

CEIS Tor Vergata

RESEARCH PAPER SERIES

Vol. 24, Issue 2, No. 622 – June 2026

Optimal Structure of Penalties with Judgment-Proof Injurers

Guillaume Pommey

OPTIMAL STRUCTURE OF PENALTIES WITH JUDGMENT-PROOF INJURERS

Guillaume Pommey ^{*†}
Tor Vergata University of Rome

June 9, 2026

Abstract

I characterize the optimal regulation of a firm constituted by potential judgment-proof agents. I investigate two cases: (i) A principal hires an agent to undertake a prevention effort on their behalf; (ii) Two agents are jointly responsible of undertaking a prevention effort. In both cases, agents are in charge of exerting an unobservable level of safety care to reduce the probability of an accident that may occur due to the firm risky activity. Agents are called judgment proof when their final wealth is not enough to pay for the monetary penalties imposed by the regulator. I show that the standard Equivalence Theorem, stating that the distribution of penalties among injurers is irrelevant, does not hold in this context. Instead, in a principal-agent firm, the optimal regulation requires to fully target the principal if the agent can be subject to judgment proofness. In a two-agent firm, the optimal regulation consists in an almost equal sharing of penalties among agents.

Keywords: *Moral Hazard, Regulation, Limited liability, Judgment Proofness.*

JEL Classification: K13, K32, G33, D86.

^{*}E-mail: guillaume.pommey@uniroma2.eu

[†]I would like to thank David Martimort, Pierre Fleckinger, Patrick Legros, Jérôme Pouyet, Julien Combe, Philippe Colo and Shaden Shabayek for discussions and critical comments on this paper.

1. INTRODUCTION

When a firm's activity may cause an accident to outside parties, how should be designed penalties to induce the firm to undertake enough safety measures, that is, what should be the optimal total amount of penalties and how should it be apportioned among the firm's members?

The earlier works of Newman and Wright (1990) and Segerson and Tietenberg (1992) suggest that only the total amount of penalties matters and not their allocation within the firm. The main argument relies on the existence of private transactions within the firm that can undo any allocation of responsibilities coming from the regulation authority. This result, known as the Equivalence Principle, had a significant influence on many works in the economic literature of tort and environmental law.

This result relies on the strong assumption that the private contract between injurers always satisfies each party's solvency towards the potential payment of fines. However, when injurers may be judgment proof (see Shavell, 1986), that is, financially insolvent when facing the penalties imposed by the regulation authority, it may be optimal not to provide them with additional resources to limit corporate liability. For instance, Ringleb and Wiggins (1990), empirically find that large firms prefer to buy inputs whose manufacture is risky to small firms with few assets. Others establish special subsidiaries or product manufacturing by contract with dedicated small-scale specialized producers. Those choices are deliberately made so as to shield the large firm's assets from potential liability in case of accident. When considering situations of long-term or large-scale hazards such as environmental degradation (hazardous waste, water pollution) or workers exposition to harmful substances (asbestos, radiation, vinyl chloride) the amount of damages are likely to be large and therefore exceed some of the injurers' financial resources if no (or small) corporate compensation exists.

In this paper, I consider a situation in which the Equivalence Principle breaks down. Under the assumption that some injurers may be judgment-proof and that the private contract between them does not guarantee solvency towards the payment of fines, the allocative role of penalties is restored. I consider a firm whose activity is risky and may cause an harm to third parties. To reduce the probability of accident, the firm has to exert costly precautionary care. I then investigate how should penalties be optimally allocated among injurers in the cases of simple and double moral hazard.

First, I consider a principal-agent firm in which the principal hires and agent to undertake an effort of prevention. The agent's effort is unobservable and the problem is a moral hazard one. Therefore, parties can contract only on outcome realizations (accident or no accident). At the contracting stage, I assume that the agent's financial resources are a random variable so that parties are unsure about the agent's ability to make monetary transfers when the outcome is realized. In the absence of regulation, the principal has no incentive to induce the agent to exert effort as the harm does not affect her. Thus, I introduce a regulation authority who can impose *ex post* penalties on the principal and on the agent when an accident occurs.

Obviously, as the agent has limited resources, he may not be able to pay for the penalties imposed on him.

In the literature, it is usually assumed that transfers from the agent to both the regulation authority and the principal are designed *ex ante* such that they never exceed the agent's resources *ex post*. I depart from this modeling by allowing unbounded transfers in the first place, which will be truncated *ex post* in case of insolvency of the agent (or equivalently "judgment-proofness"). Indeed, if one assumes that the private transaction between the principal and the agent is generally not observed by the regulation authority, then there is no reason for the principal to ensure the agent's ability to pay for the penalties. This assumption is also supported by empirical evidence (Ringleb and Wiggins, 1990) that shows that large firms intentionally choose to leave risky manufacturing to small firms with few assets.

For a given regulation policy, I derive the equilibrium contract between the principal and the agent when the former has all the bargaining power. I find that the equilibrium effort level of prevention decreases as the share of penalties imposed on the agent increases. This results stems from the fact that imposing a larger share of penalties on an *ex post* potentially insolvable agent acts as an *ex ante* decrease in the total amount of penalties imposed on the principal-agent relationship. It follows that if the regulation authority wants the firm to reduce as much as possible the probability of accident, it must impose the penalties on the principal only. Targeting the agent is always detrimental to the provision of effort of prevention.

I also investigate the design of the optimal regulation when the bargaining power varies inside the firm. I show that when the principal has most of the bargaining power, it is still optimal to impose all penalties on her. However, when the agent has most of the bargaining power, the previous optimal regulation may lead to an excess of precautionary care (with respect to the first-best level). In that case, the regulation must be such that the total expected fine paid by the firm decreases. Finally, I investigate whether we should authorize the principal to give rewards to the agent in case of accident. It appears that if the agent has most of the bargaining power, allowing the principal to reward the agent in case of accident leads to the first-best level of precautionary care.

Second, I investigate a situation in which two agents are responsible for exerting an unobservable effort of prevention and may both be potentially insolvable *ex post*. The regulation authority now faces a multiple tortfeasors problem where each injurer can be judgment-proof. The contracting problem between the two agents is now described as a double moral hazard problem in a partnership: In the first stage, agents sign a binding agreement to maximize their joint profits and, in the second stage they simultaneously choose their effort levels (similar to Cooper and Ross, 1985). In this problem, the sharing of profits among agents not only plays the traditional role of incentive provision in the partnership (due to moral hazard) but also the role of revenue concealment from the regulation authority.

When agents have symmetric initial resources distributions, the solution to this problem shows that agents sign a contract such that the profits in case of accident go to the agent who

is the least targeted by the regulation policy. This allows the partnership to escape as much as possible from the penalties and thus provide very low effort provision. This result resembles firms' strategies to create insolvent subsidiaries to escape from paying fines.

The optimal regulation policy of the partnership consists in an equal-sharing of the penalties among the two agents. Any other allocation of penalties results in a decrease in effort provision. Notice that when agents have asymmetric initial resources distribution, the optimal allocation of penalties is centered around the equal-sharing allocation while being adjusted to target more the agent with higher average initial resources.

The paper is organized as follows. In Section 2, I present the model and the benchmark case. In section 3, I develop the principal-agent problem with a judgment-proof agent. Section 4 investigates the double moral hazard problem with two judgment-proof agents. Section 6 concludes.

2. BENCHMARK MODEL

A firm undertakes a project that may cause an accident harming third parties. The owner of the firm (the principal) hires a worker (the agent) to run the firm on her behalf. Both of them are assumed to be risk-neutral. The firm's activity generates a certain surplus $\Pi \geq 0$ that accrues to the principal but also causes an environmental harm D with probability $1 - e$. The agent is responsible for preventing the harm by exerting effort $e \in [0, 1]$ at personal cost $\psi(e)$. For the problem to be well-behaved, I assume $\psi', \psi'', \psi''' > 0$, $\psi'(0) = 0$ and $\psi'(1) = +\infty$ so that effort solution is always interior.¹ The agent's effort is unobservable to both the principal and the regulator.

REGULATION. It is assumed that only third parties suffer from the harm D , leaving the firm with no natural incentives to prevent the accident. Therefore, there is room for a regulation authority (the regulator) to act in favor of third parties. Throughout the paper I assume that only ex post regulation is available to the regulator, that is, the regulator can impose fines on the firm's parties only after an harm occurred. A regulatory policy is a couple $(\alpha, F) \in [0, 1] \times [0, D]$, where F is the total amount of the fine imposed on the firm and α (resp. $1 - \alpha$) is the share of the fine charged on the agent (resp. principal).² I also assume that the regulator's objective is to reduce the probability of accident as much as possible.³ This last assumption, in addition to greatly simplify the analysis, allows me to distinguish the judgment-proof problem from other considerations.

PRIVATE CONTRACT. Due to the moral hazard problem, the principal cannot directly offer a contract contingent on the effort level exerted by the agent. Instead, she offers transfers to the agent conditional on the occurrence of an accident. Let us denote by $t_N \in \mathbb{R}$ and $t_A \in \mathbb{R}$ those

¹See for instance Chapter 5 of Laffont and Martimort (2002).

²I assume that the total fine cannot exceed total harm D caused by the firm as it is usually the case in the literature.

³This differs from the literature in which the regulator takes into account both the harm to the third parties and the profit of the firm.

transfers where the subscripts N and A stand for “no accident” and “accident”, respectively. Typically, the principal will offer $t_N \geq 0$ and $t_A \leq 0$ so that t_N and t_A act as a reward and a punishment, respectively. Therefore, the principal has the ability to punish the agent by the means of private transactions as well as the regulator has the ability to punish the agent (and the principal) by the means of the regulation policy.

THE STANDARD MODEL WITH LIMITED LIABILITY. To see how the results depart from the case with a judgment-proof agent it is useful to briefly expose the standard analysis.

For this section only, I will assume the standard limited liability constraints on transfers, that is, $t_N \geq -l$ and $t_A \geq -l + \alpha F$ where $l \in \mathbb{R}_+$ is the agent’s cash/financial resources. Those ex post constraints ensure that the agent is always endowed with enough resources to honor both the private transfer (t_A) and the regulatory transfer (αF). The limited liability constraint in case of accident $t_A \geq -l + \alpha F$ implies, *de facto*, that the principal provides the agent with the necessary resources in the private transaction. As shown below, this requirement trivially implies that the regulator’s choice of distribution of penalties is irrelevant.

