p) + w_t(1 - \pi) + \pi\gamma_t$. It is clear that W_0 is minimized upon promising zero rewards for failure and the lowest rewards for success possible that induce effort. That is, for $t < t_1$ setting $\alpha_t = \phi/(\Lambda p) + w_t(1 - \pi)$ and $\beta_t = \gamma_t = 0$ is optimal.

Hence, for $t < t_1$

$$W_t = \int_t^T e^{-\Lambda(s-t)} \Lambda \left((1-p)(1-\pi)w_s + p\alpha_s \right) ds.$$

Differentiating yields

$$\begin{aligned} \dot{W}_t &= \Lambda W_t - \Lambda((1-p)(1-\pi)w_t + p\alpha_t) \\ &= \Lambda W_t - \Lambda(w_t(1-\pi) + p(\alpha_t - w_t(1-\pi))) \end{aligned}$$

and therefore

$$\dot{W}_t - \dot{w}_t = \Lambda W_t - \Lambda(w_t(1-\pi) + p(\alpha_t - w_t(1-\pi))) - \dot{w}_t = \Lambda(W_t - w_t) + \Lambda\pi w_t,$$

where we plugged in $\alpha_t = w_t(1-\pi) + \phi/(\Lambda p)$ and $\dot{w}_t = -\phi$ for $t < t_1$. Integrating this ODE for $t < t_1$ subject to $W_{t_1} = w_{t_1}$ yields $W_t - w_t = -\int_t^{t_1} e^{-\Lambda(s-t)} \Lambda\pi w_s ds < 0$ and hence also $0 > \dot{W}_t > \dot{w}_t = -\phi$. $\square$

B.4.4 Step IV

Lemma 9. *The value function is characterized by (22) and solved subject to (20) and (21).*

Proof. Given the previous Lemmata, all that remains is to determine the optimal deadline T and the first time at which the contract incentivizes truthful disclosure of failure, $t_1 = \inf\{t \geq 0 : \beta_t > 0\}$ or equivalently $t_1 = \inf\{t \geq 0 : W_t \geq w_t\}$. These two quantities are determined to maximize the principal's ex-ante payoff F_0^* .

Note that w_t — in optimum — perfectly co-moves with time t (before completion) and therefore can be taken as state variable. Hence, we can equivalently maximize F_0^* over $w_{t_1} = w_L$ and w_0 , which uniquely pins down T and t_1 due to

$$\dot{w}_t = -\phi \quad \forall \quad t \in [0, T \wedge \tau^A] \implies w_t = w_0 - \phi t.$$

We solve the maximization problem sequentially: we first fix $T > 0$, which is equivalent to fixing $w_0 = \inf\{t \geq 0 : w_t = 0\}$ due to $\dot{w}_t = -\phi$ for all $t \in [0, T \wedge \tau^A]$, and maximize over t_1 (or equivalently w_L). We then obtain t_1 in dependence of T and thereafter maximize over T .

Let now

$$q_t = \mathbb{P}_t(\{\text{Hidden failure before } t\}) = 1 - e^{-\Lambda(1-p)(1-\pi)t}, \quad (36)$$

the principal's belief, formed over $[0, t_1)$ that the project has hiddenly failed. For $t < t_1 \wedge \tau^A$, we can rewrite the principal's payoff F_t^* (under the optimal contract) as

$$\begin{aligned}
F_t^* &= \mathbb{E}_t \left[\int_t^{T \wedge \tau^A} (\mu dS_s - \kappa ds - dc_s) \right] \\
&= \mathbb{E}_t \left[\int_t^{t_1 \wedge \tau^A} (\mu dS_s - \kappa ds - dc_s) \right] \\
&\quad + \mathbb{P}_t(\tau^A \geq t_1) \left(P_t(\{\text{No Failure before } t_1\}) F_{t_1}^* - \mathbb{P}_t(\{\text{Failure before } t_1\}) \beta_{t_1} \right) \\
&= \int_t^{t_1} e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^s (1-q_u) du} (\Lambda(1-q_s)p(\mu - \alpha_s) - \kappa) ds \\
&\quad + e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^{t_1} (1-q_u) du} \left(\mathbb{P}_t(\{\text{No Failure before } t_1\}) F_{t_1}^* - \mathbb{P}_t(\{\text{Failure before } t_1\}) \beta_{t_1} \right) \\
&= \int_t^{t_1} e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^s (1-q_u) du} (\Lambda(1-q_s)p(\mu - \alpha_s) - \kappa) ds \\
&\quad + e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^{t_1} (1-q_u) du} \left((1 - \lim_{t \uparrow t_1} q_t) F_{t_1}^* - \lim_{t \uparrow t_1} q_t \beta_{t_1} \right). \tag{37}
\end{aligned}$$

Differentiating (37) on $t \in [0, t_1)$ yields

$$\dot{F}_t^* = \Lambda(p + (1-p)(1-\pi))(1-q_t)F_t^* - \Lambda p(1-q_t)(\mu - \alpha_t) - \kappa.$$

Utilizing

$$\dot{F}_t^* = \frac{dF_t^*}{dt} = \frac{dF_t^*}{dw_t} \frac{dw_t}{dt},$$

one arrives at the ODE (19), whereby $F_t^* = f(w_t)$ for $t \in [0, t_1)$. The closed form solution — given $w_0 > w_L \geq 0$ — is derived in Lemma 4 in Appendix A.

Taking the limit $t \uparrow t_1$ in (37) yields

$$\lim_{t \uparrow t_1} F_t^* = (1 - \lim_{t \uparrow t_1} q_t) F_{t_1}^* - \lim_{t \uparrow t_1} q_t \beta_{t_1}, \tag{38}$$

which leads after substituting $\beta_{t_1} = w_{t_1} = w_L$, $\lim_{t \uparrow t_1} q_t = q(w_L)$, $\lim_{t \uparrow t_1} F_t^* = f(w_L)$, and $F_{t_1}^* = F(w_L)$ to the value matching condition (20).

As is standard for dynamic optimization problems, a necessary optimality condition is the smooth pasting condition

$$\frac{\partial}{\partial t_1} \lim_{t \uparrow t_1} F_t^* = \frac{\partial}{\partial t_1} \left((1 - \lim_{t \uparrow t_1} q_t) F_{t_1}^* - \lim_{t \uparrow t_1} q_t \beta_{t_1} \right). \tag{39}$$

If (39) did not hold, the principal could improve her payoff by increasing or decreasing t_1 , implying that optimal t_1 must adhere to (39).

Condition (39) is equivalent to condition (21), which is obtained after substituting $\beta_{t_1} = w_{t_1} = w_L$, $\lim_{t \uparrow t_1} q_t = q(w_L)$, $\lim_{t \uparrow t_1} F_t^* = f(w_L)$, and $F_{t_1}^* = F(w_L)$.

We have solved for optimal t_1 or equivalently optimal w_L , given w_0 or equivalently T , which results in a value F_0^* . That is, w_L is a function of w_0 . The optimization is complete after maximizing

F_0^* over w_0 , leading to the first order condition

$$\frac{dF_0^*}{dw_0} = \frac{\partial F_0^*}{\partial w_0} + \frac{\partial F_0^*}{\partial w_L} \frac{\partial w_L}{\partial w_0} = 0,$$

and the second order condition $\frac{d^2 F_0^*}{dw_0^2} < 0$. □

Lemma 10. $t_1 \rightarrow 0$ as $\pi \rightarrow 0$ and $t_1 \rightarrow T$ as $\pi \rightarrow 1$.

Proof. If $\pi = 0$, by Lemma 7, a full disclosure contract is optimal, leading to $t_1 = 0$. This yields the first claim by virtue of continuity.

By Proposition 1 and continuity, it must be that $T \rightarrow \infty$ and, therefore, $w_t \rightarrow \infty$ for any $t < \tau$ as $\pi \rightarrow 1$. Incentivizing the agent not to fake failure and to truthfully disclose failure at some time $t < \infty$ requires $W_t \rightarrow \infty$, which cannot be optimal. This proves the second claim. □

B.4.5 Step V

Proposition 7. *Random termination does not improve the principal's payoff derived under the contract from Proposition 3.*

Proof. To begin with, consider that the principal randomly terminates the agent's contract at time $t < \tau$ at endogenous rate $\delta_t \geq 0$ or with some atom of probability Θ_t , in that the termination time T is stochastic. Recall that upon termination at time T , the agent receives zero payoff. As a result, the agent's payoff after the project has failed at time $t < T$ (and failure is not publicly observed) is

$$\begin{aligned} w_t &:= \max_{\tau^A \geq t} \mathbb{E}_t^A[(\tau^A \wedge T - t)\phi + \mathbb{P}(\tau^A < T)\beta_{\tau^A}] \\ &= \max_{\tau^A \in [t, T]} \int_t^{\tau^A} e^{-\int_t^s \delta_u du} \prod_{t \leq u \leq s} (1 - \Theta_u) \phi ds + e^{-\int_t^{\tau^A} \delta_u du} \prod_{t \leq u \leq \tau^A} (1 - \Theta_u) \beta_{\tau^A}, \end{aligned} \quad (40)$$

where the second equality integrates out the random termination event. We can differentiate (40) with respect to time, t , and obtain

$$dw_t = (\delta_t w_t - \phi)dt + \Theta_t w_t. \quad (41)$$

Thus, if the principal terminates the project with some atom of probability $\Theta_t > 0$, then w_t increases by $\Theta_t w_t$ (i.e., $w_t = \frac{\lim_{s \uparrow t} w_s}{1 - \Theta_t}$), if the project is not terminated (which happens with probability $1 - \Theta_t$), and w_t drops to zero (i.e., $w_t = 0$), if the project is terminated at time t (which happens with probability Θ_t). In the following, we consider that random termination with some atom of probability is not optimal whenever $w_t > 0$, in that $\Theta_t = 0$ and $\dot{w}_t = \delta_t w_t - \phi$ (for $w_t > 0$). At the end of the proof, we verify that this is indeed the case.

In principle, there are two state variables: the reward for failure $w_t = w$, evolving according to (41), and the belief $q = q_t$, which is given by (18) so that $\dot{q}_t = (1 - q_t)(1 - p)(1 - \pi)\Lambda > 0$. In turn, we can express the principal's value function at time t , denoted f_t , and the agent's continuation value (before failure), denoted W_t , as functions of (w, q) , in that $f_t = f(w_t, q_t)$ and $W_t = W(w_t, q_t)$. To simplify notation, we omit time subscript whenever no confusion is likely to arise. Note that by (18), there exists a one-to-one mapping from the belief q_t to time t . Likewise, without random termination, it follows that $\delta_t = 0$ and $\dot{w}_t < 0$, so there exists a one-to-one mapping from time t to w_t . The one-to-one mapping between time t and w_t need not exist if there is the possibility of random termination (i.e., $\delta_t > 0$).

