



Research Papers



ISSN 2610-931X

CEIS Tor Vergata

RESEARCH PAPER SERIES

Vol. 24, Issue 1, No. 621 – June 2026

**Flexibility Versus Security
in Agency Contracts
with Moral Hazard**

Lorenzo Bozzoli and Guillaume Pommey

Flexibility Versus Security in Agency Contracts with Moral Hazard

Lorenzo Bozzoli¹

Toulouse School of Economics.

Guillaume Pommey²

Tor Vergata University of Rome.

June 9, 2026

We study agency contracts where a project owner (principal) privately learns her opportunity cost of continuing the relationship after the agent's effort is sunk. The principal controls the design of termination rights and faces a trade-off between preserving incentives and retaining exit flexibility. Even without legal constraints, fully flexible (at-will), fully rigid (lock-in), and intermediate contracts offering partial security and compensation can all arise at equilibrium. While some contracts may appear to protect the agent, they are always socially inefficient, justifying targeted legal constraints on termination rights. In particular, at-will contracts always under-secure the agent and under-enforce the project and can be banned on efficiency grounds. Inefficiencies also include excessive termination payments and over-securing, suggesting both too much and too little flexibility distort outcomes. Finally, we characterize a minimal mandatory termination fee, as commonly seen in employment protection laws, that partially restores efficiency. (JEL D86, J33, J65)

Keywords: *Moral Hazard, Termination rights, Severance Pay, Commitment*

¹Toulouse School of Economics, Université Toulouse 1 Capitole. Email: lorenzo.bozzoli@tse-fr.eu

²Department of Economics and Finance, Tor Vergata University of Rome. Email: guillaume.pommey@uniroma2.eu.

We would like to thank Andrea Attar, Eloisa Campioni, Matteo Escudé, Elisabetta Iossa, Alberto Iozzi, Niccolò Lomys, Thomas Mariotti, Giancarlo Spagnolo, Emanuele Tarantino, and Juha Tolvanen for helpful comments.

1. Introduction

When parties enter long-term contractual relationships, new information may emerge over time that affects their willingness to continue or terminate the agreement early. In employment relationships, for instance, a drop in business activity may lead the employer to dismiss the worker early; in partnerships, one party may find a more attractive opportunity elsewhere; in public-private collaborations, changing priorities or the emergence of a more capable private partner may lead the public entity to seek early termination.^{3,4}

Allowing for unilateral termination of long-term contracts can be valuable in such cases, as it allows parties to act on evolving information, cancel non-viable projects, and improve allocative efficiency. However, this flexibility also weakens the parties' relationship-specific investments by exposing them to the risk of making unrewarded effort. Thus, a fundamental trade-off exists between flexibility and security in long-term contracting.

Contractual forms commonly observed in various applied contexts, including labor and corporate finance, are often perceived to manage this tension inefficiently, tending to be either overly rigid or overly flexible.⁵ These considerations naturally raise two key questions: Should long-term contracts allow for early unilateral termination? And will unregulated, self-interested parties implement socially optimal termination policies?

In this paper, we develop a simple and tractable model addressing these questions within a familiar and well-established contractual theoretical framework.

Our approach builds upon the standard principal-agent problem with moral hazard, limited liability, bilateral risk neutrality, and the principal's full commitment power to design output-contingent contracts (Laffont and Martimort, 2002, Chapter 4). We assume that the principal can hire an agent to work on a *project* whose success depends on the unobservable level of effort exerted by the agent at an early stage. Our key assumption is that, during the course of the project, the principal discovers an *outside option* that she can decide to exert rather than completing the project with the agent. Importantly, we assume that the outside option value is revealed to the principal privately, after the agent has exerted effort but before the project outcome is realized. That is, the principal's new information is either her subjective assessment or cannot be verified.⁶

³For example, in 2016 Disney gave Netflix exclusive rights to stream all Disney's new releases and promote their content on their platform. Three years later, foreseeing a better future on its own, Disney was ready to launch their in-house streaming services, Disney+, and unilaterally decided to remove its content from Netflix by the end of 2019 (Source: <https://www.imdb.com/news/ni61387953/>).

⁴In 2000, the Arrow Light Rail consortium was awarded a public-private partnership contract to design, build, and operate Nottingham's first tram line. While the project was successfully completed and running, the expansion of the network under Phase Two proved too large to be integrated under the original contract. In 2011, the public authority chose a different consortium—Tramlink Nottingham—to take over the existing line and lead the expansion, effectively terminating Arrow's contract before its planned end. Despite having delivered the first phase successfully, Arrow was bought out and compensated to allow the transition to the new operator. See *Termination and Force Majeure Provisions in PPP Contracts* (2013).

⁵For instance, labor laws often aim to promote stability in employment relationships, recognizing that parties may otherwise underprovide such security. In contrast, overly protective clauses in corporate finance—such as golden parachutes—are frequently criticized. Brusa, Lee and Shook (2009), for example, find that firms adopting such measures tend to underperform.

⁶For clarity of exposition, we assume throughout the paper that the outside option corresponds to an alternative activity the principal could undertake instead of continuing the project with the agent. However, without loss of generality, the model can be relabeled to interpret the outside option as an *exogenous shock*.

These features – the timing and the private nature of the principal’s information – give rise to an intertemporal adverse selection problem for the principal: When designing the contract with the agent at the *ex ante* stage, the principal must take into account the incentives of her future (privately informed) self to terminate the project. We also assume that the principal can commit to a liquidation fee paid to the agent in case she decides to cancel the project before completion. Although conceptually simple, this departure from the standard principal-agent moral hazard problem generates a rich set of results and new insights on the design of the optimal contract.

Our findings can be organized in three main directions. First, we show that three types of contracts can emerge at equilibrium: *lock-in* contracts that bind the principal to complete the project; *partially-secured* contracts that allow the principal to cancel the project upon paying a liquidation fee to the agent; *at-will* contracts that enable the principal to withdraw without consequences. Remarkably then, contracts allowing unilateral termination without compensation, resembling those studied in limited commitment and holdup models, can emerge endogenously from complete contracting under full commitment.

Second, we provide a full characterization of the principal’s optimal choice among these three contractual forms. We show that this choice is ultimately determined by three key parameters: the expected value of the project, the agency costs associated with inducing effort, and the asymptotic hazard rate of the outside option’s distribution. When the agency costs and the asymptotic hazard rate are both large, the principal prefers a fully secured (lock-in) contract. Indeed, high agency costs mean that canceling trade imposes a significant cost on incentive provision, while the high hazard rate implies that extreme realizations of the principal’s outside option are unlikely enough that sacrificing them is not particularly costly.

Instead, for a lower level of either the agency costs or the asymptotic hazard rate, the choice between partially-secured and at-will contracts is entirely pinned down by the gross expected value of the project. At first glance, one might expect the principal to prefer full flexibility (at-will) for low-value projects—whose outcomes matter less to her—and to opt for partial flexibility (partially-secured) for higher-value projects. Instead, we find the opposite: a low project value leads the principal to impose higher liquidation fees and cancel less often. This is because the liquidation fee plays the only role of constraining the principal’s future decision to terminate the project. When the project’s gross value is high, its intrinsic return already makes continuation attractive, reducing the need for additional contractual commitment through the liquidation fee. By contrast, for lower-value projects, stronger *ex ante* security is needed to credibly commit the principal to continue, thereby supporting the agent’s effort.

Third, we compare the principal’s termination policy to the socially optimal one. We show that the latter is such that the project should be continued if and only if its expected value (gross of the *sunk* agency costs) is greater than the realized value of the outside option. While the principal is technically able to implement the socially optimal policy, it is not in her interest to do so. Our analysis reveals that the principal’s preferred contractual form balances a trade-off between minimizing allocative distortions and limiting rent-sharing with the agent. Although such trade-offs are common in incentive theory, the unusual aspect of

affecting the value of the project such as, for example, a sudden industry-wide decline in demand or an unexpected rise in implementation costs.

our environment is that the agent extracts further information rents precisely because the principal holds private information.

Interestingly, both over and under enforcement of the project can occur at equilibrium. Thus, our findings substantiate the idea that both hyper-flexible and hyper-secure contracts observed in practice may be linked to the presence of distortions. Most notably, we show that the principal optimally chooses to over enforce projects with low expected value and to under enforce those with high expected value — holding agency costs and the distribution of the outside option fixed. The principal therefore has a *perverse* incentive to over-secure bad projects and under-secure good ones. The crucial driver of this distortion is that, in our framework, the principal cannot separate the provision of incentives to the agent from her decision to terminate the project at the interim stage: the more the agent anticipates the principal will cancel the project, the less willing he is to exert effort in the first place. To make the contract more flexible—i.e., to preserve the option of terminating the project—the principal must offer stronger incentives to induce effort. But this, in turn, increases the agent's expected payoff, which further lowers the principal's perceived value of continuation and raises her temptation to cancel the project later on. This is where the use of a liquidation fee becomes critical: by reducing the value of the outside option, it helps discipline the principal's future self.

Relatedly, we also investigate whether imposing some statutory liquidation fee can mitigate the allocative distortion problem induced by the equilibrium contract. We show that the statutory liquidation fee can restore allocative efficiency in the case of high-value projects but fails to bear any equilibrium effect in the case of low-value projects. A crucial, policy-relevant implication of this finding is that the socially optimal statutory liquidation fee is strictly positive. Thus, our model provides a novel efficiency rationale for the ban of at-will contracts. This observation aligns, for example, with the existence of employment protection laws in various countries that prohibit at-will termination policies in labor contracts. While equity or risk-aversion arguments are often put forward to justify regulatory rigidity in labor contracts, our result provides a new ground for these policies, namely, allocative efficiency.

Finally, it should be noted that the flexibility-security trade-off has long been recognized in contract theory, but its causes and its analysis are usually left implicit and modeled in reduced-form. It appears most notably in the literature on limited commitment and the holdup problem which typically assumes that due to the complexity of the environment or the bounded rationality of the parties, it is too costly to write a contract that incorporates all possible contingencies.⁷ The demand for flexibility is hence usually modeled, rather than an equilibrium outcome, as an exogenous constraint on the set of available contracts, whose existence imposes additional sequential rationality requirements upon physical and incentive constraints, and engenders time-inconsistencies in the optimal contractual design.⁸

⁷See for example (Laffont and Tirole, 1988, p. 1154): “contracts are costly to write and contingencies are often hard to foresee, which gives rise to the allocation of discretion to some members of the organization”. Also, concerning the foundation of the hold-up problem as an issue of incomplete contracting, Tirole (1986, p. 239) writes: “traditional explanation of incomplete contracts relies on the existence of transaction costs. First, some contingencies may not be foreseeable. Second, it may be expensive and time consuming to write the (high number of) foreseeable ones into a contract. Third, some contingencies may be private information (as is the case here). A complete contract must then specify an “incentive-compatible” mechanism of information transmission, which makes it particularly expensive”.

⁸See, for example, (Brzustowski, Georgiadis-Harris and Szentes, 2023, p. 1338): “In other words, contracting

This approach, however, does not align with the observed practice of contracts explicitly allowing and governing the parties' rights to exit, treating the event of cancellation as an ex ante contractible contingency. In this study, instead, we show that, without relying on prior assumptions on limited commitment nor contract incompleteness, the usually exhibited time-inconsistency issues also emerge under full commitment because of the configuration of the principal's information, and that contractual forms entailing full discretion of the principal's interim self actually align with the principal's ex ante rationality and design.

Structure. We briefly review the related literature in Section 2. We expose the main framework and two relevant benchmark cases in Section 3. In Section 4 we fully characterize the equilibrium contract. Section 5 introduces our typology of contracts — lock-in, partially secured, and at-will — and analyzes their characteristics together with the conditions under which the principal prefers each type. We investigate the socially optimal regulated liquidation fee, and the agent's utility-maximizing liquidation fee in Section 6. Finally, we provide a discussion in Section 7. All proofs are relegated to the Appendix.

2. Related literature.

The seminal work by [Shavell \(1980\)](#) provides the standard framework for the study of optimal damage compensations for contractual breach, comparing the impact of different policies on ex ante investments. Our study confirms his finding that optimal contracts exist inducing optimal termination while maintaining ex ante efficient investments; however, we also show that such contracts are not featured in the principal's monopolistic solution. [Rogerson \(1984\)](#), that extend the analysis to the case of ex post renegotiation, and [Tirole \(1986\)](#), in a generalized procurement environment with renegotiation, argue that renegotiation may hinder termination fees and weaken contractual commitment. In this study, we illustrate that termination fees may also fail to emerge in equilibrium with full commitment, because of information frictions in bilateral contracting.

