

A crises-bailouts game[☆]Bruno Salcedo^a, Bruno Sultanum^b , Ruilin Zhou^c^a Western University, Canada^b University of Essex, United Kingdom^c Pennsylvania State University, United States of America

ARTICLE INFO

JEL classification:

C73

D82

Keywords:

Bailouts

Moral hazard

Repeated games

Imperfect monitoring

ABSTRACT

This paper studies the optimal design of a liability-sharing arrangement as an infinitely repeated game. We construct a noncooperative model with an active and a passive agent. The active agent can take a costly and unobservable avoidance action to reduce the incidence of a crisis, but a crisis is costly for both agents. When a crisis occurs, each agent decides how much to contribute to mitigating it. For the one-shot game, when the avoidance cost is not too high relative to the expected loss of crisis for the active agent, a no-bailout policy always achieves the first-best outcome, at which the active agent puts in effort to minimize the crisis incidence. However, the first-best is not achievable when the avoidance cost is sufficiently high. We show that, in the latter case with the same stage game, the first-best cannot be implemented as a perfect public equilibrium (PPE) of the infinitely repeated game either. Instead, at any constrained efficient PPE with avoidance, the active agent “shirks” infinitely often, though crises are always mitigated, and is bailed out infinitely often. The reason is that promises of future shirking and bailout incentivize the active player to take the costly crisis-avoidance action in the present. This result runs contrary to the typical moral hazard view that bailouts reduce incentives for agents to avoid crises. Here bailouts enhance ex-ante mitigation efforts rather than diminish them and are necessary to achieve the second-best. We use finite-state automata to approximate the constrained efficient PPE and explore some comparative statics of the repeated game numerically.

1. Introduction

Many important institutional arrangements inherently bear high incentive costs. These costs become apparent during challenging times, such as financial crises, prompting varied responses to mitigate the damage. At times, troubled institutions have been bailed out, and others have not. As a result, the fate of these troubled parties ranges from complete failure and bankruptcy to full recovery.¹ Typically, a troubled institution gets bailed out on the grounds that the alternative (failure) would have been much more costly, at least in the short run, since it might impose a big negative externality on many related parties. When this happens, economists

[☆] We thank Ed Green, V. Bhaskar, Neil Wallace, and Rishabh Kirpalani for their helpful discussion and comments.

* Corresponding author.

E-mail addresses: bsalcedo@uwo.ca (B. Salcedo), bruno@sultanum.com (B. Sultanum), rzhou@psu.edu (R. Zhou).

¹ For example, while the majority of US companies sink or float on their own, the US government has consistently bailed out large corporations in the automobile industry such as GM and Chrysler. Among financial institutions, the Federal Reserve Bank significantly assisted large banks like Citigroup and Bank of America with loans and guarantees, while it let other large financial institutions, such as Lehman Brothers and Washington Mutual, fail during the 2007–2009 Great Recession. Among sovereign countries, the US government helped Mexico to survive the 1994 Tequila crisis, while many countries suffered huge losses during the 1997 Asian financial crisis with little help from the IMF. During the Eurozone crisis that followed the Great Recession, countries such as Greece, Italy, and Spain received multiple rounds of bailouts with varying magnitudes.

are always quick to point out a fatal flaw of such rescue operation: the moral hazard — the “too big to fail” justification of bailout encourages behaviors that may lead to more failure. This incentive cost serves as the rationale for not bailing out some troubled institutions. Here, we develop a dynamic model of crises and bailouts that tackles these confronting views directly.

In this paper, we show that the decision to bail out an institution when optimally taken does not necessarily lead to more failures. Instead, it can increase mitigation efforts and reduce the likelihood of a crisis. We model this problem as a dynamic game between two players: the “crisis-inflicting” active agent, and the potential “help-to-clean-up” passive agent. The active agent can take a costly unobservable action to reduce the incidence of crisis (avoidance). Whenever a crisis occurs, both parties suffer some loss, but each agent can simultaneously contribute to mitigate the loss. It is assumed both players have nontransferable utility. They can contribute directly to reduce the loss from the crisis but not to increase each other’s consumption. This assumption rules out direct subsidies from the passive agent to the active agent to pay for his avoidance cost.²

The one-shot game exhibits various combinations of avoidance and mitigation strategies as static Nash equilibria, depending on the model parameters. Specifically, when the avoidance cost is relatively low compared to the expected loss from a crisis for the active agent, a no-bailout policy can effectively lead the active agent to take the socially desirable but costly avoidance action. However, if this cost becomes too high, there is no equilibrium where the active agent chooses the avoidance action.