In what follows, we verify that random termination at rate δ does not improve the principal's payoff derived under the contract from Proposition 3. In the general formulation with random termination, the principal's value function at time t , denoted f_t , is a function of (w, q) , in that $f_t = f(w_t, q_t)$. By the dynamic programming principle, the HJB equation becomes

$$(1 - q)\Lambda(p + (1 - p)(1 - \pi))f(w, q) = \max_{\delta \geq 0} \left\{ (1 - q)\Lambda p(\mu - \alpha(w)) - \kappa \right. \\ \left. + f_w(w, q)\phi + \delta(f_w(w, q)w - f(w, q)) + f_q(w, q)\dot{q} \right\} \quad (42)$$

where $\alpha(w)$ is described in Proposition 3 (so that $\alpha'(w) > 0$). When $\delta = 0$ in all states (w, q) , then (42) is equivalent to (19), after expressing q as function of w (i.e., $q_t = q(w_t)$). In (42), a subscript denotes the partial derivative, in that $f_x(w, q) = \frac{\partial f(w, q)}{\partial x}$ for $x \in \{w, q\}$.

We study the principal's incentives to terminate the contract at rate δ . As the contract is terminated (with certainty) once $w = 0$, it suffices to consider $w > 0$. Taking the partial derivative with respect to δ in (42) yields

$$\frac{\partial f(w, q)}{\partial \delta} \propto f_w(w, q)w - f(w, q),$$

where $\propto$ denotes proportionality. Hence, the principal's objective in (42) is linear in δ . Note that if it were $f_w(w, q)w > f(w, q)$, then — by the HJB equation (42) — setting $\delta \rightarrow \infty$ would be optimal and would yield unbounded payoff for the principal, which cannot be. Thus, it holds that

$$f_w(w, q)w - f(w, q) \leq 0,$$

with equality if $\delta > 0$ is optimal. That is, when $\delta > 0$ is optimal, the smooth pasting condition

$$f(w, q) = f_w(w, q)w \quad (43)$$

holds for $(w_t, q_t) = (w, q)$. In addition, the super contact condition

$$f_{ww}(w, q) = 0 \quad (44)$$

must hold for $(w_t, q_t) = (w, q)$ when $\delta_t > 0$ is optimal. The smooth pasting condition (43) can be interpreted as local optimality condition and the super contact condition (44) can be interpreted as global optimality condition; for a more detailed discussion, see Dumas (1991).²⁴ We study the unconditional financing stage and the disclosure stage separately.

First, we consider the disclosure stage with $t \geq t_1$, so that $\dot{q}_t = \dot{q} = 0$ and $q_t = q = 0$. Thus, the principal's value function is given by $F(w) = f(w, 0)$ where the closed-form expression for $F(w)$ is (16). As $F(w)$ is strictly concave, it follows that $F(w) < F'(w)w$ for all $w > 0$ and (43) cannot hold during the full disclosure stage (except at $w = 0$ at termination). As a result, during the disclosure stage, random termination is indeed strictly sub-optimal, in that $\delta_t = 0$ is optimal for all $t \geq t_1$.

Second, we consider the unconditional financing stage (i.e., $t < t_1$) so that $\dot{q}_t > 0$. Suppose to the contrary that during the unconditional financing stage, setting $\delta_t > 0$ is optimal and improves

²⁴To see that the super contact condition must hold, assume to the contrary that (43) holds (i.e., $f(w, q) = f_w(w, q)w$), but (44) does not hold (i.e., $f_{ww}(w, q) \neq 0$). Then, there exists $w' \neq w$ with $w' > 0$ such that $f_w(w', q)w - f(w', q) > 0$. Note that by the HJB equation (42), the condition $f_w(w', q)w - f(w', q) > 0$ implies that the principal derives unbounded payoff in state (w', q) , which cannot be.

the principal's payoff on a non-empty interval of times $t \in [s_1, s_2]$. Then, (43) and (44) hold for a non-empty interval of times $t \in [s_1, s_2]$. Note that (43) implies $f_w(w, q) \geq 0$ for times $t \in [s_1, s_2]$ with $(w, q) = (w_t, q_t)$. On $[s_1, s_2]$, we can differentiate (43) with respect to time, t , to get

$$f_w(w, q)\dot{w} + f_q(w, q)\dot{q} = f_w(w, q)\dot{w} + (f_{ww}(w, q)\dot{w} + f_{wq}(w, q)\dot{q})w.$$

Using the super-contact condition, $f_{ww}(w, q) = 0$, and simplifying yields

$$f_q(w, q) = f_{wq}w. \quad (45)$$

The left-hand-side of (45) is negative, as the value function decreases in the belief q of whether the project has failed so far. Thus, $f_{wq} \leq 0$.

Using the envelope theorem, we differentiate both sides of (42) with respect to w (under the optimal control $\delta = \delta(w)$) to obtain

$$f_{ww}(w, q)\phi = (1 - q)\Lambda(p + (1 - p)(1 - \pi))f_w(w, q) + (1 - q)\Lambda p\alpha'(w) - f_{wq}\dot{q}. \quad (46)$$

When $\delta_t > 0$, (45) and the super-contact condition $f_{ww}(w, q) = 0$ must hold for $(w_t, q_t) = (w, q)$, so that (46) simplifies to

$$f_{wq}(w, q)\dot{q} = (1 - q)\Lambda(p + (1 - p)(1 - \pi))f_w(w, q) + (1 - q)\Lambda p\alpha'(w). \quad (47)$$

As $f_w(w, q) \geq 0$ and $\alpha'(w) > 0$, the right-hand-side of (47) is positive so that $f_{wq} > 0$, a contradiction. Thus, $\delta_t = 0$ at all times during the unconditional financing stage.

Finally, we verify that random termination with some atom of probability $\Theta \in [0, 1)$ is not optimal in any state (w, q) with $w > 0$. By (41), the principal's payoff upon termination with an atom of probability Θ is

$$\mathcal{L}(\Theta) := f(\tilde{w}, q)(1 - \Theta),$$

with $\tilde{w} = \frac{w}{1 - \Theta}$. That is, if the project is not terminated, w increases times $\frac{1}{1 - \Theta}$. Note that

$$\mathcal{L}'(\Theta) = f_w(\tilde{w}, q)\tilde{w} - f(\tilde{w}, q).$$

If termination with probability Θ is optimal, the first order condition $\mathcal{L}'(\Theta) = 0$ and the second order condition

$$\mathcal{L}''(\Theta) = f_{ww}(\tilde{w}, q)\tilde{w} < 0 \iff f_{ww}(\tilde{w}, q) < 0$$

must hold. However, $\mathcal{L}'(\Theta) = 0 > \mathcal{L}''(\Theta)$ implies that there exists $w' \neq \tilde{w}$ with $f_w(w', q)\tilde{w} - f(w', q)$. As a result, by the HJB equation (42), the principal would obtain unbounded payoff in state (w', q) , which cannot be. Thus, termination with some atom of probability cannot be optimal.

Overall, we conclude that random termination does not improve the principal's payoff and is not optimal, which was to show. $\square$

B.5 Proof of Proposition 4

B.5.1 Preliminaries

We first solve the *relaxed problem*, in which the agent *cannot* fake bad outcomes and mis-report failure before it occurs, leading to the constraint in the strategy space $\tau^A \geq \tau$. We show that the optimal contract is a full disclosure contract, when the constraint $\tau^A \geq \tau$ is in place. We also show that under certain parameter conditions, the solution to the relaxed problem also solves the

full problem (without constraint $\tau^A \geq \tau$). Whenever we refer to the *relaxed problem*, we implicitly assume that the constraint $\tau^A \geq \tau$ is in place.

The proof is split in three separate parts. The first part characterizes the agent's incentives. The second part provides the principal's solution to the relaxed problem. The third part shows that the solution to the relaxed problem also solves the full problem. It then follows that the proposed contract from Proposition 4 is the optimal contract.

Throughout the following analysis, we make the Assumption that K is not prohibitively large:

Assumption 1. *Parameters satisfy the following conditions*

$$\frac{\phi(1-\pi)}{2\Lambda} \geq K \iff \bar{w} = \left(\frac{2K\phi}{\Lambda(1-\pi)} \right)^{1/2} \leq \frac{\phi}{\Lambda} \quad (48)$$

and

$$\mu p \geq \frac{3}{2} \frac{\phi(1-\pi)}{\Lambda(1-e^{-1})} + \frac{\kappa + \phi}{\Lambda} \quad (49)$$

Notably, condition (48) is met if $\pi < 1$ and $K > 0$ sufficiently small. In the limit $\pi \rightarrow 1$, monitoring becomes redundant. Condition (49) ensures that $F(0) > K$, whereby F is the (conjectured) value function, given in (26). That is, for given w_0

$$F(0|w_0) = F(0) = \mu p - \frac{\kappa + \phi}{\Lambda} - (K + w_0(1-\pi)) \left(\frac{1}{1 - e^{-\frac{-\Lambda w_0}{\phi}}} \right).$$

Here, $F(w|w_0)$ denotes the principal's value function conditional on a given (not necessarily optimal) choice of w_0 , while $F(w)$ denotes the value function under the optimal choice of w_0 . Lemma 2 implies that $w_0 < \bar{w} < \phi/\Lambda$. Because $w_0 = \phi/\Lambda$ is not (necessarily) optimal, it follows that $F(0) \geq 0 \iff F(0|\phi/\Lambda) \geq 0$ (with $F(0)$ the value under the optimal contract). Note that

$$F(0|\phi/\Lambda) \geq 0 \iff \mu p - \frac{\kappa + \phi}{\Lambda} \geq \frac{K + \phi(1-\pi)/\Lambda}{1 - e^{-1}}.$$

Utilizing condition (48) to substitute K then leads to condition (49).

Last, we define the left limit of a process x_t as $x_{t-} := \lim_{s \uparrow t} x_s$.

B.5.2 Part I — Agent's incentive compatibility

To characterize truth telling incentives, let $dM_t \in \{0, 1\}$ indicate whether the principal inspects the project over a short period of time $(t, t + dt)$, whereby $\mathbb{P}(dM_t = 1) = m_t$ is the likelihood of an inspection. In general, we can write

$$m_t = \theta_t dt + \Theta_t, \quad (50)$$

with $\theta_t \geq 0$ and $\Theta_t \in [0, 1]$. If $\Theta_t = 0$, the principal inspects the firm with infinitesimal probability $\theta_t dt$, i.e., at rate θ_t . Otherwise, if $\Theta_t > 0$, the principal inspects the firm with an atom of probability and $\Theta_t = 1$ implies a deterministic inspection at time t .