Our study can also be interpreted as *endogenizing* commitment in the form of *self-delegation* by the principal, and is thus related to the delegation literature in contract theory. Namely, we illustrate the ways in which the principal optimally delegates the choice of terminating a contract to her future, more informed self, whose sequential rationality is however misaligned to the ex ante efficient benchmark. The closest contribution to ours in this literature is [Levitt and Snyder \(1997\)](#). In fact, alike our contribution, they tackle the topic of contractual termination within the standard moral hazard setting. However, they assume that is the agent who privately observes whether the project is profitable or not for the principal, after exerting effort. In this context, they show that the principal may optimally provide the agent with the power to rescind the contract as a way to extract this private information, and that, when this is the case, the principal must compensate the agent with a fee exceeding his wage in the bad output realization. Instead, in our work, the principal extracts the information from her future self, which allows for more information transmission and an entirely

parties may refrain from signing long-term, binding contracts due to potential unforeseen or noncontractible contingencies even if such contracts were feasible [...] Although we do not model these contingencies, our assumption that the seller cannot commit not to switch off a deployed contract embodies the idea that she prefers a contract allowing for discretion in the future".

different characterization of optimal contracts, possibly involving at-will termination. In applied contexts, [Laux \(2008\)](#), [Inderst and Mueller \(2010\)](#), [Garrett and Pavan \(2012\)](#) and [Varas \(2018\)](#) assess CEO turnover in companies' boards and [Simester and Zhang \(2010\)](#) study the optimal policy for the termination of unprofitable product lines in firms in similar frameworks to [Levitt and Snyder \(1997\)](#). To conclude on the link with the delegation literature, [Halac and Yared \(2020\)](#) also analyze a trade-off between flexibility and commitment in a delegation environment of incomplete information, in which however the agent's actions is contractible and incentive-provision to exert effort is not considered.

Our work also relates to the relational contracts literature, where the principal makes subjective evaluations ([MacLeod, 2003](#)). However, we depart from this approach in two ways: we maintain verifiability and commitment in the moral hazard relationship, and subjectivity arises only with respect to the principal's outside option. Moreover, the principal's decision in our setting has real consequences—it terminates the project and nullifies the outcomes on which the incentive contract is built. The assumption that the principal privately observes her outside option also appears in the relational contracting literature ([Halac, 2012](#), [Li and Matouschek, 2013](#)). However, as we maintain the commitment assumption in our setting, our focus and characterization differ from theirs.

Finally, concerning the connections between our problem of exit design and the limited commitment literature, the principal's power to arbitrarily cancel long-term contracts has been explicitly interpreted as a form of limited commitment by [Brzustowski, Georgiadis-Harris and Szentes \(2023\)](#) in a recent contribution, within the Coase conjecture environment. However, a similar interpretation was already suggested informally by [Tirole \(1986, p. 240\)](#): *"The breach of contract argument rests on the impossibility of enforcing the initial contract [...] These features put ("individual rationality") constraints on contracts that can be enforced"*.⁹ In this contribution, we argue instead that contracts allowing unilateral cancellation, despite their adverse effects on incentive-provision, may reflect the full commitment solution to a tractable tradeoff between flexibility and incentives, and thus do not necessarily indicate the existence of constraints on contracting or rationality.¹⁰

3. The Model

3.1. Setup

We consider a contractual relationship between a principal and an agent. The principal owns a *project* whose outcome $\pi \in \{\bar{\pi}, \underline{\pi}\}$ depends on the unobservable effort of the agent. The effort decision $e \in \{H, L\}$ affects the distribution of the project's outcome $p_e := \mathbb{P}(\pi = \bar{\pi})$. We assume that $\bar{\pi} > \underline{\pi}$, and $p_H > p_L > 0$. Effort entails a nonmonetary cost ψ_e for the agent, where

⁹A similar point is made by [Laffont and Tirole \(1988, p. 1154\)](#): *"We are thus assuming that the regulator cannot commit himself not to use in the second period the information conveyed by the firm's first-period performance. This assumption, which, we believe, is reasonable in a wide range of applications, certainly deserves comment. The simplest way to motivate it is the changing principal framework. For instance, the current administration cannot bind future one"*.

¹⁰Asymmetric information explanations for incomplete contracting have already been proposed in the past, notably by [Spier \(1992\)](#) who proposes that complete contracts may be seen as negative signals about contingencies on which one party is privately informed.

$\psi_H > 0$ and we normalize $\psi_L = 0$. Both players are risk-neutral and the agent is subject to limited liability.

In addition to the project, the principal has a random *outside option* of value x , which represents an alternative opportunity that precludes the completion of the project if exercised.¹¹ The value of x is drawn from an absolutely continuous cumulative distribution function F with support on $X := [\underline{x}, +\infty)$, and density $f := F'$. Our key assumption is the following: the principal does not know the realization of x at the time the contract is offered to the agent. Instead, x is privately revealed to the principal only after the agent has exerted effort but before the project's outcome π is realized. At this stage, the principal may opt for the outside option by canceling the project. If they do so, the project's outcome does not realize and the agent's payment cannot be contingent on π .

In the following, we call a *contract with exit rights* the collection $\mathcal{C} := \{(\bar{t}, \underline{t}), z\}$, where $(\bar{t}, \underline{t}) \in \mathbb{R}_+^2$ is the *project contract* and $z \in \mathbb{R}_+ \cup +\infty$ is the *liquidation fee*.¹² The project contract specifies the agent's wage (or payment) contingent on the realization of the project outcome $\pi \in \{\underline{\pi}, \bar{\pi}\}$ while the liquidation fee z represents the monetary transfer to the agent if the project is not implemented.^{13,14} We let $Y_e := p_e \bar{\pi} + (1 - p_e) \underline{\pi}$ and $T_e := p_e \bar{t} + (1 - p_e) \underline{t}$ denote the expected gross profit and the expected transfer to the agent from implementing the project when the agent exerts effort e , respectively. The ex ante expected utility of the project is $Y_e - T_e$ for the principal and $T_e - \psi_e$ for the agent.

We assume that the principal can fully commit ex ante to the contract with exit rights, that is, can commit on the outcome-contingent monetary payments $\bar{t}$, $\underline{t}$ and z . The decision to implement the project or not, $\rho \in \{0, 1\}$, is left to the principal's discretion at the interim stage, but is administered by the conditions in the initial contract.

Formally, we study an extensive form game G in which: a principal's pure strategy specifies an offer $\mathcal{C}$ and a decision $\rho \in \{0, 1\}$ for each private history $(\mathcal{C}, x)$ she may observe; an agent's pure strategy maps any principal offer $\mathcal{C}$ into an acceptance decision and an effort decision $e \in \{H, L\}$. We study the PBE of G , and let $G(\mathcal{C})$ be the continuation game induced by a principal's offer. The timing of the game is given in Figure 1.

Importantly, as the principal takes her participation decision non-cooperatively, any rule ρ featured in a continuation equilibrium of $G(\mathcal{C})$ must be sequentially rational for her at the interim stage. It must be stressed that this last fact does not reflect limited commitment or contractual incompleteness on the principal's side: if the principal could resort to any arbitrary

¹¹As mentioned in the introduction, the model can be reinterpreted to describe a scenario in which the project is hit by an exogenous shock of size x at the interim stage—for instance, an unanticipated increase in execution costs or an economic downturn. Under this interpretation, the principal privately learns the realization of the shock before deciding whether to terminate the project.

¹²We use the notation $z = \infty$ to denote the full commitment of the principal to implementing the contract. Thus, the principal's decision to leave herself discretion for future exit is itself a choice of design.

¹³The liquidation fee can be interpreted in different ways depending if we consider the principal-agent pair to be a firm-contractor or an employer-employee relationship. In the first case, the liquidation fee represents the penalty the principal has to pay to the agent for early termination of the project. In the second case it can be interpreted both as a severance package if terminating the project means firing the agent or as the *default* wage of the agent when it is not possible to provide rewards and punishment based on an observable outcome.

¹⁴In practice, it is plausible that some new information about the project's outcome becomes available at the time the outside option is realized. The principal would then compare the updated expected value of the project with that of the outside option. For simplicity, we abstract from this feature, as explicitly modeling it would not alter our qualitative results but would complicate the exposition.

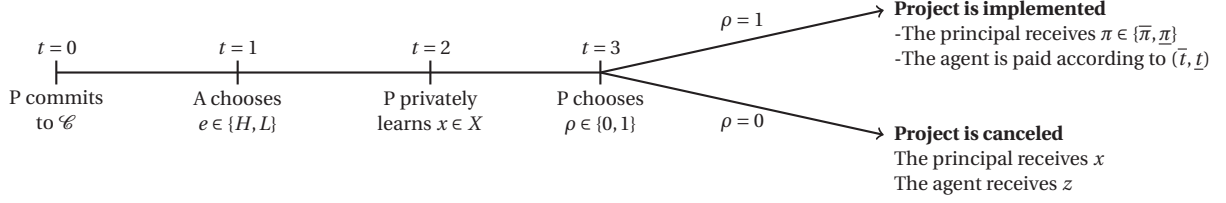


Figure 1: Timing of the game.

mechanism, it is without loss of generality to assume our simpler form of discretionary choice and noncontingent liquidation fee.

This remarkable property stems from the fact that, even if the principal has full commitment power at the ex ante stage, she faces an adverse selection problem with her future self at the interim stage. Indeed, as the value of the outside option is revealed privately to the principal at the interim stage, any mechanism that the principal designs ex ante is subject to her own incentive compatibility constraints to report truthfully at the interim stage. This last constraint renders any conditioning of the contract on the reported value of the outside option irrelevant. The following result formalizes this statement.

Lemma 1. *Assume the principal privately observes the value of the outside option at the interim stage. For any incentive-compatible allocation $\{(\bar{t}^*, \underline{t}^*), z^*, \rho^*\} : X \rightarrow [0, \infty)^3 \times \{0, 1\}$ inducing high effort from the agent, there exists a contract with exit rights $\mathcal{C} = \{(\bar{t}, \underline{t}), z\}$ inducing a payoff-equivalent continuation play. Furthermore, for any such allocation, the principal's choice to implement the project at the interim stage follows a threshold rule, $\rho(x) = \mathbb{1}\{x \leq \tilde{x}\}$, where $\tilde{x} \in X$.*

Lemma 1 guarantees that the principal cannot outperform a simple *contract with exit rights* by offering any direct, deterministic mechanism subject to her own incentive compatibility constraint. In addition to simplifying the class of contracts one has to consider, Lemma 1 also shows that the principal's equilibrium choice to implement the project is a simple threshold strategy $\rho(x) = \mathbb{1}\{x \leq \tilde{x}\}$. In other words, the project is implemented if and only if the outside option realized value is below the equilibrium threshold $\tilde{x} \in X$. Notice that the principal can choose to commit to always implement the project by setting $\rho^*(x) = 1$ for all $x \in X$. We capture this possibility in a contract with exit rights by allowing $z = \infty$, without loss of generality.

Remark 1. *Define $\phi := F(\tilde{x})$ as the ex ante probability that the project will be implemented when the threshold strategy is $\rho(x) = \mathbb{1}\{x \leq \tilde{x}\}$. Then, the principal's problem reduces to choosing:*

- A contract with exit rights $\mathcal{C}$.
- A level of security $\phi \in (0, 1]$.¹⁵
- An effort decision $e \in \{H, L\}$.

¹⁵We can exclude the case $\phi = 0$ as our assumptions guarantee that the principal always benefits from implementing the project with positive probability.

subject to the incentive-compatibility constraint that e and ϕ are featured in a continuation equilibrium of $G(\mathcal{C})$.

We henceforth refer to $\phi \in (0, 1]$ as the principal's choice level of security, that is, the ex ante probability that she implements the project. For convenience we define the quantile function of the outside option as $Q(\phi) = F^{-1}(\phi)$ and the truncated expectation of the outside option above $Q(\phi)$ as $\hat{x}(\phi) := \mathbb{E}[x \mid x \geq Q(\phi)]$.

In light of Remark 1, we find it convenient to cast the principal's problem in terms of explicitly choosing a contract with exit rights $\mathcal{C}$, a level of effort e , and an ex ante probability of implementing the contract $\phi \in (0, 1]$. As we will focus mainly on the case in which the principal wants to induce high effort, we now assume that the contract must induce $e = H$.¹⁶

In this case, the principal's ex ante expected payoff simply writes as

$$V := \phi(Y_H - T_H) + (1 - \phi)(\hat{x}(\phi) - z).$$

The ex ante expected payoff of the agent writes as

$$U := \phi T_H + (1 - \phi)z - \psi_H.$$

To ensure that the agent exerts high effort, the contract must satisfy the agent's incentive compatibility constraint $U \geq \phi T_L + (1 - \phi)z - \psi_L$. Standard computations show that this inequality can be conveniently rewritten as

$$\bar{t} - \underline{t} \geq \frac{C_H}{\phi p_H}, \quad (AIC_H)$$

where $C_H := \frac{p_H \psi_H}{p_H - p_L}$ is what we call the *agency cost* of implementing high effort.¹⁷

When $\phi = 1$, this constraint collapses to the incentive compatibility constraint in a standard moral hazard problem. Instead, when $\phi < 1$, notice that AIC_H requires the agent's payment differential within the project $\bar{t} - \underline{t}$ to weakly decrease in ϕ . That is, the more *uncertain* is the implementation of the project, the more the agent has to be put on a high-powered incentive scheme.

