



UNIVERSITÀ
DEGLI STUDI
FIRENZE

DISEI

DIPARTIMENTO DI SCIENZE
PER L'ECONOMIA E L'IMPRESA

WORKING PAPERS
QUANTITATIVE METHODS FOR SOCIAL SCIENCES

Asymmetric Interdependent Risk: Bargaining over an Insurance Externality

GIANLUCA IANNUCCI, JEAN-CHRISTOPHE PEREAU

WORKING PAPER N. 03/2026

*DISEI, Università degli Studi di Firenze
Via delle Pandette 9, 50127 Firenze (Italia) www.disei.unifi.it*

*The findings, interpretations, and conclusions expressed in the working paper series
are those of the authors alone. They do not represent the view of Dipartimento di
Scienze per l'Economia e l'Impresa*

Asymmetric Interdependent Risk: Bargaining over an Insurance Externality

Gianluca Iannucci[†] Jean-Christophe Perea[‡]

September 21, 2026

Abstract

This paper studies how private insurance affects negotiations between two firms over risk and safety. When an upstream firm buys insurance, moral hazard reduces its safety effort, creating a risk externality for a downstream firm. This outcome sets an endogenous threat point for bargaining. We evaluate three negotiation structures: a Coasean baseline, a Rigid protocol, and a Sequential protocol where insurance is chosen strategically. We show that the upstream firm intentionally over-insures to increase its bargaining leverage and extract larger payments. Counterintuitively, higher insurance prices curb this opportunistic behavior, making voluntary agreements easier to reach. Finally, even when insurance becomes too costly for risk management, a purely strategic demand for coverage persists to gain power in negotiations.

Keywords: Insurance externality, Interdependent risk, Bargaining protocols.

JEL Classifications: C78, D62, G22.

[†]Department of Economics and Management, University of Florence, Italy. E-mail: gianluca.iannucci@unifi.it.

[‡]Bordeaux School of Economics, INRAE CNRS UMR 6060, University of Bordeaux, France. E-mail: jean-christophe.pereau@u-bordeaux.fr.

1 Introduction

Consider two economic entities linked by an asymmetric risk externality, where investments by a primary (upstream) firm directly influence the vulnerability of an interconnected (downstream) counterpart (Ambec and Sprumont, 2002). The upstream firm can invest in self-protection technology—such as industrial safety equipment or cyber-resilience infrastructure—that mitigates its own risk exposure and, as a natural side effect, protects interconnected downstream operations, thereby reducing the latter’s exposure as well. In the absence of formal insurance, firm decision-making over prevention depends on private risk preferences and prudence (Ehrlich and Becker, 1972; Courbage and Loubergé, 2006). However, because individual prevention incentives are inherently distorted in environments characterized by interdependent risks (Courbage et al., 2018), the upstream firm ignores this beneficial external impact and under-invests relative to the social optimum (Lohse et al., 2012). Nevertheless, the spillover remains positive: downstream benefits unconditionally from whatever self-protection upstream chooses to maintain.

This intuitive dynamic changes fundamentally once commercial insurance becomes available. Acting rationally under risk and corporate risk management principles (Arrow, 1965; Mossin, 1968; MacMinn, 1987; Seog, 2006), and evaluating risk exposure through mean-risk or two-moment decision criteria (Meyer, 1987), the upstream firm purchases coverage to trade off expected wealth against risk exposure. Because the insurance premium is priced parametrically on private expected loss, coverage acts as a direct substitute for costly self-protection (Chiu, 2000). Once insured, the upstream firm’s incentive to maintain its self-protection effort weakens. Importantly, the resulting moral hazard does not remain confined to the policyholder; it directly erodes the spillover previously enjoyed downstream. From the downstream firm’s perspective, the upstream firm’s insurance purchase acts as an uncompensated negative risk shift, an insurance externality (Parchomovsky and Siegelman, 2022) caused entirely by a private contract to which it was never a party.

Faced with this spillover erosion—a clear manifestation of technical interdependence where an entity’s vulnerability is directly shaped by the precaution of surrounding actors (Avery and Zemsky, 1995; Heal and Kunreuther, 2007)—the nat-

ural response for the downstream firm is to seek insurance coverage itself to hedge its increased background risk (Doherty and Schlesinger, 1983; Eeckhoudt et al., 2006). While both firms purchasing coverage forms a self-enforcing Nash equilibrium of the uncoordinated insurance subgame, this market outcome merely buffers risk exposure; it does not repair the lost operational protection. Significant gains from trade remain unexploited, creating a strong rationale for direct bilateral negotiation.

Can direct bargaining allow the two firms to recover these unexploited gains? Standard Coasean logic suggests that friction-free negotiation should restore efficiency (Coase, 1960). In an interdependent risk setting, however, the availability of private insurance reshapes the baseline of negotiation: the disagreement point (Binmore et al., 1986; Crocker and Masten, 1991).

To evaluate how institutional design and contractibility shape this strategic negotiation, we analyze three distinct bargaining protocols built on this endogenous threat point. Under a first-best Coasean benchmark, prevention efforts and monetary transfers are fully contractible in a single stage, rendering private insurance obsolete under agreement. Under a Rigid protocol, prevention effort is non-contractible, but formal insurance coverage is verifiable; here, the downstream firm pays a transfer to buy back the upstream firm’s coverage, restoring baseline self-protection while retaining its own optimal insurance. Finally, under a Sequential protocol, coverages are chosen strategically in a first stage to establish threat points before bargaining over effort and transfers in a second stage, allowing the upstream firm to deliberately over-insure to extract larger bargaining transfers.

This framework reveals three key insights into the interplay between private risk markets and negotiated externalities. Our analysis demonstrates that bargaining under endogenous threat points primarily serves as a mechanism for surplus redistribution rather than pure surplus creation, as parties leverage strategic coverage decisions to claim value. Furthermore, strategic over-insurance under the Sequential protocol generates hold-up positions that can cause negotiations to stall even when positive gains from trade exist. Finally, insurance market conditions strongly influence bargaining viability: as the market hardens—such as through higher loading factors driven by reinsurance capacity constraints or market volatility—strategic over-insurance becomes increasingly costly, which counterintuitively

eases hold-up positions and facilitates voluntary agreement.

By unifying these mechanisms, this paper bridges two literature streams that have largely developed in isolation: the economics of interdependent security, risk spillovers, and insurance demand (Avery and Zemsky, 1995; Kunreuther and Heal, 2003; Heal and Kunreuther, 2007; Muermann and Kunreuther, 2008; Courbage et al., 2018; Hofmann and Ramaj, 2019; Seog, 2024) and the theory of Coasean bargaining over externalities (Deryugina and Requate, 2021). While existing studies on third-party insurance effects (Parchomovsky and Siegelman, 2022) primarily provide qualitative overviews or focus on ex-post claiming behavior, we construct a formal game-theoretic model of bilateral ex-ante negotiation under endogenous disagreement points.

The remainder of the paper is organized as follows. Section 2 sets up the theoretical framework and establishes the baseline non-cooperative equilibrium without insurance. Section 3 analyzes the upstream firm’s demand for insurance and the resulting moral hazard that erodes the physical spillover. Section 4 examines the uncoordinated subgame where both firms purchase coverage, establishing the unique self-enforcing disagreement point. Section 5 formalizes the three bilateral bargaining protocols (Coase, Rigid, and Sequential). Section 6 compares equilibrium outcomes, protocol selection, and voluntary agreement viability across insurance market conditions. Section 7 concludes.

2 The model

Let us assume two firms, i and j , endowed with initial wealth w subject to a loss function L . The loss consists of two parts: a common random shock $\tilde{z}$ and a non-random component depending on the effort e exerted to reduce the loss. To capture how much firm i ’s prevention spills over onto firm j , we introduce in the loss base of j an interdependence parameter $\gamma \in (0, 1]$. The deterministic loss base for the two agents can be written as:

$$\begin{aligned}\ell_i &= (x - e_i)\phi, \\ \ell_j &= (x - \gamma e_i - e_j)\phi,\end{aligned}\tag{1}$$

where $x \in [\underline{x}, \bar{x}]$, with $\underline{x} > 0$ and $\bar{x} < w$, represents the maximum value of the loss base if no effort is exerted, and $\phi > 0$ is an impact parameter. In (1), the risk is interdependent as the loss base of agent j is a function of both its own effort and the effort of agent i .

The total risk is, therefore, $L_i = \tilde{z}\ell_i$ and $L_j = \tilde{z}\ell_j$. We assume the random loss $\tilde{z}$ follows a probability distribution characterized by a strictly positive support and a right-skewed shape. This general framework effectively captures the common characteristics of losses arising from environmental risks, which are inherently positive and feature extreme right-tail asymmetry. Both firms manage the risk according to the same tool, the Expected Shortfall (ES). Due to the asymmetry of the loss distribution, an adequate instrument is one which calculates the mean of the losses that exceed a given threshold α , usually the 95%. For a generic continuous distribution, the ES can be written as $\theta = ES_\alpha = \frac{1}{1-\alpha} \int_\alpha^1 F^{-1}(u) du$, where F^{-1} is the quantile function (i.e., the inverse cumulative distribution function, CDF) of the random loss $\tilde{z}$.

The firms' problem is to maximize the following expected utilities:

$$\begin{aligned} \max_{e_i \in (0, x)} V_i &= w - \frac{1}{2}e_i^2 - (x - e_i)\phi\theta, \\ \max_{e_j \in (0, x - \gamma e_i)} V_j &= w - \frac{1}{2}e_j^2 - (x - \gamma e_i - e_j)\phi\theta. \end{aligned} \tag{2}$$

The first-order conditions (FOCs) of problem (2) are:

$$e_i^{NC} = e_j^{NC} = \phi\theta. \tag{3}$$

Thus, the optimal effort to reduce risk is the same for both agents. Note that the values of e_i^{NC} and e_j^{NC} lie within the interval $[0, x]$ only if $x - 2\phi\theta > 0$. The residual losses of firm i and j are strictly positive at equilibrium if and only if $x - \phi\theta > 0$ and $x - (1 + \gamma)\phi\theta > 0$. We obtain

$$V_i^{NC} = w - \left(x - \frac{1}{2}\theta\phi\right)\theta\phi, \tag{4}$$

$$V_j^{NC} = w - \left(x - \left(\frac{1}{2} + \gamma\right)\theta\phi\right)\theta\phi. \tag{5}$$

An implication of (3) is:

$$V_i^{NC} - V_j^{NC} = -\gamma\phi^2\theta^2 < 0. \quad (6)$$

Firm j is better off than firm i since j benefits from the effort realized by the latter. However, the question for firm i is: how can it improve its situation?