We characterize how monitoring provides incentives. The following Lemma derives for all times t with $dM_t = 0$ the truth telling incentive constraint

$$d\beta_t \leq m_t \beta_t - \phi dt. \quad (51)$$

Notably, with $m_t = \theta_t dt + \Theta_t$ the constraint (51) is equivalent to

$$\beta_t \leq \frac{\beta_{t-}}{1 - \Theta_t}, \text{ for } \Theta_t > 0,$$

where $d\beta_t = \beta_t - \beta_{t-}$, and

$$\dot{\beta}_t \leq \theta_t \beta_t - \phi = \theta_t \beta_{t-} - \phi, \text{ for } \Theta_t = 0.$$

Note that the last equality uses that $\beta_{t-} = \beta_t$, when β_t is continuous.

Lemma 11. *A contract $\mathcal{C}$ induces truthful disclosure of failure (i.e., $\tau^A = \tau$ with certainty) from time t' onwards if and only if (51) holds for all $t \in [t', T]$ with $dM_t = 0$ and $T \wedge \tau^M < \infty$ (almost surely) for any t , where $\tau^M = \inf\{s > t : dM_s = 1\}$. It induces full effort, $a_t = 1$ for all $t \in [0, T \wedge \tau]$, if and only if $\alpha_t \geq r_t + \phi/(\Lambda p)$.*

Proof. Without loss of generality, normalize for the proof $t' = 0$. Note that we can rewrite (51) to

$$\beta_t \leq \frac{\beta_{t-}}{1 - \Theta_t} = \frac{\lim_{s \uparrow t} \beta_s}{1 - \Theta_t} \text{ for } \Theta_t > 0$$

and

$$\dot{\beta}_t \leq \theta_t w_t - \phi \text{ for } \Theta_t = 0.$$

First, consider any $t \geq \tau$ and that the project has failed. Define $\tau^M := \inf\{s \geq t : dM_s = 1\}$ as the next time the principal inspects the project, whereby — for convenience — the notation does not make the dependence of τ^M on t explicit. We can without loss of generality assume that it is optimal to terminate financing and to fire the agent (without any severance pay), once an inspection yields that the agent has hidden failure.

Then, given a contract deadline $T \geq t$ the agent's payoff becomes

$$\begin{aligned} w_t &:= \max_{\tau^A \in [t, T]} \mathbb{E}_t^A[(\tau^A \wedge \tau^M - t)\phi + \mathbb{P}(\tau^A < \tau^M)\beta_{\tau^A}] \\ &= \max_{\tau^A \in [t, T]} \int_t^{\tau^A} e^{-\int_t^s \theta_u du} \prod_{t \leq u \leq s} (1 - \Theta_u) \phi ds + e^{-\int_t^{\tau^A} \theta_u du} \prod_{t \leq u \leq \tau^A} (1 - \Theta_u) \beta_{\tau^A}, \end{aligned}$$

where we have integrated out the random inspection event. The above expression is maximized for $\tau^A = t$ only if $\frac{\partial w_t}{\partial \tau^A} = \dot{\beta}_{\tau^A} - \theta_{\tau^A} \beta_{\tau^A} + \phi \leq 0$ for $\tau^A = t$, in case $\Theta_t = 0$, or if $\beta_t(1 - \Theta_t) \leq \beta_{t-}$, in case $\Theta_t > 0$. Notably, this is (51) and a necessary condition for truthful disclosure of failure. Since the project may complete at any time during $t \in [0, T]$, $\tau^A = \tau$ can be achieved with certainty only if (51) holds for all $t \in [0, T]$. Then, we can integrate to obtain

$$\beta_t \geq \int_t^{\tau^A} e^{-\int_t^s \theta_u du} \prod_{t \leq u \leq s} (1 - \Theta_u) \phi ds + e^{-\int_t^{\tau^A} \theta_u du} \prod_{t \leq u \leq \tau^A} (1 - \Theta_u) \beta_{\tau^A}, \quad (52)$$

for all $\tau^A \in [t, T]$. Because clearly $\beta_t < \infty$ and because the agent's limited liability requires $\beta_s \geq 0$ at any time it must be that

$$\mathbb{E}[T \wedge \tau^M] = e^{-\int_t^T \theta_u du} \prod_{t \leq u \leq T} (1 - \Theta_u) < \infty,$$

implying $T \wedge \tau^M < \infty$ almost surely.

On the other hand, if (51) holds for any $t \in [0, \tau^M \wedge T]$ and $\mathbb{E}[T \wedge \tau^M] < \infty$, it follows that (52) holds and hence that $\beta_t \geq w_t$. As a result, w_t is maximized for $\tau^A = t$ so that $\tau^A \leq \tau$ is optimal. Due to the assumed constraint $\tau^A \geq \tau$, this in fact already implies truthful disclosure of failure $\tau^A = \tau$.

Given truthful disclosure of failure, the agent chooses effort $\{a_s\}_{\{s \geq t\}}$ to maximize

$$W_t := \mathbb{E}_t \left[\int_t^{T \wedge \tau} dc_s \right] = \int_t^T e^{-\int_t^s \Lambda(s-t)} \left(\Lambda(r_s + pa_s(\alpha_s - r_s)) + \phi_t(1 - a_s) \right) ds,$$

which boils down to maximize the integrand point-wise. This implies that the incentive condition

$$\alpha_t \geq r_t + \phi/(\Lambda p)$$

must hold for any $t < T \wedge \tau$ so as to induce full effort. $\square$

B.5.3 Part II — The principal's solution to relaxed problem

Lemma 12. *Restrict the agent's strategy to $\tau^A \geq \tau$. Then, the optimal contract is a full disclosure contract, i.e., induces $\tau^A = \tau$ with certainty.*

Proof. Recall that the only reason why the optimal contract does not incentivize truthful disclosure of failure is to provide incentives not to fake failure, i.e., to relax the incentive constraint (7), that is, $W_t \geq \beta_t$. However, with the constraint $\tau^A \geq \tau$ the agent cannot fake failure anymore so that the optimal contract need not respect (7) anymore. More formal arguments follow below.

Take any time $t < \tau$ and fix a deadline T . Take a contract $\mathcal{C}^0$ that does not incentivize truthful disclosure of failure over $[t, t_1)$ with $t_1 \leq T$ but is a full disclosure continuation contract with deadline T after time t_1 , yielding continuation payoff $F(w_{t_1})$ with $w_{t_1} \geq 0$. The contract follows a monitoring strategy $\{m_s\}_{s \in [t, T \wedge \tau^A]}$ and sets rewards for public failure $\{\gamma_s\}_{s \in [0, T \wedge \tau^A]}$.

Then, for $s \in [t, t_1)$ with $dM_s = 0$:

$$dw_s \leq m_s w_s - \phi dt \quad \text{s.t.} \quad w_{t_1} = \bar{w}$$

and incentive compatibility

$$\alpha_s \geq r_s + \phi/(\Lambda p).$$

Because $\{\alpha_s, \gamma_s\}$ (i.e., $\{r_s\}$) does not affect the law of motion of w , it is clear that in optimum $\alpha_s = r_s + \phi/\Lambda p$. The agent's continuation payoff at time t reads

$$W_t = \int_t^T e^{-\Lambda(s-t)} \Lambda(r_s + p(\alpha_s - r_s)) ds.$$

Also note that conditional on the project failing at time $t' \in [t, t_1)$, the project is terminated at time $\bar{t} := \min\{t_1, \tau_{t'}^M\}$, whereby $\tau_{t'}^M = \inf\{s \geq t' : dM_s = 1\} \geq t'$ is the next inspection date after time t' . Thus, conditional on hidden project failure at time $t' \in [t, t_1)$, the principal's payoff at time t equals

$$\mathcal{F}^0(t') = -\mathbb{E} \left[\kappa(\bar{t} - t) + w_{\bar{t}} + \int_t^{\bar{t}} \theta_s K ds + \sum_{t \leq s \leq \bar{t}} \Theta_s K \right].$$

Take now a full disclosure contract $\mathcal{C}^1$ that coincides with the above contract $\mathcal{C}^0$ after time t_1 but sets over $[t, t_1)$, $\beta_s = w_s$ and $\alpha_s = r_s + \phi/(\Lambda p)$, thereby inducing full effort $a_s = 1$ and $\tau^A = \tau$, as this contract satisfies the incentive constraint (51). The full disclosure contract stipulates the same rewards for public failure $\{\gamma_s\}_{s \in [t, T \wedge \tau^A]}$ as $\mathcal{C}^0$. Also assume the full disclosure contract $\mathcal{C}^1$ employs the same monitoring strategy $\{m_s\}_{s \in [t, T \wedge \tau^A]}$ as $\mathcal{C}^0$, meaning that at any time $s < T \wedge \tau^A$ the contract $\mathcal{C}^1$ monitors with the same probability/intensity m_s as the contract $\mathcal{C}^0$. Note that both contracts $\mathcal{C}^0$ and $\mathcal{C}^1$ induce the same law of motion for w ; thus, they have the same deadline but induce potentially different τ^A .

In addition, both contracts $\mathcal{C}^0$ and $\mathcal{C}^1$ deliver the same payoff W_t to the agent at time t and stipulate the same payments to the agent in case of success. Also note that conditional on the project failing at time $t' \in [t, t_1)$, the project is terminated at time t' . Thus, the principal's payoff at time t then equals, conditional on hidden project failure at time $t' \in [t, t_1)$:

$$\mathcal{F}^1(t') = - \left(\kappa(t' - t) + (w_{t'} - w_{\bar{t}}) + w_{\bar{t}} + \int_t^{t'} \theta_s K ds + \sum_{t \leq s \leq t'} \Theta_s K \right)$$

Because of $\phi \leq \kappa$, $\bar{t} \geq t'$, it follows that $\kappa(t' - t) + (w_{t'} - w_{\bar{t}}) < \kappa(\bar{t} - t)$ for any $\bar{t} > t'$. Hence, it follows that $\mathcal{F}^1(t') \geq \mathcal{F}^0(t')$ for any $t' \in [t, t_1)$ (with strict inequality if $t' > t$).