We then study the perfect public equilibria (PPE) of the repeated game. We show that for model parameters such that the active agent *shirking* — that is, not taking the avoidance action — is the only static Nash equilibrium of the stage game, the first-best requiring him to take the avoidance action every period cannot be implemented as a PPE of the infinitely repeated game. The reason is the same as in the static game. To induce the active agent to take the costly avoidance action, the expected mitigation cost for him in case of a crisis must be higher than the avoidance cost. With high avoidance costs, the active agent is better off doing nothing.

While the first-best outcome is unattainable, agents can still improve upon the repetition of the static Nash equilibrium if they are sufficiently patient. We show that if the discount factor is above a certain threshold, any constrained efficient PPE will involve the active agent shirking infinitely often, though crises are always mitigated, and the passive agent bailing out the active agent infinitely often. As a result, crises occur more frequently compared to the first-best, but less frequently than in a repetition of the static Nash equilibrium.

The intuition behind the seemingly counterintuitive result that shirking and bailouts can incentivize good behavior lies in the dynamics of repeated interactions. In a one-shot game, the threat of a crisis and its associated costs might not be enough to induce the active agent to take costly avoidance actions if such cost is high. However, in a repeated game, future leniency (in the form of bailouts or allowing shirking) can be a reward for good behavior today. The active agent, anticipating these future benefits, may then be induced to take the costly avoidance action in the present, even if it is not in their immediate best interest. The interaction of these dynamic forces creates a delicate balance where the possibility of future moral hazard (shirking or bailouts) paradoxically leads to less moral hazard in the present, ultimately increasing overall welfare.

A natural question arises: why is it necessary to promise both future shirking and bailouts to incentivize the active agent? The promise of future bailouts is more appealing as it compensates the active agent without sacrificing efficiency. However, the passive agent’s ability to bear the mitigation cost is limited by incentive compatibility. This is because, while the active agent’s shirking imposes a cost on the passive agent, the latter has no recourse but to bear it. On the other hand, the passive agent can always choose not to mitigate the crisis if his required contribution becomes too high. As a result, after a series of successful avoidance actions and crisis-free periods, the only feasible compensation left to offer is the promise of future shirking.

The optimal mechanism to align incentives requires allowing the active agent to shirk and providing bailouts, both occurring infinitely often. This raises the question of how much welfare gains a constrained efficient PPE can achieve compared to the first-best outcome. To investigate this, we employ finite-state automata to approximate the constrained efficient PPE. Note that any PPE in the repeated game can theoretically be represented by an automaton, possibly with infinite states. We focus on finite-state automata for two reasons. One is that the PPE of a finite-state automaton gives up an implementable procedure, including both the automaton design and the equilibrium strategies, that approximates the constrained efficient PPE of the original game. The other reason is that it gives us a computationally feasible way to conduct welfare comparisons. In our numerical examples, the welfare loss relative to the first-best is relatively small, suggesting that strategic shirking and bailouts can lead to big welfare gains even in the presence of moral hazard. Furthermore, our analysis reveals that following the instructions from the approximated constrained efficient mechanism can significantly reduce costs for both the active and passive agents.

Our paper is closely related to the literature that studies incentive problems induced by bailout policies. Most papers in the literature take particular institutional design and market structures very seriously but abstract from strategic dynamic interactions. In contrast, we simplify the environment in multiple dimensions to make it manageable, while taking seriously the dynamic game played by the two parties involved. This approach enables us to elucidate the mechanism of stochastic bailout, which has some interesting economic implications. In most papers, bailouts generate bad incentives for private agents. This is true in the stage game of our model. We show that, in the repeated setting, the credible promise of conditional future bailouts can be used to generate good incentives and reduce the incidence of crises. Such strategic behavior is likely present in repeated interactions between long-lived large agents, such as members of the European Union or a government and a large corporation.

² This assumption is often not far from reality. For example, it may be politically infeasible to contribute directly to another country’s budget. Furthermore, even when transfers are feasible, they can be costly. In these situations, a combination of transfers and the mechanism we propose would likely be optimal.