3 Insurance for firm i

Assume now that firm i decides to take out an insurance policy to cover the risk. The game becomes sequential. In the first stage, firm i chooses the level of insurance coverage denoted by β_i , while in the second stage, the effort levels e_i and e_j are determined. The mean-risk objective functions become:

$$\begin{aligned} V_i^I &= w - \frac{1}{2}(e_i)^2 - (1 - \beta_i)(x - e_i)\phi\theta, \\ V_j &= w - \frac{1}{2}e_j^2 - (x - \gamma e_i - e_j)\phi\theta. \end{aligned} \quad (7)$$

where $\beta_i \in [0, 1]$ is the coinsurance rate and $1 - \beta_i$ the retention rate. Note that if $\beta_i = 0$, firm i prefers not to be insured, while $\beta_i = 1$ represents full insurance. Maximizing (7) with respect to e_i and e_j yields:

$$\begin{aligned} e_i^{I*} &= (1 - \beta_i)\phi\theta, \\ e_j^* &= \phi\theta, \end{aligned} \quad (8)$$

where the superscript I denotes the case where firm i is insured but not firm j . As expected, if $\beta_i = 0$, then e_i^{I*} in (8) coincides with e_i^{NC} in (3). Note that the effort of firm j is the same as in the non-cooperative case, $e_j^* = e_j^{NC}$. Moreover, (8) clearly illustrates the moral hazard effect of insurance coverage as an increase in β_i reduces effort. Consequently, with insurance:

$$e_i^{I*} < e_i^{NC},$$

In the first stage, firm i chooses the insurance coverage, which defines the demand for insurance. We assume that the premium is computed according to the standard deviation principle:

$$P_i = (\mu + \lambda\sigma)(x - e_i)\phi\beta_i, \quad (9)$$

where $\lambda > 0$ is the safety loading. Throughout the paper, the insurance market is competitive. The loading λ is exogenous and parametric, reflecting the insurer's price-taking position in a competitive market; correct anticipation of the realised loss base pins down the actuarially fair component of the premium, $\mu(x - e_i)\phi\beta_i$, while λ captures the (exogenous) cost of risk-bearing above that. The insurance premium is proportional to the coverage β_i and decreases with effort. This implies that with insurance, firm i 's effort is lower. Firm i 's first-stage problem is:

$$\begin{aligned} \max_{\beta_i \in [0,1]} H_i &= V_i - P_i = \\ &= w - \frac{1}{2}e_i^2 - [(1 - \beta_i)\theta + (\mu + \lambda\sigma)\beta_i](x - e_i)\phi. \end{aligned} \quad (10)$$

Problem (10) yields the optimal demand for insurance for firm i with notations $D = 2(\mu + \lambda\sigma) - \theta$ and $m = \theta - \mu - \lambda\sigma$

$$\beta_i^{I*} = \frac{(x - \phi\theta)m}{D\phi\theta}. \quad (11)$$

It gives

$$e_i^{I*} = \frac{\theta\phi(\mu + \lambda\sigma) - xm}{D}. \quad (12)$$

The second-order condition (SOC) is:

$$\frac{\partial^2 H_i}{\partial \beta_i^2} = -D\theta\phi^2,$$

which is always negative for:

$$\lambda > \hat{\lambda} \equiv \frac{\theta - 2\mu}{2\sigma}.$$

We can see that $D > 0$ for $\lambda > \hat{\lambda}$. This ensures that the denominator of β_i^{I*} is always positive. Note that $\beta_i^{I*} > 0$ for:

$$\lambda < \bar{\lambda} \equiv \frac{\theta - \mu}{\sigma}, \quad (13)$$

and $\beta_i^{I*} < 1$ for:

$$\lambda > \underline{\lambda} \equiv \frac{(\theta - \mu)x - \mu\phi\theta}{(x + \phi\theta)\sigma}. \quad (14)$$

Finally, since

$$\hat{\lambda} - \underline{\lambda} = -\frac{(x - \phi\theta)\theta}{2(x + \phi\theta)\sigma} < 0,$$

then $\lambda \in (\underline{\lambda}, \bar{\lambda})$ is a necessary and sufficient condition to obtain $\beta_i^{I*} \in (0, 1)$, i.e., partial insurance coverage. Notice that the interval $[\underline{\lambda}, \bar{\lambda}]$ is never empty.

To show that insurance increases the welfare of agent i , we first need to determine the expression of H_i^{I*} . Since $H_i(\beta_i)$ in (10), with e_i substituted by its stage-2 best response (8), is an exact quadratic in β_i , its second-order Taylor expansion around $\beta_i = 0$ is exact:

$$H_i(\beta_i) = H_i(0) + H_i'(0) \beta_i + \frac{1}{2} H_i''(0) \beta_i^2.$$

Noting that $H_i(0) = V_i^{NC}$ (no insurance reduces to the baseline problem of (2)) and that $H_i''(0) = -D\theta\phi^2$ (the SOC), and using the FOC $H_i'(\beta_i^{I*}) = 0$, i.e. $H_i'(0) = D\theta\phi^2\beta_i^{I*}$, evaluation at $\beta_i = \beta_i^{I*}$ gives

$$H_i^{I*} = V_i^{NC} + \frac{1}{2} D\theta\phi^2 (\beta_i^{I*})^2 = V_i^{NC} + \frac{(x - \phi\theta)^2 m^2}{2D\theta}. \quad (15)$$

By comparing its payoff with and without insurance, it holds

$$H_i^{I*} - V_i^{NC} = \frac{(x - \phi\theta)^2 (\theta - \mu - \lambda\sigma)^2}{2[2(\mu + \lambda\sigma) - \theta]\theta} = \frac{(x - \phi\theta)^2 m^2}{2D\theta} > 0. \quad (16)$$

Thus, taking out insurance is a rational choice for firm i . Without insurance the final wealth of i is lower than the final wealth of firm j (see (6)), but what if i is insured? We denote by V_j^* the final wealth of agent j when only agent i is insured:

$$V_j^*(\beta_i^{I*}) = w - x\phi\theta + \left(\frac{1}{2} + \gamma\right) \phi^2 \theta^2 - \frac{\gamma(x - \phi\theta)\phi\theta m}{D}. \quad (17)$$

It holds

$$\begin{aligned}\frac{\partial H_i^{I^*}}{\partial \lambda} &= -\frac{m(\mu + \lambda\sigma)(x - \phi\theta)^2\sigma}{D^2\theta} < 0, \\ \frac{\partial V_j^*}{\partial \lambda} &= \frac{\gamma(x - \phi\theta)\sigma\phi\theta^2}{D^2} > 0,\end{aligned}$$

and so $H_i^{I^*}$ and V_j^* , in the interval $[\underline{\lambda}, \bar{\lambda}]$, have at most one intersection point. Moreover, we have

$$\begin{aligned}H_i^{I^*}(\underline{\lambda}) &= w - \frac{\phi\theta x^2}{x + \phi\theta}, \\ V_j^*(\underline{\lambda}) &= w - x\phi\theta + \frac{\phi^2\theta^2}{2}, \\ H_i^{I^*}(\bar{\lambda}) &= w - x\phi\theta + \frac{\phi^2\theta^2}{2}, \\ V_j^*(\bar{\lambda}) &= w - x\phi\theta + \left(\frac{1}{2} + \gamma\right)\phi^2\theta^2.\end{aligned}$$

It is easy to check that $H_i^{I^*}(\underline{\lambda}) > V_j^*(\underline{\lambda})$ and $H_i^{I^*}(\bar{\lambda}) < V_j^*(\bar{\lambda})$, then there exists a $\lambda \equiv \lambda^*$ such that $H_i^{I^*} = V_j^*$. So, for $\lambda \in [\underline{\lambda}, \lambda^*)$ it holds $H_i^{I^*} > V_j^*$, while for $\lambda \in (\lambda^*, \bar{\lambda}]$ it holds $H_i^{I^*} < V_j^*$.

Proposition 1. *If the insurance market is soft (i.e., a relatively low loading factor is in place), the final wealth of the firm i is higher than that of firm j .*

Proposition 1 implies that if the insurance market is favorable, then being insured for i gives a higher wealth than firm j . On the contrary, if the insurance market is hard, then being insured for i gives lower wealth than firm j , but nevertheless greater than without insurance (see (16)).

What happens to agent j ? We have to compare V_j^* with V_j^{NC} , i.e. the final wealth of firm j when only firm i is insured, with its wealth when i was not insured. It holds that

$$V_j^* - V_j^{NC} = -\frac{\gamma(x - \phi\theta)m\phi\theta}{D} < 0. \quad (18)$$

Therefore, if firm i is insured, the final wealth of firm j decreases.

How can firm j react? It has two options: to seek insurance itself, or to offer a transfer to firm i to increase the latter's effort (assuming firm i remains insured).

Which strategy will yield the highest payoff?

4 Both firms are insured

We start analyzing the insurance choice of the firm j . The expected utilities in the last stage are

$$\begin{aligned} V_i^I &= w - \frac{1}{2}e_i^2 - (1 - \beta_i)(x - e_i)\phi\theta, \\ V_j^I &= w - \frac{1}{2}e_j^2 - (1 - \beta_j)(x - \gamma e_i - e_j)\phi\theta, \end{aligned} \quad (19)$$

where $\beta_j \in [0, 1]$ is the insurance coverage of firm j . Maximizing (19) yields $e_i^{I*} = (1 - \beta_i)\phi\theta$ and $e_j^{I*} = (1 - \beta_j)\phi\theta$. Interestingly, $e_i^{I*} = e_j^{I*}$ if $\beta_i = \beta_j$. We can now add the insurance premiums,

$$\begin{aligned} P_i &= (\mu + \lambda\sigma)(x - e_i)\phi\beta_i, \\ P_j &= (\mu + \lambda\sigma)(x - \gamma e_i - e_j)\phi\beta_j. \end{aligned}$$

First-stage problems become

$$\begin{aligned} \max_{\beta_i \in [0, 1]} H_i &= V_i - P_i = w - \frac{1}{2}e_i^2 - [(1 - \beta_i)\theta + (\mu + \lambda\sigma)\beta_i](x - e_i)\phi, \\ \max_{\beta_j \in [0, 1]} H_j &= V_j - P_j = w - \frac{1}{2}e_j^2 - [(1 - \beta_j)\theta + (\mu + \lambda\sigma)\beta_j](x - \gamma e_i - e_j)\phi. \end{aligned} \quad (20)$$

Problem (20) yields the optimal demand for insurance for firm i and j :

$$\begin{aligned} \beta_i^{I*} &= \frac{(x - \phi\theta)m}{D\phi\theta}, \\ \beta_j^{I*} &= \beta_i^{I*} - \frac{\gamma m}{D} (1 - \beta_i^{I*}) = \frac{m[D(x - (1 + \gamma)\phi\theta) + \gamma m(x - \phi\theta)]}{D^2\phi\theta}. \end{aligned} \quad (21)$$

Since firm i 's problem in (20) involves neither β_j nor e_j , firm j 's insurance choice has no effect on firm i : β_i^{I*} in (21) coincides with (11), and e_i^{I*} is given by (12) and H_i^{I*} by (15). Note that $\beta_i^{I*} > 0$ since $x - \phi\theta > 0$ and $\beta_j^{I*} > 0$ since $x - (1 + \gamma)\phi\theta > 0$.

Since

$$1 - \beta_j^{I*} = (1 - \beta_i^{I*}) (D + \gamma m) / D, \quad (22)$$

the corresponding effort can be rewritten as

$$e_j^{I*} = (1 - \beta_j^{I*}) \phi \theta = \frac{D + \gamma m}{D} e_i^{I*}. \quad (23)$$

The SOC's are

$$\frac{\partial^2 H_i^{I*}}{\partial (\beta_i^{I*})^2} = \frac{\partial^2 H_j^{I*}}{\partial (\beta_j^{I*})^2} = -D\theta\phi^2 < 0.$$

Who asks for more insurance? (22) gives the difference immediately:

$$\beta_i^{I*} - \beta_j^{I*} = \frac{\gamma m}{D} (1 - \beta_i^{I*}) = \frac{\gamma m e_i^{I*}}{D\phi\theta} > 0, \quad (24)$$

for $\lambda > \underline{\lambda}$. So firm i , since it has a higher loss base, demands more insurance, and the gap is exactly proportional to i 's retention $1 - \beta_i^{I*}$. We can infer

$$\beta_i^{I*} > \beta_j^{I*} \longrightarrow e_i^{I*} < e_j^{I*}.$$

Payoff for firm j . The same exact-quadratic argument behind (15) applies to $H_j(\beta_j)$, since $H_j''(0) = -D\theta\phi^2$ (the same SOC as above) and, at $\beta_j = 0$, firm j is back to the passive role of Section 3, facing i already insured at β_i^{I*} : $H_j^{I*}(0) = V_j^*(\beta_i^{I*})$ of (17). The vertex formula then gives

$$H_j^{I*} = V_j^*(\beta_i^{I*}) + \frac{1}{2} D\theta\phi^2 (\beta_j^{I*})^2 = V_j^*(\beta_i^{I*}) + \frac{m^2 [D(x - (1 + \gamma)\phi\theta) + \gamma m(x - \phi\theta)]^2}{2D^3\theta}. \quad (25)$$

The partial insurance increases the wealth of the firm j ,

$$H_j^{I*}(\beta_j^{I*}) - H_j^{I*}(0) = \frac{m^2 [D(x - (1 + \gamma)\phi\theta) + \gamma m(x - \phi\theta)]^2}{2D^3\theta} > 0, \quad (26)$$

which is immediate from (25).