Thus, if the project does not complete over $[t, t_1)$ both contracts $\mathcal{C}^1$ and $\mathcal{C}^0$ yield the same payoff to the principal. Likewise, if there is project success or public project failure over $[t, t_1)$, both contracts $\mathcal{C}^1$ and $\mathcal{C}^0$ yield the same payoff to the principal. But, if the project hiddenly fails over $[t, t_1)$, contract $\mathcal{C}^1$ always yields strictly higher payoff to the principal than contract $\mathcal{C}^0$ (in expectation), because of $\mathcal{F}^0(t') \leq \mathcal{F}^1(t')$ for all $t' \in [t, t_1)$. Because of the arbitrariness of t, t_1 , the deadline T , rewards for public failure $\{\gamma_s\}_{s \in [t, T \wedge \tau^A]}$, and the monitoring strategy $\{m_s\}_{s \in [t, T \wedge \tau^A]}$, it follows that any contract $\mathcal{C}^0$, that does not always incentivize truthful disclosure of failure, cannot be optimal. This concludes the proof. $\square$

Lemma 13. *Under the optimal contract, the principal's value function solves (25) with the proposed controls.*

Proof. We provide a verification argument in the spirit of Sannikov (2008) and recall the left limit of a process x_t as $x_{t-} := \lim_{s \uparrow t} x_{t-}$. Let $F(w)$ the value function under the contract, proposed by the Proposition 4:

$$F(w) = \mu p - \frac{\kappa + \phi}{\Lambda} - w(1 - \pi) - (K + w_0(1 - \pi)) \left(\frac{e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} \right).$$

solving (25) by Lemma 2. It is easy to see that $F''(w) < 0$.

Suppose the principal follows another strategy up to time $t < T \wedge \tau$ and then switches to the strategy proposed by the optimal contract at time t , whereby the alternative strategy induces truthful disclosure of failure. Her payoff then equals

$$G_t = \int_0^t e^{-\Lambda s} (\Lambda p(\mu - \alpha_s) - \Lambda(1 - p)r_{s-} - \kappa - \theta_s K) ds - \sum_{s \leq t} e^{-\Lambda s} \Theta_s K + e^{-\Lambda t} F(w_{t-}).$$

subject to (11) and (51) for $w_t = \beta_t$:

$$dw_t = (\Theta_t + \theta_t dt)w_t - \phi dt + (w_t^* - w_{t-})dM_t + d\ell_t,$$

whereby $d\ell_t \leq 0$ almost surely and w_t^* endogeneous. Note that whenever there is no inspection, i.e., for times t with $dM_t = 0$, this reduces to

$$dw_t = (\Theta_t + \theta_t dt)w_t - \phi dt + d\ell_t,$$

while at an inspection (i.e., for $dM_t = 1$), w_t is set to an endogeneous value w_t^* , that can be freely chosen without affecting local incentive compatibility.

Differentiating G_t with respect to t yields

$$\begin{aligned} e^{\Lambda t} dG_t = & \left\{ -\Lambda F w_{t-} dt + (\Lambda p\mu - \kappa) dt - \Lambda(r_{t-} + p(\alpha_t - r_{t-})) dt + F'(w_{t-})(\theta_t w_{t-} - \phi) dt \right. \\ & + (\Theta_t + \theta_t dt)(F(w_t^*) - F(w_{t-}) - K) + (1 - \Theta_t) \left(F\left(\frac{w_{t-}}{1 - \Theta_t}\right) - F(w_{t-}) \right) \\ & \left. + F(w_{t-} + d\ell_t) - F(w_{t-}) \right\} + (F(w_t^*) - F(w_{t-}) - K)(dM_t - \Theta_t - \theta_t dt) \\ & = dU_t + (F(w_t^*) - F(w_{t-}) - K)(dM_t - \Theta_t - \theta_t dt), \end{aligned} \quad (53)$$

where the term in curly brackets equals dU_t and the second term has zero expectation in that

$$\mathbb{E}[(F(w_t^*) - F(w_{t-}) - K)(dM_t - \Theta_t - \theta_t dt)] = 0.$$

First, note that because of $F'(w_t) > 0 > F''(w_t)$ with $F'(w_0) = 0$, it follows that dU_t is maximized only if $d\ell_t = \max\{w_{t-} - w_0, 0\}$ and $w_t^* = w_0$.

Second, taking the derivative w.r.t. Θ_t yields

$$\frac{dU_t}{d\Theta_t} =: \mathcal{L}(w_{t-}, \Theta_t) := F(w_0) - K - F\left(\min\left\{\frac{w_{t-}}{1 - \Theta_t}, w_0\right\}\right) + F'\left(\min\left\{\frac{w_{t-}}{1 - \Theta_t}, w_0\right\}\right) \frac{w_{t-}}{1 - \Theta_t}.$$

Note that if $w_{t-} > (1 - \Theta_t)w_0$, then $d\ell_t = w_0 - w_{t-}/(1 - \Theta_t)$ so that w_{t-} is never set above w_0 . This implies that $w_t = \min\{w_{t-}/(1 - \Theta_t), w_0\}$ at any time t with $\Theta_t > 0$.

Strict concavity implies that

$$\mathcal{L}^1(w_{t-}, \Theta_t) = F'\left(\min\left\{\frac{w_{t-}}{1 - \Theta_t}, w_0\right\}\right) \frac{w_{t-}}{1 - \Theta_t} - F\left(\min\left\{\frac{w_{t-}}{1 - \Theta_t}, w_0\right\}\right) < 0 \quad \text{for } w_t > 0,$$

and that $\mathcal{L}^1(w_{t-}, \Theta_t)$ decreases in both of its arguments, i.e., in w_{t-} and Θ_t . Hence:

$$\begin{aligned} \mathcal{L}(w_{t-}, \Theta_t) &= F(w_0) - K + \mathcal{L}^1(w_{t-}, \Theta_t) \leq F(w_0) - K + \mathcal{L}^1(w_{t-}, 0) \\ &= F(w_0) - F(w_{t-}) - K \leq F(w_0) - F(0) - K = 0, \end{aligned}$$

where the last inequality is strict for $w_{t-} > 0$. Hence, dU_t is maximized only if $\Theta_t = 0$, whenever $w_{t-} > 0$.

Third, taking the derivative w.r.t. to θ_t yields

$$\frac{dU_t}{d\theta_t} = F(w_0) - K + F'(w_{t-})w_{t-} - F(w_{t-}) \leq F(w_0) - K - F(0) = 0$$

where the inequality follows from strict concavity and is strict for $w_{t-} > 0$. Hence, dU_t is maximized only if $\theta_t = 0$, whenever $w_{t-} > 0$.

Fourth, note that dU_t is maximized only if $\alpha_t = r_{t-} + \phi/(\Lambda p)$ and $r_{t-} = w_{t-}(1 - \pi)$, i.e., $\alpha_t = (1 - \pi)w_{t-} + \phi/(\Lambda p)$

On the other hand, setting $\Theta_t = \mathbf{1}_{\{w_{t-}=0\}}$, $\theta_t = \gamma_t = 0$, $r_{t-} = w_{t-}(1 - \pi)$, and $\alpha_t = r_{t-} + \phi/(\Lambda p) = (1 - \pi)w_{t-} + \phi/(\Lambda p)$ implies $dU_t = 0$, since dU_t reduces to the HJB equation (25) (holding in equality). It follows that $dU_t \leq 0$ for all $t < T \wedge \tau$. Hence, G_t is a super-martingale, implying that

$$F(w_0) = G_0 \geq \mathbb{E}G_t.$$

That is, the contract, that is proposed by the Proposition 4 and yields payoff $F(w_0)$, is indeed optimal among all incentive compatible, full disclosure contracts. $\square$

B.5.4 Part III — Relaxed problem solves full problem

The following Lemma completes the proof of Proposition 4.

Lemma 14. *The solution to the relaxed problem solves the full problem in which the agent can fake failure through reporting $\tau^A < \tau$. In addition, $\phi < \dot{W}_t$ and $W_t \geq w_t$ for all $t \leq \tau$.*

Proof. Let the agent's continuation utility under truthful reporting before time t be W_t , given in (53). We must show that under the proposed solution $W_t \geq w_t$ for all t or equivalently $W(w) \geq w$ in the state space $[0, w_0]$, in which case the agent is not tempted to fake failure, i.e., is not tempted to set $\tau^A < \tau$. Define $\Delta^W(w) = W(w) - w$, i.e., $\Delta_t^W = W_t - w_t$.

Let $\tau^M = \inf\{t \geq 0 : dM_t = 1\}$ the first inspection time. Given the shape of the derived contract and given that stages repeat, it suffices to show that $W_t \geq w_t$ for all $t < \tau^M$. Note that for $t < \tau^M$ the agent's continuation value is given by

$$W_t = \int_t^T e^{-\Lambda(s-t)} (\Lambda(r_s + p(\alpha_s - r_s))) ds,$$

so that

$$\dot{\Delta}_t^W = \dot{W}_t - \dot{w}_t = \Lambda W_t - \Lambda(r_t - p(\alpha_t - r_t)) + \phi = \Lambda \Delta_t^W + \Lambda \pi(w_t - \gamma_t) = \Lambda \Delta_t^W + \Lambda \pi w_t, \quad (54)$$

where it was used that $\gamma_t = 0$ and $\alpha_t = (1 - \pi)w_t + \phi/(\Lambda p)$. That is, $\dot{\Delta}_t^W > 0$ if $\Delta_t^W > 0$. Hence, it follows that $\Delta_t^W \geq 0$ for all $t \geq 0$ if and only if $\Delta_0^W \geq 0$. The proof is therefore complete, if we can show that $\Delta_0^W \geq 0$.

To do so, we transform (54) into a ODE of w rather than of time t , using $\frac{d\Delta_t^W}{dt} = \frac{d\Delta_t^W}{dw_t} \frac{dw_t}{dt}$:

$$(\Delta^W)'(w) \dot{w} = \Lambda \Delta^W(w) + \Lambda \pi w.$$

This ODE is solved subject to $\Delta^W(w_0) + w_0 = \Delta^W(0)$ and admits the unique solution (see Appendix A and Lemma 3 for details):

$$\Delta^W(w) = \frac{(1 - \pi)w_0 e^{-\frac{\Lambda w}{\phi}}}{1 - e^{-\frac{\Lambda w_0}{\phi}}} + \pi \left(\frac{\phi}{\Lambda} - w \right).$$

The first term is always positive. The second term is positive for $w \leq \phi/\Lambda$. Because of $\bar{w} = \left(\frac{2K\phi}{\Lambda(1-\pi)} \right)^{1/2} > w_0$, it follows then that $\Delta^W(w_0) \geq 0$ if $\Delta^W(\bar{w}) \geq 0$. A sufficient condition is given

by

$$\frac{\phi}{\Lambda} \geq \bar{w} \iff \frac{\phi(1-\pi)}{2\Lambda} \geq K,$$

which — by Assumption (1) — is met. Hence, $\Delta_0^W \geq 0$. To conclude the proof, note that for all $t < \tau^M$, it follows that $\dot{\Delta}_t > 0$ so that $\dot{W}_t > \dot{w}_t = -\phi$. $\square$

C Solution details for Section 4

We sketch how to solve the model under the specification presented in Section 4 under the assumption that full effort is efficient (i.e., $\mu - \mu^f$ is sufficiently large). Without loss of generality, we can assume that parameters satisfy $\mu^f < 0$, in that a failed project is inefficient to continue. Otherwise, equations are merely shifted by a constant. Notably, the proof of Corollary 2 does not make use of the heuristic solution presented in this Section. We also assume already that the agent is only paid at time τ^A , in that (28) holds.²⁵

Incentives. Like in the baseline version of the model, it is optimal to not pay the agent for observed failure. The incentive compatibility condition w.r.t. effort is simply

$$\alpha_t \geq w_t(1-\pi) + \frac{\phi}{\Lambda p}. \quad (55)$$

To incentivize the agent not to fake bad outcomes, it must be that $W_t \geq \beta_t$, in that (7) is met like in the baseline version of the model.