Some examples of papers highlighting the negative incentives generated by bailout policies are [Chari and Kehoe \(2016\)](#), [Farhi and Tirole \(2012, 2018\)](#).³ [Chari and Kehoe \(2016\)](#) study the time inconsistency problem of bailouts. The paper focuses on the dynamic policy decision of a *bailout authority* who cannot commit to future actions (like the two players in our model). [Farhi and Tirole \(2012, 2018\)](#) consider a commitment problem from the government side, and focus on the strategic complementarity of firms' risk-taking behavior. The last two papers study a finite-horizon stage game and the former assumes bailout policies to be noncontingent in agents' identities.

[Green \(2010\)](#) and [Keister \(2016\)](#) highlight that bailouts can be welfare-enhancing, but they do so without relying on dynamic incentives. [Keister \(2016\)](#) studies a version of [Diamond and Dybvig \(1983\)](#) that allows the government to divert expenditures from public goods to bailout banks and highlights two important implications of bailouts. On the one hand, bailouts induce bad behavior in banks, making them less cautious and more illiquid. On the other hand, bailouts in their environment also provide insurance to depositors. [Keister \(2016\)](#) demonstrates that the insurance effect prevails when the likelihood of a crisis is low. Similarly, [Green \(2010\)](#) suggests that the welfare-enhancing benefits of bailouts stem from the necessity of such measures in a regime with limited-liability firms. In this case, bailouts enable firms to offer perfect risk-sharing. The mechanism that makes bailouts welfare-enhancing in these models differs from others in the literature, including ours. In comparison to our model, they lack the dynamic incentive properties we highlight.

The constrained efficient PPE we study is characterized by recurrent crises, as in [Green and Porter \(1984\)](#). Such a phenomenon reflects an equilibrium that passes through several distinct states, rather than independent randomization by individual agents. Unlike in [Green and Porter \(1984\)](#), however, the recurrent shirking by the active agent is not punished by a severe ex-post inefficient outcome. Instead, all crises are mitigated regardless of the cause (hence ex-post efficient). The solution here is a deliberate arrangement of interwoven occasional shirking and bailout as mechanisms to incentivize the good behavior of the active agent as often as possible. Such a solution is more related to [Rubinstein \(1979\)](#), which shows in the context of criminal proceedings that, with repeated interactions, it is optimal to be lenient on offenders with good records.

The design of optimal incentive contracts in dynamic settings has also been a central focus in contract theory. A substantial body of literature follows [Spear and Srivastava \(1987\)](#) in studying a discrete-time framework and applying recursive methods with the agent's continuation value as a state variable. We build on this literature, specifically on subsequent developments by [Abreu et al. \(1986, 1990\)](#), to derive our results. A subsequent literature, notably [Sannikov \(2008\)](#) and [Williams \(2015\)](#), has extended these insights to continuous time. Our study returns to the discrete-time setting and departs from this literature in two other ways. The principal, our passive agent, cannot commit to the contract, and there is no transferable utility — instead, the game participants decide how to split the mitigation cost. The lack of commitment introduces a limit on how much utility the passive agent can promise to the active agent to induce effort. This difference put this paper closer to the relational contract theory (principal–agent dynamic moral hazard problems without commitment) as reviewed in [Watson \(2021\)](#). The absence of transferable utility restricts the range of possible contractual arrangements. Our application to a crisis-bailout game also differs from most papers on contract theory, offering new insights into bailout policies.

The paper is structured as follows. Section 2 introduces the stage game, while Section 3 describes the repeated game. Section 4 contains our main theoretical results. Section 5 uses numerical exercises to illustrate how the incentive mechanism works and to explore some comparative statics. Section 6 discusses alternative mechanisms under different assumptions. Finally, Section 7 concludes.

2. The stage game

There are two agents, agent 1 and agent 2, and two subperiods. In the first subperiod, agent 1 either takes an avoidance action to avert a crisis, $a = 1$, or not, $a = 0$. The cost of taking the avoidance action is $d > 0$, and the cost of not taking the avoidance action is normalized to zero. Agent 1's action a is unobservable to agent 2.