Combining (15) and (25) with (22), the comparison collapses to

$$H_i^{I*} - H_j^{I*} = \frac{\gamma\phi e_i^{I*}}{2} [m(\beta_i^{I*} + \beta_j^{I*}) - 2\theta]. \quad (27)$$

The bracket is always strictly negative: $\beta_i^{I*}, \beta_j^{I*} \leq 1$ gives $m(\beta_i^{I*} + \beta_j^{I*}) \leq \max(2m, 0)$, while $2\theta > 2m$ since $\mu + \lambda\sigma > 0$; either way the bracket is < 0 . Since $e_i^{I*} \geq 0$, with equality iff $\beta_i^{I*} = 1$,

$$H_i^{I*} \leq H_j^{I*}, \quad \text{equality iff } \lambda = \underline{\lambda}. \quad (28)$$

Firm j is therefore always (weakly) wealthier than firm i once both insure, equality holding only at the boundary where i insures fully ($e_i^{I*} = 0$). This pins down the boundary values already noted:

$$\begin{aligned} H_i^{I*}(\underline{\lambda}) &= H_j^{I*}(\underline{\lambda}) = w - \frac{\phi\theta x^2}{x + \phi\theta}, \\ H_i^{I*}(\bar{\lambda}) &= w - x\phi\theta + \frac{\phi^2\theta^2}{2} < w - x\phi\theta + \left(\frac{1}{2} + \gamma\right)\phi^2\theta^2 = H_j^{I*}(\bar{\lambda}). \end{aligned}$$

The following lemma shows that the outcome when both firms are insured is the unique equilibrium of the uncoordinated insurance sub-game itself. It will serve as the disagreement point for every bargaining protocol in the next section.

Lemma 1 (Self-enforcing status quo).

- *Starting from the non-cooperative outcome, $(\beta_i, \beta_j) = (0, 0)$ is not a Nash equilibrium of the uncoordinated insurance sub-game. Firm i strictly gains by insuring at β_i^{I*} regardless of β_j , and firm j 's best response to an insured i is to insure at $\beta_j^{I*}(\gamma)$.*
- *Conversely $(\beta_i^{I*}, \beta_j^{I*})$ is a Nash equilibrium: neither firm can gain by unilaterally dropping coverage. Both firms insuring is therefore the unique self-enforcing outcome of the uncoordinated insurance sub-game.*

Firm i strictly gains by insuring at β_i^{I*} regardless of β_j , since its problem (10) does not involve β_j , and firm j 's best response to an insured i is to insure at $\beta_j^{I*}(\gamma)$, since $H_j^{I*}(\gamma) > V_j^*(\beta_i^{I*})$ ((26)). Neither firm can prevent the other from insuring, so neither can hold the threat of “staying at NC” over the other, and there is no self-enforcing configuration in which i is insured while j is not. This is

why the bargaining protocols considered next take the both-insured outcome, not the non-cooperative outcome, as their disagreement point.

5 Bargaining

Instead of acting unilaterally as in the previous sections, firms i and j can engage in direct negotiation. We examine three alternative bargaining protocols that differ in what the two firms can contractually agree upon and how insurance choices interact with the negotiation.

Crucially, in all three protocols, the relevant disagreement point reflects the uncoordinated insurance outcome derived in [Lemma 1](#) (or its strategic equivalent), as neither firm can be credibly held at the non-cooperative baseline once private insurance is available.

The three protocols capture three distinct institutional environments regarding contractibility and coverage renegotiation:

- i) **Coase (First-Best Benchmark):** Effort levels (e_i, e_j) and a monetary transfer T are fully contractible in a single stage. Because effort can be perfectly specified by contract, private insurance loses its risk-reduction and incentive roles; under agreement, both firms drop coverage entirely ($\beta_i = \beta_j = 0$).
- ii) **Rigid Protocol:** Precautionary effort is unobservable or non-contractible, but the formal insurance policy β_i is verifiable. Firm j pays a lump-sum transfer T in exchange for firm i cancelling its coverage ($\beta_i = 0$), thereby restoring i 's effort to its uninsured level e_i^{NC} . Firm j 's coverage remains untouched at its individually optimal level β_j^{I*} .
- iii) **Sequential Protocol:** A two-stage game where both firms non-cooperatively set their insurance coverages (β_i, β_j) in Stage 1, strategically anticipating a Stage 2 Nash bargain over effort (e_i, e_j) and a transfer T . Under agreement, the strategically chosen coverages remain active while effort is contractually adjusted.

These protocols thus reflect a spectrum of contractibility: Coase eliminates all private insurance under agreement, Rigid renegotiates only firm i 's coverage, and Sequential retains both insurance policies as strategic baseline positions.

5.1 Coase bargaining: a first-best benchmark

Insurance plays no role in this benchmark once an agreement is reached: $\beta_i = \beta_j = 0$ under agreement. Firms i and j bargain directly over the pair (e_i, e_j) and a monetary transfer $T \geq 0$ paid by j to i . Payoffs are quasi-linear in T , information is complete, and $\tau \in (0, 1)$ represents firm i 's bargaining power. The Nash bargaining problem is given by

$$\max_{e_i, e_j, T} \left[u_i(e_i, T) - d_i \right]^\tau \left[u_j(e_i, e_j, T; \gamma) - d_j \right]^{1-\tau}, \quad (29)$$

with $u_i(e_i, T) = V_i(e_i) + T$ and $u_j(e_i, e_j, T; \gamma) = V_j(e_i, e_j; \gamma) - T$. Following [Lemma 1](#), the disagreement payoffs are $d_i = H_i^{I*}$ and $d_j = H_j^{I*}(\gamma)$, corresponding to the mutually-insured outcome.

The first-order condition with respect to T yields $\tau/(u_i - d_i) = (1 - \tau)/(u_j - d_j) \equiv \kappa$. Substituting κ into the first-order conditions with respect to e_i and e_j cancels it entirely, leaving

$$\frac{\partial V_i}{\partial e_i} + \frac{\partial V_j}{\partial e_i} = 0, \quad \frac{\partial V_j}{\partial e_j} = 0,$$

independently of τ , d_i , d_j , and T . Thus, the efficient effort levels are chosen *prior* to splitting the surplus and do not depend on the specific disagreement point used.

Solving this system yields

$$e_i^C(\gamma) = (1 + \gamma)\phi\theta, \quad e_j^C = \phi\theta. \quad (30)$$

At $\gamma = 1$, this recovers $e_i^C = 2\phi\theta$. For $\gamma < 1$, firm i 's effort increase is strictly smaller, as a weaker externality calls for less correction. Firm j 's effort remains unchanged relative to [\(3\)](#) for any γ .

The loss base of firm j at the bargained profile is $x - \gamma e_i^C(\gamma) - e_j^C = x - \phi\theta(1 +$

$\gamma + \gamma^2$). Maintaining a strictly positive loss base requires

$$x > \phi\theta(1 + \gamma + \gamma^2). \quad (31)$$

Relative to the mutually-insured disagreement point of [Lemma 1](#), the efficient profile [\(30\)](#) generates the bargaining surplus

$$S^C(\gamma) \equiv (V_i^C + V_j^C) - (H_i^{I*} + H_j^{I*}(\gamma)), \quad (32)$$

which, defining

$$B = \frac{(\phi\theta - e_i^{I*})(\phi\theta^2 + m e_i^{I*})}{\theta} > 0, \quad C = \frac{m^2(e_i^{I*})^2}{2D\theta} > 0, \quad (33)$$

can be expressed as a quadratic function of γ :

$$S^C(\gamma) = \left(\frac{\phi^2\theta^2}{2} - C \right) \gamma^2 + B\gamma - 2\Delta H_i, \quad (34)$$

where $\Delta H_i \equiv H_i^{I*} - V_i^{NC}$ ([\(16\)](#)).

From [\(34\)](#), the three terms admit a clear economic interpretation. The constant term $-2\Delta H_i$ represents a fixed cost: reaching an agreement requires both firms to forgo the individual insurance gains they have already secured (ΔH_i each). The linear term $B\gamma$ captures a direct benefit channel: as γ increases, firm j gains more from firm i 's efficiently increased effort. The quadratic term $\left(\frac{\phi^2\theta^2}{2} - C \right) \gamma^2$ nets out two competing forces: the efficiency gain from internalizing the externality, since $e_i^C(\gamma) = (1 + \gamma)\phi\theta$ increases with γ — against the erosion of bargaining surplus caused by firm j 's disagreement payoff $H_j^{I*}(\gamma)$ rising with γ .

Consequently, $S^C(\gamma)$ can be negative for low values of γ . Bargaining from the mutually-insured disagreement point is viable if and only if $\gamma \geq \gamma^C(\lambda)$, where $\gamma^C(\lambda)$ is the unique positive root of $S^C(\gamma) = 0$:

$$\gamma^C(\lambda) = \frac{-B + \sqrt{B^2 + 8 \left(\frac{\phi^2\theta^2}{2} - C \right) \Delta H_i}}{2 \left(\frac{\phi^2\theta^2}{2} - C \right)}. \quad (35)$$

When an agreement is viable ($\gamma \geq \gamma^C(\lambda)$), the resulting payoffs are

$$u_i^C = H_i^{I*} + \tau S^C(\gamma), \quad u_j^C = H_j^{I*}(\gamma) + (1 - \tau) S^C(\gamma). \quad (36)$$

Each firm receives its disagreement payoff plus its proportional share of the generated surplus.

Had bargaining instead originated from the standard non-cooperative disagreement point $e_i^{NC} = e_j^{NC} = \phi\theta$ (with payoffs V_i^{NC} and $V_j^{NC}(\gamma)$), the same efficient action profile would yield the total surplus

$$TS(\gamma) \equiv (V_i^C(e_i^C) + V_j^C(e_i^C, e_j^C; \gamma)) - (V_i^{NC} + V_j^{NC}(\gamma)) = \frac{\gamma^2 \phi^2 \theta^2}{2}. \quad (37)$$

Unlike $S^C(\gamma)$, this surplus is strictly positive for any $\gamma > 0$ and monotonically increases with the strength of the spillover. The wedge between the two measures stems solely from the baseline subtracted: $S^C(\gamma) = TS(\gamma) - \Delta H_i - \Delta_j(\gamma)$, where $\Delta_j(\gamma) \equiv H_j^{I*}(\gamma) - V_j^{NC}(\gamma)$. The terms ΔH_i and $\Delta_j(\gamma)$ reflect the gains already captured unilaterally via private insurance prior to negotiation; what the firms claim individually can exceed the surplus created by coordination itself.

5.2 The Rigid protocol: buying back the externality

Unlike the Sequential protocol of [Section 5.3](#), this is a single-stage negotiation over T alone: $(\beta_i^{I*}, \beta_j^{I*})$ are inherited from [Sections 3](#) and [4](#), chosen before, and independently of, this protocol – not a first stage that anticipates it. In this protocol, effort is not directly contractible. Firm i has chosen to insure at β_i^{I*} , reducing its effort to $e_i^{I*} < \phi\theta$. The sole instrument negotiated is a transfer $T \geq 0$ paid by j to i in exchange for i – and only i – dropping its coverage to $\beta_i = 0$, restoring its effort to $e_i^{NC} = \phi\theta$; j 's coverage $\beta_j^{I*}(\gamma)$ is never renegotiated, whether or not agreement is reached. Payoffs are quasi-linear in T , information is complete, and $\tau \in (0, 1)$ is i 's bargaining power.

Failing an agreement, i remains insured at β_i^{I*} with effort e_i^{I*} , and j insures optimally at

$$\beta_j^{I*}(\gamma) = \frac{(x - \phi\theta - \gamma e_i^{I*}) m}{D\phi\theta}, \quad (38)$$

yielding the disagreement payoffs $d_i = H_i^{I*}$, $d_j = H_j^{I*}(\gamma)$ of [Lemma 1](#).