As is already derived in Section 4, the agent prefers to disclose failure if and only if

$$\dot{\beta}_t \leq -\phi - \lambda(\alpha_t - \beta_t),$$

which implies in optimum

$$\dot{w}_t = -\phi - \lambda(\alpha_t - w_t). \quad (56)$$

Specifically, one can write

$$w_t := \max_{\tau^A \in [t, T]} \left[\int_t^{T \wedge \tau^A} e^{-\lambda(s-t)} (\phi + \lambda(\alpha_s - w_s)) ds + \beta_{\tau^A} \right],$$

whereby $\tau^A = t = \tau$ under a full disclosure contract, i.e., whenever the principal incentivizes disclosure of failure. By contrast, if the contract features an unconditional financing phase $[0, t_1)$, then $\tau^A = t_1$ for any $\tau = t < t_1$. Also note that $T = \inf\{t \geq 0 : w_t = 0\}$.

Solution: full disclosure contract. A full disclosure contract with deadline T implies that $\beta_t = w_t$, $w_T = \beta_T = 0$, and, optimally, $W_t = w_t$ for all t . Hence, $\dot{W}_t = \dot{w}_t = 0$ for all t , which

²⁵In fact, delaying payments for failure β_t up to time τ^λ could be optimal to provide incentives to the agent not to fake failure. However, so as not to complicate the analysis, we assume that this is not done, e.g., because it is too costly due to high κ .

becomes

$$\begin{aligned}
0 &= \dot{W}_t - \dot{w}_t = \Lambda(W_t - w_t(1 - \pi) - p(\alpha_t - w_t(1 - \pi))) - \dot{w}_t \\
&= -\Lambda p(\alpha_t - w_t(1 - \pi)) + \Lambda \pi w_t + \phi + \lambda(\alpha_t - w_t) \\
&= -(\Lambda p - \lambda)(\alpha_t - w_t(1 - \pi)) + \pi(\Lambda - \lambda)w_t - \phi.
\end{aligned}$$

Hence:

$$\alpha_t = w_t(1 - \pi) + \frac{\phi + \pi(\Lambda - \lambda)w_t}{\Lambda p - \lambda}, \quad (57)$$

implying that the incentive condition w.r.t. effort (55) is slack if $\lambda > 0$ or $\pi > 0$.

Optimal contract. In order to reduce rewards for success α_t and hence the agent's stake W_t and agency costs, the optimal contract features an unconditional financing stage $[0, t_1)$, during which $\alpha_t = w_t(1 - \pi) + \phi/\Lambda$ and $\beta_t = 0$. During that stage, the contract does not incentivize disclosure of failure. Thereafter, the optimal contract is a full disclosure contract with some deadline T , stipulating $\beta_t = w_t$ and $\alpha_t = w_t(1 - \pi) + \frac{\phi + \pi(\Lambda - \lambda)w_t}{\Lambda p - \lambda}$. The further solution steps look like those presented in Section 2 and are therefore omitted.

C.1 Proof of Corollary 2

Proof. Let $\mathcal{T}$ the endogenous set of times, during which the principal incentivizes truthfull disclosure of failure. That is, $\tau \in \mathcal{T} \implies \tau^A = \tau$. A contract $\mathcal{C} = (c, T)$ induces the set $\mathcal{T}$.

Let the time of failure τ^F and define a jump process F such that $t = \tau^F \Leftrightarrow dF_t = 1$ with $F_0 = 0$. Likewise, let the time of success τ^S and define a jump process S such that $t = \tau^S \Leftrightarrow dS_t = 1$ with $S_0 = 0$.

We already impose that the agent is optimally not paid for observed failure. The principal's problem can be written as

$$F_0 := \max_{\mathcal{C}} \mathbb{E} \left[\int_0^{T \wedge \tau^A} (\mu dS_t + \mu^f dF_t - \kappa dt - dc_t) \right]. \quad (58)$$

subject to

$$\begin{aligned}
dc_t &= \left(\alpha_t \mathbf{1}_{\{t=\tau^S\}} + \beta_t \mathbf{1}_{\{t=\tau^F\}} \right) \mathbf{1}_{\{t \leq \tau^A\}} \geq 0 \quad \text{for all } t \in [0, T] \\
\alpha_t &\geq w_t(1 - \pi) + \phi/\Lambda p \quad \text{for all } t \in [0, T] \\
\beta_t &\geq 0 \quad \text{for all } t \in [0, T]
\end{aligned} \quad (59)$$

$$\begin{aligned}
W_t &= \int_t^T e^{-\Lambda(s-t)} (w_s(1 - \pi) + p(\alpha_s - w_s(1 - \pi))) ds \geq \beta_t \quad \text{for all } t \in [0, T] \\
w_t &= \max_{\tau^A \in [t, T]} \left[\int_t^{T \wedge \tau^A} e^{-\lambda(s-t)} (\phi + \lambda(\alpha_s - w_s)) ds + \beta_{\tau^A} \right] = \beta_t \quad \text{for all } t \in \mathcal{T}.
\end{aligned}$$

It can be seen that — all else equal, i.e., holding the deadline T and pay schedules $\{dc_s\}_{s \geq t}$ fixed — w_t increases in λ . This tightens incentive compatibility. Formally, denote the set of contracts $\mathcal{C}$ satisfying (59) by $\mathbf{C}_\lambda$. Then:

$$\mathbf{C}_{\lambda_1} \subset \mathbf{C}_{\lambda_2} \quad \text{for } \lambda_1 > \lambda_2.$$

This is because an increase in λ increases w_t (i.e., $\partial w_t / \partial \lambda > 0$), which tightens incentive compatibility w.r.t. effort (i.e., $\partial \alpha_t / \partial w_t > 0$) and incentive compatibility w.r.t. reporting (i.e., $\partial \beta_t / \partial w_t > 0$ for $t \in \mathcal{T}$). Because in addition $\partial w_t / \partial \alpha_s > 0$ and $\partial w_t / \partial \beta_s \geq 0$ for $s \geq t$, it necessarily must be that $\mathbf{C}_\lambda$ is an increasing set in λ .

The principal's problem can be written as

$$F_0 := \max_{\mathcal{C} \in \mathbf{C}_\lambda} \left(e^{-\Lambda T} (p\mu + (1-p)\mu^F) - \mathbb{E} \left[\int_0^{T \wedge \tau^A} (\kappa dt + dc_t) \right] \right). \quad (60)$$

Notably, by definition a mean preserving spread in (λ, μ) leaves $p\mu + (1-p)\mu^F$ unaffected/constant. As a result, a mean preserving spread in (λ, μ) affects the principal's payoff only via the optimization constraint $\mathcal{C} \in \mathbf{C}_\lambda$. As the set $\mathbf{C}_\lambda$ decreases in λ , it follows that the principal's payoff increases, when a mean preserving spread increases μ but decreases λ and μ^f . $\square$

D Additional results

D.1 Implementation with debt and equity

We provide the details for the implementation of the optimal contract with debt and equity. For the implementation, we consider a slightly different version of the optimal contract that differs from the contract in Proposition 3 only in the values of (α_t, γ_t) during the disclosure stage (but this contract features the same deadlines (t_1, T) and the same unconditional financing stage). In this contract, $\beta_t = \gamma_t = w_t$ during the disclosure stage and $\alpha_t = w_t + \phi/(\Lambda p)$. Proposition 6 demonstrates that this alternative contract is optimal and yields the same payoffs for principal and agent as the contract from Proposition 3. The only difference between the alternative contract and the contract from Proposition 3 is that to incentivize the agent not to fake failure during the disclosure stage, the principal boosts the agent's rents by stipulating rewards for publicly observed failure $\gamma_t = \beta_t$ rather than by stipulating excessive rewards for success.

D.1.1 Unconditional financing stage

We start by characterizing the implementation of the unconditional financing stage, $[0, t_1)$. At the beginning at time $t = 0$, the principal allocates κt_1 dollars to the project, which is sufficient for project development until time t_1 . Thus, at time $t > 0$, there are $\kappa(t - t_1)$ dollars left (within the firm).

We consider that during the unconditional financing stage, the project is financed with a mixture of debt and equity. There is one unit of debt outstanding that has face value $D := \kappa t_1$ and accrues non-compounding interest at rate R . Thus, at time $t \geq 0$, the face value of debt including interest is $D(1 + Rt)$. Interest is paid only when the face value (principal) is paid back. Debt including interest is paid back during the unconditional financing stage only when the project succeeds at time $t \in [0, t_1)$. Otherwise, if the project fails at some time $t \in [0, t_1)$, debt defaults and there are zero repayments. If neither failure nor success occurs over $[0, t_1)$, debt remains outstanding until the next financing stage. As a result, note that the form of debt we consider in this implementation resembles a credit line the principal grants to the agent.

There is one unit of equity outstanding. During the unconditional financing stage, the agent holds e units of equity and the principal holds the remainder. If the project succeeds at time

$t \in [0, t_1)$, equity pays dividends

$$\mu + \kappa(t_1 - t) - D(1 + Rt),$$

which is the sum of the project payoff μ and the remaining project development funds $\kappa(t_1 - t)$ net debt repayments $D_t(1 + R_t)$. Equity does not pay dividends if the agent discloses failure at time t_1 .