Given a contract $\mathcal{C}$, the ex ante probability that the principal implements the project is given by the probability that the net value of exerting the outside option is lower than the net value of the project, that is, $\phi = \mathbb{P}(x - z \leq Y_H - T_H) = F(Y_H - T_H + z)$. We call this constraint the principal's incentive constraint and we rewrite it as follows:

$$Y_H - T_H = Q(\phi) - z. \quad (PIC)$$

The *PIC* constraint simply ensures that the principal's choice of probability of implementation is consistent with her interim decision to implement or terminate the project. It also makes it clear that to increase the probability of trade ϕ , the principal must either increase her expected profit from implementing the project $Y_H - T_H$ or increase the liquidation fee z .

¹⁶We briefly explore the case $e = L$ in Section 6.3.

¹⁷This cost C_H is the cost incurred by the principal in the standard moral hazard with limited liability. See for instance Laffont and Martimort (2002), Chapter 4, p. 155-157.

As usual, we impose limited liability and participation constraint on the agent's side. That is, we assume that the agent transfers must be such that $\bar{t} \geq 0$, $\underline{t} \geq 0$, $z \geq 0$, and participation requires $U \geq 0$.

3.2. Benchmark results

To fix ideas, we first derive the socially optimal contract when a utilitarian planner can observe the value of the outside option and only the moral hazard problem is present. We make the following assumption to ensure that inducing high effort is socially optimal:

Assumption 1. $\int_{Y_L}^{Y_H} F(x) dx \geq C_H$.

We further discuss this assumption in Section 6.3 when we show how our framework can be reinterpreted as an option to sell for the principal.

The social planner chooses $\phi \in (0, 1]$ and $\mathcal{C}$ to maximize

$$U + V = \phi Y_H + (1 - \phi) \hat{x}(\phi) - \psi_H,$$

subject to AIC_H , $U \geq 0$, $V \geq \mathbb{E}[x]$ and limited liability constraints on $\bar{t}$, $\underline{t}$, and z . The principal's constraint $V \geq \mathbb{E}[x]$ stems from the fact that she could always refuse to participate and consume her outside option.

Proposition 1. *Under assumption (1), the socially optimal implementation policy is such that $Q(\phi^{FB}) = Y_H$ and $e = H$.*

In other words, it is socially optimal to implement the project when the realized value of the outside option is greater than the expected gross value of the project, i.e. when $x \leq Y_H$, and to terminate it otherwise. Notice also that the agency cost C_H is irrelevant in the socially optimal termination policy as it is a sunk cost at the time at which the decision to implement the project is taken.

Importantly, the liquidation fee z plays no role in implementing the socially optimal implementation policy and could only serve as a tool to redistribute rent towards the agent if set to a positive value.

Our second benchmark is the one in which the principal can commit to any deterministic mechanism $\{(\bar{t}^*, \underline{t}^*), z^*, \rho^*\}$ and where we assume that the realized value of the outside option is *verifiable* at the interim stage.

Proposition 2. *Assume that the principal can offer an augmented contract $\{(\bar{t}^*, \underline{t}^*), z^*, \rho^*\}$ and that the realized value of the outside option $x \in X$ is verifiable at the interim stage. Then, under assumption (1), the principal wants to induce $e = H$, and the equilibrium contract is such that $\rho^*(x) = \mathbb{1}\{x \leq Y_H\}$, $z^*(x) = 0$, $\bar{t}^*(x) = \frac{C_H}{\phi^{FB} p_H}$, and $\underline{t}^*(x) = 0$ for all $x \in X$.*

It is immediate that the contract offered by the principal implements the socially optimal implementation policy as $\mathbb{P}(\rho^*(x) = 1) = F(Y_H) = \phi^{FB}$. Moreover, as in the socially optimal case, the principal does not make any use of the liquidation fee. Verifiability of the outside option – and full commitment at the ex ante stage – allows the principal to separate the

problem of inducing the agent to exert effort from that of efficiently terminating the project when it is worth less than the outside option.

Proposition 2 reveals that, when the value of the outside option is verifiable, it is ex ante optimal for the principal to achieve the socially optimal implementation policy. Doing so, the principal first maximizes the joint surplus from trade and then extracts as much rent as possible from the agent. In fact, as shown in the proof, any benefit from distorting the socially optimal implementation policy is outweighed by the additional costs to provide incentives to the agent.

Yet, it is important to stress that under the equilibrium contract characterized in Proposition 2, $\rho^*(x) = \mathbb{1}\{x \leq Y_H\}$ is not sequentially optimal for the principal. At the interim stage, the expected value of the project for the principal is $Y_H - p_H \bar{t}^* = Y_H - \frac{C_H}{\phi^{FB}}$ as payments $\bar{t}^*$ to the agent are conditional on the implementation of the project. The principal would therefore prefer, at the interim stage, to implement the project *less often* than the socially optimal policy, that is, only when $Y_H - \frac{C_H}{\phi^{FB}} \geq x$.

4. Optimal Contract Design: Equilibrium

We now turn to the main analysis in which we characterize the equilibrium contract and termination policy when the value of the outside option is privately revealed to the principal *after* the agent has exerted effort and *before* the realization of the project outcome as presented in Figure 1. In this setting, the principal cannot separate the decision to implement or terminate the project from the incentive provision problem with the agent and new trade-offs appear.

For illustration, assume that the principal sets the project contract and the liquidation fee as in Proposition 2, i.e., $\bar{t} = \frac{C_H}{\phi^{FB} p_H}$, $\underline{t} = 0$, and $z = 0$. These payments align the agent's incentives to exert high effort with the socially optimal probability to implement the project ϕ^{FB} but the principal will find it optimal to implement the project less often than the socially optimal policy, $\phi < \phi^{FB}$, as we have already seen above. Anticipating this, the agent will not exert high effort in the first place.¹⁸

In this context, the principal will have to either increase the agent's rent within the project, distort the probability of implementing the project, make use of the liquidation fee to lower the temptation of her future self to terminate the project, or a combination of these. As we show below, the optimal contract design of the principal will now have to solve a trade-off between flexibility to terminate the project and security of the agent in order to provide effort incentives.

¹⁸It is easy that if the principal implements the project with probability $\phi < \phi^{FB}$, then AIC_H is violated as $\bar{t} - \underline{t} = \frac{C_H}{\phi^{FB} p_H} < \frac{C_H}{\phi p_H}$.

4.1. The principal's problem

As mentioned earlier, we focus on the case in which the principal wants to implement high effort. The principal's problem is as follows:

$$\begin{aligned}
& \max_{(\phi, \bar{t}, \underline{t}, z)} \quad \phi(Y_H - T_H) + (1 - \phi)(\hat{x}(\phi) - z) \\
& \text{s.t.} \quad \phi T_H + (1 - \phi)z - \psi_H \geq 0 \quad (IR_H) \\
& \quad \bar{t} - \underline{t} \geq \frac{C_H}{\phi p_H} \quad (AIC_H) \\
& \quad Y_H - T_H = Q(\phi) - z \quad (PIC) \\
& \quad z \geq 0, \bar{t} \geq 0, \underline{t} \geq 0, \quad ((LL_z), (LL_{\bar{t}}), (LL_{\underline{t}}))
\end{aligned}$$

where recall that $T_H = p_H \bar{t} + (1 - p_H) \underline{t}$.

As in the standard moral hazard framework with limited liability, it is immediate that we can ignore $LL_{\bar{t}}$ and IR_H as they are implied by the other constraints.¹⁹ Moreover, the usual argument that AIC_H must bind applies so that $\bar{t} = \underline{t} + \frac{C_H}{\phi p_H}$ and thus $T_H = \frac{C_H}{\phi} + \underline{t}$.

However, contrary to the standard framework, it is not straightforward whether the limited liability constraint on $\underline{t}$ should bind as this transfer enters the PIC constraint and could potentially be used as a tool to enforce some desired value of ϕ at equilibrium. Indeed, using $T_H = \frac{C_H}{\phi} + \underline{t}$, we can rewrite PIC as follows:

$$Y_H - \frac{C_H}{\phi} - \underline{t} = Q(\phi) - z. \quad (1)$$

Hence, a priori, both the transfer $\underline{t}$ and the liquidation fee z can serve as instruments to adjust each side of the equation, by lowering the net value of the project or that of the outside option, respectively. Although, $\underline{t}$ and z seem to play a symmetric role in the principal's problem, we will show that this is not the case at equilibrium. It is also useful to notice that the left-hand side of equation (1) – the expected value of the project for the principal – is increasing in ϕ as implementing the project more often decreases the total cost of providing incentives to the agent. It follows that this equation may have multiple solutions in ϕ as the right-hand side is also increasing in ϕ .

For convenience and as it proves to be an important object of the analysis we define

$$\zeta(\phi) := Q(\phi) - Y_H + \frac{C_H}{\phi},$$

and use PIC to express transfer $\underline{t}$ as follows:

$$\underline{t} = z - \zeta(\phi). \quad (PIC)$$

Substituting $T_H = \frac{C_H}{\phi} + \underline{t}$ in the principal's objective function and using the expression of $\underline{t}$

¹⁹More precisely, AIC_H and $LL_{\underline{t}}$ together imply $LL_{\bar{t}}$. And IR_H holds whenever AIC_H , LL_z , $LL_{\bar{t}}$, and $LL_{\underline{t}}$ simultaneously hold.

given by *PIC*, we can reformulate the principal's program as follows:

$$\begin{aligned}
 & \max_{\phi, z} \quad \phi Q(\phi) + (1 - \phi) \hat{x}(\phi) - z \\
 (\mathcal{P}) \quad & \text{s.t.} \quad z \geq 0 \\
 & \quad \quad z \geq \zeta(\phi),
 \end{aligned}$$

where the second constraint corresponds to LL_t .

Interestingly, the term $\phi Q(\phi) + (1 - \phi) \hat{x}(\phi)$ can be interpreted as the principal's expected value of the *option to sell* her outside option at *strike price* $Q(\phi)$ at the interim stage.²⁰ Similarly, the liquidation fee $z \geq \max\{0, \zeta(\phi)\}$ represents the price of buying the option to sell. Therefore, the project can be seen as a way for the principal to secure profits in case of a bad realization of her outside option.

Although $\mathcal{P}$ appears to be a significantly easier problem than the original one, it may not be concave or even admit a unique solution, notably due to the behavior of ζ . The following result makes a first step in further simplifying the problem.

Proposition 3. *Let $\phi^* \in (0, 1]$ denote a solution to problem $\mathcal{P}$ and define $\underline{\phi} := \sup\{\phi \in [0, 1] \mid \zeta(\phi) = 0\}$ if it exists, and $\underline{\phi} := 0$ otherwise. Then, any ϕ^* must be such that $\phi^* \in [\underline{\phi}, 1]$. Moreover, $\zeta(\phi) > 0$ for all $\phi \in (\underline{\phi}, 1]$ so that the equilibrium termination fee always satisfies $z^* = \zeta(\phi^*)$.*

An immediate consequence of this result is that at equilibrium LL_t is always binding as $\underline{t}^* = z^* - \zeta(\phi^*) = 0$. It shows that there is no benefit for the principal to leave the agent with any additional rent within the project on top of the standard limited liability rent $\frac{C_H}{\phi}$ even if $\underline{t}$ may serve as an instrument to adjust the probability of implementation in *PIC*. Hence, the liquidation fee z is the only instrument used by the principal to adjust the *PIC* constraint. Secondly, Proposition 3 shows that, contrary to the socially optimal and the verifiable outside option cases, the principal always (weakly) benefits from the possibility of using nonnull liquidation fees.

4.2. Equilibrium contract

To fully characterize the solution to the principal's problem, we assume the following.

Assumption 2 (Single-crossing). *For any C_H , there exists at most one $\tilde{\phi} \in (0, 1]$ such that $I(Q(\tilde{\phi})) = \frac{C_H}{\tilde{\phi}^2}$.*

This assumption simply ensures some regularity of the distribution of the outside option such that program $\mathcal{P}$ is well-behaved. Notice that any distribution with nonincreasing hazard rate immediately satisfies Assumption 2.²¹

Lemma 2. *Under the single-crossing assumption, $(\phi Q(\phi) + (1 - \phi) \hat{x}(\phi) - \zeta(\phi))$ is strictly quasi-concave over $\phi \in [0, 1]$. It attains its unique maximum at $\phi^u \in (0, 1]$, defined as $I(Q(\phi^u)) = \frac{C_H}{(\phi^u)^2}$ if it exists, and $\phi^u = 1$ otherwise.*

²⁰Indeed, with probability ϕ the realized value will be such that $x < Q(\phi)$ and the principal exercises the option at strike price $Q(\phi)$, yielding an expected payoff $\phi Q(\phi)$. With probability $1 - \phi$, we have $x > Q(\phi)$ so that the principal optimally keeps the outside option and enjoys its expected value $\hat{x}(\phi)$, yielding an expected payoff $(1 - \phi) \hat{x}(\phi)$. We further discuss this interpretation in Section 6.3.