In the second subperiod, one of two things happens: either there is a crisis, denoted by $\xi = 1$, or there is no crisis, denoted by $\xi = 0$. Define $X \equiv \{0, 1\}$, $\xi \in X$. The probability of a crisis, conditional on agent 1's action in the first subperiod, $a \in \{0, 1\} \equiv A$, is $\pi^a \in (0, 1)$. We assume $\pi^1 < \pi^0$ so that taking the avoidance action reduces the probability of crisis. Agent 2 cannot infer agent 1's action from observing whether there is a crisis. Throughout the text, we refer to agent 1 as the active agent given that his action affects the probability of a crisis, and agent 2 as the passive agent since he is forced to face the consequences of a crisis but does not influence its occurrence.

In the event of a crisis, $\xi = 1$, the two agents can jointly mitigate the crisis. Let $m_i \geq 0$ denote agent i 's contribution to mitigation. The crisis is mitigated if the total contribution of the two agents, $m_1 + m_2$ is no less than one. If the crisis is mitigated ($m_1 + m_2 \geq 1$), the cost to agent i is only his mitigation contribution m_i . If the crisis is not mitigated ($m_1 + m_2 < 1$), agent i suffers a loss $c_i > 0$ due to the crisis and his contribution m_i . If there is no crisis, $\xi = 0$, nothing needs to be mitigated and agents do not suffer any loss. It is implicit in the payoff structure that there is no transferable utility; contribution $m_1 + m_2$ is made only to mitigate a crisis. Neither party can consume it once it is made. The two agents cannot make payments to each other in the first subperiod. In Section 6, we show that relaxing this assumption would greatly reduce the difficulty of achieving a better allocation in equilibrium. Fig. 1 summarizes the stage-game structure.

We are interested in studying the case where mitigation after a crisis is ex-post efficient. Therefore, we make the following assumptions on the model parameters.

³ Other examples of papers highlighting the negative incentives generated by bailout policies are [Schneider and Tornell \(2004\)](#), [Ennis and Keister \(2009\)](#) and [Brunnermeier et al. \(2016\)](#).

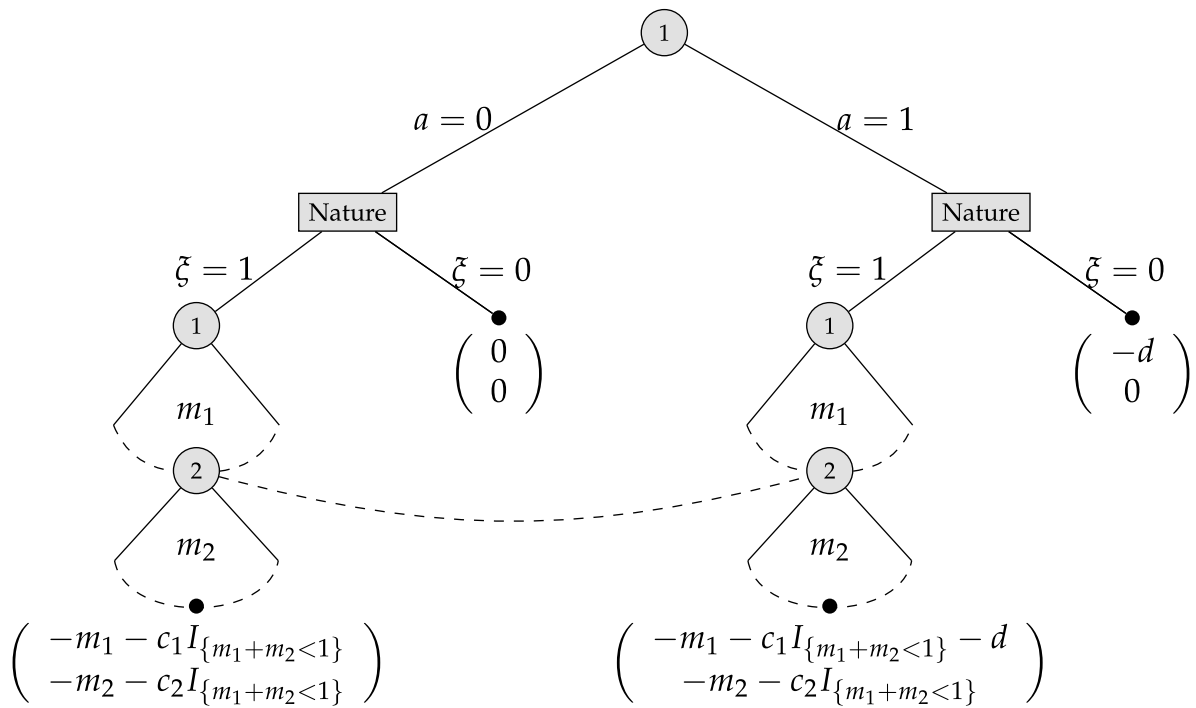


Fig. 1. The stage game.