Under agreement, firm i drops its insurance ($\beta_i = 0$) and its effort returns to $e_i^{NC} = \phi\theta$. Firm j 's coverage and effort are unaffected – $\beta_j^{I*}(\gamma)$, $e_j^{I*}(\gamma)$ remain exactly as at the disagreement point – since j 's optimal effort $e_j^*(\beta_j) = (1 - \beta_j)\phi\theta$ depends on β_j alone, never on e_i . What changes for j is purely the loss base it faces, now smaller because i 's effort has risen. A transfer $T \geq 0$ is paid by j to i . The payoffs under agreement are

$$u_i(T) = V_i^{NC} + T, \quad u_j(T) = \tilde{V}_j(\gamma) - T, \quad (39)$$

where $\tilde{V}_j(\gamma)$ is j 's payoff at its own unchanged $(\beta_j^{I*}(\gamma), e_j^{I*}(\gamma))$, evaluated at i 's new effort $\phi\theta$ instead of e_i^{I*} :

$$\tilde{V}_j(\gamma) = H_j^{I*}(\gamma) + K_j(\gamma) \phi \gamma (\phi\theta - e_i^{I*}), \quad K_j(\gamma) \equiv \theta - \beta_j^{I*}(\gamma) m. \quad (40)$$

The Nash program is

$$\max_T \left[u_i(T) - d_i \right]^\tau \left[u_j(T) - d_j \right]^{1-\tau} \quad \text{s.t. } u_i(T) \geq d_i, \quad u_j(T) \geq d_j. \quad (41)$$

The two participation constraints require respectively $T \geq \Delta H_i$ and $T \leq \tilde{V}_j(\gamma) - H_j^{I*}(\gamma)$. They are simultaneously satisfied if and only if $S^R(\gamma) \geq 0$, where

$$S^R(\gamma) \equiv \left[\tilde{V}_j(\gamma) - H_j^{I*}(\gamma) \right] - \Delta H_i = C' \gamma^2 + B' \gamma - \Delta H_i, \quad (42)$$

with

$$B' = \frac{(\phi\theta - e_i^{I*}) (D\phi\theta^2 - (x - \phi\theta)m^2)}{D\theta}, \quad C' = \frac{(\phi\theta - e_i^{I*}) m^2 e_i^{I*}}{D\theta}. \quad (43)$$

From this expression, the three terms admit a clear interpretation. The constant term $-\Delta H_i$ is a fixed cost: i must be compensated for the individual insurance gain it gives up, while j gains nothing from the deal at $\gamma = 0$ – with no spillover, i 's effort is irrelevant to j . The linear term $B' \gamma$ is the direct benefit channel: as γ increases, j benefits more from i 's restored effort. The quadratic term $C' \gamma^2$ is a second, reinforcing channel: a stronger spillover already lowers j 's disagreement-

point loss base, so j needs less coverage at the disagreement point ($\beta_j^{I*}(\gamma)$ falls in γ), which raises j 's marginal sensitivity $K_j(\gamma)$ to further changes in its loss base – amplifying the direct benefit rather than eroding it.

Since $C' > 0$ and $S^R(0) = -\Delta H_i < 0$, $S^R(\gamma)$ has exactly one positive root, defining a threshold $\gamma^R(\lambda)$ such that a deal exists if and only if $\gamma \geq \gamma^R(\lambda)$,

$$\gamma^R(\lambda) = \frac{-B' + \sqrt{(B')^2 + 4C'\Delta H_i}}{2C'}. \quad (44)$$

For $\gamma \geq \gamma^R(\lambda)$, the Nash Bargaining Solution (NBS) splits $S^R(\gamma)$ in proportions $(\tau, 1 - \tau)$. The equilibrium transfer follows from $u_i^R(T^{R*}) - d_i = \tau S^R(\gamma)$:

$$T^{R*} = \Delta H_i + \tau S^R(\gamma). \quad (45)$$

The equilibrium payoffs are

$$u_i^{R*} = H_i^{I*} + \tau S^R(\gamma), \quad u_j^{R*} = H_j^{I*}(\gamma) + (1 - \tau) S^R(\gamma). \quad (46)$$

Each firm obtains its disagreement payoff plus its share of the bargaining surplus. This yields the following Proposition.

Proposition 2 (Rigid protocol viability). *An agreement is reached if and only if $\gamma \geq \gamma^R(\lambda)$, given by (44). For $\gamma < \gamma^R(\lambda)$, no agreement is reached and both firms remain at the disagreement point $(d_i, d_j) = (H_i^{I*}, H_j^{I*}(\gamma))$ of Lemma 1. For $\gamma \geq \gamma^R(\lambda)$, the joint payoff at agreement is*

$$u_i^{R*} + u_j^{R*} = V_i^{NC} + \tilde{V}_j(\gamma), \quad (47)$$

independently of the disagreement point, so the entire effect of τ is distributional, not allocative.

Proof. Existence follows from $S^R(\gamma) \geq 0$, established above. The joint payoff identity (47) holds because $u_i(T) + u_j(T) = V_i^{NC} + \tilde{V}_j(\gamma)$ for any T , independently of the disagreement point, and the NBS only reallocates this fixed total between i and j via T^{R*} . $\square$

5.3 Sequential bargaining

Like the Rigid protocol, the disagreement point here has both firms insured, never the non-cooperative outcome. Unlike Rigid, however, it is not fixed at the point $(\beta_i^{I*}, \beta_j^{I*}(\gamma))$ of [Lemma 1](#): coverage is chosen strategically in stage 1, in anticipation of the bargaining that follows, so the disagreement payoffs d_i, d_j are themselves endogenous to that choice. Effort, moreover, is directly contractible in the stage-2 bargaining – neither restriction of the Rigid protocol applies here.

In this protocol the coverages are chosen *before* any negotiation and are not renegotiated afterwards. The game has two stages. In *stage 1*, firms i and j choose their coverages $\beta_i, \beta_j \in [0, 1]$ simultaneously and non-cooperatively, each anticipating the outcome of stage 2. In *stage 2*, taking (β_i, β_j) as given, the two firms bargain à la Nash (with bargaining power $\tau \in (0, 1)$ for i) over the effort pair (e_i, e_j) and a transfer $T \geq 0$.

We solve by backward induction, stage 2 first. Denote $L \equiv \mu + \lambda\sigma$, so that $m = \theta - L$ and $D = 2L - \theta$.

5.3.1 Stage 2

Taking (β_i, β_j) as given, the stage-2 Nash program is

$$\max_{e_i, e_j, T} \left[u_i(e_i, T) - d_i \right]^\tau \left[u_j(e_i, e_j, T) - d_j \right]^{1-\tau}, \quad (48)$$

where $u_i(e_i, T) = V_i(e_i; \beta_i) - P_i + T$ and $u_j(e_i, e_j, T) = V_j(e_i, e_j; \beta_j) - P_j - T$, with V_i, V_j the gross mean-risk payoffs of [\(19\)](#) (both firms insured) and P_i, P_j the premiums, computed under correct anticipation of the effort levels (e_i, e_j) chosen in this stage.

As in the previous protocols, T enters u_i, u_j linearly with slopes $+1, -1$, so the FOC with respect to T gives the same proportional-split condition $\tau[u_j - d_j] = (1-\tau)[u_i - d_i]$, and since T cancels from the sum, $[u_i(e_i, T) - d_i] + [u_j(e_i, e_j, T) - d_j] \equiv S^S(e_i, e_j; \beta_i, \beta_j)$ is independent of T . Substituting the resulting split back into [\(48\)](#) shows that the maximised value of the Nash product, for any (e_i, e_j) , is $\tau^\tau (1-\tau)^{1-\tau} S^S(e_i, e_j; \beta_i, \beta_j)$. Maximizing [\(48\)](#) over all three variables is therefore equivalent to first maximizing $S^S(e_i, e_j; \beta_i, \beta_j)$ over (e_i, e_j) alone, and only then

splitting the resulting total in proportions $(\tau, 1 - \tau)$ via T .

Since d_i, d_j, P_i, P_j do not depend on (e_i, e_j) , maximizing $S^S(e_i, e_j; \beta_i, \beta_j)$ over (e_i, e_j) is equivalent to maximizing the joint gross payoff $V_i + V_j$. The FOCs

$$\frac{\partial(V_i + V_j)}{\partial e_i} = -e_i + (1 - \beta_i)\phi\theta + \gamma(1 - \beta_j)\phi\theta = 0, \quad \frac{\partial(V_i + V_j)}{\partial e_j} = -e_j + (1 - \beta_j)\phi\theta = 0,$$

yield

$$e_j^S(\beta_j) = (1 - \beta_j)\phi\theta, \quad e_i^S(\beta_i, \beta_j; \gamma) = (1 - \beta_i)\phi\theta + \gamma(1 - \beta_j)\phi\theta. \quad (49)$$

Only i raises its effort under agreement, by exactly γe_j^S , internalizing its own externality; j 's effort is unchanged from disagreement.

Substituting (49) into $S^S(e_i, e_j; \beta_i, \beta_j)$ gives the bargaining surplus

$$S^S(\beta_i, \beta_j; \gamma) = \frac{\gamma\phi^2\theta(1 - \beta_j)}{2} \left[\gamma\theta(1 - \beta_j) + 2L(\beta_i + \gamma\beta_j) \right]. \quad (50)$$

It is non-negative on the whole admissible region and strictly positive whenever $\gamma > 0$ and $\beta_j < 1$. Unlike the previous protocol, agreement in the Sequential protocol *never* breaks down. At $\beta_i = \beta_j = 0$ it reduces to the Coase surplus $\gamma^2\phi^2\theta^2/2$ of (37); at $\beta_j = 1$ it vanishes. The NBS splits S^S in proportions $(\tau, 1 - \tau)$, so

$$u_i^S = d_i + \tau S^S, \quad u_j^S = d_j + (1 - \tau) S^S, \quad (51)$$

where d_i, d_j are the disagreement payoffs (individual insurance at (β_i, β_j)).

5.3.2 Stage 1

Anticipating the stage-2 outcome (51), firm i chooses β_i to maximize its own payoff, taking β_j as given:

$$\max_{\beta_i \in [0,1]} u_i^S(\beta_i, \beta_j) = d_i(\beta_i) + \tau S^S(\beta_i, \beta_j; \gamma), \quad (52)$$

and, symmetrically, firm j solves

$$\max_{\beta_j \in [0,1]} u_j^S(\beta_i, \beta_j) = d_j(\beta_i, \beta_j; \gamma) + (1 - \tau) S^S(\beta_i, \beta_j; \gamma), \quad (53)$$

taking β_i as given. Here $d_i(\beta_i)$ is i 's individual-insurance payoff of [Section 3](#), while $d_j(\beta_i, \beta_j; \gamma)$ is j 's individual-insurance payoff of [Section 4](#), which also depends on β_i through i 's disagreement effort $e_i(\beta_i)$.

Differentiating [\(52\)](#) with respect to β_i and using $d'_i(\beta_i) = D\theta\phi^2(\beta_i^{I*} - \beta_i)$ together with $\partial S/\partial\beta_i = L\gamma\phi^2\theta(1 - \beta_j)$ (from [\(50\)](#)) gives firm i 's best response

$$\beta_i^{BR}(\beta_j) = \beta_i^{I*} + \underbrace{\frac{\gamma\tau L(1 - \beta_j)}{D}}_{\text{strategic mark-up}}. \quad (54)$$

By insuring more, i lowers its disagreement effort, which widens the effort gap it can later sell, hence enlarges S and the transfer it receives. The mark-up decomposes into three effects. First, a *surplus-size effect*: $\partial S/\partial\beta_i > 0$ is the pure effect of β_i on the negotiable pie, prior to any sharing. Second, a *bargaining-power weight*: i only captures a fraction τ of that enlarged surplus, so the incentive to grow S is scaled by τ ; at $\tau = 0$ this channel vanishes entirely. Third, a *deviation-cost effect*: moving β_i away from its private optimum β_i^{I*} costs i in terms of d_i , at rate $D\theta\phi^2$ (the curvature of d_i). The mark-up equates the bargaining-weighted surplus gain to this private cost,

$$\frac{\gamma\tau L(1 - \beta_j)}{D} = \frac{\tau \partial S/\partial\beta_i}{D\theta\phi^2},$$

so i trades off a private loss for a larger share of a bigger pie. The mark-up increases in the externality γ , in i 's bargaining power τ , and in the premium level L , and decreases in β_j . As $\tau \rightarrow 0$, $\beta_i^{BR} \rightarrow \beta_i^{I*}$. With no bargaining power, i maximizes its disagreement payoff and the protocol collapses to the individual insurance problem of [Section 3](#).