Notice that the amount of resources l may not necessarily represent the full range of the agent’s resources but may represent a reasonable lower bound on the “collectible assets” that can easily be observable and seized when an accident occurs. I will therefore abstract from the situation in which the agent can engage in strategies to hide the real value of his assets.⁴

The timing of the game is as follows. First the regulator publicly commits to a regulatory policy (α, F) . Second the principal offers a contract (t_N, t_A) to the agent. Third, the agent chooses his effort level e .

For a given regulation policy (α, F) and private contract (t_N, t_A) the agent’s expected utility is given by

$$U^B = et_N + (1 - e)[t_A - \alpha F] - \psi(e),$$

and the expected profit of the principal writes

$$V^B = \Pi - et_N - (1 - e)[t_A + (1 - \alpha)F].$$

To induce a particular level of effort e , the principal must choose (t_N, t_A) such that the agent has the proper incentives to do so. Relying on the *first-order approach*, I replace the set of the agent’s incentive constraints by the first-order condition of his utility maximization problem with respect to effort:

$$t_N - t_A + \alpha F = \psi'(e), \tag{1}$$

which is a necessary condition when the effort level is interior.⁵ This equation reveals that

⁴Hiriart and Martimort (2006) follow the same approach.

⁵From the initial assumption it also a sufficient condition for the agent’s problem. Indeed, from convexity of ψ the agent’s maximization problem is concave in e . Moreover, ψ' is an increasing function of e and (1) defines a unique solution for any given t_N , t_A and αF .

the agent's incentives to exert effort depends on monetary incentives from both private transactions and regulation policy.

The principal must also ensure the participation of the agent ($U^B \geq 0$) and accounts for ex post limited liability. Her maximization problem then writes

$$\max_{\{e, t^N, t^A\}} V^B = \Pi - [et_N + (1-e)t_A] - (1-e)(1-\alpha)F,$$

subject to equation (1), $U^B \geq 0$, $t_N \geq -l$ and $t_A \geq -l + \alpha F$ for all $l \in [0, \bar{l}]$.

Examining equation (1) immediately reveals that the ex post limited liability constraint on t_N is always satisfied when the one on t_A holds. It is therefore possible to drop condition $t_N \geq -l$ from the principal's problem.

Using (1), the agent's expected utility becomes

$$U^B = e\psi'(e) - \psi(e) + t_A - \alpha F = R(e) + t_A - \alpha F, \quad (2)$$

where $R(e) := e\psi'(e) - \psi(e)$ is nonnegative, increasing and convex in e . Using (1) and (2) it is useful to rewrite the principal's objective in terms in the plane (e, U^B) as follows

$$V^B = \Pi - \psi(e) - (1-e)F - U^B. \quad (3)$$

This formulation clearly shows that the principal's objective takes into account the total amount of the fine and must also leave some rent to the agent. The problem of the principal rewrites

$$\begin{aligned} \max_{\{e, U^B\}} \quad & \Pi - \psi(e) - (1-e)F - U^B \\ \text{s.t.} \quad & U^B \geq 0 \\ & U^B \geq R(e) - l, \end{aligned}$$

where the first constraint is the agent's participation constraints and the second one the limited liability constraints in the plane (e, U^B) .

From linearity of U^B it is clear that at least one constraint must bind at the optimum. Let us focus on the case where only the limited liability constraint binds, $U^B = R(e) - l$.⁶ Substituting the binding constraints in the principal's objective gives

$$\max_e \Pi - e\psi'(e) - (1-e)F.$$

Notice that the principal's maximization problem is independent of α and that the principal takes into account the entire value of the fine F . This leads us to the following statement.

Proposition 1 (Equivalence Principle ⁷) *In a principal-agent relationship with moral hazard, when the principal has to ensure the agent's ability to pay the regulator in any states of the world, the way*

⁶The case where $U^B = 0$ at optimum implies that the level of effort attains the first best, i.e. $\psi'(e^{FB}) = F$. This occurs when the level of liability l is large enough, namely, when $l \geq R(e^{FB})$. This case presents no interest in the analysis as the moral hazard problem entails no distortion. Therefore, assume $l < R(e^{FB})$.

the regulator allocates responsibilities within the firm does not affect the equilibrium level of effort to prevent the harm to occur.

This result is not new and generally known as the Equivalence principle. It states that, even in the presence of moral hazard and agent's limited liability, the regulator cannot increase the level of effort by targeting more the principal or the agent.

The equilibrium level of effort, e^B solves the following first-order condition:

$$\psi'(e^B) + e^B \psi''(e^B) = F. \quad (4)$$

As the left-hand side of the equation is increasing in e , the optimal regulation policy sets $F = D$ according to the *maximum punishment principle*. The α can take any value in $[0, 1]$ without changing the optimal level of effort nor the distribution of revenues between the principal and the agent.

3. JUDGMENT-PROOF AGENT

In this section, I develop the problem of designing a regulatory policy (α, F) for a judgment-proof agent, that is, when he may be unable to pay his share of the fine for some levels of wealth. I argue that the formal treatment of limited liability constraints in the regulation literature is generally too restrictive as they require the principal to fully ensure the agent's ability to pay his share of the penalty.

In the standard treatment of moral hazard problem, limited liability constraints on transfers formally put a limit on the monetary punishment the agent may have to pay to the principal. When the relationship generates only private benefits and no negative externality (such as an accident), the limited liability constraints only applies on the transfers between the principal and the agent and do not involve a third party such as a regulator. In that case, those constraints reasonably assume that the principal must offer a contract whose transfers do not exceed some lower bound.

In the presence of a regulator, however, the assumption of ex post limited liability means that the principal must ensure the agent's ability to pay both private transfers and potential fines to the regulator. This would require (i) that the principal is legally bound to ensure agent's liability in any scenario and (ii) sufficient observability of the private contract that takes place within the firm. When at least one of these conditions fails it is unlikely that the principal will ensure agent's liability towards the regulator. Indeed, assume that the principal does not provide the agent with enough cash to pay for the fine for some realizations of the agent's wealth. When the accident occurs, the agent may simply not be able to pay for the entire fine and will pay at most with his disposable cash. The fact that the agent cannot pay for punishments coming from the regulator does not hurt the principal. On the contrary, the agent's inability to pay the fine to the regulator reduces the total amount of the firm paid by the firm making the principal better off.

To model a judgment-proof agent, I do not assume ex post limited liability constraints as in the previous section. In other words, the private contract offered by the principal does not

need to provide the agent with some minimal level of resources in case of accident. Instead, the principal can impose unlimited punishment on the agent. Obviously, as the agent is still resource-constrained, he will never pay above his level of resources and his transfer will be naturally bounded below. It is then necessary to modify the agent's expected utility accordingly.

AGENT'S PAYOFF. Assume that an accident occurs. For a given level of resources $l \in \mathbb{R}_+$, the agent now faces the private transfer $t_A < 0$ to the principal and the regulatory transfer αF to the regulator. If $t_A - \alpha F$ is larger than the agent's resource $-l$, the agent can fully pay the principal and the regulator. On the contrary, if $t_A - \alpha F < -l$ then the agent has not enough financial resources to cover both transfers. In this section, I assume that the regulator has the ability to collect money before the principal so that when $t_A - \alpha F < -l$, the agent pays the regulator first and gives the remaining resources, if any, to the principal. Let $\tilde{m} := \tilde{c}_P(t_A, \alpha, F; l) - \tilde{c}_R(\alpha, F; l)$ denote the transfer the agent receives when an accident occurs where $\tilde{c}_P(\cdot)$ and $\tilde{c}_R(\cdot)$ denote the agent's transfers to the principal and the regulator, respectively. Formally,

$$\begin{aligned}\tilde{c}_P(t_A, \alpha, F; l) &= \mathbb{1}_{\{l \geq \alpha F - t_A\}} t_A - \mathbb{1}_{\{l \in [\alpha F, \alpha F - t_A]\}} (l - \alpha F), \\ \tilde{c}_R(\alpha, F; l) &= \mathbb{1}_{\{l \geq \alpha F\}} \alpha F + \mathbb{1}_{\{l \leq \alpha F\}} l.\end{aligned}$$

Notice that the transfer to the regulator, $\tilde{c}_R(\cdot)$ is independent of t_A . This stems from the fact that the regulator has priority over the agent's wealth so that the principal cannot reduce the agent's payment to the regulator by punishing him more through t_A .

In order to greatly simplify the mathematical reasoning I now make the following assumption.

Assumption 1 *The agent's wealth l is a random variable drawn from an absolutely continuous cumulative distribution function $H(l)$ over the support $[0, \bar{l}]$ where $\bar{l} > D$ so that the maximal agent's wealth exceeds the monetary damage caused by an accident. The realization of the agent's wealth is known to all players only at the end of the game.*

As all players now view l as a random variable, they will evaluate their payoff by taking expectations of $\tilde{c}_P(\cdot)$ and $\tilde{c}_R(\cdot)$ over l . Define $c_P(t_A, \alpha, F) := \mathbb{E}_l \tilde{c}_P(t_A, \alpha, F; l)$ and $c_R(\alpha, F) := \mathbb{E}_l \tilde{c}_R(\alpha, F; l)$. This makes the analysis easier as now $c_P(t_A, \alpha, F)$ and $c_R(\alpha, F)$ are differentiable in each argument. No particular conditions are imposed on $H(\cdot)$ so that the randomness assumption of l is quite mild and could simply represent that there is always a small uncertainty about the agent's wealth when the regulator decides to enforce penalties. Furthermore, I show below that the framework with a judgment-proof agent and random wealth can be equivalently rewritten in a model with certain wealth that resembles the model presented in section 2.