Assumption 1. For $i = 1, 2$, $c_i \in (0, 1)$, and $c_1 + c_2 > 1$.

Assumption 1 implies that neither agent alone is willing to mitigate the crisis, but together they should. Since the total cost of mitigation is less than that when there is no mitigation — that is, $1 < c_1 + c_2$ — mitigation is efficient.

2.1. Equilibrium of the stage game

The structure of the game allows us to restrict attention to pure and public strategies without loss of generality. We thus solve for (pure-strategy, public-perfect) Nash equilibria of the two-subperiod normal-form game.

Denote the strategy for the active agent 1 by (a, m_1) , and for the passive agent 2 by m_2 , where a is agent 1's avoidance action, and m_i is agent $i = 1, 2$ mitigation contribution. Given the strategy profile (a, m_1, m_2) , agent 1's expected payoff is

$$u_1(a, m_1, m_2) = -ad - \pi^a(m_1 + c_1 I_{\{m_1+m_2 \geq 1\}}), \quad (1)$$

and agent 2's expected payoff is

$$u_2(a, m_1, m_2) = -\pi^a(m_2 + c_2 I_{\{m_1+m_2 \geq 1\}}). \quad (2)$$

As usual, a Nash equilibrium of the stage game is given by a strategy profile (a, m_1, m_2) such that (a, m_1) is a best response for agent 1 given the strategy m_2 , and vice versa.

The stage game has many equilibria depending on the parameters, and we are interested in the parameter region where the efficient outcome cannot be supported as an equilibrium outcome. There are multiple best responses in the mitigation stage after a crisis. In particular, by Assumption 1, after a crisis, not contributing to mitigation $m_i = 0$, is always the best response if the other agent is doing the same — although this is ex-post inefficient. If agent 1 contributes m_1 , and agent 2 contributes the remaining $1 - m_1$, one can infer that $m_1 \leq c_1$ and $m_2 = 1 - m_1 \leq c_2$. That is, for $(m_1, 1 - m_1)$ to be both agents' best mitigation response to each other, it must be that $1 - c_2 \leq m_1 \leq c_1$. This mitigation outcome is ex-post efficient. The cases with $m_1 + m_2 > 1$ or $m_1 + m_2 < 1$ with either $m_1 > 0$ or $m_2 > 0$ can be easily ruled out since at least one agent would be strictly better off by decreasing his mitigation contribution.

When deciding whether to take the avoidance action, agent 1 weighs the cost d against the expected gain of taking the action and successfully avoiding a crisis, which is either the change in his expected contribution $(\pi^0 - \pi^1)m_1$ or expected loss $(\pi^0 - \pi^1)c_1$. Define a composite parameter

$$\hat{d} \equiv \frac{d}{\pi^0 - \pi^1}$$

which is the cost of avoidance adjusted by its impact on the probability of a crisis. If $\hat{d} \leq c_1$, the efficient equilibrium where agent 1 takes the avoidance action and crisis is mitigated is an equilibrium, and dominates all other equilibria. This scenario exemplifies the prevailing concern in the literature that bailouts can encourage moral hazard by diminishing the extent to which the active agent internalizes the cost of a crisis.

Definition 1. A *bailout* in the stage game is a situation where a crisis occurs, the agents jointly contribute sufficient resources to mitigate it, and the contribution of the active agent is less than his private crisis loss, i.e., $m_1 + m_2 \geq 1$, and $m_1 < c_1$.