The same steps applied to [\(53\)](#), differentiating $d_j(\beta_i, \beta_j; \gamma) + (1 - \tau)S^S(\beta_i, \beta_j; \gamma)$ with respect to β_j , give firm j 's best response

$$\beta_j^{BR}(\beta_i) = \frac{m[x - \phi\theta(1 + \gamma + \gamma^2(1 - \tau))]}{D\phi\theta\Omega} - \frac{\gamma(D - L\tau)}{D\Omega}\beta_i, \quad \Omega \equiv 1 + \gamma^2(1 - \tau). \quad (55)$$

At the upper corner $\beta_i = 1$, [\(55\)](#) simplifies to

$$\beta_j^{BR}(1) = \frac{m(x - \phi\theta) - \gamma\phi\theta(1 - \tau)(L + \gamma m)}{D\phi\theta[1 + \gamma^2(1 - \tau)]}, \quad (56)$$

whose numerator sets the pure insurance motive $m(x - \phi\theta)$ against a strategic motive $\gamma\phi\theta(1 - \tau)(L + \gamma m)$ pushing j towards no insurance.

Firm i 's objective is concave: $\partial^2 u_i^S / \partial \beta_i^2 = -\phi^2 \theta D < 0$. Firm j 's objective satisfies

$$\frac{\partial^2 u_j^S}{\partial \beta_j^2} = -\phi^2 \theta D [1 + \gamma^2(1 - \tau)],$$

which is strictly negative unconditionally, since $D > 0$ and $1 + \gamma^2(1 - \tau) > 0$ throughout the admissible region. Both problems are therefore concave and the best responses are linear.

Writing the best responses (54) and (55) in slope-intercept form, $\beta_i^{BR}(\beta_j) = A_i - k\beta_j$ and $\beta_j^{BR}(\beta_i) = A_j - q\beta_i$, with slopes $k = \gamma\tau L/D$, $q = \gamma(D - L\tau)/(D\Omega)$ and intercepts $A_i = \beta_i^{I*} + k$, $A_j = m[x - \phi\theta(1 + \gamma + \gamma^2(1 - \tau))]/(D\phi\theta\Omega)$, the interior equilibrium solves the linear 2×2 system

$$\beta_i^S = \frac{A_i - kA_j}{1 - kq}, \quad \beta_j^S = \frac{A_j - qA_i}{1 - kq}. \quad (57)$$

The system is invertible whenever $1 - kq \neq 0$; clearing denominators shows $1 - kq = \Delta/(D^2\Omega)$, where

$$\Delta = D^2 [1 + \gamma^2(1 - \tau)^2] + D\gamma^2 m\tau(2\tau - 1) + \gamma^2 m^2 \tau^2,$$

is strictly positive on the whole admissible region, so (57) is well defined and the interior equilibrium, when it exists, is unique. It exhibits strategic over-insurance by i and withdrawal by j , namely $\beta_i^S \geq \beta_i^{I*}$ and $\beta_j^S \leq \beta_j^{I*}$, with β_i^S strictly increasing and β_j^S strictly decreasing in γ .

Because the strategic mark-up drives β_i^S up and pushes β_j^S down, the equilibrium is generically at a corner for moderate-to-high γ . The domain partitions into four regimes: (1) interior/interior for low γ , given by (57); (2) $\beta_i^S = 1$ with $\beta_j^S = \beta_j^{BR}(1)$ interior, from (56), for high γ and high τ ; (3) $\beta_j^S = 0$ with $\beta_i^S = \beta_i^{I*} + \gamma\tau L/D$ interior, for intermediate γ and low τ ; (4) corner/corner $(1, 0)$ for high γ , where both conditions below bind. Regimes (2) and (3) are mutually exclusive along any path in γ at fixed τ : depending on whether τ exceeds a critical value $\tilde{\tau}$ (defined below), only one of the two can be crossed before the corner $(1, 0)$ is

reached. [Proposition 3](#) formalises the two resulting trajectories. The two corner conditions and the thresholds separating the four regimes are given in closed form in [Appendix A.1](#) ((61) to (63)). In words: γ_i^c and γ_j^c are the values of γ at which the corners $\beta_i^S = 1$ and $\beta_j^S = 0$ first bind – the roots of $\beta_i^{BR}(0) = 1$ and $\beta_j^{BR}(1) = 0$; γ_i^o, γ_j^o are the thresholds at which the interior equilibrium (57) itself first violates $\beta_i^S \leq 1$ or $\beta_j^S \geq 0$.

Proposition 3. (i) *The Sequential protocol never breaks down: $S^S(\beta_i, \beta_j; \gamma) > 0$ whenever $\beta_j < 1$, so participation always holds strictly.* (ii) *The interior equilibrium (57) is unique whenever it exists, and exhibits strategic over-insurance by i and under-insurance by j , monotone in γ .* (iii) *There exists a critical $\tilde{\tau}$, defined by $\gamma_i^o = \gamma_j^o$, such that as γ increases:*

- *if $\tau > \tilde{\tau}$: interior/interior for $\gamma < \gamma_i^o$; then $\beta_i^S = 1$ with β_j^S interior for $\gamma_i^o \leq \gamma < \gamma_j^c$; then the corner $(1, 0)$ for $\gamma \geq \gamma_j^c$;*
- *if $\tau < \tilde{\tau}$: interior/interior for $\gamma < \gamma_j^o$; then $\beta_j^S = 0$ with β_i^S interior for $\gamma_j^o \leq \gamma < \gamma_i^c$; then the corner $(1, 0)$ for $\gamma \geq \gamma_i^c$.*

The corner $(1, 0)$ is reached within the admissible range only if the relevant threshold (γ_j^c , resp. γ_i^c) does not exceed it; otherwise the mixed regime persists up to $\gamma = 1$.

The Sequential protocol *reverses* the configuration of the previous one. There, j 's ability to insure was the credible threat and i dropped its coverage under agreement. Here it is i who over-insures strategically – up to full coverage – to strengthen its bargaining position, while j strategically withholds insurance, because additional coverage by j shrinks the negotiable surplus (50) without improving its disagreement position enough to compensate. The asymmetry traces directly to the fact that only i carries the externality: i 's coverage has strategic value in the negotiation that j 's does not.

We check the admissibility at the corner $\beta_i^S = 1$. In the Sequential protocol the binding loss base is that of j at the bargained profile, $\ell_j^S/\phi = x - \gamma(1 - \beta_i)\phi\theta - (1 - \beta_j)(1 + \gamma^2)\phi\theta$, which over $[0, 1]^2$ is minimized at $(\beta_i, \beta_j) = (0, 0)$ and

there equals $x - \phi\theta(1 + \gamma + \gamma^2)$. The maintained condition (31) thus keeps every loss base non-negative on the whole domain. At the corner $\beta_i^S = 1$, where i 's disagreement effort $e_i(1) = 0$ and its bargained effort reduces to the spillover term $e_i^S = \gamma(1 - \beta_j)\phi\theta$, the loss base becomes $\ell_j^S/\phi = x - (1 - \beta_j)(1 + \gamma^2)\phi\theta$, minimized at $\beta_j = 0$ (Regime 4). Admissibility there requires only $x \geq \phi\theta(1 + \gamma^2)$, strictly weaker than (31). Full insurance by i suppresses its private effort and hence relaxes j 's loss base, so the corner is a fortiori admissible.

To illustrate the relationship between the critical bargaining power $\tilde{\tau}$ and the loading factor λ , we consider the following baseline calibration ($x = 3$, $\phi = \theta = 1$, $\mu = 0.1$, $\sigma = 0.5$). At the baseline calibration, $\underline{\lambda} = 1.30$ and $\bar{\lambda} = 1.80$ ((13) and (14)). The grid above covers this admissible interval from just above $\underline{\lambda}$ – where β_i^{I*} reaches 1, a degenerate corner – up to $\lambda = 1.70$, comfortably short of $\bar{\lambda}$. As $\tilde{\tau}$ is increasing in λ and, numerically, approaches its maximum value 1 as $\lambda \rightarrow 1.76$, i.e. before $\bar{\lambda}$ is reached, we therefore stop the table at $\lambda = 1.70$ rather than at $\bar{\lambda} = 1.80$.

λ	1.31	1.35	1.40	1.45	1.50	1.55	1.60	1.65	1.70
$\tilde{\tau}$	0.0462	0.2078	0.3690	0.4975	0.6027	0.6910	0.7669	0.8336	0.8936

Table 1: Critical bargaining power $\tilde{\tau}$ as a function of the loading factor λ , at baseline calibration ($x = 3$, $\phi = \theta = 1$, $\mu = 0.1$, $\sigma = 0.5$).

The relationship between λ and $\tilde{\tau}$ shows that a harder market (higher λ) lowers both firms' individually optimal coverage, β_i^{I*} and β_j^{I*} . This widens firm i 's room to over-insure strategically before hitting its corner $\beta_i^S = 1$, while narrowing firm j 's room before it withdraws entirely to its corner $\beta_j^S = 0$. Because the γ -threshold at which i reaches its corner is decreasing in τ while the threshold at which j reaches its corner is increasing in τ , raising λ shifts the former up and the latter down – and their crossing point, $\tilde{\tau}$, is pushed to a higher value of τ .

Three patterns are worth spelling out. Concerning the coverage, in both panels β_i^S rises and β_j^S falls in γ , the strategic over-insurance/under-insurance of Proposition 3; only in panel (a), where $\tau > \tilde{\tau}$, does β_i^S reach its corner at 1 within the grid, consistently with the regime ordering of Proposition 3(iii). Concerning the effort, $e_j^S = (1 - \beta_j^S)\phi\theta$ mechanically mirrors β_j^S , while e_i^S climbs sharply once β_i^S is pinned

Table 2: Equilibrium coverages, efforts and payoffs in the Sequential protocol as γ varies, at baseline calibration ($x = 3$, $\phi = \theta = 1$, $\mu = 0.1$, $\sigma = 0.5$, $w = 4$), $\lambda = 1.40$, for τ above and below the critical $\tilde{\tau} = 0.3969$. Recall $e_i, e_j \in (0, x)$ are not capped at 1 the way $\beta_i, \beta_j \in [0, 1]$ are.