²¹Notably, the exponential and the Pareto distributions have constant and decreasing hazard rate, respectively.

Notice that $\phi^u = \arg \max_{\phi \in (0,1]} (\phi Q(\phi) + (1-\phi)\hat{x}(\phi) - \zeta(\phi))$ is simply a solution to problem $\mathcal{P}$ when the limited liability constraint on the liquidation fee, $z \geq 0$, is omitted. Naturally, if $\phi^u \in [\underline{\phi}, 1]$, then $z^u = \zeta(\phi^u) \geq 0$ and the two problems coincide. In the general case, the following proposition fully determines the features of the optimal contract proposed by the principal in the case of unverifiable outside option.

Proposition 4. *Under the single crossing condition, the equilibrium probability of implementation of the project is unique, strictly positive, and satisfies $\phi^* = \max\{\underline{\phi}, \phi^u\}$. We also have,*

- $\phi^* = \phi^u$, and $z^* > 0$ if and only if $\zeta(\phi^u) > 0$;
- $\phi^* = \underline{\phi}$, and $z^* = 0$ if and only if $\zeta(\phi^u) \leq 0$.

Collecting our previous results, the equilibrium contract offered by the principal is unique and satisfies $\phi^* = \max\{\underline{\phi}, \phi^u\}$ and $\mathcal{C}^* := \{(\bar{t}^*, \underline{t}^*), z^*\} = \{(\frac{C_H}{\phi^*}, 0), \zeta(\phi^*)\}$.

5. The Flexibility-Security Trade-Off

In this section we begin by classifying equilibrium contract types, then study how their structure responds to changes in the environment.

5.1. Contract typology

Relying on Propositions 3 and 4 it is useful to classify the equilibrium contract according to the following typology.

Remark 2. *Under the single crossing condition, the equilibrium contract is one of the three following contractual forms.*

- **Lock-in contract:** $\phi^* = 1, z^* = +\infty$;
- **Partially-secured contract:** $\phi^* < 1, z^* \in (0, \infty)$;
- **At-will contract:** $\phi^* < 1, z^* = 0$.

As its name suggests, a lock-in contract is such that the principal commits to *always* implement the project, regardless of the value of the outside option.²² A partially secured contract leaves the principal with the freedom of canceling the project together with a positive liquidation fee so as to provide the agent with some guarantee that cancellation will not happen *too often*. Finally, in an at-will contract, the principal can freely withdraw the original offer without having to compensate the agent.

While the lock-in contracts are exactly equivalent to the standard principal-agent problem with moral hazard, the other two contractual forms exhibit new trade-offs for the principal. Consider the case of partially-secured contracts, that is when $\phi^* = \phi^u < 1$ where ϕ^u solves

$$I(Q(\phi^u)) = \frac{C_H}{(\phi^u)^2}.$$

²²Recall that $z^* = +\infty$ is only a convention that represents full commitment to implement the project. In the general contract space considered in Lemma 1 it corresponds to the case in which $\rho^*(x) = 1$ for all $x \in X$.

On the left-hand side, the inverse hazard rate of the outside option, $I(Q(\phi))$, can be interpreted as the principal's ex ante marginal cost to incentivize her future interim self to implement the project with probability ϕ . That is, the inverse hazard rate represents the information cost for the principal to satisfy her own future incentive constraints. In other words, $I(Q(\phi))$ represents the *virtual marginal cost* associated with the principal with type $x = Q(\phi)$. On the right-hand side, $\frac{C_H}{\phi^2}$ is the marginal benefit of increasing the probability of implementing the project on the limited liability rent of the agent. Indeed, at equilibrium the principal's cost of providing effort incentives is $T_H^* = \frac{C_H}{\phi}$, and $\frac{\partial}{\partial \phi} T_H^* = -\frac{C_H}{\phi^2}$ represents the marginal benefit of increasing the agent's security level. The security level ϕ^u simply equalizes the marginal virtual cost of the information rent on the principal's side with the marginal benefit of the limited liability rent on the agent's side. The positive liquidation fee only serves as an instrument to ensure that the PIC constraint is satisfied at ϕ^u . As we will see later, the security level ϕ^u is generically distorted away from the first-best level ϕ^{FB} as it captures the trade-off between flexibility and security, absent from the fully verifiable outside option case.

Instead, at-will contracts correspond to the case $\phi^* = \phi < 1$ and $z^* = 0$. This happens when satisfying PIC at ϕ^u would require a negative liquidation fee. Hence, the limited liability constraint on the liquidation fee LL_z binds and forces the principal to further distort the security level. From Proposition 4, we have $\phi > \phi^u$ so that the quasiconcavity of the unconstrained objective function implies that $I(Q(\phi)) > \frac{C_H}{\phi^2}$. In other words, the security level ϕ is such that the principal provides *over security* to the agent in order to maintain a reasonable cost of providing effort incentives. Notice that in the unconstrained problem (i.e. without LL_z) the principal would have chosen ϕ^u , inducing a large cost of providing effort incentives and use and $z < 0$ as a tool to extract this rent from the agent.

5.2. Comparative statics

We now investigate how the equilibrium contract varies with the features of the environment. In order to obtain meaningful comparative static results, we focus on changes in the reduced-form parameters Y_H and C_H rather than on the primitive parameters from which they are derived. In what follows, when we say "consider an increase in Y_H " we mean "consider changes in the primitive parameters such that Y_H increases, while C_H remains constant", and similarly for C_H .²³ This approach is motivated by the fact that, as in the standard moral hazard second best analysis, our equilibrium characterization only relies on these reduced-form parameters.

It is useful to restate Proposition 4 in terms of the reduced-form parameters. First define, $\overline{Y}_H$ as the unique value of the expected outcome of the project such that $\zeta(\phi^u) = 0$, that is,

$$\overline{Y}_H = Q(\phi^u) + \frac{C_H}{\phi^u}.$$

For clarity, we write $\overline{Y}_H = \overline{Y}_H(C_H)$ to make the dependence on C_H explicit.

²³Recall that $Y_H = p_H \overline{\pi} + (1 - p_H) \underline{\pi}$ and $C_H = \frac{p_H \psi_H}{p_H - p_L}$. Each reduced-form parameter can always be increased or decreased while keeping the other constant as they do not share the exact same subset of primitive parameters.

Second, let $C^* := \lim_{x \rightarrow +\infty} I(x)$ denote the limit of the inverse hazard rate. Whether C^* is finite or not depends on the thickness of the upper tail of the outside option distribution. Notice that when $C_H < C^* < \infty$, then $I(Q(\phi)) = C_H/\phi^2$ admits a unique solution at ϕ^u , otherwise $\phi^u = 1$ maximizes the unconstrained version of problem $\mathcal{P}$ when $C_H > C^*$.

Lemma 3. *The function $\overline{Y}_H : (0, C^*) \rightarrow (\underline{x}, \infty)$ is continuous, strictly increasing, and satisfies $\lim_{C_H \rightarrow 0} \overline{Y}_H(C_H) = \underline{x}$ and $\lim_{C_H \rightarrow +\infty} \overline{Y}_H(C_H) = +\infty$.*

Proposition 5 (Proposition 4 bis). *Under the single crossing condition, the equilibrium probability of implementation of the project is unique and the principal offers a*

- **Lock-in contract** if and only if $C_H \geq C^*$;
- **Partially-secured contract** if and only if $C_H < C^*$ and $Y_H < \overline{Y}_H(C_H)$;
- **At-will contract** if and only if $C_H < C^*$ and $Y_H \geq \overline{Y}_H(C_H)$.

Lock-in contracts emerge at equilibrium only when the agency cost C_H is larger than the limit of the inverse hazard rate C^* , regardless of the expected value of the project Y_H . This can occur only when C^* is finite, that is, when the distribution of the outside option has a relatively thin upper tail. Therefore, the standard second-best contract occurs when both the moral hazard problem is very costly to the principal and the outside option is relatively unattractive from an ex ante perspective.

When instead the agency cost is relatively low $C_H < C^*$ or when the outside option is very attractive such that C^* is infinite, the principal offers a partially-secured or an at-will contract depending on the expected value of the project Y_H . At first glance, one could intuitively think that the worse is the value of the project Y_H , the less the principal chooses to secure the agent so as to leave herself with the more freedom to enjoy the outside option. Proposition 4 suggests that the opposite is true: projects with relatively low value ($Y_H < \overline{Y}_H$) induce partially-secured contracts while projects with larger value ($Y_H \geq \overline{Y}_H$) induce at-will contracts at equilibrium.

Together with Proposition 5, the following result helps to further interpret partially-secured and at-will contracts.

Proposition 6. *Assume $C_H < C^*$. If $Y_H < \overline{Y}_H(C_H)$ the equilibrium security level ϕ^* is constant in Y_H and the liquidation fee z^* is decreasing in Y_H . If instead $Y_H \geq \overline{Y}_H(C_H)$, the equilibrium security level ϕ^* is increasing in Y_H and the liquidation fee z^* is constant in Y_H .*

Figure 2 illustrates the results of Propositions 5 and 6.

When the project is not very valuable to the principal compared to the outside option ($Y_H < \overline{Y}_H$), the agent expects a high cancellation probability and requires a large rent to be incentivized to exert high effort. From the point of view of the principal, it is beneficial to set a positive liquidation fee to discipline her future self to cancel the project less often and decrease the agent's rent. As Y_H increases towards the threshold, the project becomes more and more valuable which lowers the need for the liquidation fee as a self-commitment device. The partially-secured contract solves the principal's trade-off between granting herself enough flexibility to cancel the project and curb the agent's effort rent.

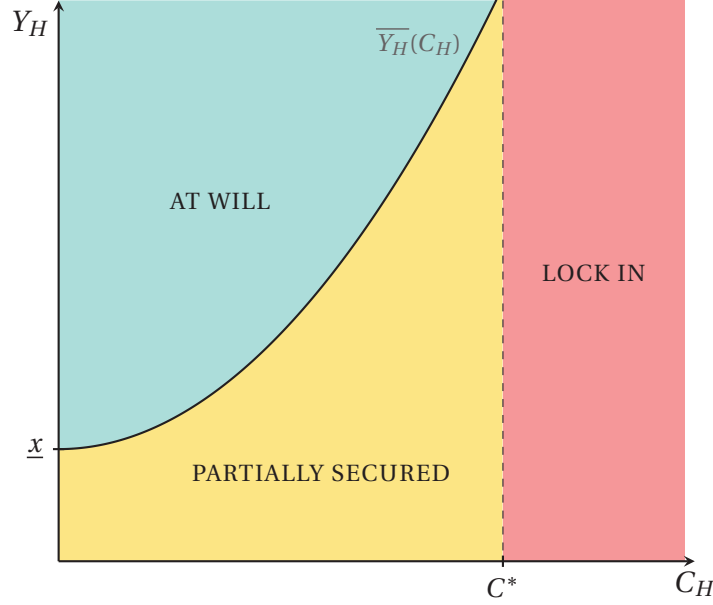


Figure 2: Characterizations of the types of contracts (lock in, partially secured, and at will) in the (C_H, Y_H) space.

Instead, when the project becomes much more valuable compared to the outside option ($Y_H \geq \overline{Y}_H$), the expects a higher probability of implementation of the project. In other words, projects with high expected outcomes are naturally more self-secured in their likelihood to be implemented by the principal. In that case, Proposition 6 states that the higher is Y_H , the higher will be the probability of enforcement $\phi^* = \underline{\phi}$ while the liquidation fee is constant (and null). This behavior is the result of the limited liability constraints on the liquidation fee: In the absence of LL_z , the principal would find this equilibrium probability *too large* and would be better-off implementing $\phi^u < \underline{\phi}$ together with a negative liquidation so as to self-commit to cancel the project more often at the interim stage. The negative liquidation fee would also serve as a way to compensate for the agent's effort rent at a low security level ϕ^u . As LL_z prevents the principal from doing so, she is forced to distort the security level upwards as the project becomes more valuable.

At this stage, it is clear that in general the principal distorts the equilibrium probability of implementation of the project away from the socially optimal one. This raises the question of whether the principal tends to under or over enforce the project with respect to its socially optimal level.

Proposition 7. *There exists a unique threshold $\overline{Y}_H^{FB}(C_H) \in [\underline{x}, \overline{Y}_H(C_H))$ such that*

- *If $Y_H < \overline{Y}_H^{FB}$, then $\phi^* = \phi^u > \phi^{FB}$;*
- *If $\overline{Y}_H^{FB} \leq Y_H < \overline{Y}_H$, then $\phi^* = \phi^u \leq \phi^{FB}$;*
- *If $Y_H \geq \overline{Y}_H$, then $\phi^* = \underline{\phi} < \phi^{FB}$.*

Figure 3 summarizes the findings of Propositions 5, 6 and 7.