In a bailout, the active agent fails to internalize both the social and private costs associated with a crisis.⁴ The condition $m_1 + m_2 \geq 1$ ensures that the crisis is mitigated, implying that the social cost of the crisis (the necessity for mitigation) is borne by both agents. The condition $m_1 < c_1$ highlights that the active agent's contribution to the mitigation is less than their private cost of the crisis c_1 . This indicates that the active agent is not fully bearing the consequences of his (in)action, even at the individual level, which can lead to *shirking* (not taking the avoidance action) and a higher frequency of crisis. In an equilibrium without bailouts, on the other hand, the active agent's cost during a crisis must be c_1 . If the adjusted avoidance cost $\hat{d}$ is less than c_1 , the active agent's best response is then to take the avoidance action.

Proposition 1. In the state game, if $\hat{d} < c_1$, the active agent takes the avoidance action in any Nash equilibrium without bailouts. Moreover, if $\hat{d} > c_1$, the active agent does not take the avoidance action in all equilibria.

Proposition 1 demonstrates that a no-bailouts policy can effectively prevent the moral hazard associated with the active agent neglecting the avoidance action. However, this holds true only when the avoidance cost is not excessively high. The question then arises: what if the cost is substantial, rendering a no-bailout policy ineffective? In other words, what if $c_1 < \hat{d} < 1$? In such a scenario, no static Nash equilibrium can achieve the first-best outcome. The subsequent analysis will focus on this case.

Assumption 2. $c_1 < \hat{d} < 1$.

Assumption 2 implies that $a = 0$ is always agent 1's optimal action. With agent 1 never taking the avoidance action, there are two types of equilibrium: a nonmitigation equilibrium where $(a, m_1, m_2) = (0, 0, 0)$, and a continuum of mitigation equilibria where $(a, m_1, m_2) = (0, m_1, 1 - m_1)$ where $m_1 \in [1 - c_2, c_1]$. These equilibria are ex-ante inefficient when $\hat{d} < 1$ since the cost of action d is less than the expected social gain of avoiding a crisis $\pi^0 - \pi^1$. Yet, the continuum of mitigating equilibrium is ex-post efficient given that all crises are mitigated. In the next section, we investigate what can be achieved in this region for the repeated game.

3. The repeated game

In the *repeated game*, time is discrete and is indexed by $t \in \{1, 2, \dots\}$. The two agents live forever and discount future payoffs with the same discount factor $\delta \in (0, 1)$. At the beginning of each period t , agents observe a payoff-irrelevant public signal $\theta_t \sim \mathcal{U}[0, 1]$, which is i.i.d. across periods. After observing the public signal, the agents play the stage game described in the previous section. The public signal allows agents to take correlated actions in each period. This serves a technical purpose — it convexifies the payoff set without explicitly considering the randomized strategy for agent 1's action a_t . The public information at the beginning of date t is denoted by $h_t \in H_t$. It consists of the realizations of all past and current public signals, the history of all past crises, and the history of all past contributions.

We focus on perfect Bayesian equilibria where both agents play *pure* and *public* strategies. Proposition A.2 establishes that this restriction is without loss of generality (see Appendix A for details). A public strategy for the active agent 1 is a sequence of measurable functions $\sigma_1 = (\alpha_t, \mu_{1t})_{t=1}^\infty$, where $\alpha_t(h_t) \in \{0, 1\}$ is whether to take the avoidance action in period t given the public information h_t , and $\mu_{1t}(h_t) \geq 0$ is his contribution to the mitigation in case of a crisis. A public strategy for the passive agent 2 is a sequence of functions $\sigma_2 = (\mu_{2t})_{t=1}^\infty$, where $\mu_{2t}(h_t) \geq 0$ is his contribution to the mitigation in case of a crisis. We denote a public strategy profile by $\sigma = (\sigma_1, \sigma_2)$.

The expected discounted utility for agent i from date t onward, given the strategy profile σ and public history h_t is

$$v_{it}(\sigma, h_t) = (1 - \delta) \mathbb{E} \left[\sum_{\tau=t}^{\infty} \delta^{\tau-t} u_i(\alpha_\tau(h_\tau), \mu_{1\tau}(h_\tau), \mu_{2\tau}(h_\tau)) \mid h_t \right], \quad (3)$$

where the expectation is taken with respect to the crisis realizations from period t onward and the public signal from period $t + 1$ onward. With slight abuse of notation, the average expected discounted utility for agent i at the beginning of the game is denoted by $v_i(\sigma) = \mathbb{E}[v_{i1}(\sigma, h_1)]$.