(a) $\tau = 0.5 > \tilde{\tau}$: i 's corner ($\beta_i^S = 1$) reached first						
γ	β_i^S	β_j^S	e_i^S	e_j^S	u_i^S	u_j^S
0.1	0.6930	0.6055	0.3465	0.3945	1.6454	1.6715
0.2	0.7288	0.5341	0.3644	0.4659	1.6655	1.7098
0.3	0.7759	0.4537	0.3880	0.5463	1.6963	1.7474
0.4	0.8357	0.3660	0.4179	0.6340	1.7405	1.7828
0.5	0.9091	0.2727	0.4545	0.7273	1.8008	1.8149
0.6	0.9964	0.1758	0.4982	0.8242	1.8798	1.8422
0.7	1.0000	0.0950	0.6335	0.9050	1.9706	1.9059
0.8	1.0000	0.0202	0.7838	0.9798	2.0722	1.9802
0.9	1.0000	0.0000	0.9000	1.0000	2.1625	2.0625
1.0	1.0000	0.0000	1.0000	1.0000	2.2500	2.1500

(b) $\tau = 0.2 < \tilde{\tau}$: j 's corner ($\beta_j^S = 0$) reached first						
γ	β_i^S	β_j^S	e_i^S	e_j^S	u_i^S	u_j^S
0.1	0.6780	0.5763	0.3644	0.4237	1.6385	1.6800
0.2	0.6948	0.4723	0.4107	0.5277	1.6475	1.7350
0.3	0.7180	0.3589	0.4744	0.6411	1.6617	1.7999
0.4	0.7477	0.2406	0.5561	0.7594	1.6816	1.8759
0.5	0.7838	0.1216	0.6554	0.8784	1.7079	1.9642
0.6	0.8258	0.0057	0.7708	0.9943	1.7405	2.0656
0.7	0.8533	0.0000	0.8467	1.0000	1.7675	2.1810
0.8	0.8800	0.0000	0.9200	1.0000	1.7963	2.3026
0.9	0.9067	0.0000	0.9933	1.0000	1.8276	2.4302
1.0	0.9333	0.0000	1.0667	1.0000	1.8613	2.5640

at 1 (panel (a), $\gamma \geq 0.7$), where it reduces to the pure spillover term $\gamma(1 - \beta_j^S)\phi\theta$. Recall that e_i, e_j live in $(0, x) = (0, 3)$, not $[0, 1]$: some values slightly exceed 1 (e.g. $e_i^S = 1.0667$ in panel (b) at $\gamma = 1$) purely because $\phi\theta = 1$ at this calibration, well inside the admissible domain. Concerning the payoffs, we show that in panel (a), u_i^S overtakes u_j^S around $\gamma \approx 0.6$. Firm i 's strategic over-insurance eventually pays off enough, at this bargaining power, to reverse its initial disadvantage. In panel (b), no such reversal occurs. Firm j retains the advantage throughout and the gap widens sharply as $\gamma \rightarrow 1$, since firm i 's lower bargaining power limits how much of the growing surplus it can capture through over-insurance.

Remark 1 (Hard-market limit). As $\lambda \rightarrow \bar{\lambda}$ ((13)), individual insurance demand vanishes, $\beta_i^{I*}, \beta_j^{I*}(\gamma) \rightarrow 0$: insurance becomes too expensive to be worth buying for its own, risk-transfer purpose. One might expect the Sequential protocol to degenerate along with it. It does not. In the best responses (54) and (55), $m \rightarrow 0$ forces $L \rightarrow \theta$ and $D \rightarrow \theta$, so

$$\beta_i^{BR}(\beta_j) \rightarrow \gamma\tau(1 - \beta_j), \quad \beta_j^{BR}(\beta_i) \rightarrow 0, \quad (58)$$

whose fixed point is

$$(\beta_i^S, \beta_j^S) \longrightarrow (\gamma\tau, 0) \quad \text{as } \lambda \rightarrow \bar{\lambda}. \quad (59)$$

Firm j stops insuring altogether, consistent with the vanishing pure-insurance motive. Firm i does not: it keeps a strictly positive coverage, purely as a bargaining device – stripped of any risk-transfer content and pinned down solely by the externality γ and its bargaining power τ , independently of $x, \phi, \theta, \mu, \sigma$. The Rigid protocol, by contrast, genuinely degenerates at this limit: with nothing left to buy back, $S^R(\gamma) \rightarrow 0$ for every γ . So even where insurance is too expensive to be worth buying for its own sake, the Sequential protocol still generates a purely strategic insurance demand that the Rigid protocol cannot replicate.

6 Comparing insurance and bargaining protocols

Firms i and j face a menu of arrangements – no insurance (NC), individual insurance for i alone, both insured and the three implementable bargaining protocols (Coase, Rigid and Sequential). We assume that neither firm can unilaterally impose a choice on the other. We ask which arrangement each firm prefers, taken separately, and then which arrangement they can agree on jointly.

6.1 The full landscape

[Table 3](#) reports every payoff at a reference calibration, $\lambda = 1.40$, $\tau = 0.5$; the full dependence on λ and τ is examined in [Section 6.2](#). Two very different stories emerge for i and for j .

Firm i . Once i insures at all, the outcome for i never depends on β_j (the SOC and vertex formula of [Section 4](#) apply regardless of firm j 's choice), so the I (i only) and I (i and j) columns coincide. NC is dominated by every other arrangement throughout the table. Among the bargained arrangements, the Sequential protocol dominates for i throughout: it is the first to become viable ($\gamma = 0.1$) and remains i 's best option at every γ in the table, so i never has reason to prefer Coase or Rigid once Sequential is on the table. Behind it, the ranking between Coase and Rigid is not fixed. Rigid becomes viable first, crossing between the $\gamma = 0.2$ and $\gamma = 0.3$ rows, where γ crosses $\gamma^R(1.40) = 0.2294$ ([Table 4](#)); Coase becomes viable next, between $\gamma = 0.3$ and $\gamma = 0.4$, crossing $\gamma^C(1.40) = 0.3086$. Because both share the exact same disagreement point of [Lemma 1](#), their ranking at any γ is settled entirely by the sign of $S^C(\gamma) - S^R(\gamma)$ – the same sign for both firms, since $u_i^C - u_i^R = \tau[S^C(\gamma) - S^R(\gamma)]$ and $u_j^C - u_j^R = (1 - \tau)[S^C(\gamma) - S^R(\gamma)]$. At this calibration that sign is negative up to $\gamma \approx 0.41$: Rigid, not Coase, is the better arrangement over the range where both are viable but γ is not yet too large – visible at $\gamma = 0.4$, where $u_i^R = 1.6834$ narrowly exceeds $u_i^C = 1.6819$. Only beyond $\gamma \approx 0.41$ does Coase overtake Rigid, remaining ahead for the rest of the table. SO — the social optimum without insurance, splitting $TS(\gamma)$ from the

Table 3: Payoffs across all arrangements, baseline calibration ($x = 3$, $\phi = \theta = 1$, $\mu = 0.1$, $\sigma = 0.5$, $w = 4$), $\lambda = 1.40$, $\tau = 0.5$.

Firm i							
γ	NC	SO	I (i only)	I (i and j)	Coase	Rigid	Seq
0.1	1.5000	1.5025	1.6333	1.6333	—	—	1.6454
0.2	1.5000	1.5100	1.6333	1.6333	—	—	1.6655
0.3	1.5000	1.5225	1.6333	1.6333	—	1.6540	1.6963
0.4	1.5000	1.5400	1.6333	1.6333	1.6819	1.6834	1.7405
0.5	1.5000	1.5625	1.6333	1.6333	1.7398	1.7130	1.8008
0.6	1.5000	1.5900	1.6333	1.6333	1.8027	1.7427	1.8798
0.7	1.5000	1.6225	1.6333	1.6333	1.8705	1.7725	1.9706
0.8	1.5000	1.6600	1.6333	1.6333	1.9433	1.8025	2.0722
0.9	1.5000	1.7025	1.6333	1.6333	2.0210	1.8327	2.1625
1.0	1.5000	1.7500	1.6333	1.6333	2.1037	1.8630	2.2500
Firm j							
γ	NC	SO	I (i only)	I (i and j)	Coase	Rigid	Seq
0.1	1.6000	1.6025	1.5333	1.6623	—	—	1.6715
0.2	1.7000	1.7100	1.5667	1.6913	—	—	1.7098
0.3	1.8000	1.8225	1.6000	1.7203	—	1.7410	1.7474
0.4	1.9000	1.9400	1.6333	1.7495	1.7981	1.7996	1.7828
0.5	2.0000	2.0625	1.6667	1.7787	1.8852	1.8583	1.8149
0.6	2.1000	2.1900	1.7000	1.8080	1.9773	1.9173	1.8422
0.7	2.2000	2.3225	1.7333	1.8374	2.0745	1.9766	1.9059
0.8	2.3000	2.4600	1.7667	1.8668	2.1767	2.0360	1.9802
0.9	2.4000	2.6025	1.8000	1.8963	2.2840	2.0957	2.0625
1.0	2.5000	2.7500	1.8333	1.9259	2.3963	2.1463	2.1500

NC disagreement point – sits between I and NC for low γ , but crosses I around $\gamma \approx 0.73$ (the root of $V_i^{NC} + \tau TS(\gamma) = H_i^{I*}$) and remains above it for the rest of the table – a reminder that $TS(\gamma)$ grows in γ^2 while the gain from individual insurance, ΔH_i , does not grow with γ at all. SO is not itself an arrangement the two firms can reach, however: it presupposes bargaining from the NC disagreement point, which [Lemma 1](#) rules out as self-enforcing – i always insures unilaterally first.

Remark 2 (Why not bargain from NC?). *A natural objection is that i should never have insured in the first place: had the two firms bargained directly from NC, would either be better off? The SO column of [Table 3](#) answers this precisely – it is the same efficient action of (30), split instead from the NC disagreement point. Because $V_i^C + V_j^C$ is the same fixed total whichever disagreement point the split starts from, what i gains by starting from the both-insured point instead of NC is exactly what j loses:*

$$u_i^C(\tau) - u_i^{SO}(\tau) = u_j^{SO}(\tau) - u_j^C(\tau) = \Delta H_i(1 - \tau) - \tau \Delta_j(\gamma). \quad (60)$$

Since $\Delta H_i > 0$ always ((16)) and $\gamma^C(\lambda) > \gamma_j^{NC}(\lambda)$ throughout the admissible range, $\Delta_j(\gamma) < 0$ everywhere Coase is viable, so (60) is strictly positive for every $\tau \in (0, 1)$: i strictly prefers bargaining from the both-insured point to bargaining from NC, at any bargaining power. Insuring first is a dominant move for i , not merely an assumption. j would indeed prefer to bargain from NC – but [Lemma 1](#) already rules this out: i insures unilaterally regardless of what j does or negotiates, so j can never hold i at NC long enough to bargain from there.

Firm j . The picture is qualitatively different. Individual insurance by i alone hurts j relative to NC ((18)); insurance by both firms repairs part of the damage but, once γ exceeds the threshold $\gamma_j^{NC}(\lambda)$, defined by $H_j^{I*}(\gamma) = V_j^{NC}(\gamma)$, does not fully restore NC (closed form in [Appendix A.2](#)). At $\lambda = 1.40$, $\gamma_j^{NC}(1.40) = 0.188$: for all but the very first row of the table, $H_j^{I*}(\gamma) < V_j^{NC}(\gamma)$. Coase and Rigid inherit this shortfall directly, since both share the exact disagreement point of [Lemma 1](#); the Sequential protocol shares only the same *type* of disagreement point (both insured, never NC), evaluated instead at firms’ own strategically chosen

coverage – yet its payoffs for j remain below NC throughout the table as well. So for $\gamma \geq 0.2$, j does *strictly better under NC, and under SO, than under any of the three bargained arrangements*. Among the three, the same reversal found for i appears for j around the same $\gamma \approx 0.41$: Rigid gives j more than Coase up to that point ($u_j^R = 1.7996 > u_j^C = 1.7981$ at $\gamma = 0.4$), after which Coase takes over and stays ahead for the rest of the table. The Sequential protocol, despite being i 's best arrangement throughout and the only one viable at low γ , is typically the *worst* of the three for j once more than one is viable ($\gamma \geq 0.4$) – the gap to Rigid narrowing again as $\gamma \rightarrow 1$, where the two are nearly tied. Firm j is the residual claimant on an externality it cannot control on its own. Every implementable arrangement lets j claw back *some* of what i 's (privately rational) insurance choice cost it, but none restores what j had before i started insuring – not even Coase, the arrangement i itself prefers once γ is large enough.

6.2 Bargaining outcomes and voluntary agreements

Section 6.1 established this ordering at the reference calibration $\lambda = 1.40$: i never prefers Rigid to Sequential, and Rigid dominates Coase, whenever both are viable, up to $\gamma \approx 0.41$. We now characterize the relevant thresholds – $\gamma^R(\lambda)$, $\gamma^C(\lambda)$, and the crossing point $\gamma^{RC}(\lambda)$ where $S^C(\gamma) = S^R(\gamma)$ – as functions of λ , and use them to establish which arrangement, if any, the two firms voluntarily adopt.

Table 4: Viability thresholds $\gamma^R(\lambda)$ (Rigid), $\gamma^C(\lambda)$ (Coase), and the crossing point $\gamma^{RC}(\lambda)$ where $S^C(\gamma) = S^R(\gamma)$, as a function of the loading factor λ , at baseline calibration ($x = 3$, $\phi = \theta = 1$, $\mu = 0.1$, $\sigma = 0.5$).