At equilibrium, both under an over enforcement of the project with respect to the socially optimal level ϕ^{FB} are possible. As before, fix the level of agency cost and consider a variation

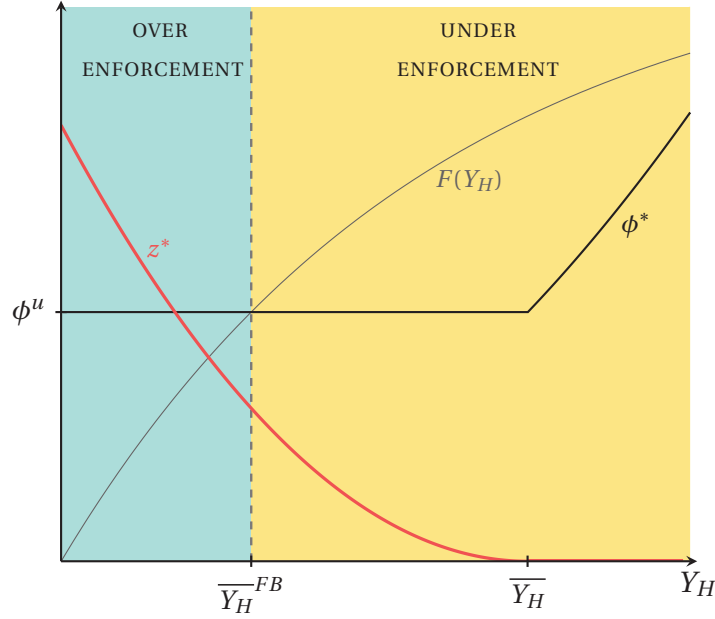


Figure 3: Case $C_H < C^*$. The equilibrium level of security ϕ^* , liquidation fee z^* , and socially optimal level of security ϕ^{FB} are represented by the black, red, and gray curves, respectively. Over and under enforcement regions corresponds to the blue-green and yellow filled areas, respectively.

in Y_H . Perhaps surprisingly, the principal will optimally over enforce projects with very low expected value compared to the agency cost. This arises because Y_H is such that $I(Q(\phi^{FB})) < \frac{C_H}{(\phi^{FB})^2}$, that is, increasing ϕ above the socially optimal level yields larger effort incentive cost savings than the virtual marginal cost of the principal's incentive constraint. Instead for larger project values ($Y_H \geq \bar{Y}_H^{FB}$), the principal will optimally under enforce the project as she perceives the cost to self-discipline herself higher than the one of providing extra rent to the agent. Importantly, even when the principal is forced to distort the probability of enforcement upwards due to LL_z (when $Y_H \geq \bar{Y}_H$), she is still underenforcing the project with respect to the socially optimal level.

6. Welfare and Normative Implications

In this section we analyze the principal's optimal contracting policy from the point of view of both society's and the agent's equilibrium welfare. Furthermore, we assess the viability of some remedies to mitigate the allocative distortions induced by the principal's monopolistic power.

In particular, in subsection 6.1 we assess the principal's solution on the ground of utilitarian social welfare and discuss the effects of a minimum mandatory liquidation fee; in subsection 6.2 we assess the principal's solution in terms of the agent's associated level of welfare and study how a statutory liquidation fee impacts the agent's equilibrium payoff; in section 6.3 we explore the misalignment between society's and the principal's objective in deciding which level of effort to induce from the agent.

6.1. Social welfare

Let us consider the regulatory problem of a social planner that can constrain the set of contracts that the parties are allowed to conclude. We assume in this subsection that the social planner has a utilitarian social welfare function. This implies, as already noted in Proposition 1, that any allocation such that $Q(\phi^{FB}) = Y_H$ maximizes the social planner's objective function.

In principle, the planner can easily induce the first-best level of security by banning any contract with exit rights that induces the principal to implement the project with a different frequency than ϕ^{FB} . This, however, implies heavy restrictions on the parties' freedom to contract.

A less demanding and more realistic solution is to impose a statutory liquidation fee $z_0 \in (0, +\infty)$ on the set of the admissible contracts: this is equivalent to imposing a ban on all contracts such that $z \in [0, z_0]$. The following result examines the role of such fees in restoring efficiency.

Lemma 4. *Suppose the planner imposes a statutory liquidation fee $z_0 \in (0, \infty)$, so that, the principal must now solve the constrained design problem $\mathcal{P}^{z_0}$ obtained from $\mathcal{P}$ by substituting to LL_z the requirement that $z \geq z_0$. Then, if $z_0 \leq z^*$, $\mathcal{P}^{z_0}$ and $\mathcal{P}$ have identical solutions; if $z_0 > z^*$, the principal implements the project with the frequency $\phi^{z_0} = \zeta^{-1}(z_0)$ where $\zeta^{-1}(z_0) : (z^*, +\infty) \rightarrow (0, 1)$ is defined, continuous, strictly increasing, and such that*

$$\lim_{z_0 \rightarrow z^*} \zeta^{-1}(z_0) = \phi^* \quad \text{and} \quad \lim_{z_0 \rightarrow +\infty} \zeta^{-1}(z_0) = 1.$$

Lemma 4 shows that the principal's preferred level of security is irresponsive to z_0 if it is lower than $z^* = \zeta(\phi^*)$, and otherwise increases in z_0 within the continuum identified in Proposition 3. This has an interesting implication.

Proposition 8. *The planner cannot do better than imposing the statutory liquidation fee $z_0^{FB} = \zeta(\phi^{FB})$. This policy achieves the first-best social surplus when there is under-enforcement in the baseline case ($\phi^* < \phi^{FB}$), but is completely ineffective if there is over-enforcement ($\phi^* > \phi^{FB}$). Furthermore, the optimal statutory level z_0^{FB} of the liquidation fee is strictly positive and bounded away from zero.*

The last statement in Proposition 8 stems from the fact that, as noted in Proposition 7, a contract such that $z^* = 0$ can only emerge in equilibrium if the project is implemented with frequency $\phi < \phi^{FB}$, implying under-enforcement with respect to the first-best level. Remarkably, Proposition 8 provides a rationale for banning at-will termination in contracts, a provision that, for example, is observed in employment protection legislation across various judiciary systems.²⁴ While equity concerns and the agent's risk aversion are commonly cited as justifications for such policies, our framework offers an alternative rationale that is grounded in economic efficiency, without relying on assumptions about the parties' attitude to risk.

²⁴Several countries prohibit at-will termination while allowing dismissal without just cause, provided compensation is paid. For instance, Italy and Spain require employers to compensate unfairly dismissed workers with severance based on tenure and salary, while the UK mandates statutory redundancy payments in similar circumstances. See e.g. <https://cms.law/en/int/expert-guides/cms-expert-guide-to-dismissals/italy>, <https://www.lawhub.es/severance-pay-for-employee-dismissal-rules-in-spain/>, and <https://www.gov.uk/redundancy-your-rights/redundancy-pay>.

6.2. Agent's utility

As already noted in Section 4.1, the inspection of the incentive compatibility constraint and limited liability constraints of program $\mathcal{P}$ reveals that the agent's utility from any feasible contract cannot be less than $C_H - \psi_H$, which equates, the limited liability rent featured in the textbook case of moral hazard (see for example [Laffont and Martimort, 2002](#)). However, the rent featured in the optimal contract with exit rights is likely to exceed that value. The following remark clarifies the point.²⁵

Remark 3. *The agent's utility in the principal's optimal contract with exit rights inducing the effort decision $e = H$ is*

$$U^* := (C_H - \psi_H) + (1 - \phi^*)z^*.$$

The agent obtains an *excess* rent in addition to the standard limited liability one. Curiously, the agent's rent has the character of an *information rent* even if the incomplete information component within the contractual relationship lies on the principal's side.

To briefly comment on why is that so, note that the principal's design problem can be expressed as follows:

$$\max_{\phi \in (0,1]} \phi(Y_H - T_H) + (1 - \phi)\hat{x}(\phi) - (1 - \phi)z \quad \text{s.t. } (IR_H, AIC_H, PC, LL_z, LL_{\bar{t}}, LL_{\underline{t}}).$$

Hence the principal maximizes the difference between her expected (gross) profits, within or outside the project, and the expected transfer $(1 - \phi)z$, which is needed to make the project contract consistent with her own truth-telling requirements.

Now, while in principle the transfer could be paid to a third party, the *strict* budget-balance requirement that we implicitly assume throughout the analysis entails that the mechanism cannot run a surplus, and all its transfers must accrue to the agent. Then, remarkably, the principal's program, in the usual fashion to standard mechanism design problems, can be casted as one of maximizing the difference between the total surplus from trade and the agent's expected *excess* rent $(1 - \phi)z$, although the agent's excess rent is obtained indirectly because of the principal's own incentive-compatibility constraints.

Let us briefly discuss the policy relevance of this observation. In the previous subsection, we have analyzed the problem of the social planner who wants to impose a statutory liquidation fee to restore first-best efficiency in terms of utilitarian social surplus. Let us now take a look to the planner's optimal policy under the objective of maximizing the agent's welfare.

Let (ϕ^*, z^*) denote the security level and the penalty of the principal's favorite contract, in the baseline case without public intervention. As shown in the previous subsection, an appropriate statutory liquidation fee $z_0 = \zeta(\phi)$ can be used to induce any level of security $\phi \in [\phi^*, 1]$, while such a policy cannot be exploited to induce any $\phi \in (0, \phi^*)$.

This has the implication that the level of the statutory liquidation fee that maximizes the

²⁵Indeed, the agent obtains with probability ϕ^* the expected payment $\frac{C_H}{\phi^*}$ from the project contract, and, with probability ϕ^* , the optimal liquidation fee z^* .

agent's equilibrium expected payoff is $\zeta(\phi^A) \in [z^*, \infty)$ which solves the following program:

$$\phi^A := \arg \max_{\phi \in [\phi^*, 1]} (1 - \phi) \zeta(\phi).$$

Notably, both under full enforcement and under at-will termination the agent has zero excess rent. The derivation of a general solution is not easy even under the single-crossing assumption, because the objective function may exhibit non-concavities, and is postponed to future research.²⁶

6.3. Agent's effort

We argue in Section 3.2 that, from the point of view of a utilitarian social planner, inducing $e = H$ is optimal over $e = L$ when

$$\int_{Y_L}^{Y_H} F(x) dx > \psi_H. \quad (2)$$

To gain further insight into this requirement, consider the function $v : X \rightarrow \mathbb{R}$ where

$$v(x) := xF(x) + \mathbb{E}[s \mid s > x](1 - F(x)).$$

This function represents the expected payoff from the option to sell, at a strike price $x \in [x, \infty)$, an asset whose value follows the distribution of the principal's outside option.²⁷

Now, one can note that $v(x) = \mathbb{E}(x)$ and $\frac{\partial v(x)}{\partial x} = F(x)$.²⁸ This, in turn implies that $\int_{Y_L}^{Y_H} F(x) dx = v(Y_H) - v(Y_L)$ and (2) rewrites as

$$v(Y_H) > v(Y_L) + \psi_H.$$

That is, the social planner is better off inducing $e = H$ if the value of the option with strike price Y_H exceeds that of the option with strike price Y_L by at least the direct cost of the agent exerting high effort. The economic intuition behind this result is straightforward: conditionally on the effort level $e \in \{H, L\}$ exerted by the agent, at the interim stage, the planner can always secure the level of output Y_e by implementing the project (which is analogue to *selling* the option at the strike price Y_e), while it is better off canceling the project (which is analogue to retaining the option) when the outside payoff exceeds Y_e . If the value of the equivalent option for high effort is greater than the one for low effort by more than the direct cost of $e = H$, the latter induces a higher level of social welfare than $e = L$.

²⁶To see it, note that

$$\frac{\partial(1 - \phi)\zeta(\phi)}{\partial \phi} = I(Q(\phi)) - \frac{C_H}{\phi^2} - Q(\phi) + Y_H,$$

where the term $I(Q(\phi)) - \frac{C_H}{\phi^2}$ is strictly increasing, while $-Q(\phi)$ is strictly decreasing. Numerical evaluations for the case of x being exponentially and Pareto distributed are available at request from the authors.

²⁷Specifically, the holder of such an option would exercise it, securing the strike price x , if, with probability $F(x)$, the asset's value falls below x . Otherwise, she would retain the asset.

²⁸To see it, observe that, by Leibniz' rule, $\frac{\partial \mathbb{E}[s \mid s > x]}{\partial x} = \frac{f(x)}{1 - F(x)} (\mathbb{E}[s \mid s > x] - x)$. The rest follows from elementary algebraic manipulations.

The following remark highlights the conditions under which the principal is better off inducing $e = H$ than $e = L$.