Definition 2. A public strategy profile σ^* is a *perfect public equilibrium* (PPE) if and only if $v_{it}(\sigma^*, h_t) \geq v_{it}(\sigma'_i, \sigma_{-i}^*, h_t)$ for any agent i , any public strategy σ'_i any period $t \geq 1$, and any public history h_t .

⁴ An alternative definition could focus on the relative contributions of each agent to the mitigation effort, and define a bailout as a situation where the active agent's contribution as a proportion of the total mitigation cost is less than their private crisis loss as a proportion of the total potential loss, i.e. $\frac{m_1}{m_1 + m_2} < \frac{c_1}{c_1 + c_2}$. While this definition has its merits, it does not capture the concept we aim to emphasize: that a no-bailout policy ensures the active agent internalizes the benefit of the avoidance action.

We denote the set of PPE payoffs by $\mathcal{V}^* = \{v(\sigma^*) \mid \sigma^* \text{ is a PPE}\}$. A PPE always exists because unconditional repetition of a static Nash equilibrium of the stage game is a PPE. In Appendix A, using a version of the recursive decomposition introduced in [Abreu et al. \(1990\)](#), we show that the set of PPE can be characterized recursively, and establish technical properties of equilibria used in the proofs.

4. Optimal level of crises and bailouts

Under [Assumptions 1](#) and [2](#), the first-best requires that in every period the active agent takes the avoidance action, and both agents mitigate a crisis if it happens. However, this is not achievable in an equilibrium of the stage game because the active agent never takes the avoidance action in such circumstances. In this section, we study how, and to what extent, welfare can be improved in the repeated setting. Our first finding is that, in the infinitely repeated game, the first-best outcome is also not achievable. This is a strong impossibility result: it holds for games with agents with any discount factor $\delta \in (0, 1)$.

Given that the first-best is not achievable, we then turn our attention to constrained-efficient allocations by investigating the properties of the constrained-efficient PPEs. In any constrained-efficient PPE, crises are always mitigated and the welfare loss arises from the avoidance action not being taken every period. Furthermore, the frequency of avoidance action depends on the agents' discount factor. For low enough discount factors, the active agent never takes the avoidance action at equilibrium. Once the discount factor is above a certain threshold, the active agent takes the avoidance action infinitely often in any constrained-efficient PPE. Moreover, the passive agent has to *bailout* the active agent infinitely often, where “*bailout*” is defined more precisely later. The optimal frequency of avoidance and bailouts is determined endogenously.

In what follows, we discuss the logical reasoning and intuition of these results. The details of all proofs are in Appendix B.

4.1. The impossibility of implementing the first-best

Suppose that in an equilibrium the active agent takes the avoidance action with positive probability. When he takes the action, his expected discounted payoff (v_1) is a convex combination of the expected discounted payoff conditional on the event of a crisis (w_1^1), and that conditional on no crisis (w_1^0). Because taking the avoidance action is costly, it must be the case that w_1^0 is strictly greater than w_1^1 , so that the active agent finds it optimal to incur the cost. Moreover, w_1^1 cannot be too negative because of individual rationality. Using these facts, we show that there is a fixed positive constant γ , such that the ex-post value at the good state for the active agent w_1^0 has to be strictly greater than that of the ex-ante value v_1 by γ ; $w_1^0 > v_1 + \gamma$. That is, whenever the active agent takes the avoidance action as part of a PPE and there is no crisis, his continuation value must increase by at least a fixed amount. Therefore, when there is no crisis for a sufficiently long time, the implied continuation value required for the active agent to be willing to take the avoidance action value stops being feasible.⁵ We thus obtain the following result.

Proposition 2. *There is no PPE in which the active agent takes the avoidance action almost surely at every period along the equilibrium path.*

4.2. Efficient mitigation

It is not possible to have avoidance action played on every period at equilibrium. However, when the discount factor is not too low, a PPE exists in which the active agent sometimes takes the avoidance action (see Lemma B.4 in the appendix for detail). Such “good deed” requires incentives, for instance, punishing the active agent after a crisis, or rewarding him if there is no crisis. Two possible ways to punish the active agent after a crisis are to let him suffer the cost of the crisis (no mitigation), or to ask him to contribute more than necessary to mitigate the crisis (money burning). Our second result is that neither of these forms of punishment schemes is optimal. In every constrained efficient PPE, agents contribute exactly as much as needed to mitigate a crisis when it happens.