λ	1.31	1.35	1.40	1.45	1.50	1.55	1.60	1.65	1.70
γ^R	0.3198	0.2742	0.2294	0.1923	0.1598	0.1302	0.1025	0.0760	0.0503
γ^C	0.4020	0.3575	0.3086	0.2644	0.2233	0.1842	0.1464	0.1095	0.0729
γ^{RC}	0.4915	0.4549	0.4057	0.3547	0.3033	0.2520	0.2011	0.1504	0.1001

$\gamma^R(\lambda)$, $\gamma^C(\lambda)$ and $\gamma^{RC}(\lambda)$ are all strictly decreasing in λ , for the same reason: as the loading grows, firm i 's individual coverage β_i^{I*} falls, both firms give up less at the disagreement point by reaching an agreement, and a smaller spillover suffices to make bargaining – or switching from Rigid to Coase – worth it. The ordering $\gamma^R(\lambda) < \gamma^C(\lambda) < \gamma^{RC}(\lambda)$ holds throughout the admissible range.

For γ in the range where Rigid is the relevant alternative to Sequential, j switches at a threshold $\tau_j^{*,R}(\gamma, \lambda)$, defined by $u_j^S(\gamma, \tau_j^{*,R}; \lambda) = u_j^R(\gamma, \tau_j^{*,R}; \lambda)$. For γ in the range where Coase is the relevant alternative, j switches at $\tau_j^{*,C}(\gamma, \lambda)$, defined by $u_j^S(\gamma, \tau_j^{*,C}; \lambda) = u_j^C(\gamma, \tau_j^{*,C}; \lambda)$, and i switches at $\tau_i^*(\gamma, \lambda)$, defined by $u_i^S(\gamma, \tau_i^*; \lambda) = u_i^C(\gamma, \tau_i^*; \lambda)$. Firm i never switches to preferring Rigid over Sequential, at any γ or τ : no analogous threshold exists on that side. Whenever the two firms prefer different arrangements, neither can impose its choice on the other, so the arrangement actually implemented is the disagreement point itself. None of the three thresholds admits a closed form; all are obtained numerically, by bisection on $\tau \in (0, 1)$, throughout the paper.

Proposition 4 (Voluntary protocol choice). *(i) For $\gamma < \gamma^R(\lambda)$: neither Rigid nor Coase is viable at any τ : the Sequential protocol is adopted for lack of an alternative.*

(ii) For $\gamma^R(\lambda) \leq \gamma < \gamma^{RC}(\lambda)$: i always prefers Sequential. For $\tau < \tau_j^{,R}(\gamma, \lambda)$: both firms prefer Sequential, and agree on it. For $\tau \geq \tau_j^{*,R}(\gamma, \lambda)$: j switches to preferring Rigid, while i still prefers Sequential.*

(iii) For $\gamma \geq \gamma^{RC}(\lambda)$: the same three-region structure applies, with Coase in place of Rigid. For $\tau < \tau_j^{,C}(\gamma, \lambda)$: both prefer Sequential, agree on it. For $\tau_j^{*,C}(\gamma, \lambda) \leq \tau < \tau_i^*(\gamma, \lambda)$: j prefers Coase, i still prefers Sequential. For $\tau \geq \tau_i^*(\gamma, \lambda)$ (when it exists): both prefer Coase, agree on it.*

Firm i 's switching point. Under Coase, i 's payoff is its disagreement payoff H_i^{I*} plus a τ -share of $S^C(\gamma)$. Under the Sequential protocol, i has one additional lever: it can strategically over-insure at stage 1 ($\beta_i^S \geq \beta_i^{I*}$) to enlarge the negotiable surplus and extract a larger transfer. At $\tau = 0$ this lever is worthless – i 's stage-1 problem collapses to maximizing its own disagreement payoff, giving $\beta_i^S = \beta_i^{I*}$ exactly and $u_i^S(\tau = 0) = H_i^{I*} = u_i^C(\tau = 0)$. Firm i is exactly indifferent between the two protocols at $\tau = 0$, for every (γ, λ) – the same identity that held against Rigid, since both share the disagreement point. As τ rises from 0, i strictly prefers the Sequential protocol, but the lever is bounded by $\beta_i \leq 1$: near $\underline{\lambda}$, β_i^{I*} is already close to 1 (e.g. 0.96 at $\lambda = 1.31$), so the strategic room $1 - \beta_i^{I*}$ is thin, the corner is reached early, and i 's edge over Coase erodes fast; near $\bar{\lambda}$, the room is ample and i 's preference for Sequential persists to much higher τ . Past a threshold $\tau_i^*(\gamma, \lambda)$

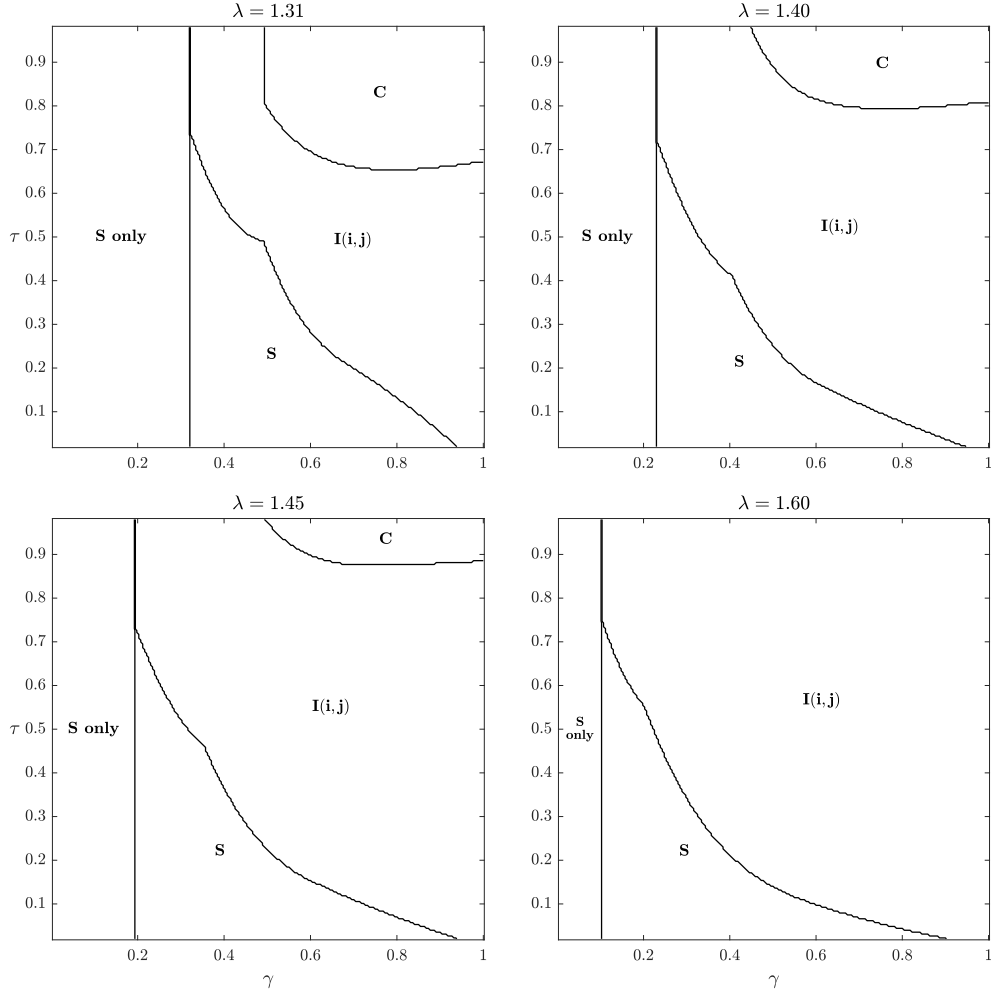


Figure 1: Actual outcome of voluntary protocol choice between i and j , as a function of (γ, τ) , at four values of λ across the admissible range. Region “S only”: $\gamma < \gamma^R(\lambda)$, no alternative viable, Sequential adopted by default. Region “S”: both firms prefer Sequential and agree on it. Region “I (i,j)”: the two firms prefer different arrangements: neither can impose its choice, so the disagreement point of [Lemma 1](#) is implemented instead – merging what are, in [Proposition 4](#), two distinct conflicts (Sequential versus Rigid below $\gamma^{RC}(\lambda)$, Sequential versus Coase above it) into a single observed outcome. Region “C”: both firms prefer Coase and agree on it.

– which fails to exist once λ is large enough that the lever never exhausts within $\tau \in (0, 1)$ – i switches to preferring Coase.

Firm j 's switching point. Since i is exactly indifferent at $\tau = 0$, j there captures the entire surplus in whichever protocol is being compared, and simply prefers whichever creates more of it. Where Rigid is the relevant alternative, the Sequential surplus S^S exceeds $S^R(\gamma)$ throughout our grid at $\tau = 0$, so j strictly prefers Sequential there; where Coase is the relevant alternative, likewise against $S^C(\gamma)$. As τ rises, i 's strategic over-insurance erodes its own disagreement effort, which erodes j 's disagreement payoff through the spillover channel; j switches to preferring the relevant alternative past $\tau_j^{*,R}(\gamma, \lambda)$ or $\tau_j^{*,C}(\gamma, \lambda)$, and i switches to Coase past $\tau_i^*(\gamma, \lambda)$ in turn, when it exists. None of the three reduces to a tractable closed form – each depends on which of the four regimes of [Proposition 3](#) the Sequential equilibrium is in – so we locate them numerically. j always switches before i does, whenever both thresholds exist, giving the region structure of [Proposition 4](#).

Table 5: Thresholds $\tau_j^{*,R}(\gamma, \lambda)$, $\tau_j^{*,C}(\gamma, \lambda)$ and $\tau_i^*(\gamma, \lambda)$ at $\lambda = 1.40$ ($\gamma^R(1.40) = 0.2294$, $\gamma^{RC}(1.40) = 0.4057$).

$\gamma^R(1.40) \leq \gamma < \gamma^{RC}(1.40)$ – Rigid is the relevant alternative							
γ	0.25	0.28	0.30	0.32	0.35	0.38	0.40
$\tau_j^{*,R}$	0.6704	0.5968	0.5531	0.5150	0.4682	0.4329	0.4148
$\gamma \geq \gamma^{RC}(1.40)$ – Coase is the relevant alternative							
γ	0.45	0.5	0.6	0.7	0.8	0.9	1.0
$\tau_j^{*,C}$	0.3235	0.2515	0.1661	0.1188	0.0763	0.0373	—
τ_i^*	0.9771	0.8897	0.8152	0.7960	0.7943	0.8001	0.8090

Both $\tau_j^{*,R}(\gamma, \lambda)$ and $\tau_j^{*,C}(\gamma, \lambda)$ fall monotonically in γ within their respective ranges, meeting at $\gamma^{RC}(1.40)$ where, by construction, $S^C(\gamma) = S^R(\gamma)$ makes the two comparisons coincide; the combined boundary is continuous there but kinks, since the two sides have different slopes. $\tau_j^{*,C}(\gamma, \lambda)$ vanishes (into the region $\tau \in (0, 1)$) by $\gamma \approx 1$: j switches to Coase almost immediately once it is the relevant alternative, for high enough γ . $\tau_i^*(\gamma, \lambda)$, by contrast, is comparatively flat, around 0.79–0.98 – consistent with i 's strategic lever depending on β_i^{I*} and λ , not directly on γ .

Same outcome, different reasons. For $\gamma < \gamma^R(\lambda)$ (*S only*), and for $\tau < \tau_j^{*,R}(\gamma, \lambda)$ or $\tau < \tau_j^{*,C}(\gamma, \lambda)$ in the higher- γ ranges (*S*), both end up on the Sequential protocol – but for different reasons: below $\gamma^R(\lambda)$, no alternative is viable at all; above it, an alternative is viable but both firms prefer Sequential anyway, a genuine revealed preference rather than the absence of a benchmark to compare against.