Formally, the total expected transfer of the agent in case of accident is given by

$$m(t_A, \alpha, F) = \int_{\alpha F - t_A}^{\bar{l}} (t_A - \alpha F) dH(l) + \int_0^{\alpha F - t_A} (-l) dH(l). \quad (5)$$

Taking partial derivatives of $m(t_A, \alpha, F)$ with respect to t_A and α gives:

$$\begin{aligned}\frac{\partial m(t_A, \alpha, F)}{\partial t_A} &= \int_{\alpha F - t_A}^{\bar{l}} dH(l) \geq 0, \\ \frac{\partial m(t_A, \alpha, F)}{\partial \alpha} &= - \int_{\alpha F - t_A}^{\bar{l}} F dH(l) \leq 0.\end{aligned}$$

Thus, the transfer of the agent is nondecreasing in t_A and nonincreasing in α . More precisely, $m(t_A, \alpha)$ increases in t_A when t_A is high enough, but as soon as t_A becomes too low the agent's payment becomes flat and equal to $-l$. The same thing happens with α : the agent's transfer is decreasing in α as long as α is not too high and then becomes flat when the regulator asks too much money.

It is now possible to write the agent's expected utility as

$$U = e t_N + (1 - e) [c_P(t_A, \alpha, F) - c_R(\alpha, F)] - \psi(e),$$

where the only difference with the benchmark case is that t_A is replaced by $c_P(t_A, \alpha, F) - c_R(\alpha, F)$. For a given private contract (t_N, t_A) the induced level of effort e is given by

$$t_N - [c_P(t_A, \alpha, F) - c_R(\alpha, F)] = \psi'(e). \quad (6)$$

Therefore, when the principal offers t_A in case of accident, the agent now considers $c_P(t_A, \alpha, F) - c_R(\alpha, F)$ rather than directly t_A . Using (6), the agent's expected utility becomes

$$U = R(e) + c_P(t_A, \alpha, F) - c_R(\alpha, F) \quad (7)$$

where $R(e) = e\psi'(e) - \psi(e) \geq 0$ is defined as in section 2. Equation (7) shows that when the principal wants to induce level of effort e , she has to give a positive rent $R(e)$ to the agent. Then the principal can extract this rent through the side-payment $c_P(t_A, \alpha, F) - c_R(\alpha, F)$. Notice that for any $t_A \leq \alpha F - \bar{l}$, $c_P(t_A, \alpha, F) - c_R(\alpha, F)$ is bounded below by $-\mathbb{E}[l]$. That is, even when the principal sets a very low t_A , she cannot extract more than $\mathbb{E}[l]$ from the agent through the side-payment.

Again, when the principal sets t_A she only expects to receive $c_P(t_A, \alpha, F)$ from the agent. Thus, her expected profit is given by

$$V = \Pi - e t_N - (1 - e) c_P(t_A, \alpha, F) - (1 - e)(1 - \alpha)F.$$

Notice that

$$\begin{aligned}\frac{\partial c_P(t_A, \alpha, F)}{\partial t_A} &= \int_{\alpha F - t_A}^{\bar{l}} dH(l) \geq 0, \\ \frac{\partial c_P(t_A, \alpha, F)}{\partial \alpha} &= \int_{\alpha F}^{\alpha F - t_A} F dH(l) \geq 0.\end{aligned}$$

Naturally, the more the principal increases the punishment in case of accident (reduces t_A)

the more she can expect to collect on the agent. More interestingly, let us investigate what happens for the principal when the regulator targets the agent more, *i.e.* when α increases. Two effects are at play: On the one hand, as the agent becomes more targeted by the regulation, the principal expects to collect less on the agent as the regulator has priority over the agent's wealth which is illustrated by $\partial c_P(\cdot)/\partial \alpha \leq 0$. On the other hand, an increase in α increases the principal's expected payoff through a decrease in his share of the fines. The overall effect, however, is positive on the principal's expected payoff.⁸

Using (6) and (7), it is useful to rewrite the principal's profit as

$$V = \Pi - \psi(e) - (1 - e)[c_R(\alpha, F) + (1 - \alpha)F] - U. \quad (8)$$

It is instructive to compare the objective of the principal with a judgment-proof injurer with the one obtained in the benchmark model, namely, equation (3). The principal still receives benefits from production Π , has to pay $\psi(e)$ as if she were exerting effort herself and leaves a rent U to the agent. However, the principal considers the threat of the penalty differently. Notice that $c_R(\alpha, F) + (1 - \alpha)F \leq \alpha F + (1 - \alpha)F = F$, that is, the principal does not consider the total amount of fines F as some of it is imposed on a potentially insolvent agent. The assumption of judgment-proofness on the agent's side formalizes the idea that potential insolvency of an injurer creates a discrepancy between the sanction and the way it is perceived by the firm.

When the regulator targets more the agent (α increases), the firm then faces an overall lower sanction than F . At the same time, the burden of payments rests more upon the agent whose participation must be ensured.

The principal's problem now consists in maximizing her objective subject to the agent's participation. It is worth to stress once again that the principal does not ensure a limited liability constraints on transfers so that t_A is here unconstrained. For the sake of clarity, let us first consider the following formulation of the principal's problem

$$\begin{aligned} \max_{e, t_A} \quad & V = \Pi - e t_N - (1 - e)c_P(t_A, \alpha, F) - (1 - e)(1 - \alpha)F \\ \text{s.t.} \quad & U = R(e) + c_P(t_A, \alpha, F) - c_R(\alpha, F) \geq 0. \end{aligned}$$

As in section 2, let us make a change of variable so that the principal chooses e and U instead of e and t_A . However, this requires to take into account that even if t_A is unbounded below, $c_P(t_A, \alpha, F) - c_R(\alpha, F)$ is bounded by $-\mathbb{E}[U]$. This limits how much the principal can take collect on the agent through the side-payment. Using (7), this condition can be written as $c_P(t_A, \alpha, F) - c_R(\alpha, F) = U - R(e) \geq -\mathbb{E}[U]$. Notice that this constraint resembles the standard

⁸This can be easily seen by differentiating the principal's expected payoff in case of accident with respect to α . Indeed, $\frac{\partial}{\partial \alpha}(-c_P(t_A, \alpha, F) - (1 - \alpha)F) = -\int_{\alpha_F}^{\alpha_F - t_A} F dH(l) + F \geq 0$.

limited liability constraint of section 2. The principal's problem rewrites

$$\begin{aligned}
 (\text{JP}) : \max_{e, \mathcal{U}} \quad & \Pi - \psi(e) - (1 - e)[c_R(\alpha, F) + (1 - \alpha)F] - \mathcal{U} \\
 \text{s.t.} \quad & \mathcal{U} \geq 0 \\
 & \mathcal{U} \geq R(e) - \mathbb{E}[\mathcal{U}].
 \end{aligned}$$

The following proposition summarizes the solution to the optimization problem.

Proposition 2 *Assume the agent may be judgment proof for some realizations of his wealth. Then, the equilibrium effort level, $e(\alpha, F)$, induced by the principal is nondecreasing in α . More precisely,*

- For $\alpha \in [0, \tilde{\alpha}_1)$, $c_P(t_A, \alpha, F) - c_R(\alpha, F) = -\mathbb{E}[\mathcal{U}]$, the agent has a positive utility, $\mathcal{U} > 0$, and the optimal effort level is given by

$$\psi'(e) + e\psi''(e) = c_R(\alpha, F) + (1 - \alpha)F. \quad (9)$$

- For $\alpha \in [\tilde{\alpha}_1, \tilde{\alpha}_2]$, $c_P(t_A, \alpha, F) - c_R(\alpha, F) = -\mathbb{E}[\mathcal{U}]$ and the optimal effort level is given by the binding participation constraint of the agent, $\mathcal{U} = 0$,

$$e\psi'(e) - \psi(e) = \mathbb{E}[\mathcal{U}]. \quad (10)$$

- For $\alpha \in (\tilde{\alpha}_2, 1]$, $c_P(t_A, \alpha, F) - c_R(\alpha, F) > -\mathbb{E}[\mathcal{U}]$, the agent's participation constraint is binding, $\mathcal{U} = 0$, and the optimal effort level is given by

$$\psi'(e) = c_R(\alpha, F) + (1 - \alpha)F. \quad (11)$$

Proof. The proof and the characterization of thresholds $\tilde{\alpha}_1$ and $\tilde{\alpha}_2$ are in the appendix. ■

DISTRIBUTION OF PENALTIES. When the agent may be judgment proof, the equilibrium solution crucially depends on the distribution of liabilities α and $1 - \alpha$ within the firm. Proposition 2 states that the optimal effort level $e(\alpha, F)$ is nonincreasing in α , that is, the firm exerts less and less precautionary effort to prevent an accident as the regulator increases the liability of the agent to pay for the damage. The intuitive explanation for this results is as follows. From the agent's point of view, only the size of the punishment matters and not to whom it is due. Therefore, whether the monetary punishment comes from the principal or the regulator is irrelevant for the agent's decision to exert the precautionary effort. From the principal's point of view, however, a shift in regulation that targets more the agent induces a lower expected total fine on the firm through the agent's inability to pay in some states of the world. As a result, the principal perceives the fine less and less as a threat and does not want to induce a high effort level.

Therefore, the Equivalence Principle does not hold anymore in this context. In other words, the structure of penalties designed by the regulator affects the equilibrium level of effort chosen by the firm. This naturally raises the question of determining the optimal regulatory

policy (α, F) with a judgment-proof agent. In the simple case in which the regulator is only concerned about reducing the probability of accident, the optimal distribution of fines is as follows.

Corollary 1 *With a judgment-proof agent, the optimal regulatory policy (reduce the probability of accident) consists in targeting only the principal of the firm, that is, $\alpha = 0$ so that the principal faces the total amount of the fine F .*

This result contrasts with Segerson and Tietenberg (1992). As soon as the agent may be unable to pay the regulator in some states of the world, the Equivalence Principle fails and the distribution of penalties within the firm is no more neutral. Considering the standard formulation presented in section 2, this result should not be surprising at all. Indeed, the ex post limited liability constraint $t_A \geq -l + \alpha F$ in the standard model artificially assumes that the principal must ensure the agent's ability to pay his share fines for any value of α . This implicitly amounts to saying that the whole burden of penalties lies on the principal, which is equivalent to consider that $\alpha = 0$ in the judgment-proof case.