Remark 4. *If the outside option is verifiable, the principal induces $e = H$ if and only if*

$$\int_{Y_L}^{Y_H} F(x) dx > C_H,$$

and, if the outside option is her private information, if and only if

$$\int_{Y_L}^{Q(\phi^*)} F(x) dx > \zeta(\phi^*).$$

Consider first the scenario in which x is contractible. In this case, as already argued in Proposition 2, the principal is able to incentivize $e = H$ while inducing a socially optimal allocation, but she must also leave the standard limited liability rent $C_H - \psi_H$ to the agent. The only difference between the principal's and the planner's effort choice in this case is thus that the cost of high effort is larger for the principal, as she does not internalize the agent's rent. Apart from this, the option interpretation remains, meaning that, gross of the agent's compensation, the principal's maximal profit equals $v(Y_H) = \mathbb{E}[x] + \int_{\underline{x}}^{Y_H} F(x) dx$ when $e = H$ and $v(Y_L) = \mathbb{E}[x] + \int_{\underline{x}}^{Y_L} F(x) dx$ when $e = L$.²⁹ As a result, the principal is less likely than the social planner to induce high effort.

Concerning the private information case, it has already been shown in the proof of Proposition 3 that the principal's design problem can be reduced to

$$\max_{\phi \in [\underline{\phi}, 1]} \phi Q(\phi) + (1 - \phi) \hat{x}(\phi) - \zeta(\phi) = v(Q(\phi^*)) - \zeta(\phi^*).$$

In other words, the principal's decision can be interpreted as selecting from a continuum of options indexed by the probability $\phi \in [\underline{\phi}, 1]$ with which she expects to sell them, where each option entails an upfront cost equal to $\zeta(\phi)$. The principal's optimal implementation probability ϕ^* indexes her preferred option under this interpretation. Then, the value of the principal's optimal option can be expressed as

$$v(Q(\phi^*)) - \zeta(\phi^*) = \mathbb{E}[x] + \int_{\underline{x}}^{Q(\phi^*)} F(x) dx - \zeta(\phi^*).$$

Hence, inducing $e = H$ is optimal if such a value exceeds $v(Y_L)$, that is the principal's optimal profit when inducing $e = L$. This yields the second condition in Remark 4.

To conclude, a final remark can be made.

Remark 5. *The principal is less likely to induce $e = H$ when x is private than observable.*

Remark 5 follows since the principal's optimal payoff when inducing $e = L$ is the same whether x is observable or not, while, when inducing $e = H$, the principal obtains a strictly lower payoff than when the outside option is observable, since the optimal contract exhibited

²⁹Indeed, the principal can straightforwardly appropriate the entire social surplus associated to the low effort decision $e = L$ by deploying the trivial project contract $\bar{t} = \underline{t} = z = 0$.

in Proposition 2 and inducing the first-best level of security without imposing any liquidation fee fails to be incentive-compatible when x is the principal's private information. Therefore, the presence of private information on x may lead the principal to distort the agent's provision of effort with respect to both the social optimum and the public information benchmark.

7. Discussion

This study has reviewed the conditions in which parties fail to address the flexibility security trade-off efficiently, establishing an efficiency-based rationale for regulatory concern on this subject. In this discussion, we first revisit the two core elements driving this inefficiency: asymmetric information about the value of the outside option and the distribution of bargaining power. We then discuss additional model assumptions and outline both descriptive and theoretical implications.

Nature of the outside option. We assume that the outside option only realizes when effort is sunk, and that it is privately observed by the principal. These two elements are both crucial to drive the efficiency concerns of the model.

Regarding timing, the assumption captures the natural tension between long-term commitment to sunk incentives and the need for flexibility in response to evolving information. It also aligns with the recurring features of limited commitment models that address similar dynamics from a different angle (e.g., [Netzer and Scheuer, 2010](#)). As for private observability, it realistically reflects the fact that the information affecting the principal's outside option is often difficult to verify and inherently subjective, typically based on the principal's assessment of market conditions or future prospects. Also, while some aspects may be observable in principle, incorporating them into enforceable contracts as verifiable contingencies is likely impractical.³⁰

Bargaining power and inefficiencies. The other driver of allocative inefficiency in our model is the principal's full bargaining power; which may be interpreted, to some extent, as the principal's monopolistic position in contracting. On this respect, while we abstract from modeling competition among principals explicitly, it is possible to show that the agent's utility-maximizing contract, interpretable as a benchmark for competitive settings, remains efficient even under stringent limited liability constraints, in contrast to the principal's optimal contract.³¹

On a related note, our findings offer a novel, efficiency-based justification for policy interventions that prohibit at-will termination clauses in labor contracts, since, in our model, such clauses are a telltale sign of asymmetric bargaining power, and cannot arise in the agent's or the planner's solution. Thus, our results support these interventions not on equity or risk-sharing grounds, as commonly argued, but on the basis of allocative efficiency as mitigating "monopolistic" distortions.

³⁰The idea that parties possess evolving private, unverifiable or "soft" information about a project's viability is also present in the literature on delegation ([Levitt and Snyder, 1997](#)), subjective evaluation ([MacLeod, 2003](#)), and relational contracts ([Halac, 2012](#)).

³¹Proof available from the authors upon request.

Our analysis also suggests that large termination payments to agents—such as those found in golden parachute clauses—need not reflect strong bargaining power on the agent’s side. Instead, they can emerge from the principal’s own incentive and informational constraints. In fact, we show that even when the agent has no hidden information and no bargaining power, the principal may still end up offering substantial information rents due to the interaction of incentive compatibility and budget balance requirements.

Project outcome observability. Throughout the paper, we assume that if the principal terminates the contract, the project outcome does not materialize and any information about the agent’s action is lost.³² This assumption serves two purposes. First, it allows us to abstract from cancellation as a disciplinary tool and instead isolate the core trade-off between flexibility and security in termination clauses. Second, it reflects many real-world situations—such as employment contracts—where the agent’s compensation upon cancellation is predetermined and independent of performance, having been negotiated ex ante. In conclusion, even allowing for partial observability and contractibility of the project outcome would not substantially alter our results: in case of cancellation, the principal would still be limited in how much of the original incentive structure could be preserved.

Intensity of incentives. Contrary to the standard moral hazard framework, high- and low-powered incentive schemes are not in one-to-one correspondence with the level of agency costs. In our model, the wedge between payments reflects not only agency costs but also the agent’s level of security. The least secured agents are those who receive the strongest incentives, effectively compensating for the fragility of project implementation.

Importantly, because the agent’s security depends on the expected value of the project, the strength of their incentive scheme is also indirectly shaped by project value. This link is absent in the standard model, where the agent’s compensation is typically independent of the value they generate.

The way in which the intensity of incentives responds to such factors is non-trivial and non-monotonic, as it depends on the complex interplay of their direct and indirect effects channeled via those on the security of trade. In addition, such effects are sensitive to the distribution of the principal’s outside option, and in particular to the weight of extremely large valuations, as reflected by its asymptotic hazard rate.

Limited commitment. Our analysis provides a novel explanation for contractual forms frequently observed in applied contracting settings and commonly attributed to limited commitment in the thematic literature, in which one party is allowed to terminate a contract mid-execution without compensating the other.³³ Our explanation, by addressing cancella-

³²Levitt and Snyder (1997) make a similar assumption, where cancellation eliminates information about contractible variables.

³³Brzustowski, Georgiadis-Harris and Szentes (2023) explicitly model limited commitment as the principal’s capacity to replace inter-temporal mechanisms mid-execution; on the subject, Netzer and Scheuer (2010, p. 1087) also write: “One could think of stage 1 being preceded by an additional stage in which the government announces an initial policy. Then the agents choose an effort level in anticipation of some final policy, possibly different from the announcement. If the government is not committed to its initial announcement, it is free to change the policy ex post, after effort choice has taken place (but remains unobservable). The initial

tion rules as part of contract design, aligns with the observation that actual contracts often include detailed cancellation clauses, treating termination as a foreseeable and contractible event, and, in some cases, explicitly permit termination at will.

The notion that parties may intentionally incorporate unilateral termination clauses into contracts, to preserve flexibility, is well established in the limited commitment literature. However, this interpretation is usually presented in a reduced form fashion, primarily as a way to motivate the framework, and is typically related to unmodeled factors such as nonverifiability, bounded rationality or incomplete contracting.³⁴ We argue instead that cancellation policies of this kind may align with full rationality and commitment ex ante, reflecting the principal's desire to *delegate* the cancellation decision to her future, better informed self. We examine this idea in detail, identifying the economic forces that lead the principal to adopt such contractual arrangements. Our analysis shows in particular that incentive and time-inconsistency frictions, typically associated with limited commitment, can also emerge under full commitment, arising naturally from simple asymmetric information frictions.

Appendix

Proof of Lemma 1. We first show that (i) any incentive-compatible allocation must exhibit a threshold rule for ρ ; and (ii) for any fixed incentive-compatible allocation, we construct an equivalent contract with exit option $\mathcal{C} = (\bar{t}, \underline{t}, z)$ that replicates its payoffs in a sequentially rational continuation play while inducing high effort from the agent.

(i) Let $(\bar{t}^*, \underline{t}^*, z^*, \rho^*)$ be an incentive-compatible allocation and define $X^k = \{x \in X : \rho^*(x) = k\}$ for $k = 0, 1$. We first show that $z^*(x)$ must be constant on X^0 . Suppose instead that there exist $x_a, x_b \in X^0$ such that $z^*(x_a) > z^*(x_b)$. Then a principal of type x_a could misreport x_b and obtain the utility $x_a - z^*(x_b) > x_a - z^*(x_a)$, violating incentive compatibility. Similarly, $T_H^*(x)$ must be constant on X^1 : otherwise, if $T_H^*(x_a) > T_H^*(x_b)$ for $x_a, x_b \in X^1$, then the principal of type x_a would strictly prefer to misreport x_b and obtain $Y_H - T_H^*(x_b) > Y_H - T_H^*(x_a)$.

Now fix $z^*(x) = z^*$ on X^0 and $T_H^*(x) = T_H^*$ on X^1 . Under these conditions, the principal's condition for truthful reporting is that:

$$x - z^* \geq Y_H - T_H^* \quad \text{if and only if} \quad \rho(x) = 0.$$

Hence, truthful reporting requires $\rho^*(x) = \mathbb{1}\{x \leq \tilde{x}\}$, where $\tilde{x} := Y_H - T_H^* + z^*$. Therefore, any incentive-compatible allocation must feature a threshold rule in x for ρ^* .

(ii) Now consider any allocation such that $z^*(x) = z^*$ for all $x \in X^0$, $T_H^*(x) = T_H^*$ for all $x \in X^1$, and $\rho^*(x) = \mathbb{1}\{x \leq \tilde{x}\}$ for $\tilde{x} = Y_H - T_H^* + z^*$, as established in (i). Let us construct the contract with exit rights $(\bar{t}, \underline{t}, z) = \left(\frac{T_H^*}{p_H}, 0, z^*\right)$. In this contract, the principal chooses to implement the

announcement is then irrelevant, as captured by the reduced time structure."

³⁴See, for example, (Brzustowski, Georgiadis-Harris and Szentes, 2023, p. 1338): "In other words, contracting parties may refrain from signing long-term, binding contracts due to potential unforeseen or noncontractible contingencies even if such contracts were feasible [...] Although we do not model these contingencies, our assumption that the seller cannot commit not to switch off a deployed contract embodies the idea that she prefers a contract allowing for discretion in the future".

project if and only if $x \leq \tilde{x}$, as in the original allocation, and the agent receives the same net utilities $T_H^* - \psi_H$ if the project is implemented, and $z^* - \psi_H$ if the project is not implemented.

It follows that this contract satisfies the agent's participation constraint if and only if the original allocation does. Furthermore, it induces high effort from the agent whenever:

$$T_H^* \geq \frac{p_H}{p_H - p_L} \psi_H.$$

The latter condition, however, is required for high effort to be incentive-compatible in the original allocation. If this condition fails, then the agent earns less than the limited liability rent, which means that the limited liability constraints and the agent's incentive-compatibility constraint cannot be respected at the same time.

Thus, the constructed contract with exit rights is payoff-equivalent with the original allocation, and it induces high effort and participation from the agent only if the original allocation is incentive-compatible. ■

Proof of Proposition 1. As the objective function of the social planner only depends on ϕ , let us first ignore the constraints and solve $\max_{\phi} \phi Y_H + (1 - \phi) \hat{x}(\phi) - \psi_H$. The first-order condition imply that $Y_H - Q(\phi)Q'(\phi)f(Q'(\phi)) = 0$, where we immediately notice that $Q'(\phi) = 1/f(Q'(\phi))$ from the inverse function rule. Hence, we must have that $Y_H = Q(\phi)$. This condition is also sufficient for optimality as $Q(\cdot)$ is an increasing function.

Let ϕ^{FB} denote this solution and assume it is feasible. Then the planner obtains

$$\begin{aligned} V^{FB} + U^{FB} &= \phi^{FB} Q(\phi^{FB}) + (1 - \phi^{FB}) \hat{x}(\phi^{FB}) - \psi_H \\ &= F(Y_H) Y_H + (1 - F(Y_H)) \mathbb{E}[x \mid x \geq Y_H] - \psi_H. \end{aligned}$$

Instead if the planner were to induce effort $e = L$, she would choose ϕ_L such that $Q(\phi_L) = Y_L$ and obtain $V_L + U_L = \phi_L Q(\phi_L) + (1 - \phi_L) \hat{x}(\phi_L) = F(Y_L) Y_L + (1 - F(Y_L)) \mathbb{E}[x \mid x \geq Y_L]$. It is easy to show (by integrating by part the right-hand side) that :

$$(V^{FB} + U^{FB}) - (V_L + U_L) = \int_{Y_L}^{Y_H} F(x) dx - \psi_H.$$

Assumption 1 immediately implies that the planner prefers to induce high effort.