During the unconditional financing stage, the agent must receive

$$e[\mu + \kappa(t_1 - t) - D(1 + Rt)] = \alpha_t \iff e = \frac{\alpha_t}{\mu + \kappa(t_1 - t) - D(1 + Rt)}, \quad (61)$$

dollars for success upon time t . In addition, we know that $\dot{\alpha}_t = -\phi(1 - \pi)$ (see Propositions 3 and 6), so we can differentiate (61) with respect to time, t , and obtain the interest rate

$$R = \frac{\phi(1 - \pi) - e\kappa}{eD} \quad (62)$$

We can solve (61) and (62) to get closed-form expressions for e and R . Specifically, evaluate (61) at time $t = 0$ and use $D = \kappa t_1$ to get

$$e = \frac{\alpha_0}{\mu} \quad (63)$$

and

$$R = \frac{\phi(1 - \pi)\mu/\alpha_0 - \kappa}{D}. \quad (64)$$

Also note that if the agent reports failure at some time $t \in [0, t_1)$, debt holders enjoy seniority and seize all funds $\kappa(t_1 - t)$ that are left within the firm. Thus, the agent finds it never optimal to disclose failure during $[0, t_1)$.

D.1.2 Disclosure stage

Next, we look at the implementation of the disclosure stage, $[t_1, T]$. At the beginning of the disclosure stage at time t_1 , the principal allocates $\kappa(T - t_1) + \xi$ dollars to the project. Thus, there are $\kappa(T - t) + \xi$ dollars left within the firm at any time $t \in [t_1, T]$. In exchange for the funds she contributes, the principal receives additional equity, but no new debt is issued at the beginning of the disclosure stage. Thus, the face value of outstanding debt including interest is $D^* := D(1 + Rt_1)$ at the beginning of the disclosure stage. During the disclosure stage, debt accrues non-compounding interest at rate R^* , so the face value of outstanding debt including interest is $D^*(1 + R^*(t - t_1))$. Interest is paid only when the face value (principal) is paid back. Debt (including interest) is paid back upon project completion or at the deadline T . If the project fails, the agent optimally reports failure and financing is terminated. Then, debt has seniority over equity and debt holders are paid first from the remaining funds $\kappa(T - t) + \xi$. What is left after repaying debt is distributed as dividends to equity holders.

During the disclosure stage, the agent owns e^* units of equity and there is one unit of equity outstanding. Equity pays

$$\mu + \kappa(T - t) + \xi - D^*[1 + R^*(t - t_1)]$$

dollars as dividends in case of success at time $t \in [t_1, T]$ and

$$\max\{0, \kappa(T - t) + \xi - D^*[1 + R^*(t - t_1)]\}$$

in case of failure at time $t \in [t_1, T]$. We implement the contract such that $\kappa(T - t) + \xi - D^*[1 + R^*(t - t_1)] \geq 0$, so there is always cash left within the project after failure and after paying back debt including interest. Thus, equity pays a dividend in the event of failure. That is, during the disclosure stage (but not during the unconditional financing stage), funds allocated to the project exceed the face value of debt (plus accrued interest), so termination of financing due to failure leads to repayment of debt and dividend payouts to equity generating rewards for failure for the agent.

According to the optimal contract, the agent must receive

$$e^* [\mu + \kappa(T - t) + \xi - D^*[1 + R^*(t - t_1)]] = \alpha_t \quad (65)$$

dollars in case of success at time $t \in [t_1, T]$ and

$$e^* [\kappa(T - t) + \xi - D^*[1 + R^*(t - t_1)]] = \beta_t \quad (66)$$

in case of failure at time $t \in [t_1, T]$.

We subtract (66) from (65) to get

$$e^* \mu = \alpha_t - \beta_t = \frac{\phi}{\Lambda p} \iff e^* = \frac{\phi}{\Lambda p \mu}. \quad (67)$$

Next, we evaluate (66) at $t = t_1$ to solve

$$D^* = \xi + \kappa(T - t_1) - \frac{\beta_{t_1}}{e^*}. \quad (68)$$

Next, we choose ξ to achieve $D^* = D(1 + R t_1)$. In summary, debt (raised at inception at time $t = 0$) is paid back either i) when the project succeeds during the unconditional financing stage, ii) when the project is completed during the disclosure stage, or iii) at the deadline T . Otherwise, if the project fails during the unconditional financing stage, debt defaults.

Next, we differentiate (66) with respect to time t and note that $\dot{\beta}_t = \dot{\alpha}_t = -\phi$, leading to

$$R^* = \frac{\phi - e^* \kappa}{e^* D^*}. \quad (69)$$

Note that

$$e = \frac{\alpha_0}{\mu} > \frac{\alpha_0 - (1 - \pi)w_0}{\mu} = \frac{\phi}{\Lambda p \mu} > e^*,$$

so the agent's equity stake is lower during the disclosure stage than during the unconditional financing stage. Using that $D^* = D$, it follows that

$$R^* = \frac{\phi - e^* \kappa}{e^* D} > \frac{\phi(1 - \pi) - e \kappa}{e D} = R. \quad (70)$$

Finally, note that the form of debt we consider in this type of implementation resembles credit line debt. Our implementation therefore rationalizes the use of credit lines in venture capital financing.

D.1.3 Implementation with convertible debt and convertible equity

Generally, the implementation of the optimal contract is not unique. Note that the optimal contract can alternatively be implemented, also using convertible debt and convertible equity. In what follows, we demonstrate how to integrate convertible debt and convertible equity into the financ-

ing of the project, as convertible debt and convertible equity are widely used in venture capital financing.

First, to integrate convertible equity, note that it is always possible to grant the principal additional convertible equity at time $t = 0$ that converts into equity at time t_1 , as long as the pre-specified terms and conversion rate ensure that the agent holds fraction e of project equity before time t_1 and fraction e^* after time t_1 .

Second, to integrate convertible debt, it is possible to stipulate that some of the debt raised at time $t = 0$ converts at time t_1 into equity, as long as the pre-specified terms and conversion rate ensure that the agent holds fraction e of equity before time t_1 and fraction e^* after time t_1 . To see this, observe that the implementation of the optimal contract requires (68) to hold, but the actual values of D^* and ξ in (68) do not matter. Recall that the implementation with debt and equity sets $D^* = D(1 + Rt_1)$, so that at time t_1 , no debt is paid back or is converted into equity. However, by stipulating that $D^* < D(1 + Rt_1)$, one could implement the decrease in debt outstanding from $D(1 + Rt_1)$ to D^* at time t_1 as part of the debt being converted into equity, i.e., debt is partially convertible debt. That is, by stipulating $D^* < D(1 + Rt_1)$, one finances the project, (partially) using convertible debt.

D.2 Solution when agent affects completion timing

We present an alternative formulation of the moral hazard problem similar to [Mason and Välimäki \(2015\)](#), [Green and Taylor \(2016\)](#), or [Varas \(2017\)](#). In this alternative model, the agent controls project completion, while the project is subject to failure risk during its development phase. The agent affects the project completion timing $\tau = \inf\{t \geq 0 : dN_t = 1\}$ with his effort $a_t \in \{0, 1\}$, whereas project failure at time $\tau^\delta := \inf\{t \geq 0 : dN_t^\delta = 1\}$ occurs at for simplicity exogenous rate $\delta > 0$. Specifically, assume that N and N^δ are jump processes with $\mathbb{E}dN_t = \Lambda a_t dt \cdot \mathbf{1}_{\{t < \tau^\delta\}}$ and $\mathbb{E}dN_t^\delta = \delta dt$.

The agent derives private benefits $\phi(1 - a_t)$ as long as the project receives sufficient financing, i.e., before time $T_0 = T \wedge \tau^A$. Completion at time τ always results into success and yields terminal payoff μ to the principal. Failure (during project development) occurs at time τ^δ , in which case the project becomes worthless and produces zero payoffs for all times $t \geq \tau^\delta$. In particular, due to $\mathbb{E}dN_t = \Lambda a_t dt \cdot \mathbf{1}_{\{t < \tau^\delta\}}$, the project cannot be (successfully) completed anymore after failure has occurred during the project development phase. Also note that by exerting effort to complete the project, the agent accelerates completion and hence reduces the risk of failure during project development, in that the agent effectively controls project failure (risk) too. Like in the baseline model, failure is publicly observed with probability $\pi \in [0, 1]$. Otherwise, failure is privately observed by the agent and reported failure is not verifiable.

Incentive Compatibility. The agent is not paid for observed failure but obtains payoff w_t upon privately observed failure. We take the agent's continuation value for $t < T \wedge \tau \wedge \tau^\delta$:

$$W_t = \int_t^T e^{-(\Lambda + \delta)(s-t)} (\Lambda \alpha_s + \delta(1 - \pi)w_s) ds,$$

so that

$$\dot{W}_t = (\Lambda + \delta)W_t - \delta(1 - \pi)w_t - \Lambda \alpha_t.$$

Like in [Green and Taylor \(2016\)](#), the incentive condition w.r.t. effort a_t becomes

$$\alpha_t - W_t \geq \frac{\phi}{\Lambda}. \quad (71)$$

The intuition is as follows. If the agent exerts effort a_t , the project succeeds with probability Λdt in which case the agent receives a reward for success α_t but loses his continuation payoff W_t . On the other hand, shirking over a time interval of length dt yields private benefits ϕdt but then the project does not complete for sure.

Like in our baseline model, the incentive condition w.r.t. disclosure of failure is

$$\dot{\beta}_t \leq -\phi,$$

which in optimum leads to $\dot{w}_t = -\phi$. In addition, the agent must not find it optimal to fake bad outcomes, which requires

$$W_t \geq \beta_t.$$

Full disclosure contract. In a (optimal) full disclosure contract, it holds that $W_t = w_t = \beta_t$ and $\dot{w}_t = \dot{\beta}_t = -\phi$ for all t . Due to $W_T = w_T = 0$, this requires

$$0 = \dot{W}_t - \dot{w}_t = \Lambda(W_t - \alpha_t) + \delta\pi w_t + \phi.$$

Hence:

$$\alpha_t = W_t + \frac{\phi + \delta\pi w_t}{\Lambda},$$

so that the incentive condition w.r.t. effort (71) is slack.

Optimal contract. In order to reduce rewards for success α_t and hence the agent's stake W_t and agency costs, the optimal contract features an unconditional financing stage $[0, t_1)$, during which $\alpha_t = W_t + \phi/\Lambda$ and $\beta_t = 0$. During that stage, the optimal contract does not incentivize disclosure of failure. Thereafter, the optimal contract is a full disclosure contract with some deadline T , stipulating $\beta_t = w_t$ and $\alpha_t - W_t = \frac{\phi + \delta\pi w_t}{\Lambda}$. The further solution steps look like those presented in Section 2 and are therefore omitted.