It is worth noting that when the regulator chooses $\alpha = 0$, the agent's transfer to the regulator is naturally $c_R(0, F) = 0$. In that case, Proposition 2 gives that the equilibrium effort level is uniquely defined by $\psi'(e) + e\psi''(e) = c_R(0, F) + F = F$.⁹ This effort level is the same as the one obtained in section 2 (equation (4)) with the standard limited liability constraints. However, as soon as $\alpha > 0$, the equilibrium effort level in the judgment proof case decreases.

EQUILIBRIUM CHARACTERIZATION. The equilibrium characterization along the value of α is also worth to investigate. When the principal is mostly targeted, *i.e.* $\alpha \in [0, \tilde{\alpha}_1)$, the equilibrium effort level is given by (9). As the left-hand side of the equation is strictly decreasing in α , so is the equilibrium level of effort. Moreover, the principal chooses t_A such that $c_P(t_A, \alpha, F) - c_R(\alpha, F) = -\mathbb{E}[U]$ in order to extract as much as possible from the agent through the side-payment. Even if the principal is not as concerned by the threat of the sanction as in the benchmark case, she is still inducing a quite large effort level which forces her to leave a positive rent to the agent, $U > 0$. Indeed, to induce this effort level, she must leave a rent $R(e)$ to the agent that is greater than what she can extract from the agent through the side-payment. When $\alpha \in [\tilde{\alpha}_1, \tilde{\alpha}_2]$, the effort level is given by equation (10), which is simply the binding participation constraint of the agent. The principal still chooses t_A so as to extract all the rent from the agent. She can therefore implement the effort level at no informational cost. The equilibrium effort level is constant over the region $\alpha \in [\tilde{\alpha}_1, \tilde{\alpha}_2]$ as the principal is still concerned by the threat of sanction and does not face the trade-off between increasing the effort level and minimizing the rent left to the agent. However, when the agent becomes mainly targeted, *i.e.* when $\alpha \in (\tilde{\alpha}_2, 1]$, the principal is almost not anymore concerned by the threat of the sanction and the effort level decreases once again with respect to α . As effort level becomes low, the rent $R(e)$ left to the agent becomes low as well and the principal chooses higher values of t_A so that the agent still wants to participate. Therefore, we

⁹To see that the effort level is uniquely defined by the equation, simply notice that $\frac{\partial}{\partial e}(\psi'(e) + e\psi''(e)) = 2\psi''(e) + e\psi'''(e) > 0$ as ψ'' , $\psi''' > 0$ by assumption.

can also see that the more the principal is targeted by the regulation the larger the punishment t_A she imposes on the agent.

Consider finally the case where the regulator mainly targets the agent, that is, $\alpha \in (\tilde{\alpha}_2, 1]$. Recall that the first-best solution to the moral hazard problem is defined by $\psi'(e^{FB}) = F$. Then, it is clear that it resembles equation (11) that determines the equilibrium level of effort for $\alpha \in (\tilde{\alpha}_2, 1]$. When the agent is mainly targeted, his incentives to exert effort comes mainly from the regulator and that from the principal decreases as she is less and less concerned by the risk of accident. Equation (11) resembles the first-best solution as the agent internalizes the risk of accident but the effort level is much lower as the fine the agent's expect to pay is less and less important.

TOTAL AMOUNT OF PENALTIES. The second choice of the regulator is the total amount of fine F that is imposed on the firm. Recall that $F \in [0, D]$ as tort law generally precludes fines to exceed monetary damages caused by the accident. The optimal level of fine is defined as follows.

Corollary 2 *When the regulator optimally allocates fines within the firm ($\alpha = 0$), the optimal level of fine that maximizes the equilibrium agent's effort is $F = D$, that is, the maximum punishment principle applies with a judgment-proof agent.*

If Corollary 1 challenges the view that allocation of penalties within the firm is relevant when the agent may be judgment-proof, Corollary 2 re-establishes a very well known result: the so-called Becker's maximum punishment principle. This last result is not surprising: potential judgment proofness of the agent reduces the expected total fine perceived by the principal. Thus, in essence, Corollary 1 tells us that fully targeting the principal is the only way to maximize the "perceived" total fine which is a maximum punishment principle in itself. Intuitively, choosing the maximal amount of total fines ensures that the firm internalizes the risk of accident at most. In particular, when $\alpha = 0$ and $F = D$, the effort level is given by $\psi'(e) + e\psi''(e) = D$ which corresponds to the equilibrium effort in the standard case (when the distribution of fines does not matter) when the regulator chooses the optimal regulation. It is also interesting to notice that choosing $F = D$ is optimal even when $\alpha > 0$.

UNCERTAINTY OF PENALTIES. So far, I have assumed that the regulation was nonrandom so that the firm can perfectly anticipate the amount of fines and its distribution between the principal and the agent. In practice, however, it is likely that there is some uncertainty about the exact amount and who will be accountable for it. The analysis of the judgment-proof case shows that what matters for regulation is how the firm perceives the threat of paying fines. As soon as the fine is not at its maximum or if there is a chance that the agent cannot fully pay his share, the threat perceived by the firm becomes lower. As a result, any uncertainty in the total amount of penalties or about its distribution will makes the firm less concerned about the accident. Whenever possible, the regulator should announce and commit to a certain regulation policy to ensure that the firm better internalizes the risk of accident.

AN EQUIVALENT FORMULATION. The judgment-proof problem can be simply rewritten in a formulation very similar to the standard model with ex post limited liability constraints.

This reformulation is useful to interpret and compare the judgment-proof problem with the standard formulation as well as providing a simpler and more tractable model.

Simply consider the following expected payoffs and limited liability constraints:

$$\begin{aligned} U^E &= e t_N + (1 - e) [t_A - c_R(\alpha, F)] - \psi(e) \\ V^E &= \Pi - e t_N - (1 - e) [t_A + (1 - \alpha)F] \\ t_A &\geq -\mathbb{E}[l] + c_R(\alpha, F) \end{aligned}$$

where the only difference with the standard model of section 2 is that the agent's payment to the regulator is $c_R(\alpha, F)$ instead of αF and l replaced by $\mathbb{E}[l]$. The incentive constraints immediately writes $t_N - [t_A - c_R(\alpha, F)] = \psi'(e)$. Therefore, $U^E = R(e) + t_A - c_R(\alpha, F)$ and the limited liability constraint can be rewritten as $U^E - R(e) + c_R(\alpha, F) \geq -\mathbb{E}[l] + c_R(\alpha, F) \Leftrightarrow U^E \geq R(e) - \mathbb{E}[l]$. The participation constraint of the agent still writes $U^E \geq 0$ and the principal's expected payoff is $V^E = \Pi - \psi(e) - (1 - e) [c_R(\alpha, F) + (1 - \alpha)F] - U$. The principal's problem therefore writes exactly as problem (JP) and the equilibrium effort level is characterized by proposition 2.

This equivalent formulation shows that it is as if ex post limited liability constraints in the standard model were replaced by interim limited liability constraints. In other words, the private contract ensures that the agent has enough resources to pay the regulator but only in average and not ex post. This suggests that even if the principal is legally bound on the size of the punishment is limited the nonneutrality of the distribution of fines still hold when the principal is constrained on the size of punishment

LIMITS OF THE OPTIMAL REGULATION. Although the optimal regulation with a judgment-proof agent seems to be very clear, it can be difficult to implement as it if for various reasons. First, they might be legal restriction on the choice of the distribution of liabilities within the firm. The question of whether the principal can be vicariously liable for the acts of the agent is likely to depend on the nature of their relationship. Vicarious liability generally applies when the agent commits negligent acts in the course of employment. In the case of the present paper, the agent can be seen as an employee or an independent contractor whose responsibility is to ensure some level of precautionary care on behalf of the principal. Although the principal is responsible to incentivize the agent, the latter can still be seen as responsible for his choice of precautionary care. Unobservability of the effort level or the precise private contract between the two makes it difficult to assess each injurer's responsibility in the accident.

Second, if the regulator decides to target only the principal two problems can arise: (i) The principal may simply not engage in production as she expects that the benefits from production Π do not cover the expected fines and costs of precaution. In other words, the principal's participation constraint must also be taken into account in the design of the optimal regulation. Or (ii) If the fine is large, the principal may also have insufficient financial resources to cover the whole payment by herself. In that case, also relying on the agent's wealth may improve the amount that the regulator is able to collect.

DISTRIBUTION OF REVENUES WITHIN THE FIRM. One interesting aspect of the framework with a judgment-proof agent is that α also affects the distribution of revenues within the firm. Let $e = e(\alpha, F)$ the equilibrium effort level defined by proposition 2. First, the agent's expected utility is $U = R(e(\alpha, F)) - \mathbb{E}[l] > 0$ for all $\alpha \in [0, \hat{\alpha}_1)$ and then $U = 0$ for all $\alpha \in (\hat{\alpha}_1, 1]$. On the principal's side, let $V(\alpha)$ denote the value function of program (JP). From the envelope theorem, $V'(\alpha) = -(1 - e) \frac{\partial}{\partial \alpha} (c_R(\alpha, F) + (1 - \alpha)F) > 0$ so that the principal's expected payoff is increasing in α . Therefore, an increase in α benefits the principal but decreases the rent of the agent. It may also appear surprising that the optimal regulation $(\alpha, F) = (0, D)$ is such that the agent gets the highest possible rent although he is the one who may cause the accident by taking too little precautionary care. In fact, leaving the agent with a high rent is the best possible way to incentivize him to maximize the effort level.

3.1. Participation of the Principal

As mentioned above, when the regulator imposes the optimal regulation $(\alpha, F) = (0, D)$ it may occur that the principal makes negative profits $V < 0$. Even if the regulator wants to avoid as much as possible the occurrence of an accident, it does not mean that the economic activity must be prohibited. Ensuring the participation of the principal is therefore important if the regulator values the production of the firm (though not modeled here).

For simplicity, first assume that the regulator sets $F = D$ by default and can only use α as an instrument policy. It follows that the optimal regulation $(\alpha, F) = (0, D)$ ensures the principal's participation as long as

$$\Pi - \psi(e^*) - (1 - e^*)[c_R(0, D) + (1 - \alpha)D] - R(e^*) + \mathbb{E}[l] \geq 0,$$

where e^* solves equation (9) for $\alpha = 0$ and $F = D$. This condition can be violated when Π or $\mathbb{E}[l]$ are low. If this occurs, assuming that the regulator only considers $\alpha \in [0, \hat{\alpha}_1]$, the new optimal choice of allocation of responsibilities, α^V within the firm must solve

$$\Pi - \psi(e(\alpha^V, D)) - (1 - e(\alpha^V, D))[c_R(\alpha^V, D) + (1 - \alpha)D] - R(e(\alpha^V, D)) + \mathbb{E}[l] = 0$$

This is indeed possible as the profit of the principal is increasing in α so that there exists a $\alpha^V > 0$ such that the principal's profit is nonnegative.