Let now show that ϕ^{FB} is feasible. Choose, for instance, $\bar{t}^{FB} = \frac{C_H}{\phi^{FB} p_H}$, $\underline{t}^{FB} = 0$, and $z^{FB} = 0$. It follows that AIC_H , $LL_{\bar{t}}$, $LL_{\underline{t}}$, and LL_z are satisfied. It is immediate that $U^{FB} = C_H - \psi_H > 0$. We also have that

$$\begin{aligned} V^{FB} &= \phi^{FB} Q(\phi^{FB}) + (1 - \phi^{FB}) \hat{x}(\phi^{FB}) - C_H \\ &= F(Y_H) Y_H + (1 - F(Y_H)) \mathbb{E}[x \mid x \geq Y_H] - C_H \\ &= \int_{\underline{x}}^{Y_H} F(x) dx + \mathbb{E}[x] - C_H. \end{aligned}$$

It immediately follows from Assumption 1 that $V^{FB} \geq \mathbb{E}[x]$. Hence, ϕ^{FB} , $\bar{t}^{FB}$, $\underline{t}^{FB}$, and z^{FB} satisfy all the constraints.

■

Proof of Proposition 2. The principal chooses $\{\bar{t}^*, \underline{t}^*, z^*, \rho^*\}$ to maximize her profit subject to AIC_H , $U \geq 0$, and the three limited liability constraints.

$$\begin{aligned}
& \max_{(\bar{t}, \underline{t}, z, \rho)} \mathbb{E}[\rho(x)(Y_H - T_H(x)) + (1 - \rho(x))(x - z(x))] \\
& \text{s.t. } \mathbb{E}[\rho(x)T_H(x) + (1 - \rho(x))z(x)] - \psi_H \geq 0 \quad (IR_H) \\
& \mathbb{E}[\rho(x)(\bar{t}(x) - \underline{t}(x))] \geq \frac{C_H}{p_H} \quad (AIC_H) \\
& z(x) \geq 0, \bar{t}(x) \geq 0, \underline{t}(x) \geq 0, \quad ((LL_z), (LL_{\bar{t}}), (LL_{\underline{t}}))
\end{aligned}$$

It is easy to see that AIC_H and the limited liability constraint imply IR_H so that we can ignore it. It then becomes clear that $z(x) = 0$ for all $x \in X$ as the liquidation fee serves no purposes and only penalizes the principal's objective.

Similarly, notice that $\mathbb{E}[\rho(x)T_H(x)] = p_H \mathbb{E}[\rho(x)(\bar{t}(x) - \underline{t}(x))] + \mathbb{E}[\rho(x)\underline{t}(x)]$. Hence, to minimize this quantity while satisfying AIC_H the principal can simply choose $\underline{t}(x) = 0$ for all $x \in X$ and $\bar{t}(x)$ such that $\mathbb{E}[\rho(x)\bar{t}(x)] = C_H/p_H$. For every choice of ρ , it is always possible to find $\bar{t}$ such that this last equation is satisfied (for instance simply let $\bar{t}(x) = \frac{C_H}{p_H \mathbb{E}[\rho(x)]}$ for all $x \in X$).

Hence, the principal's objective simply rewrites as $\mathbb{E}[\rho(x)Y_H + (1 - \rho(x))x] - C_H$ when AIC_H is binding. By linearity of ρ it immediately follows that $\rho(x) = \mathbb{1}\{Y_H \geq x\}$ which corresponds to implementing the first-best security level ϕ^{FB} .

■

Proof of Proposition 3. Define $\Phi^- := \{\phi \in [0, 1] \mid \zeta(\phi) \leq 0\}$ and $\Phi^+ := \{\phi \in [0, 1] \mid \zeta(\phi) \geq 0\}$. Notice that $\Phi^- \cap \Phi^+$ contains the "zeros" of the zeta function.

Consider first the case $\Phi^- \cap \Phi^+ = \emptyset$. By assumption, $\zeta(1) = \bar{x} - (Y_H - C_H) > 0$ so that by continuity of ζ we must have $\Phi^- = \emptyset$ and $\Phi^+ = [0, 1]$. Hence, from our definition $\underline{\phi} = 0$ and it trivially follows that $\phi^* \in [\underline{\phi}, 1]$, $\zeta(\phi) > 0$ for all $\phi \in (\underline{\phi}, 1]$ and $z^* = \zeta(\phi^*)$.

Consider now the case $\Phi^- \cap \Phi^+ \neq \emptyset$. Using once again that $\zeta(1) = \bar{x} - (Y_H - C_H) > 0$ and by continuity of ζ , there exists a $\underline{\phi} = \sup\{\phi \in [0, 1] \mid \zeta(\phi) = 0\}$ such that $[\underline{\phi}, 1] \subseteq \Phi^+$. It also immediately follows that $\sup(\Phi^-) = \underline{\phi}$. For convenience let $g(\phi) := \phi Q(\phi) + (1 - \phi)\hat{x}(\phi)$ and notice that solving problem $\mathcal{P}$ on Φ^- is equivalent to solve $\max_{\phi \in \Phi^-} g(\phi)$ as $\zeta(\phi) \leq 0$ on Φ^- . It is easy to show that g is strictly increasing in ϕ so that the maximum must be attained at $\sup(\Phi^-) = \underline{\phi}$. It also means that any $\phi \in \Phi^+$ such that $\phi < \underline{\phi}$ cannot be optimal as $g(\phi) - \zeta(\phi) < g(\underline{\phi})$ follows from $\zeta(\phi) \geq 0$ on Φ^+ and g increasing. Hence, the only admissible candidates are $\phi^* \in [\underline{\phi}, 1]$. We also have that $\zeta(\phi) > 0$ for all $\phi \in (\underline{\phi}, 1]$ as $(\underline{\phi}, 1] \subseteq \Phi^+$ and $z^* = \zeta(\phi^*)$.

■

Proof of Lemma 2. Let $g(\phi) := \phi Q(\phi) + (1 - \phi)\hat{x}(\phi)$ for convenience. Notice that the inverse hazard rate satisfies $I(Q(\phi)) = (1 - \phi)Q'(\phi)$. Then, we have that

$$(g - \zeta)'(\phi) = \frac{C_H}{\phi^2} - I(Q(\phi)).$$

Notice that $\lim_{\phi \rightarrow 0} (g - \zeta)'(\phi) = +\infty$. Hence, if there exists no crossing point $\tilde{\phi}$ as defined in Assumption 2, $(g - \zeta)'(\phi)$ is strictly increasing on $[0, 1]$ and $g - \zeta$ is therefore quasiconcave. If otherwise such a point $\tilde{\phi}$ exists, then $(g - \zeta)' > 0$ on $\phi \in [0, \tilde{\phi})$ and $(g - \zeta)' < 0$ on $\phi \in [\tilde{\phi}, 1]$ which is a sufficient condition for quasiconcavity of $g - \zeta$.

The solution $\phi^u \in \arg\max_{\phi \in (0, 1]} g(\phi) - \zeta(\phi)$ and its uniqueness straightforwardly stem from the above argument. ■

Proof of Proposition 4. Once again, let $g(\phi) := \phi Q(\phi) + (1 - \phi)\hat{x}(\phi)$. Assume first that $\phi^u \leq \underline{\phi}$. From Lemma 2 it follows that $g - \zeta$ is decreasing on $\phi \in [\phi^u, \underline{\phi}]$ so that $\phi^* = \underline{\phi}$ is the best the principal can do without violating LL_z . If instead, $\phi^u \geq \underline{\phi}$, then $\zeta(\phi^u) \geq 0$ and thus the solution to the unconstrained problem coincides with that of the full problem, i.e., $\phi^* = \phi^u$. Hence $\phi^* = \max\{\underline{\phi}, \phi^u\}$.

Now notice that $\phi^u \geq \underline{\phi} \Leftrightarrow \zeta(\phi^u) \geq 0$. By definition of $\underline{\phi}$, we already showed that $\phi^u \geq \underline{\phi} \Rightarrow \zeta(\phi^u) \geq 0$. Assume instead that $\zeta(\phi^u) \geq 0$ and $\phi^u < \underline{\phi}$. It follows that $g(\phi^u) - \zeta(\phi^u) \leq g(\underline{\phi})$ as g is increasing but this implies $\phi^u = \underline{\phi}$ which contradicts $\phi^u < \underline{\phi}$. Hence, $\phi^u < \underline{\phi} \Leftrightarrow \zeta(\phi^u) < 0$.

The results on z^* immediately follow. Indeed we know from Proposition 3 that $z^* = \zeta(\phi^*) = \zeta(\max\{\underline{\phi}, \phi^u\})$ which is strictly positive when $\phi^* = \phi^u$ as $\zeta(\phi^u) > 0$ and null when $\phi^* = \underline{\phi}$ as $\zeta(\underline{\phi}) = 0$. ■

Proof of Lemma 3. Recall that $\overline{Y}_H(C_H)$ is implicitly defined by $\overline{Y}_H(C_H) = Q(\phi^u) + \frac{C_H}{\phi^u}$. Consider first the case $C_H < C^*$ such that ϕ^u solves $I(Q(\phi^u)) = \frac{C_H}{(\phi^u)^2}$. Continuity of $\overline{Y}_H(C_H)$ is immediate. Differentiating the first equality on both sides gives:

$$\frac{\partial \overline{Y}_H}{\partial C_H} = \frac{\partial \phi^u}{\partial C_H} \left[Q'(\phi^u) - \frac{C_H}{(\phi^u)^2} \right] + \frac{1}{\phi^u}.$$

From Assumption 2, it is clear that ϕ^u is strictly increasing (and differentiable) in C_H so that $\frac{\partial \phi^u}{\partial C_H} > 0$. Notice that for all $\phi \in (0, 1)$, $Q' = \frac{I}{1-\phi} > I$ so that $Q'(\phi^u) - \frac{C_H}{(\phi^u)^2} > I(\phi^u) - \frac{C_H}{(\phi^u)^2} = 0$. It follows that $\overline{Y}_H(C_H)$ is strictly increasing in C_H on $(0, C^*)$. At the limit $C_H \rightarrow 0$, it is clear that $\phi^u \rightarrow 0$. Notice that by definition of ϕ^u , $\lim_{C_H \rightarrow 0} \frac{C_H}{\phi^u} = \lim_{C_H \rightarrow 0} \phi^u I(Q(\phi^u)) = 0$ as $I(Q(\phi^u))$ is finite at 0. It immediately follows that $\lim_{C_H \rightarrow 0} \overline{Y}_H(C_H) = \underline{x}$ as $\lim_{C_H \rightarrow 0} Q(\phi^u) = \underline{x}$.

Now, consider the case $C_H \geq C^*$ so that $\phi^u = 1$. It is immediate that $\overline{Y}_H(C_H)$ is continuous and strictly increasing over $[C^*, +\infty)$ and that $\lim_{C_H \rightarrow +\infty} \overline{Y}_H(C_H) = +\infty$. ■

Proof of Proposition 5. The proof is simply a restatement of Proposition 4. When $C_H \geq C^*$, $\phi^u = 1$ and $\zeta(\phi^u) = +\infty > 0$ so that the principal offers a lock-in contract. Instead when $C_H < C^* < \infty$, $\phi^u < 1$ solves $I(Q(\phi^u)) = C_H/(\phi^u)^2$. The cases $Y_H < \overline{Y}_H(C_H)$ and $Y_H > \overline{Y}_H(C_H)$ correspond to $\zeta(\phi^u) > 0$ and $\zeta(\phi^u) \leq 0$, respectively. That is, to the cases of partially-secured and at-will contract, respectively. ■

Proof of Proposition 6. Assume $C_H < C^*$, then ϕ^u solves $I(Q(\phi^u)) = C_H/(\phi^u)^2$. If $Y_H < \overline{Y}_H(C_H)$, then $\phi^* = \phi^u$ from Proposition 5 which is constant in Y_H by definition of ϕ^u (independent of Y_H). The equilibrium liquidation fee, $z^* = \zeta(\phi^u) = Q(\phi^u) - Y_H + \frac{C_H}{\phi^u}$ is clearly decreasing in Y_H .