D.3 Solution when success is unobservable

Suppose the principal observes success — just like failure — only with probability π . Otherwise, success is privately observed by the agent. Crucially, reported success is — unlike failure — verifiable. Also note that to obtain a non-trivial solution, at least one of the two possible outcomes, success and failure, must be verifiable. We consider that success is verifiable as it is arguably more difficult to fake good outcomes rather than bad outcomes.

We start with some notation. Denote the pay for (publicly) observed success by ω_t and denote the pay for reported success by α_t . Then, the agent discloses success truthfully if and only if $\dot{\alpha}_t \leq -\phi$. Likewise, the agent discloses failure truthfully if and only if $\dot{\beta}_t \leq -\phi$.

Full disclosure contract. Let us start by looking at the disclosure stage or — equivalently — at a full disclosure contract (recall that the disclosure stage is a full disclosure continuation contract). We illustrate that it is not consequential whether success is observable or not during the disclosure

stage (i.e., during a full disclosure contract). To start with, note that the full disclosure contract from Proposition 2 stipulates

$$\alpha_t = w_t \left(1 - \pi + \frac{\pi}{p}\right) + \frac{\phi}{\Lambda p} \implies \dot{\alpha}_t < \dot{w}_t = -\phi.$$

Hence, within the optimal full disclosure contract from Proposition 2 the agent always strictly prefers to disclose success truthfully. That is, for this contract the observability of success does not matter (one can set $\alpha_t = \omega_t$).

In the following, we take the optimal full disclosure contract (see Proposition 5) that sets $\gamma_t = w_t$ and $\alpha_t = w_t + \frac{\phi}{\Lambda p}$, thereby “minimizing” rewards for success. In this contract, $\beta_t = w_t$ and

$$\dot{\beta}_t = \dot{\alpha}_t = -\phi,$$

so that the agent possesses sufficient incentives to disclose failure and success truthfully. Rewards for observed and reported success (failure) are equal, in that $\omega_t = \alpha_t$. In this full disclosure contract, the agent’s stake, capturing agency costs, is given by $W_t = w_t$.

Optimal contract. We construct the optimal contract using heuristic arguments. A formal proof conjectures the shape of the contract — which is discussed here — and then provides a (formal) verification argument — which is omitted here. The optimal contract features an unconditional financing stage $[0, t_1]$ during which it does not incentivize disclosure of failure and success. Thereafter, after t_1 , the contract becomes a full disclosure contract with time to deadline $T - t_1$, stipulating $\beta_t = \gamma_t = w_t$ and $\alpha_t = \omega_t = w_t + \phi/(\Lambda p)$. It therefore holds — by construction — that $W_{t_1} = w_{t_1}$.

If the project fails (succeeds) at time $t < t_1$, the agent reports failure (success) at time t_1 . Hence, privately observed project completion at time t yields payoff $\beta_{t_1} + \phi(t_1 - t)$, in case of failure, and payoff $\alpha_{t_1} + \phi(t_1 - t) = \beta_{t_1} + \phi(t_1 - t) + \phi/(\Lambda p)$ in case of success. In addition, the agent is not paid for observed failure, i.e., $\gamma_t = 0$ for $t < t_1$, and receives pay $\omega_t = \phi/(\Lambda p)$ for observed success, when $t < t_1$.

Let us look at the agent’s incentives to exert effort over $[t, t + dt)$ with $t < t_1$. Shirking entails benefits ϕdt and, if the project completes, the project fails. Failure is observable with probability π , in which case the agent receives zero payoff, and otherwise with probability $1 - \pi$ the agent’s payoff is $\beta_{t_1} + \phi(t_1 - t)$. If the agent works and the project succeeds, the agent receives pay $\phi/(\Lambda p)$, when success is observed (with probability π), and otherwise with probability $1 - \pi$ pay $\alpha_{t_1} + \phi(t_1 - t) = \beta_{t_1} + \phi(t_1 - t) + \phi/(\Lambda p)$. Therefore, the agent prefers to work (i.e., to exert full effort $a_t = 1$) if and only if

$$\begin{aligned} & \Lambda dt \cdot \left(p[(1 - \pi)(\alpha_{t_1} + \phi(t_1 - t)) + \pi\phi/(\Lambda p)] + (1 - p)(1 - \pi)[\beta_{t_1} + \phi(t_1 - t)] \right) \\ & \geq \phi dt + \Lambda dt \cdot (1 - \pi)[\beta_{t_1} + \phi(t_1 - t)]. \end{aligned}$$

The first line depicts the agent’s (expected) pay upon exerting effort and the second line depicts the agent’s pay upon shirking. Using simple algebra, it can be verified that the above condition is — by construction of the proposed contract — satisfied. More straightforwardly, the proposed contract motivates effort because the agent’s pay for success in each state of the world exceeds his pay for failure by $\phi/(\Lambda p)$, which outweighs the disutility of effort.

Note that the agent's continuation utility during the unconditional financing stage is given by

$$W_t = \int_t^T e^{-\Lambda(s-t)} \left[\Lambda \left((1-p)(1-\pi)w_s + p[\pi\omega_s + (1-\pi)(w_s + \phi/(\Lambda p))] \right) \right] ds,$$

so that

$$\begin{aligned} \dot{W}_t &= \Lambda W_t - \Lambda \left((1-p)(1-\pi)w_t + p[\pi\omega_t + (1-\pi)(w_t + \phi/(\Lambda p))] \right) \\ &= \Lambda(W_t - w_t) - \phi + \pi w_t. \end{aligned}$$

The second equality follows after plugging in the previously derived expressions for ω_t and α_t . Hence, due to $\dot{w}_t = -\phi$:

$$\dot{W}_t - \dot{w}_t = \Lambda(W_t - w_t) + \pi w_t,$$

and because of $W_{t_1} = w_{t_1}$ it follows that $W_t < w_t$ and $0 > \dot{W}_t > \dot{w}_t = -\phi$ for $t < t_1$. Hence, the provision unconditional financing over some period $[0, t_1)$ reduces the agent's stake W_t and hence agency costs relative to a full disclosure contract. Thus, the provision of unconditional financing over $[0, t_1)$ is optimal and the optimal contract (likely) takes the conjectured shape. A rigorous proof can be constructed along the lines of the proof of Proposition 6.

D.4 Calculating the average financing horizon

Take $\mathcal{E} = \mathbb{E}[T \wedge \tau^A]$ the average length of the financing period. Define

$$\mathcal{E}_t := \mathbb{E}_t[T \wedge \tau^A | t < T \wedge \tau^A]$$

and note that $\mathcal{E} = \mathcal{E}_0$ and $\mathcal{E}_T = T$, by definition.

On $[t_1, T]$ for $t < T \wedge \tau^A$, we can write

$$\mathcal{E}_t = \mathbb{E}_t[T \wedge \tau^A] = \mathbb{E}_t[T \wedge \tau] = \mathbb{E}_t \left[\int_t^{T \wedge \tau} s ds + \mathbf{1}_{\{T \geq \tau\}} T \right] = \int_t^T e^{-\Lambda s} \Lambda s ds + e^{-\Lambda(T-t)} T,$$

where the second equality uses truthful disclosure of failure, $\tau = \tau^A$ and the third equality uses integration by parts. Thus, differentiating the above expression w.r.t. t implies that $\mathcal{E}_t$ solves on $[t_1, T]$ the time ODE

$$\Lambda \mathcal{E}_t - \dot{\mathcal{E}}_t = \Lambda t, \tag{72}$$

subject to $\mathcal{E}_T = T$. We obtain the closed-form solution

$$\mathcal{E}_t = t + \frac{1 - e^{-\Lambda(T-t)}}{\Lambda}.$$

Taking this solution, we proceed and solve backwards in time.

Recall that

$$q_t = 1 - e^{-\Lambda(1-p)(1-\pi)t}.$$

At $t = t_1$, the contract elicits a progress report. If the project has failed already, which is with probability $\lim_{t \uparrow t_1} q_t$, the financing deadline is $T \wedge \tau^A = \tau^A = t_1$, since the agent reports failure at t_1 . Otherwise, with probability $1 - \lim_{t \uparrow t_1} q_t$, the project is not complete at time t_1 , in which case

the expected financing date is given by $\mathcal{E}_{t_1}$. This leads to the value matching condition

$$\lim_{t \uparrow t_1} \mathcal{E}_t = \mathcal{E}_{t_1} (1 - \lim_{t \uparrow t_1} q_t) + t_1 \lim_{t \uparrow t_1} q_t. \quad (73)$$

On $[0, t_1)$, the financing is only terminated if the project succeeds, which happens at rate $\Lambda(1 - q_t)p$. Thus, for $t < T \wedge \tau^A$ we can write

$$\begin{aligned} \mathcal{E}_t &= \mathbb{E}_t \left[\int_t^{T \wedge \tau^A} s ds + \mathbf{1}_{\{T \geq \tau\}} T \right] \\ &= \int_t^{t_1} e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^s (1-q_u) du} \Lambda(p + (1-p)(1-\pi)) s ds \\ &\quad + e^{-\Lambda(p+(1-p)(1-\pi)) \int_t^{t_1} (1-q_u) du} \lim_{t \uparrow t_1} \mathcal{E}_t. \end{aligned}$$

Differentiating w.r.t. t implies that $\mathcal{E}_t$ solves the time ODE

$$\Lambda(p + (1-p)(1-\pi))(1 - q_t)\mathcal{E}_t - \dot{\mathcal{E}}_t = \Lambda(p + (1-p)(1-\pi))(1 - q_t)t, \quad (74)$$

subject to (73). This is a first order linear ODE and can be solved in closed form.