A more complete analysis of the optimal regulation subject to the principal's participation would allow for change in both α and F with respect to Corollaries 1 and 2. Let $\Omega(\alpha, F) := c_R(\alpha, F) + (1 - \alpha)F$ and $V(\Omega(\alpha, F))$ denote the total amount of fines imposed on the firm and the value function of problem (JP), respectively. From the Envelope Theorem, it is straightforward to see that the principal's expected payoff decreases in $\Omega(\cdot)$ as $\partial V(\Omega(\alpha, F))/\partial \Omega = -(1 - e(\alpha, F)) < 0$. Furthermore, proposition 2 states that the equilibrium effort level is (weakly) increasing in $\Omega(\cdot)$. Therefore, to ensure the principal's participation constraint and maximize the equilibrium effort level, the regulator must choose Ω such that

When both α and F are close to 0 and D , respectively, the equilibrium is given by Proposition 2 (a). Therefore, when the regulator seeks to maximize precautionary care it is equivalent to choose α and F so that

$$\begin{aligned} \max_{\Omega \in [\mathbb{E}[U], D]} \quad & \Omega \\ \text{s.t.} \quad & V(\Omega) \geq 0 \end{aligned}$$

It is clear that the constraint must bind so that the optimal choice of Ω satisfies $V(\Omega^*) = 0$. Uniqueness of Ω^* is guaranteed by $V'(\Omega) < 0$. Then, the regulator must choose (α^*, F^*) so that $c_R(\alpha^*, F^*) + (1 - \alpha^*)F^* = \Omega^*$. Observe, however, that the choice of (α^*, F^*) is not unique so that the optimal regulation can be achieved with various combinations of the regulation instruments. More precisely, if (α^*, F^*) implements Ω^* then it is also possible to find $(\hat{\alpha}, \hat{F})$ with $\hat{\alpha} > \alpha^*$ and $\hat{F} > F^*$, that is, another regulation scheme in which the agent is more targeted but the total amount of fines imposed on the firm increases.¹⁰

Proposition 3 *When the optimal regulation $(\alpha^*, F^*) = (0, D)$ is such that the principal does not want to participate ex ante, the regulation policy (α, F) must be such that the total perceived fine solves $V(\Omega(\alpha, F)) = 0$, that is, the principal's expected profit is null. The choice of (α, F) is nonunique.*

Participation of the principal crucially depends on the benefits of the productive activity Π and on the agent's expected wealth $\mathbb{E}[U]$. Indeed, let $e = e(\Omega)$, then the optimal regulation subject to the principal's participation sets Ω to solve

$$\Pi - \psi(e(\Omega)) - (1 - e(\Omega))\Omega - R(e(\Omega)) + \mathbb{E}[U] = 0.$$

Totally differentiating this expression yields $\frac{d\Omega}{d\Pi} = \frac{d\Omega}{d\mathbb{E}[U]} = -\frac{1}{V'(\Omega)} > 0$. Hence, an increase in either the productive activity of the expected agent's wealth allows the regulator to set a higher total perceived fine Ω which, in turn, leads to higher level of precautionary care. Quite intuitively, it seems therefore easier to regulate profitable businesses with wealthy members rather than small returns activities with very financially constrained agents.

3.2. Bargaining Power and Optimal Regulation

So far, I have assumed that the principal had all the bargaining power in the choice of the private contract. It seems important to investigate the importance of this assumption for the choice of the equilibrium effort level. The results for the optimal regulation suggest that imposing the total amount of the fine on the principal forces her to fully internalize the sanction. But this result holds because the principal is the one with the bargaining power in the relationship. Indeed, if the agent has more bargaining power in the private relationship, he will try to impose his terms to the principal and will obviously benefit from a heavily targeted principal.

To investigate this issue, assume that $b \in [0, 1]$ and $1 - b$ denote the bargaining power of the principal and the agent, respectively. Let us now assume that the principal and the agent

¹⁰Naturally, I assume that $\Omega^* < \Omega(0, D)$ and $(\alpha^*, F^*) \in (0, 1) \times (0, D)$ so that there exists $\hat{\alpha} \in (\alpha^*, 1]$ and $\hat{F} \in (F^*, D]$.

have the following objective

$$bV + (1 - b)U = b \left[\Pi - et_N - (1 - e)c_P(t_A, \alpha, F) - (1 - e)(1 - \alpha)F \right] \\ + (1 - b) \left[et_N + (1 - e)[c_P(t_A, \alpha, F) - c_R(\alpha, F)] - \psi(e) \right]$$

The agent's local incentive constraint is still characterized by (7). Substituting it into the objective of the coalition and expressing everything in the plane (e, U) as before, the objective can be rewritten as

$$T := (\Pi - \psi(e) - (1 - e)[c_R(\alpha, F) + (1 - \alpha)F] - U) + \beta U,$$

where $\beta := \frac{1-b}{b}$ represents the agent's relative bargaining power.¹¹ When the agent has bargaining power, he will try to extract rent from the principal. However, I still assume that the principal cannot reward the agent in case of accident, that is, $t_A \leq 0$, so that there is an upper bound on how much rent the agent can extract from the principal. Formally, if $t_A \leq 0$ then $c_P(t_A, \alpha, F) \leq 0$. This constraint can be rewritten as $c_P(t_A, \alpha, F) = U - R(e) + c_R(\alpha, F) \geq 0$. The maximization problem writes

$$\begin{aligned} \max_{e, U} \quad & (\Pi - \psi(e) - (1 - e)[c_R(\alpha, F) + (1 - \alpha)F] - U) + \beta U \\ \text{s.t.} \quad & U \geq 0 \\ & U \geq R(e) - \mathbb{E}[U] \\ & U \leq R(e) - c_R(\alpha, F) \\ & \Pi - \psi(e) - (1 - e)[c_R(\alpha, F) + (1 - \alpha)F] - U \geq 0, \end{aligned}$$

Notice that the value of β plays a crucial part in the solution of this problem. Let us first consider the case $\beta \in [0, 1)$. It immediately follows that U enters the objective negatively. Therefore, as in the case in which the principal has all the bargaining power (special case $\beta = 0$ here), at least one of the agent's constraint must bind and the principal's participation constraint is relaxed as U decreases.

Let μ and ν be the Lagrange multipliers associated with the two first constraints. Ignoring the principal's participation constraint, the first-order conditions of the problem write¹²

$$\begin{aligned} \psi'(e) + \nu e\psi''(e) &= c_R(\alpha, F) + (1 - \alpha)F \\ \mu + \nu &= 1 - \beta. \end{aligned}$$

For simplicity let us focus on the case where α is low (similarly to Proposition 2, case (a)). Only the second constraint binds and thus $\mu = 0$. It follows that the equilibrium effort level is given by $\psi'(e) + (1 - \beta)e\psi''(e) = c_R(\alpha, F) + (1 - \alpha)F$.

¹¹After plugging (7) into T and changing variables the actual objective of the coalition is bT . As it is equivalent to maximize bT and T , I choose the latter for convenience.

¹²As long as Π is large enough, the principal's participation constraint will not be a problem here as she has most of the bargaining power.

Proposition 4 *When the principal is dominant, $\beta < 1$, the equilibrium effort level is increasing in β and the optimal regulation policy still satisfies $(\alpha, F) = (0, D)$.*

This result simply confirms results obtained in section 3 when the principal has most of the bargaining power. It shows, however, that an increase in the agent's bargaining power mitigates the judgment-proof problem as the equilibrium effort level increases in β .

When $\beta > 1$, the agent has most of the bargaining power and U now enters the objective positively. It is therefore clear that either the third or the fourth constraint is now binding. For simplicity, I assume that $\Pi > D$ so that the participation constraint of the principal is never binding (as shown below). Let $U = R(e) - c_R(\alpha, F)$, then maximizing the objective with respect to e gives the following first-order condition¹³

$$\psi'(e) + (1 - \beta)e\psi''(e) = c_R(\alpha, F) + (1 - \alpha)F.$$

Notice that when $\beta = 1$ and $\alpha = 0$, the equilibrium effort level solves $\psi'(e) = F$. Hence, if the regulator sets $F = D$, it is possible to achieve the first-best effort level (characterized by $\psi'(e) = D$). This occurs as $\beta = 1$ is equivalent to $b = 1/2$, that is, the principal-agent coalition puts the same weights on each member's payoff. However, as equilibrium effort level is increasing β it follows that if $\beta > 1$ and $(\alpha, F) = (0, D)$ then the equilibrium effort level is higher than the first-best one. The following proposition summarizes those results.

Proposition 5 *When the agent is dominant, $\beta \geq 1$, the equilibrium effort level is increasing in β . If the regulator chooses $(\alpha, F) = (0, D)$, the equilibrium effort level attains the first-best level for $\beta = 1$ but exceeds the first-best level when $\beta > 1$.*

When the agent is *dominant*, setting the regulation policy $(\alpha, F) = (0, D)$ might lead to an excessive precautionary effort level. The intuition behind this finding is the following. The *dominant* agent tries to extract as much rent as possible from the principal through increases in payments t_N and t_A . However, recall that only punishments are available in case of accident ($t_A \leq 0$) so that once the agent has set $t_A = 0$ (from the binding constraint $U = R(e) - c_R(\alpha, F)$), the only to extract more rent from the principal is through an increase in t_N , the payment in the absence of accident. This makes the agent even more incentivized that no accident occurs and in that situation the effort level becomes higher than the first-best one. The optimal regulation in that case might therefore surprisingly be milder and a decrease in the total amount of perceived fines $c_R(\alpha, F) + (1 - \alpha)F$ becomes desirable. Once again, several combinations of (α, F) can achieve a lower total amount of perceived fines.