If instead $Y_H \geq \overline{Y}_H(C_H)$ then $\phi^* = \underline{\phi}$ from Proposition 5. Recall that by definition $\underline{\phi}$ solves $\zeta(\underline{\phi}) = 0$. Consider the implicit solution $\underline{\phi}(Y_H)$ and differentiate $\zeta(\underline{\phi}) = 0$ with respect to Y_H . After simple rearrangements, we obtain:

$$\frac{\partial \underline{\phi}}{\partial C_H} \left[Q'(\underline{\phi}) - \frac{C_H}{\underline{\phi}^2} \right] = 1.$$

Recall from Lemma 2 that the first-order derivative of the objective of the unconstrained problem is equal to $(g - \zeta)'(\phi) = \frac{C_H}{\phi^2} - (1 - \phi)Q'(\phi) > \frac{C_H}{\phi^2} - Q'(\phi)$ for all $\phi \in (0, 1]$. Furthermore, from Proposition 4 we can deduce that $\phi \geq \phi^u$ as $\phi^* = \max\{\phi, \phi^u\} = \underline{\phi}$. Hence by quasiconcavity of $g - \zeta$ that $(g - \zeta)'(\phi) < 0$ so that $\frac{C_H}{\phi^2} - Q'(\phi) < 0$ as well. It immediately follows that $\underline{\phi}$ is increasing in Y_H . Finally, $z^* = \zeta(\underline{\phi}) = 0$ is trivially constant in Y_H . ■

Proof of Proposition 7. Define $\overline{Y}_H^{FB}(C_H)$ as the solution to $\phi^u = F(Y_H)$, that is the Y_H such that ϕ^u and ϕ^{FB} coincide at a given C_H . Hence, we can rewrite it as $\overline{Y}_H^{FB}(C_H) = Q(\phi^u)$ and $\overline{Y}_H^{FB}(0) = \underline{x}$. Notice also that $\overline{Y}_H^{FB}(C_H) = Q(\phi^u) < Q(\phi^u) + \frac{C_H}{\phi^u} = \overline{Y}_H(C_H)$.

Assume first that $Y_H < \overline{Y}_H^{FB}(C_H)$. From the above argument and Proposition 5 we therefore have that $\phi^* = \phi^u$. Furthermore, $Y_H < \overline{Y}_H^{FB}(C_H)$ imply that $F(Y_H) < F(\overline{Y}_H^{FB}(C_H)) \Leftrightarrow \phi^{FB} < \phi^u$. If instead, $\overline{Y}_H^{FB}(C_H) < Y_H \leq \overline{Y}_H(C_H)$, then $\phi^* = \phi^u$ and the inequality implies $F(\overline{Y}_H^{FB}(C_H)) \leq F(Y_H) \Leftrightarrow \phi^u \leq \phi^{FB}$.

Finally, the case $Y_H \geq \overline{Y}_H(C_H) > \overline{Y}_H^{FB}(C_H)$ yields $\phi^* = \underline{\phi}$. By definition $\underline{\phi}$ is such that $Q(\underline{\phi}) - Y_H + \frac{C_H}{\underline{\phi}} = 0$, that is, $Q(\underline{\phi}) - Q(\phi^{FB}) + \frac{C_H}{\underline{\phi}} = 0$. It immediately follows that $Q(\phi^{FB}) - Q(\underline{\phi}) > 0$ and thus that $\underline{\phi} < \phi^{FB}$. ■

Proof of Lemma 4. If $z_0 < z^*$, then the solution to $\mathcal{P}$ remains feasible in $\mathcal{P}^{z_0}$, so the principal's optimal contract does not change. This holds in particular when $\phi^u = \phi^* = 1$.

Now assume $z_0 > z^*$. Since z_0 is finite, we must have $z^* \in [0, \infty)$, and hence $\phi^u \in (0, 1)$. Consider $\zeta(\phi) = Q(\phi) - Y_H + \frac{C_H}{\phi}$ and note that it is continuously differentiable over $(0, 1]$. Then, from the definition of $I(Q(\phi)) = \frac{1-\phi}{f(Q(\phi))}$, it follows that

$$\zeta'(\phi) = \frac{1}{f(Q(\phi))} - \frac{C_H}{\phi^2} = \frac{I(Q(\phi))}{1-\phi} - \frac{C_H}{\phi^2}.$$

By Assumption 2 (single-crossing), $I(Q(\phi)) > \frac{C_H}{\phi^2}$ for all $\phi \in (\phi^u, 1]$, so $\zeta'(\phi) > 0$ over $[\phi^u, 1]$, implying ζ is strictly increasing on $[\phi^*, 1]$ – since, as established in Proposition 4, $\phi^* = \max\{\phi^u, \underline{\phi}\} \geq \phi^u$.

Now consider problem $\mathcal{P}^{z_0}$, which can be rewritten (following the steps in Proposition 3)

as:³⁵

$$\max_{\phi \in (0,1]} \phi Q(\phi) + (1 - \phi) \hat{x}(\phi) - z \quad \text{s.t.} \quad z \geq z_0 \text{ and } z \geq \zeta(\phi).$$

But then, none of the $\phi \in (0, \zeta^{-1}(z_0))$ can solve $\mathcal{P}$: the term $\phi Q(\phi) + (1 - \phi) \hat{x}(\phi)$ is increasing in ϕ , while the penalty term is minimized at $\phi = \zeta^{-1}(z_0)$, where it equals z_0 . Hence, the principal's problem reduces to solving the unconstrained program:

$$\max_{\phi \in [\zeta^{-1}(z_0), 1]} \phi Q(\phi) + (1 - \phi) \hat{x}(\phi) - \zeta(\phi). \quad (3)$$

Since $z_0 > \zeta(\phi^*)$ and ζ is strictly increasing on $[\phi^*, 1]$, it follows that $\zeta^{-1}(z_0) > \phi^*$. Moreover, Assumption 2 implies that the objective of program 3 is strictly decreasing over $[\phi^*, 1]$, and therefore also over $[\zeta^{-1}(z_0), 1] \subset [\phi^*, 1]$. It follows that the unique maximizer of problem 3 is $\phi^{z_0} = \zeta^{-1}(z_0)$. ■

Proof of Proposition 8. Lemma 4 establishes that the planner, by imposing a statutory liquidation fee $z_0 \in [z^*, \infty]$, can induce any implementation probability $\phi^{z_0} \in [\phi^*, 1]$. Thus a utilitarian planner that must select an optimal value of the liquidation fee that solves

$$\max_{\phi \in [\phi^*, 1]} \phi Y_H + (1 - \phi) \hat{x}(\phi).$$

As shown in Proposition 1, the planner's objective function attains its unique maximum at $\phi^{FB} = F(Y_H)$, is strictly increasing on $[0, \phi^{FB}]$ and strictly decreasing on $[\phi^{FB}, 1]$. Hence, two cases can arise:

- If $\phi^* \in (0, \phi^{FB})$ (under-enforcement), the planner can restore the first-best by imposing $z_0 = z_0^{FB} := \zeta(\phi^{FB})$;
- If $\phi^* \in [\phi^{FB}, 1]$, the objective is strictly decreasing in the feasible interval and the planner cannot do better than the principal's solution by setting $z_0 = z^*$.

To conclude, observe that, since $Q(\phi^{FB}) = Y_H$,

$$z_0^{FB} = \zeta(\phi^{FB}) = Q(\phi^{FB}) - Y_H + \frac{C_H}{\phi^{FB}} = \frac{C_H}{\phi^{FB}} > C_H > 0,$$

where the last inequality holds for every $\phi^{FB} \in (0, 1)$. Hence, the first-best liquidation fee z_0^{FB} is strictly positive and bounded away from zero by C_H . ■

References

Brusa, Jorge, Wayne L Lee and Carole Shook. 2009. "Golden parachutes, managerial incentives and shareholders' wealth." *Managerial Finance* 35(4):346–356.

³⁵In particular, the principal's incentive-compatibility constraint requires that $\underline{t} = z - \zeta(\phi)$. Thus, $LL_{\underline{t}}$ rewrites as $z \geq \zeta(\phi)$.

- Brzustowski, Thomas, Alkis Georgiadis-Harris and Balázs Szentes. 2023. "Smart contracts and the coase conjecture." *American Economic Review* 113(5):1334–1359.
- Garrett, Daniel F and Alessandro Pavan. 2012. "Managerial turnover in a changing world." *Journal of Political Economy* 120(5):879–925.
- Halac, Marina. 2012. "Relational contracts and the value of relationships." *American Economic Review* 102(2):750–779.
- Halac, Marina and Pierre Yared. 2020. "Commitment versus flexibility with costly verification." *Journal of Political Economy* 128(12):4523–4573.
- Inderst, Roman and Holger M Mueller. 2010. "CEO replacement under private information." *The Review of Financial Studies* 23(8):2935–2969.
- Laffont, Jean-Jacques and David Martimort. 2002. *The Theory of Incentives: The Principal-Agent Model*. Princeton University Press.
- Laffont, Jean-Jacques and Jean Tirole. 1988. "The dynamics of incentive contracts." *Econometrica: Journal of the Econometric Society* pp. 1153–1175.
- Laux, Volker. 2008. "Board independence and CEO turnover." *Journal of Accounting Research* 46(1):137–171.
- Levitt, Steven D and Christopher M Snyder. 1997. "Is no news bad news? Information transmission and the role of "early warning" in the principal-agent model." *The RAND Journal of Economics* pp. 641–661.
- Li, Jin and Niko Matouschek. 2013. "Managing conflicts in relational contracts." *American Economic Review* 103(6):2328–2351.
- MacLeod, W Bentley. 2003. "Optimal contracting with subjective evaluation." *American Economic Review* 93(1):216–240.
- Netzer, Nick and Florian Scheuer. 2010. "Competitive markets without commitment." *Journal of political economy* 118(6):1079–1109.
- Rogerson, William P. 1984. "Efficient reliance and damage measures for breach of contract." *The Rand Journal of Economics* pp. 39–53.
- Shavell, Steven. 1980. "Damage measures for breach of contract." *The Bell Journal of Economics* pp. 466–490.
- Simester, Duncan and Juanjuan Zhang. 2010. "Why are bad products so hard to kill?" *Management Science* 56(7):1161–1179.
- Spier, Kathryn E. 1992. "Incomplete contracts and signalling." *The RAND Journal of Economics* pp. 432–443.
- Termination and Force Majeure Provisions in PPP Contracts*. 2013. *Review of current European practice and guidance*.

- Tirole, Jean. 1986. "Procurement and renegotiation." *Journal of Political Economy* 94(2):235–259.
- Varas, Felipe. 2018. "Managerial short-termism, turnover policy, and the dynamics of incentives." *The Review of Financial Studies* 31(9):3409–3451.

RECENT PUBLICATIONS BY *CEIS Tor Vergata*

The Misallocation Costs of Inflation: A Sufficient Statistics Approach

Klaus Adam, Andrey Alexandrov and Henning Weber

CEIS Research Paper, 620 April 2026

Sensitivity of the Euro OIS Term Structure to ECB Policy Rate Surprises

Stefano Herzel and Marco Nicolosi

CEIS Research Paper, 619 January 2026

Pink queue and Double Standard: what Markov Chains uncover in academic career progressions

Marianna Brunetti and Annalisa Fabretti

CEIS Research Paper, 618 December 2025

Spending Multipliers and the Party in Power: Evidence from United States Political Cycles

Francesco Morelli

CEIS Research Paper, 617 November 2025

Banks' Maturity Choices and the Transmission of Interest-Rate Risk

Paolo Varraso

CEIS Research Paper, 616 November 2025

Keeping the Agents in the Dark: Competing Mechanisms, Private Disclosures, and the Revelation Principle

Andrea Attar, Eloisa Campioni, Thomas Mariotti and Alessandro Pavan

CEIS Research Paper, 615 October 2025

Optimal Capital Taxation with Borrowing Constraints and Entrepreneurial Heterogeneity

Lorenzo Carbonari, Fabrizio Mattesini and Alberto Petrucci

CEIS Research Paper, 614 October 2025

Modeling Euro Area Benchmark Rates After the End of LIBOR

Flavio Angelini, Stefano Herzel, and Marco Nicolosi

CEIS Research Paper, 613 October 2025

From Knowledge to Allocation of Sustainable Assets: Results From An In-Field Survey In Italy

Beatrice Bertelli, Marianna Brunetti, Costanza Torricelli and Mariangela Zoli

CEIS Research Paper, 612 October 2025

VAR Models With An Index Structure: A Survey With New Results

Gianluca Cubadda

CEIS Research Paper, 611 September 2025

DISTRIBUTION

Our publications are available online at www.ceistorvergata.it

DISCLAIMER

The opinions expressed in these publications are the authors' alone and therefore do not necessarily reflect the opinions of the supporters, staff, or boards of CEIS Tor Vergata.

COPYRIGHT

Copyright © 2026 by authors. All rights reserved. No part of this publication may be reproduced in any manner whatsoever without written permission except in the case of brief passages quoted in critical articles and reviews.

MEDIA INQUIRIES AND INFORMATION

For media inquiries, please contact Barbara Piazzzi at +39 06 72595584 or by e-mail at barbara.piazzzi@ceis.uniroma2.it. Our web site, www.ceistorvergata.it, contains more information about Center's events, publications, and staff.

DEVELOPMENT AND SUPPORT

For information about contributing to CEIS Tor Vergata, please contact at +39 06 72595587 or by e-mail at sagr.ceis@economia.uniroma2.it