Trend across λ . Across the four panels, as λ rises from 1.31 to 1.60, $\gamma^R(\lambda)$ and $\gamma^{RC}(\lambda)$ both fall (Table 4), so the region where Sequential is the only option (*S only*) shrinks, while the region of disagreement ($I(i, j)$) expands sharply at the expense of the region where both firms agree on Sequential (*S*). The region where both agree on Coase (*C*) also grows, but more modestly. A single mechanism drives all three movements: a harder market pushes β_i^{I*} down, widening *i*’s strategic room $1 - \beta_i^{I*}$ and making *j* switch away from Sequential – toward Rigid, then Coase – at ever lower τ .

The overall picture qualifies the reading of Propositions 2 and 3 in an important way: those propositions characterise what each protocol delivers *conditional on being adopted*, but adoption itself is not automatic. The Sequential protocol, despite dominating both Rigid and, over part of the parameter space, even Coase, is not always implementable, because the firm that stands to lose from the reallocation of surplus it induces – generically *j*, once τ is large enough – can simply decline it in favor of the disagreement point, or of Coase, that already guarantees it more.

7 Conclusion

The objective of this paper was to ask whether two firms exposed to an interdependent risk, one insured and the other affected by that choice, can resolve the resulting externality through direct negotiation. Our main results characterize whether bargaining is even attempted, and on what terms. Based on the three protocols we have considered, our results show that at low bargaining power for the upstream firm *i*, both firms prefer the Sequential protocol. There, *i*’s strategic option to over-insure lets it extract a share of a larger surplus. As *i*’s bargaining power rises, *j*’s disagreement payoff is eroded by *i*’s own over-insurance. *j* comes to prefer whichever alternative is viable – Rigid first, the Coase benchmark beyond

a further threshold – while i still favors the Sequential protocol.

For a substantial range of parameters, these preferences are strictly opposed. Neither arrangement can be adopted by mutual consent. The two firms remain stuck at the self-enforcing, both-insured outcome. This is worse for j than the situation that prevailed before i started insuring. It is reached not because bargaining failed on its own terms, but because the two parties could not agree on the terms on which to bargain in the first place. This is a Coasian failure of a different kind than the usual suspects – private information, transaction costs, ill-defined property rights. Here, information is complete and gains from trade generically exist. The obstacle is a coordination failure over the bargaining protocol itself. Only once i 's bargaining power is high enough, and the insurance market leaves it little strategic room to begin with, does i too switch to preferring Coase.

Taken together, these results trace a full arc from a positive externality to its resolution, or lack thereof. When the spillover is weak, little is at stake. Neither Rigid nor Coase is viable, and the Sequential protocol prevails by default, delivering an outcome close to each firm's individual-insurance payoff. When the spillover is strong and the insurance market leaves firm i meaningful room to manoeuvre, the outcome bifurcates sharply with the balance of bargaining power: cooperation on the Sequential protocol at low bargaining power for i , outright conflict once j 's share grows large enough to prefer an alternative – Rigid, then Coase – that i still does not. Only once the insurance market is soft enough that firm i 's strategic room is exhausted does agreement return, this time on Coase itself, precisely because firm i no longer has anything left to protect by insisting on the Sequential alternative.

These results speak to a trend already visible in several insurance markets exposed to climate risk. As the underlying hazard intensifies, insurers raise loadings or withdraw coverage altogether. The market hardens rather than softens. Our model suggests this hardening is not benign for the resolution of the resulting insurance externality. A harder market – flood insurance premiums rising as basins prove less reliable, wildfire coverage retreating as fire risk grows – leaves the upstream firm more, not less, strategic room to over-insure and extract surplus in bargaining. It widens the region in which the two parties cannot agree on how to resolve the externality between them. The conventional expectation is that costlier insurance simply forces both parties back toward self-protection and co-

operation. That expectation does not hold here. It is precisely when insurance becomes expensive that direct negotiation is least likely to succeed.

Two extensions seem natural. The first would consider more than two firms, keeping the upstream-downstream structure but allowing a chain or a network of interdependent risks, and adopting a Nash-in-Nash bargaining protocol across the resulting bilateral relationships ([Collard-Wexler, 2019](#)). This would raise new questions about disagreement points. Some firms in the chain might insure, others might not, and the disagreement payoff of each pairwise negotiation would then depend on the insurance choices made elsewhere in the network. The second would introduce a third strategic agent, a monopoly insurer, into the bargaining itself rather than treating the loading factor as parametric ([Hofmann, 2007](#)). This would let the insurer's own pricing respond to the externality and to the bargaining outcome, closing the loop between insurance market power and the resolution of the externality it helps create.

A Threshold characterizations

This appendix collects the expressions of the thresholds used throughout the paper.

A.1 The Sequential protocol's thresholds: γ_i^c , γ_j^c , γ_j° , and γ_i°

The two corner conditions are

$$\beta_i^S = 1 \iff m \geq \frac{\phi L \theta (1 - \gamma \tau)}{x}, \quad \beta_j^S = 0 \iff m \leq \frac{\gamma \phi \theta (1 - \tau) L}{x - \phi \theta \Omega}. \quad (61)$$

The exact thresholds in γ separating the four regimes of [Section 5.3](#) follow directly:

γ_i^c , γ_j^c (roots of $\beta_i^{BR}(0) = 1$ and $\beta_j^{BR}(1) = 0$), in closed form

$$\gamma_i^c = \frac{\phi L \theta - m x}{\tau \phi L \theta}, \quad \gamma_j^c = \frac{-L(1 - \tau)\phi\theta + \sqrt{L^2(1 - \tau)^2\phi^2\theta^2 + 4m^2(1 - \tau)\phi\theta(x - \phi\theta)}}{2m(1 - \tau)\phi\theta}, \quad (62)$$

and γ_j° , the unique root in $(0, 1)$ of $P_j(\gamma) = 0$, where P_j is the numerator of $\beta_j^S(\gamma)$

$$\phi\theta \left(L^2\tau^2 - D^2\tau - Dm \right) \gamma^2 - m(Dx - L\tau(x - \phi\theta))\gamma + Dm(x - \phi\theta) = 0 \quad (\beta_j^S = 0). \quad (63)$$

The symmetric threshold γ_i° is the unique root in $(0, 1)$ of $P_i(\gamma) = 0$, where P_i is the numerator of $\beta_i^S(\gamma) - 1$ (equilibrium [\(57\)](#) cleared of its common denominator).

Writing $B \equiv \phi L \theta - m x$ for the numerator appearing in γ_i^c above and

$$Q_i \equiv m^2(2\phi\theta - \phi\tau\theta + 2x) + m(2\phi\tau\theta^2 - 3\phi\theta^2 - \theta x) + \phi\theta^3(1 - \tau),$$

the cubic is

$$P_i(\gamma) = \phi\tau\theta(1 - \tau)L^2\gamma^3 - (1 - \tau)Q_i\gamma^2 + B(L\tau\gamma - D). \quad (64)$$

The linear-and-constant part factors through the same B that determines γ_i^c , but the quadratic and cubic coefficients do not reduce further, so no comparably simple closed form is available for γ_i° ; it is obtained numerically as the unique root of [\(64\)](#) in $(0, 1)$.

A.2 The threshold $\gamma_j^{NC}(\lambda)$

The threshold $\gamma_j^{NC}(\lambda)$ solves $H_j^{I*}(\gamma) = V_j^{NC}(\gamma)$. Using the identity $V_j^*(\beta_i^{I*}) - V_j^{NC} = -\frac{\gamma(x-\phi\theta)m\phi\theta}{D}$ from (18), this equality expands to:

$$2D^2\theta\gamma(x-\phi\theta)\phi\theta = m\left[D\left(x-(1+\gamma)\phi\theta\right) + \gamma m(x-\phi\theta)\right]^2. \quad (65)$$

Solving (65) for its unique positive root yields $\gamma_j^{NC}(\lambda)$.

References

- Ambec, S. and Sprumont, Y. (2002). Sharing a river. *Journal of Economic Theory*, 107(2), 453–462.
- Arrow, K. J. (1965). *Aspects of the Theory of Risk-Bearing*. Yrjö Jahnssonin Säätiö, Helsinki.
- Avery, C. and Zemsky, P. (1995). Multidimensional uncertainty and technical interdependence in financial markets. *American Economic Review*, 88(4), 724–748.
- Binmore, K., Rubinstein, A., and Wolinsky, A. (1986). The Nash bargaining solution in economic modelling. *RAND Journal of Economics*, 17(2), 176–188.
- Chiu, W. H. (2000). Degree of risk aversion and the demand for insurance. *Journal of Risk and Insurance*, 67(2), 263–277.
- Coase, R. H. (1960). The problem of social cost. *Journal of Law and Economics*, 3, 1–44.
- Collard-Wexler, A., Gowrisankaran, G., and Lee, R. (2019). “Nash-in-Nash” bargaining: A microfoundation for applied work. *Journal of Political Economy*, 127(1), 163–195.
- Courbage, C. and Loubergé, H. (2006). On the demand for insurance with prudence. *Journal of Risk and Insurance*, 73(4), 679–688.
- Courbage, C., Loubergé, H., and Rey, B. (2018). On the properties of high-order non-monetary measures for risks. *The Geneva Risk and Insurance Review*, 43(1), 77–94.
- Crocker, K. J. and Masten, S. E. (1991). Pretia ex Machina? Prices and Process in Long-Term Contracts *Journal of Law, Economics, & Organization*, 34(1), 69–99.
- Deryugina, T., Moore, F. and Tol, R. (2021). Environmental applications of the Coase Theorem. *Environmental Science and Policy*, 120, 81–88.

- Doherty, N. A. and Schlesinger, H. (1983). Optimal insurance in the presence of uninsurable background risk. *Journal of Political Economy*, 91(6), 1045–1054.
- Eeckhoudt, L., Gollier, C., and Schlesinger, H. (2006). *Economic and Financial Decisions under Risk*. Princeton University Press, Princeton, NJ.
- Ehrlich, I. and Becker, G. S. (1972). Market insurance, self-insurance, and self-protection. *Journal of Political Economy*, 80(4), 623–648.
- Heal, G. and Kunreuther, H. (2007). You can only die once: Interdependent security in an uncertain world. In *Managing Strategic Surprise: Lessons from Risk Management and Risk Assessment*, pages 85–104. Cambridge University Press.
- Hofmann, A. (2007). Internalizing externalities of loss prevention through insurance monopoly: an analysis of interdependent risks. *The Geneva Risk and Insurance Review*, 32(1), 91–111.
- Hofmann, A. and Ramaj, H. (2019). Interdependent risk networks: the threat of cyber attack. *The Geneva Risk and Insurance Review*, 44(2), 207–221.
- Kunreuther, H. and Heal, G. (2003). Interdependent security. *Journal of Risk and Uncertainty*, 26(2), 231–249.
- Lohse, T., Robledo, J. R., and Schmidt, U. (2012). Self-insurance and self-protection as public goods. *Journal of Risk and Insurance*, 79(1), 57–76.
- MacMinn, R. D. (1987). Corporate insurance and the underinvestment problem. *Journal of Risk and Insurance*, 54(1), 19–31.
- Meyer, J. (1987). Two-moment decision models and expected utility maximization. *American Economic Review*, 77(3), 421–430.
- Mossin, J. (1968). Aspects of rational insurance purchasing. *Journal of Political Economy*, 76(4), 553–568.
- Muermann, A. and Kunreuther, H. (2008). Self-protection and insurance with interdependent risks. *Journal of Risk and Uncertainty*, 36(3), 201–223.

- Parchomovsky, G. and Siegelman, P. (2022). Third-party moral hazard and the problem of insurance externalities. *Iowa Law Review*, 107(5), 2115–2160.
- Seog, S. H. (2006). Strategic demand for insurance. *Journal of Risk and Insurance*, 73(2), 279–295.
- Seog, S. H. (2024). On the efficiency of insurance institutions under interdependent risks. *Journal of Risk and Insurance*, 91(4), 1025–1048.