REWARDS IN CASE OF ACCIDENT. So far, I have assumed that the principal could not offer rewards ($t_A \geq 0$) in case of accident but only a punishment ($t_A \leq 0$). Although it may be difficult to regulate private contracts between a principal and an agent, it may be possible to make sure that an agent does not receive bonuses when accident occurs. This legal restriction can be imposed for ethical reasons or following the reasoning of incentive theory. Furthermore,

¹³The second-order condition requires that $-\psi''(e) - (1 - \beta)[\psi''(e) + e\psi'''(e)] \leq 0$ or, equivalently, that β is not too large. To ensure that the effort level is always interior, I assume that β is such that the second-order condition always hold.

when the principal has all the bargaining power, it is intuitive that she would not offer rewards to the agent as it would both reduce incentive provision and the possibility of extracting rent from him.

When the agent is dominant, however, we have seen that the private contract is such that $t_A = 0$ when only punishments are available. This suggests that a dominant agent would like to set positive t_A if possible. Assume now that rewards in case of accident, *i.e.* $t_A \geq 0$, are allowed. The agent's payoff in case of accident can still be written as $m(t_A, \alpha, F) = c_P(t_A, \alpha, F) - c_R(\alpha, F)$ so that neither his expected utility nor his incentive constraint changes. For the principal, however, the payoff in case of accident simply becomes t_A (instead of $c_P(t_A, \alpha, F)$) as the reward she pays the agent is independent from the regulation policy. Her expected payoff therefore rewrites $V^R = \Pi - e t_N - (1 - e)t_A - (1 - e)(1 - \alpha)F$.

For simplicity, assume that the agent has all the bargaining power. It follows that the maximization problem writes

$$\begin{aligned} \max_{e, t_A} \quad & R(e) + m(t_A, \alpha, F) \\ \text{s.t.} \quad & \Pi - e\psi'(e) - (1 - e)(1 - \alpha)F - e m(t_A, \alpha, F) - (1 - e)t_A \geq 0, \end{aligned}$$

where it is clear that we can ignore the agent's participation constraint. As $\frac{\partial m(t_A, \alpha, F)}{\partial t_A} \geq 0$, the objective is increasing in t_A whereas the principal's profit is decreasing in t_A . Then, the principal's participation constraint is binding. Notice that if $t_A \geq \alpha F$, *i.e.* in case of accident the agent receives a reward that covers his share of fines, then $m(t_A, \alpha, F) = t_A - \alpha F$. Assume that the equilibrium t_A is greater than αF , then it follows that the binding participation constraint of the principal gives $t_A = \Pi - e\psi'(e) - (1 - e)(1 - \alpha)F + e\alpha F$. Plugging this expression into the objective reduces the problem to $\max_e \Pi - \psi(e) - (1 - e)F$. This program is simply the social objective and it yields the first-best level of effort $\psi'(e) = F$.

This result, although not surprising, stresses two important insights for the regulation.¹⁴ First, this finding shows that the allocation of penalties between the principal and the agent is irrelevant when the agent has all the bargaining power. Intuition could have suggested that members with more bargaining power should be more targeted but this is not the case. Second, it shows that rewards should be allowed in case of accident. When they are, a fully dominant agent can extract the whole surplus from the principal and the first-best level of effort is attained. This contrasts with the result of proposition 5 in which only punishments are allowed and the agent chooses an excessively large effort level.

4. TWO-SIDED MORAL HAZARD

In many situations, the responsibility of preventing an accident lies with more than one agent. The probability of accident may then depend upon the action of several agents and it is not possible to hire a single agent in charge of safety. In that case, what should be the optimal

¹⁴As in standard principal-agent models, giving the bargaining power to the informed party makes the moral hazard problem disappear as the rent extraction-efficiency trade-off vanishes.

targeting policy if some of the tortfeasors are judgment-proof? Agents may differ in their wealth characteristics and cost of effort. Should the regulation change with respect to agents' efficiency and wealth?

To capture the fact that the probability of accident depends on more than one agent, I now model a two-agent partnership with double moral hazard. Assume that each partner receives a non-contractible benefit b only when no accident occurs whereas they receive $\Pi < b$ independently of the accident. Therefore, their contract consists in choosing a sharing of Π for the two state of the world. The environmental harm is denoted by $D < \Pi$. Agent 1 and agent 2 exert efforts $e \in [0, 1]$ and $a \in [0, 1]$, respectively. The probability of accident is determined by $(1 - p(e, a))$ where $p(e, a)$ is the joint production function. Individual cost functions of effort are $\psi(e)$ and $C(a)$ that are both increasing and convex functions. To obtain a tractable model I further assume the following specific functional forms: $p(e, a) = e + a$, $\psi(e) = \gamma_1 \frac{e^2}{2}$ and $C(a) = \gamma_2 \frac{a^2}{2}$.

As in the one-sided moral hazard case, the regulator chooses the level of total fine $F \in [0, D]$ and a distribution of fines (α_1, α_2) among the two agents, where $\alpha_1 + \alpha_2 = 1$ and $\alpha_i F$ is the share of the fine imposed on agent $i = 1, 2$. For simplicity, I will assume that $F = D$ so that the regulator only has to determine the distribution of fines.

Each agent $i = 1, 2$ has a liability l_i , uniformly distributed over $[0, \bar{l}_i]$. Neither the regulator nor the agents know the value of l_1 and l_2 until the end of the game. The final wealth of agent i is given by the realization of his liability and his share of the total profit Π . An agent $i = 1, 2$ is said to be judgment proof when his final wealth is lower than the amount of the fine $\alpha_i D$ he has to pay to the regulator. I will further assume that $D > \max\{\mathbb{E}[l_1], \mathbb{E}[l_2]\}$ so that in expectation, none of the agent has enough resources to pay the fine in full.

The timing of the game is the following. First, the regulator announces an ex post regulation policy (D, α) . Second, the agents observe the regulation policy and contract upon profit. The contracting stage is close to Cooper and Ross (1995). Agents play a two-stage game in which they first agree on a binding contract with respect to their respective share of profit in the two possible states of the world ("no accident" and "accident"). In the second stage, taking the terms of the contract as given, they simultaneously choose their effort level e and a .

The design of the private transactions among the two agents has a particular interest in the context of judgment proofness. Here, I assume that agents agree on a contract $(t_N, t_A) \in [0, \Pi]^2$ where t_k is the share of agent 1 in state $k = N, A$ and thus $\Pi - t_k$ is the share of agent 2 in state $k = N, A$. As usual in double moral hazard problems, the sharing of the profit ensures distribution of incentives in the partnership. With judgment-proof agents, however, the sharing of profit plays the additional role of concealing revenue to the regulation authority.

Following the same line as the previous principal-agent model, the potential agents' inability to pay the regulator modify their ex post transfer in case of accident. Let $m_i(t_A)$ denote the transfer of agent $i = 1, 2$ in case of accident. Again, this transfer incorporates both

the private and the public transactions. For uniformly distributed level of liability $l_i \in [0, \bar{l}_i]$, I define

$$m_1(t_A) := \int_{\alpha_1 D - t_A}^{\bar{l}_1} \frac{(t_A - \alpha_1 D)}{\bar{l}_1} dl + \int_0^{\alpha_1 D - t_A} \frac{(-l)}{\bar{l}_1} dl,$$

$$m_2(t_A) := \int_{\alpha_2 D - \Pi + t_A}^{\bar{l}_2} \frac{(\Pi - t_A - \alpha_2 D)}{\bar{l}_2} dl + \int_0^{\alpha_2 D - \Pi + t_A} \frac{(-l)}{\bar{l}_2} dl.$$

Naturally, the payoff of agent 1 is increasing in t_A while the one of agent 2 is decreasing in t_A . Let $M(t_A) := m_1(t_A) + m_2(t_A)$ be the total profit of the partnership when an accident occurs. Notice that it depends upon t_A , that is, the way profit is split up in case of accident. This stems directly from the assumption that agent may be judgment proof: internal distribution of profits matter now.

I solve the perfect Bayesian equilibrium of the contracting game by backward induction. For a given contract $(t_N, t_A) \in [0, \Pi]^2$, agents' utility functions write

$$U_1 = p(e, a)(b + t_N) + [1 - p(e, a)]m_1(t_A) - \psi(e),$$

$$U_2 = p(e, a)(b + \Pi - t_N) + [1 - p(e, a)]m_2(t_A) - C(a).$$

At the second stage of the contracting game, agents simultaneously choose their effort level. Differentiating the utility of each agent with respect to own effort and equating to zero gives the two incentive constraints:

$$b + t_N - m_1(t_A) = \gamma_1 e, \quad (12)$$

$$b + \Pi - t_N - m_2(t_A) = \gamma_2 a. \quad (13)$$

Given that functions $m_i(\cdot)$ are monotonic, the equilibrium effort levels e and a are uniquely defined by a contract (t_N, t_A) in the subgame. For simplicity, I ignore both agents' participation constraints. This approach would therefore fit with a situation in which agents are already engaged in production and cannot decide to quit ex ante like for instance if agents run an established nuclear power plant that cannot be stopped overnight.¹⁵

Therefore, assume agents choose a sharing of profit and effort levels to maximize joint profits as follows

$$\max_{\{e, a, t_N, t_A\}} (e + a)(2b + \Pi) + (1 - (e + a))M(t_A) - \psi(e) - C(a)$$

subject to constraints (12), (13) and $(t_N, t_A) \in [0, \Pi]^2$.

Notice that t_A enters directly into the objective function of the firm. Usually, without judgment-proofness, profit sharing only serves as a way to distribute incentives within the

¹⁵In the one-sided moral hazard case, the agent can be seen as an employee or an independent contractor who is in charge of take precautionary care on behalf of the principal. In that case, it is crucial to take into account his participation constraint as he may simply refuse to take part in a risky activity.