Take

$$B_t = -\frac{p + (1-p)(1-\pi)}{(1-p)(1-\pi)} e^{-\Lambda(1-p)(1-\pi)t},$$

which is the anti-derivative of $\Lambda(p+(1-p)(1-\pi))(1-q_t)$ and define $a_t = -\Lambda(p+(1-p)(1-\pi))(1-q_t)t$. The ODE (74) can be rewritten in the form $\dot{\mathcal{E}}_t = \dot{B}_t \mathcal{E}_t + a_t$ and it is well known that such an ODE possesses general solution on $[0, t_1)$

$$\mathcal{E}_t = C e^{B_t} + e^{B_t} \int_t^{t_1} e^{-B_s} a_s ds.$$

The constant is determined using the boundary condition (73):

$$C = e^{-B_{t_1}} \left((1 - \lim_{t \uparrow t_1} q_t) + t_1 \lim_{t \uparrow t_1} q_t \right),$$

yielding the solution

$$\mathcal{E}_t = e^{-B_{t_1}} \left((1 - \lim_{t \uparrow t_1} q_t) + t_1 \lim_{t \uparrow t_1} q_t \right) e^{B_t} + e^{B_t} - \int_t^{t_1} e^{-B_s} a_s ds.$$

References

- Acharya, V. V., P. DeMarzo, and I. Kremer (2011). Endogenous information flows and the clustering of announcements. *American Economic Review* 101(7), 2955–79.
- Aghion, P. and J. Tirole (1994). The management of innovation. *The Quarterly Journal of Economics* 109(4), 1185–1209.
- Bergemann, D. and U. Hege (1998). Venture capital financing, moral hazard, and learning. *Journal of Banking & Finance* 22(6-8), 703–735.
- Bergemann, D. and U. Hege (2005). The financing of innovation: Learning and stopping. *RAND Journal of Economics*, 719–752.
- Bertomeu, J., I. Marinovic, S. J. Terry, and F. Varas (2019). The dynamics of concealment. *Working Paper*.
- Beyer, A., D. A. Cohen, T. Z. Lys, and B. R. Walther (2010). The financial reporting environment: Review of the recent literature. *Journal of Accounting and Economics* 50(2-3), 296–343.
- Biais, B., T. Mariotti, G. Plantin, and J.-C. Rochet (2007). Dynamic security design: Convergence to continuous time and asset pricing implications. *The Review of Economic Studies* 74(2), 345–390.
- Biais, B., T. Mariotti, J.-C. Rochet, and S. Villeneuve (2010). Large risks, limited liability, and dynamic moral hazard. *Econometrica* 78(1), 73–118.
- Board, S. and M. Meyer-ter Vehn (2013). Reputation for quality. *Econometrica* 81(6), 2381–2462.
- Daley, B. and B. Green (2012). Waiting for news in the market for lemons. *Econometrica* 80(4), 1433–1504.
- Daley, B. and B. Green (2016). An information-based theory of time-varying liquidity. *The Journal of Finance* 71(2), 809–870.
- Daley, B. and B. S. Green (2019). Bargaining and news. *American Economic Review* (forthcoming).
- Davis, J., A. Morse, and X. Wang (2020). The leveraging of silicon valley. *Working Paper*.
- DeMarzo, P. M. and M. J. Fishman (2007a). Agency and optimal investment dynamics. *The Review of Financial Studies* 20(1), 151–188.
- DeMarzo, P. M. and M. J. Fishman (2007b). Optimal long-term financial contracting. *The Review of Financial Studies* 20(6), 2079–2128.
- DeMarzo, P. M., M. J. Fishman, Z. He, and N. Wang (2012). Dynamic agency and the q theory of investment. *The Journal of Finance* 67(6), 2295–2340.
- DeMarzo, P. M. and Y. Sannikov (2006). Optimal security design and dynamic capital structure in a continuous-time agency model. *The Journal of Finance* 61(6), 2681–2724.
- DeMarzo, P. M. and Y. Sannikov (2016). Learning, termination, and payout policy in dynamic incentive contracts. *The Review of Economic Studies* 84(1), 182–236.

- Diamond, D. W. and R. E. Verrecchia (1991). Disclosure, liquidity, and the cost of capital. *The Journal of Finance* 46(4), 1325–1359.
- Duffie, D. and D. Lando (2001). Term structures of credit spreads with incomplete accounting information. *Econometrica* 69(3), 633–664.
- Dumas, B. (1991). Super contact and related optimality conditions. *Journal of Economic Dynamics and Control* 15(4), 675–685.
- Edmans, A. and X. Gabaix (2016). Executive compensation: A modern primer. *Journal of Economic literature* 54(4), 1232–87.
- Fishman, M. J. and K. M. Hagerty (1989). Disclosure decisions by firms and the competition for price efficiency. *The Journal of Finance* 44(3), 633–646.
- Fishman, M. J. and K. M. Hagerty (1990). The optimal amount of discretion to allow in disclosure. *The Quarterly Journal of Economics* 105(2), 427–444.
- Gompers, P. and J. Lerner (2001). The venture capital revolution. *Journal of economic perspectives* 15(2), 145–168.
- Gompers, P. A., W. Gornall, S. N. Kaplan, and I. A. Strebulaev (2020). How do venture capitalists make decisions? *Journal of Financial Economics* 135(1), 169–190.
- González-Uribe, J. and W. Mann (2020). New evidence on venture loans.
- Gornall, W. and I. A. Strebulaev (2020). Squaring venture capital valuations with reality. *Journal of Financial Economics* 135(1), 120–143.
- Green, B. and C. R. Taylor (2016). Breakthroughs, deadlines, and self-reported progress: contracting for multistage projects. *American Economic Review* 106(12), 3660–99.
- Grenadier, S. R. and A. Malenko (2011). Real options signaling games with applications to corporate finance. *The Review of Financial Studies* 24(12), 3993–4036.
- Gryglewicz, S., S. Mayer, and E. Morellec (2019). Agency conflicts and short-versus long-termism in corporate policies. *Journal of Financial Economics* (forthcoming).
- Guttman, I., I. Kremer, and A. Skrzypacz (2014). Not only what but also when: A theory of dynamic voluntary disclosure. *American Economic Review* 104(8), 2400–2420.
- Halac, a. and I. Kremer (2019). Experimenting with career concerns. *American Economic Journal Microeconomics* (forthcoming).
- Halac, M. and A. Prat (2016, October). Managerial attention and worker performance. *American Economic Review* 106(10), 3104–32.
- Hall, B. H. and J. Lerner (2010). The financing of r&d and innovation. In *Handbook of the Economics of Innovation*, Volume 1, pp. 609–639. Elsevier.
- Hartman-Glaser, B., T. Piskorski, and A. Tchisti (2012). Optimal securitization with moral hazard. *Journal of Financial Economics* 104(1), 186–202.
- He, Z. (2009). Optimal executive compensation when firm size follows geometric brownian motion. *The Review of Financial Studies* 22(2), 859–892.

- He, Z. (2011). A model of dynamic compensation and capital structure. *Journal of Financial Economics* 100(2), 351–366.
- He, Z. (2012). Dynamic compensation contracts with private savings. *The Review of Financial Studies* 25(5), 1494–1549.
- Hochberg, Y. V., C. J. Serrano, and R. H. Ziedonis (2018). Patent collateral, investor commitment, and the market for venture lending. *Journal of Financial Economics* 130(1), 74–94.
- Hoffmann, F., R. Inderst, and M. Opp (2019). Only time will tell: A theory of deferred compensation and its regulation. *Working Paper*.
- Hoffmann, F. and S. Pfeil (2010). Reward for luck in a dynamic agency model. *The Review of Financial Studies* 23(9), 3329–3345.
- Hoffmann, F. and S. Pfeil (2019). Dynamic multitasking and managerial investment incentives. *Working paper*.
- Holmstrom, B. (1989). Agency costs and innovation. *Journal of Economic Behavior & Organization* 12(3), 305–327.
- Innes, R. D. (1990). Limited liability and incentive contracting with ex-ante action choices. *Journal of Economic Theory* 52(1), 45–67.
- Jovanovic, B. (1982). Truthful disclosure of information. *The Bell Journal of Economics*.
- Kaplan, S. N. and P. Strömberg (2003). Financial contracting theory meets the real world: An empirical analysis of venture capital contracts. *The review of economic studies* 70(2), 281–315.
- Kaplan, S. N. and P. E. Strömberg (2004). Characteristics, contracts, and actions: Evidence from venture capitalist analyses. *The Journal of Finance* 59(5), 2177–2210.
- Kerr, W. R. and R. Nanda (2015). Financing innovation. *Annual Review of Financial Economics* 7, 445–462.
- Kerr, W. R., R. Nanda, and M. Rhodes-Kropf (2014). Entrepreneurship as experimentation. *Journal of Economic Perspectives* 28(3), 25–48.
- Lerner, J. (1995). Venture capitalists and the oversight of private firms. *Journal of Finance* 50(1), 301–318.
- Lerner, J. and R. Nanda (2020). Venture capital’s role in financing innovation: What we know and how much we still need to learn. *Journal of Economic Perspectives* 34(3), 237–61.
- Malenko, A. (2019, 08). Optimal dynamic capital budgeting. *Review of Economic Studies* 86(4), 1747–1778.
- Manso, G. (2011). Motivating innovation. *The Journal of Finance* 66(5), 1823–1860.
- Marinovic, I. and M. Szydlowski (2019). Monitor reputation and transparency. *Working paper*.
- Marinovic, I. and F. Varas (2016). No news is good news: Voluntary disclosure in the face of litigation. *The RAND Journal of Economics* 47(4), 822–856.

- Marinovic, I. and F. Varas (2019). Ceo horizon, optimal pay duration, and the escalation of short-termism. *The Journal of Finance* 74(4), 2011–2053.
- Mason, R. and J. Välimäki (2015). Getting it done: dynamic incentives to complete a project. *Journal of the European Economic Association* 13(1), 62–97.
- Morellec, E. and N. Schürhoff (2011). Corporate investment and financing under asymmetric information. *Journal of financial Economics* 99(2), 262–288.
- Nanda, R. and M. Rhodes-Kropf (2013). Investment cycles and startup innovation. *Journal of Financial Economics* 110(2), 403–418.
- Nanda, R. and M. Rhodes-Kropf (2017). Financing risk and innovation. *Management Science* 63(4), 901–918.
- Nanda, R., K. Younge, and L. Fleming (2014). Innovation and entrepreneurship in renewable energy. In *The changing frontier: Rethinking science and innovation policy*, pp. 199–232. University of Chicago Press.
- Orlov, D. (2018). Optimal design of internal disclosure. *Working Paper*.
- Piskorski, T. and M. M. Westerfield (2016). Optimal dynamic contracts with moral hazard and costly monitoring. *Journal of Economic Theory* 166, 242–281.
- Robb, A. M. and D. T. Robinson (2014). The capital structure decisions of new firms. *The Review of Financial Studies* 27(1), 153–179.
- Sannikov, Y. (2008). A continuous-time version of the principal-agent problem. *The Review of Economic Studies* 75(3), 957–984.
- Szydlowski, M. (2019). Optimal financing and disclosure. *Management Science (forthcoming)*.
- Townsend, R. M. (1979). Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory* 21(2), 265–293.
- Varas, F. (2017). Managerial short-termism, turnover policy, and the dynamics of incentives. *The Review of Financial Studies*.
- Varas, F., I. Marinovic, and A. Skrzypacz (2020). Random inspections and periodic reviews: Optimal dynamic monitoring. *Review of Economic Studies (forthcoming)*.
- Zhu, J. Y. (2012). Optimal contracts with shirking. *Review of Economic Studies* 80(2), 812–839.