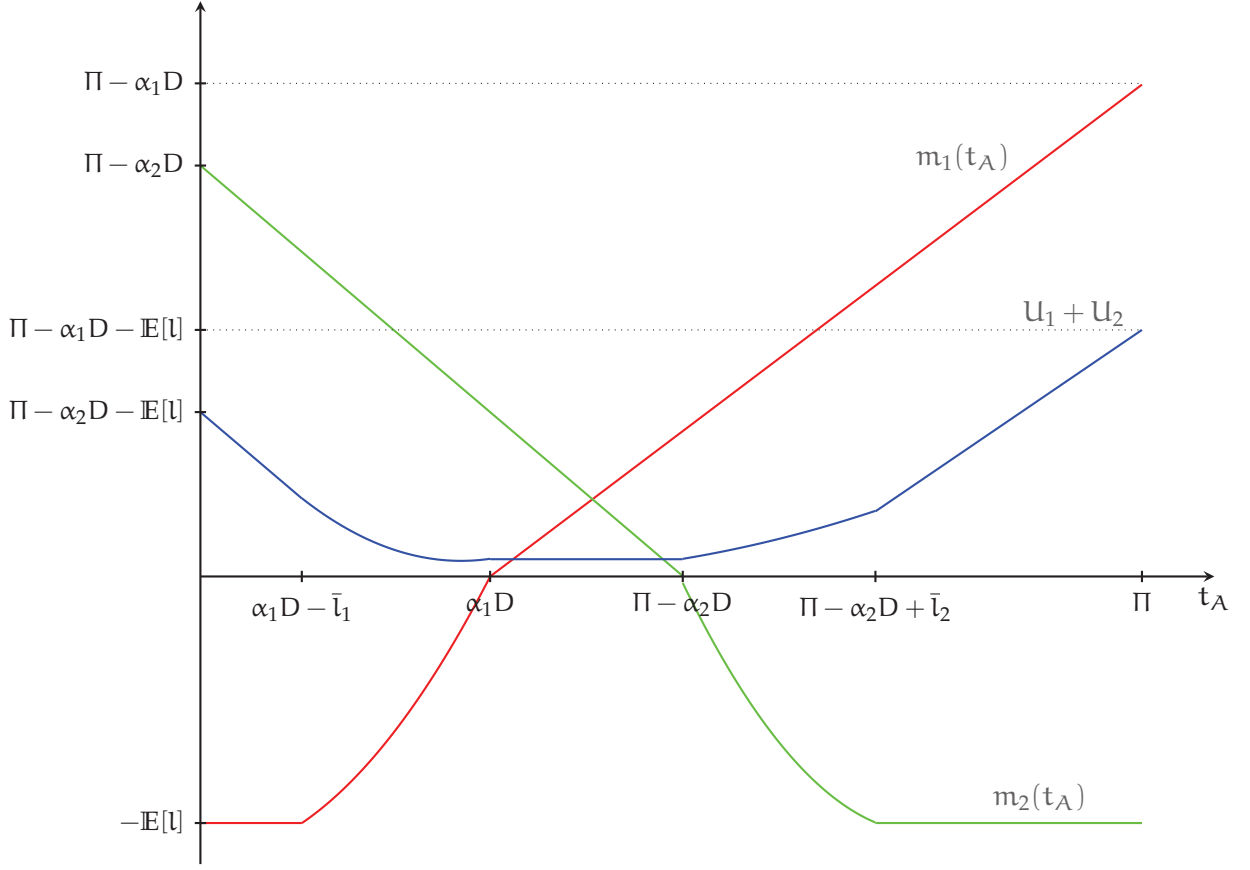


Figure 1: Total profits in case of accident, where $\bar{l}_1 = \bar{l}_2$ and $\alpha_1 \leq 0.5$

firm and affects joint profits only indirectly through changes in equilibrium effort level. In the judgment-proof case, profit sharing also directly affects profits in case of accident as shifting monetary resources from one agent to the other also serves as a way of concealing profits to the regulator. This can easily be seen on Figure 1. For intermediate values of α_1 , the total profit of the firms in case of accident is a U-shape function of t_A . The minimal total profit in case of accident is attained for intermediate values of t_A and is exactly $\Pi - D$ for $t_A \in [\alpha_1 D, \Pi - D + \alpha_1 D]$. In other words, for intermediate values of t_A , the firm pays the whole fine whereas it can increase total profit by shifting resources more extremely to one or another agent. Consider first the case of an interior solution in the sense $(t_N, t_A) \in (0, \Pi)^2$. Plugging (12) into (13), the first-order conditions with respect to e , a and t_A write:

$$2b + \Pi - M(t_A) - \gamma_1 e - \lambda \gamma_1 = 0 \quad (14)$$

$$2b + \Pi - M(t_A) - \gamma_2 a - \lambda \gamma_2 = 0 \quad (15)$$

$$(1 - (e + a) - \lambda)M'(t_A) = 0. \quad (16)$$

where λ is the Lagrange multiplier associated with (13). Analyzing those first-order conditions (details in the appendix) shows that there exists a local maximum for which $M'(t_A^I) = 0$ so

that $t_A^I \in [\alpha_1 D, \Pi - D + \alpha_1 D]$ and

$$e^I = \frac{\gamma_2(2b + D)}{\gamma_1(\gamma_1 + \gamma_2)}$$

$$a^I = \frac{\gamma_1(2b + D)}{\gamma_2(\gamma_1 + \gamma_2)}$$

Notice that only t_A^I depends upon α_1 whereas e^I and a^I only depend upon D and marginal costs of efforts. More importantly, as the solution requires $M'(t_A^I) = 0$, it means that total profits of the firm are at their lowest possible value, namely $M(t_A^I) = \Pi - D$. Therefore, the firm faces the whole amount of fines and agents have to provide quite high effort levels.¹⁶

Intuition suggests that this candidate is, in some cases, the worst possible scenario for the firm as it has to pay the fine in full and exert effort levels accordingly. As mentioned above, this candidate to the maximization problem is only a local maximum and it may not be a global one.

This is indeed the case when the regulator heavily targets one agent. Let us assume that α_1 is close to zero, that is, agent 1 is almost not targeted by the regulation while agent 2 faces almost the whole fine D . Then the following may arise.¹⁷

Proposition 6 *When the regulation heavily targets agent 1 (resp. agent 2), the equilibrium contract may exhibit an extreme sharing out of the profit, namely, $t_N = t_A = \Pi$ (resp. $t_N = t_A = 0$).*

That is, when an accident occurs, the two-agent partnership secures profit by giving it to the least targeted agent. Notice that when $t_A = \Pi$ and α_1 , then $M(t_A) = \Pi - \alpha_1 D - \mathbb{E}[l_2] > \Pi - D$ if α_1 is small enough as $D > \mathbb{E}[l_2]$. As the firm is able to secure profits in case of accident by moving resources to the least targeted agent, it will also exert a lower total effort level as the threat of an accident is also lower. In the case where $t_N = t_A = \Pi$, equilibrium effort levels are given by

$$e = \frac{b + \alpha_1 D}{\gamma_1},$$

$$a = \frac{b + \mathbb{E}[l_2]}{\gamma_2}.$$

The interpretation is as follows. Each agent has natural incentives to exert effort as $b > 0$ is obtained only when no accident occurs. Agent 1 has all the contractible profit Π whether an accident occurs or not so that an increase in α_1 gives him additional incentives to exert effort. Agent 2, however, receives nothing whether an accident occurs or not so that his incentives to exert effort are unchanged with respect to α_1 . However, agent 2 faces a large share of the fine (as α_1 is small) and therefore expects to pay $\mathbb{E}[l_2]$ if an accident occurs.

¹⁶As a way of comparison, the first-best levels of effort solve $\max_{e,a} (e + a)(2b + \Pi) + (1 - (e + a))(\Pi - D) - \psi(e) - C(a)$ and thus they write $e^{FB} = (2b + D)/\gamma_1$ and $a^{FB} = (2b + D)/\gamma_2$. Therefore $e^I = \frac{\gamma_2}{\gamma_1 + \gamma_2} e^{FB} < e^{FB}$ and $a^I = \frac{\gamma_1}{\gamma_1 + \gamma_2} a^{FB} < a^{FB}$. When $\gamma_1 = \gamma_2$, the total level of effort $e^I + a^I$ is exactly twice as less as the first-best total of effort $e^{FB} + a^{FB}$.

¹⁷Whether the interior equilibrium is a global maximum crucially depends on the size of D . For low D , the interior is a global maximum, whereas it is always dominated by extreme sharing when D becomes higher.

For larger α_1 , the extreme equilibrium $t_N = t_A = \Pi$ may not hold anymore. In that case, we have $t_A = \Pi$ but $t_N < \Pi$ so that $e = \frac{\gamma_2(2b+\alpha_1 D+E[l_2])}{\gamma_1(\gamma_1+\gamma_2)}$ and $a = \frac{\gamma_1(2b+\alpha_1 D+E[l_2])}{\gamma_2(\gamma_1+\gamma_2)}$. By symmetry, when α_1 is close to 1, we have $t_N = t_A = 0$, $e = \frac{b+E[l_1]}{\gamma_1}$ and $a = \frac{b+(1-\alpha_1)D}{\gamma_2}$ and for lower α_1 , $t_A = 0$, $t_N > 0$, $e = \frac{\gamma_2(2b+(1-\alpha_1)D+E[l_1])}{\gamma_1(\gamma_1+\gamma_2)}$ and $a = \frac{\gamma_1(2b+(1-\alpha_1)D+E[l_1])}{\gamma_2(\gamma_1+\gamma_2)}$.

When agents both share the same marginal cost of effort, that is, $\gamma_1 = \gamma_2$, the optimal regulation takes a very simple form.

Proposition 7 *When $\gamma_1 = \gamma_2$, the optimal regulation (minimize probability of accident) requires:*

$$\alpha_1^* = \frac{1}{2} + \frac{E[l_1] - E[l_2]}{2D}.$$

Thus, the optimal regulation policy is centered around 1/2 and targets more the agent with more liability. More importantly, this optimal regulation is unique and thus the way it is designed matter for incentivizing agents to exert effort.

Two important things stems from Proposition 7. First, the Equivalence Principle fails to apply due to the judgment-proofness possibility. Agents anticipate the regulation and write contracts accordingly. Second, in a double moral hazard setting, the optimal regulation allocates the total fines evenly among injurers. This is in sharp contrast with the optimal regulation in the case of one-sided moral hazard in which targeting the principal only was optimal.

5. CONCLUSION

In this paper, I have proposed a theoretical foundation for the optimal distribution of penalties among several potential injurers. Assuming that injurers have limited financial resources and therefore sometimes declared judgment proof, I show that the usual Equivalence Principle does not hold anymore. On the contrary, both in the principal-agent firm and in the two-agent partnership firm, the optimal regulation distributes fines toward the injurers with available cash resources. On their side, firms anticipate the regulation and try to prevent paying the fines as much as possible. This requires the optimal contract to solve a trade-off between allocating incentives (to avoid the accident) and sharing profits in case of accident to avoid paying penalties. My result stems from relaxing the modeling assumption that the individuals must contract ex ante to avoid ex post insolvency. Instead, I assume that there is no need to contract ex ante on that matter as insolvency simply implies not paying the fines.

REFERENCES

- [1] Balkenborg, Dieter. 2001. "How liable should a lender be? The case of judgment-proof firms and environmental risk: Comment." *American Economic Review* 91(3):731–738.
- [2] Boyd, James and Daniel E Ingberman. 1997. "The search for deep pockets: Is "extended liability" expensive liability?" *The Journal of Law, Economics, and Organization* 13(1):232–258.

- [3] Boyer, Marcel and Jean-Jacques Laffont. 1997. "Environmental risks and bank liability." *European Economic Review* 41(8):1427–1459.
- [4] Che, Yeon-Koo and Kathryn E Spier. 2008. "Strategic judgment proofing." *The RAND Journal of Economics* 39(4):926–948.
- [5] Cooper, Russell and Thomas W Ross. 1985. "Product warranties and double moral hazard." *The RAND Journal of Economics* pp. 103–113.
- [6] Feess, Eberhard and Ulrich Hege. 1998. "Efficient liability rules for multi-party accidents with moral hazard." *Journal of Institutional and Theoretical Economics (JITE)/Zeitschrift für die gesamte Staatswissenschaft* pp. 422–450.
- [7] Ganuza, Juan José and Fernando Gomez. 2008. "Realistic standards: optimal negligence with limited liability." *The Journal of Legal Studies* 37(2):577–594.
- [8] Hiriart, Yolande and David Martimort. 2006. "The benefits of extended liability." *The RAND Journal of Economics* 37(3):562–582.
- [9] Newman, Harry A and David W Wright. 1990. "Strict liability in a principal-agent model."
- [10] Pitchford, Rohan. 1995. "How liable should a lender be? The case of judgment-proof firms and environmental risk." *The American Economic Review* pp. 1171–1186.
- [11] Ringleb, Al H and Steven N Wiggins. 1990. "Liability and large-scale, long-term hazards." *Journal of Political Economy* 98(3):574–595.
- [12] Segerson, Kathleen and Tom Tietenberg. 1992. "The structure of penalties in environmental enforcement: an economic analysis." *Journal of Environmental Economics and Management* 23(2):179–200.
- [13] Shavell, Steven. 1986. "The judgment proof problem." *International review of law and economics* 6(1):45–58.
- [14] Sykes, Alan O. 1984. "The economics of vicarious liability." *The Yale Law Journal* 93(7):1231–1280.

APPENDIX

Proof of Proposition 2. Let λ and μ be the Lagrange multipliers associated with the participation constraint $U \geq 0$ and the constraint $U \geq R(e) - \mathbb{E}[U]$, respectively. First-order conditions write

$$\begin{aligned}\mu e\psi''(e) + \psi'(e) &= c_R(\alpha, F) + (1 - \alpha)F \\ \lambda + \mu &= 1.\end{aligned}$$

It is clear that $\lambda = \mu = 0$ is impossible. Consider first that $\lambda = 0$, so that $\mu = 1$. It follows that the equilibrium effort level e_1 is given by

$$e_1\psi''(e_1) + \psi'(e_1) = c_R(\alpha, F) + (1 - \alpha)F,$$

which is equation (9). Let $e_1(\alpha)$ be the implicit solution to this equation. Notice that the LHS is strictly increasing in e and the RHS is strictly decreasing in α so that $e_1(\alpha)$ is decreasing in α . As $\mu = 1$, the second constraint is binding so that $U = R(e_1(\alpha)) - \mathbb{E}[l]$. This holds as long as the agent's participation constraint is also satisfied, that is $U = R(e_1(\alpha)) - \mathbb{E}[l] \geq 0$. As $e_1(\alpha)$ decreases in α , there exists a threshold $\hat{\alpha}_1$ such that $U = R(e_1(\hat{\alpha}_1)) - \mathbb{E}[l] = 0$. The principal sets $t_A = \alpha F - \hat{l}$ to extract as much rent as possible from the agent.

Consider now the case where both $\mu > 0$ and $\lambda > 0$. Complementary slackness implies that $U = R(e) - \mathbb{E}[l] = 0$. For this case, the effort level $e_2 = e_1(\hat{\alpha}_1)$ is constant with respect to α . Notice that the first-order condition with respect to e gives that $\mu = \frac{c_R(\alpha, F) + (1-\alpha)F - \psi'(e_2)}{e_2 \psi''(e_2)}$. It is easy to see that for $\alpha = \hat{\alpha}_1$ we have $c_R(\hat{\alpha}_1, F) + (1-\alpha)F = e_2 \psi''(e_2) + \psi'(e_2) > \psi'(e_2)$ so that $\mu > 0$. This solution is then valid until α reaches the second threshold $\hat{\alpha}_2$ defined by $c_R(\hat{\alpha}_2, F) + (1-\alpha)F = \psi'(e_2)$.

Finally, when $\mu = 0$ it implies $\lambda = 1$ and $U = 0$. The equilibrium effort level is defined by $\psi'(e_3) = c_R(\alpha, F) + (1-\alpha)F$ and is decreasing in α . Recall that $U = R(e) + c_P(t_A, \alpha, F) - c_R(\alpha, F)$ so that $U = 0$ requires that $c_P(t_A, \alpha, F) - c_R(\alpha, F) > -\mathbb{E}[l]$ in that situation as $R(e_3(\alpha)) < R(e_2) = \mathbb{E}[l]$ for all $\alpha > \hat{\alpha}_2$. ■

5.1. Interior solution for double moral hazard problem.

Equation (16) directly rewrites $M'(\underline{t}) [1 - (e + a) - \lambda] = 0$.

• Let us assume that $M'(\underline{t}) \neq 0$. Thus, we must have $1 - (e + a) - \lambda = 0$. The bordered Hessian of the problem writes:

$$H = \begin{bmatrix} 0 & -\gamma_1 & -\gamma_2 & -M'(\underline{t}) \\ -\gamma_1 & -\gamma_1 & 0 & -M'(\underline{t}) \\ -\gamma_2 & 0 & -\gamma_2 & -M'(\underline{t}) \\ -M'(\underline{t}) & -M'(\underline{t}) & -M'(\underline{t}) & 0 \end{bmatrix}$$

Using Sydsæte, Hammond & Seierstad (2005) we have to compute two determinants:

$$B_2 := \begin{vmatrix} 0 & -\gamma_1 & -\gamma_2 \\ -\gamma_1 & -\gamma_1 & 0 \\ -\gamma_2 & 0 & -\gamma_2 \end{vmatrix} = \gamma_1 \gamma_2 (\gamma_1 + \gamma_2) > 0$$

$$B_3 := \begin{vmatrix} 0 & -\gamma_1 & -\gamma_2 & -M'(\underline{t}) \\ -\gamma_1 & -\gamma_1 & 0 & -M'(\underline{t}) \\ -\gamma_2 & 0 & -\gamma_2 & -M'(\underline{t}) \\ -M'(\underline{t}) & -M'(\underline{t}) & -M'(\underline{t}) & 0 \end{vmatrix} = [M'(\underline{t})]^2 [\gamma_1^2 + \gamma_1 \gamma_2 + \gamma_2^2] > 0$$

These determinants are neither both negative (local min) nor of alternate sign (local max). Thus, this stationary point is a saddle point.

• Let us assume that $M'(\underline{t}) = 0$. Solving $M'(\underline{t}) = 0$ for $\underline{t}$ we get:

$$\int_{\alpha_1 D - \underline{t}}^{\bar{l}_1} dH(l) - \int_{\alpha_2 D - \Pi + \underline{t}}^{\bar{l}_2} dH(l) = 0$$

Then, the solutions t^F lie in the interval $[\alpha_1 D, \Pi - D + \alpha_1 D]$. We simply obtain that ¹⁸

$$M(\underline{t}^F) = \Pi - D.$$

In other words, if this were a solution to the partnership problem, the agents would pay all the damages without trying to escape and would exert maximal levels of efforts.

¹⁸Be careful, here other $\underline{t}$ might solve $M'(\underline{t}) = 0$. Check what happens.

Equation (14) and (15) write

$$2b + D - \lambda\gamma_1 = \gamma_1 e,$$

$$2b + D - \lambda\gamma_2 = \gamma_2 a.$$

Summing this two constraints and plugging this into the constraint gives:

$$\lambda = \frac{2b + D}{(\gamma_1 + \gamma_2)},$$

and then

$$e = \frac{\gamma_2(2b + D)}{\gamma_1(\gamma_1 + \gamma_2)}$$

$$a = \frac{\gamma_1((2b + D))}{\gamma_2(\gamma_1 + \gamma_2)}$$

The partnership's total revenue at this equilibrium writes:

$$V^F = \Pi - D + \frac{\gamma_1^2 + \gamma_1\gamma_2 + \gamma_2^2}{2\gamma_1\gamma_2(\gamma_1 + \gamma_2)}(2b + D)D^2.$$

Let L^* be the Hessian of the Lagrangian at the candidate solution:

$$L^* := \begin{bmatrix} -\gamma_1 & 0 & 0 \\ 0 & -\gamma_2 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

From Luenberger & Ye (2008, p. 335) we have a local maximum if $y^T L^* y < 0$ for all $y \in M = \{y = (y_1, y_2, y_3) : (-\gamma_1, -\gamma_2, 0)y = 0\}$. Here, we have

$$y^T L^* y = -\gamma_1 y_1^2 - \gamma_2 y_2^2 < 0.$$

Then we have a local maximum.

RECENT PUBLICATIONS BY *CEIS Tor Vergata*

Flexibility Versus Security in Agency Contracts with Moral Hazard

Lorenzo Bozzoli and Guillaume Pommey

CEIS Research Paper, 621 June 2026

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

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