Proposition 3. *In any constrained-efficient PPE, crises are efficiently mitigated, that is, $\mu_{1t}(h_t) + \mu_{2t}(h_t) = 1$ almost surely along the equilibrium path.*

This result is very natural since both of these forms of punishment are ex-post inefficient. However, the proof is far from trivial because, given that there is imperfect monitoring, some degree of inefficiency ex-post could be necessary to generate incentives ex-ante. This is a common feature of models with imperfect monitoring that can be traced back to [Green and Porter \(1984\)](#). We obtain the result because we show that there are always more efficient ways to punish the active agent. As it turns out, any incentive scheme that can be generated in equilibrium via no-mitigation or money-burning can also be generated by adjusting the shares of the mitigation cost in the future without incurring any efficiency losses due to either insufficient or excessive mitigation.

⁵ It is crucial for this proposition that the avoidance action is not observable, see Section [6.1](#). Hence, this is a result of moral hazard and not of the structure of the payoffs.

4.3. Bailouts as an incentive mechanism for avoidance

The difficulty in inducing the active agent to take the avoidance action is that the cost is too high for him to pay on his own. The solution seems to be that the passive agent should help pay part of it. In a world with perfectly transferable utility, we could consider schemes where the passive agent directly subsidizes the active agent.⁶ However, we have assumed that the agents' contributions can only be used to mitigate crises. In our environment, the only way for the passive agent to compensate the active agent is by sometimes paying more in mitigation costs after a crisis has occurred. When this happens, we call it a bailout.

Definition 3. A *bailout* in the dynamic game is a situation where a crisis occurs, the agents jointly contribute sufficient resources to mitigate it, and the contribution of the active agent is less than his private crisis loss, i.e., $\mu_{1t}(h_t) + \mu_{2t}(h_t) \geq 1$, and $\mu_{1t}(h_t) < c_1$.

We can show that bailouts are the only form of compensation available, and if the active agent is not compensated, then he has no reason to choose avoidance. It follows that bailouts are not only sufficient to induce the avoidance action, but also necessary.

Proposition 4. In any PPE where the avoidance action is taken with positive probability, bailouts occur with positive probability.

Proposition 4 shows that bailouts are necessary to support avoidance actions, but it says nothing about sufficiency or efficiency. When is it possible to support any avoidance at all? When is it efficient to do so? If the active agent expects to be bailed out in the future as a form of compensation, he may be willing to take the avoidance action, at least in some instances, and such an arrangement is necessary for efficiency when feasible. The following proposition formalizes these results.

Proposition 5. There exists $\tilde{\delta} \in (0, 1)$ such that:

1. If $\delta < \tilde{\delta}$, every PPE (and therefore every constrained-efficient PPE) has avoidance played with probability zero at all periods.
2. If $\delta > \tilde{\delta}$, in every constrained-efficient PPE the avoidance action is played infinitely often, and bailouts take place infinitely often.

Proposition 5 indicates that, for low discount factors, it is *not* possible to induce the active agent to take any avoidance actions, and the set of efficient PPE essentially reduces to a repetition of static Nash equilibria of the stage game. For higher discount factors, avoidance is not only possible, but it is also necessary for constrained efficiency. Propositions 2 and 5 combined imply that when $\delta > \tilde{\delta}$, in any constrained-efficient PPE the active agent takes the avoidance action infinitely often, takes the non-avoidance action infinitely often, and is bailed out infinitely often.

The proof of this result involves two key steps. First, we show that when agents are sufficiently patient, any equilibrium without avoidance is Pareto dominated by one with avoidance and bailouts. The intuition behind this is that a patient active agent can be motivated to bear the immediate cost of avoidance in exchange for the long-term benefits of reduced crisis frequency and future bailouts. At the same time, a patient passive agent is willing to provide bailouts to obtain the long-term benefits of reduced crises. Once we establish that constrained efficient equilibria must include at least one instance of avoidance, the next step is to prove that avoidance must occur infinitely often. This is possible because we can incentivize the active agent to choose avoidance without altering their expected payoff. Although playing avoidance is costly and reduces the active agent's immediate payoff, this can be offset by the decreased probability of a crisis and the prospect of future bailouts. Consequently, if an equilibrium does not feature avoidance after a certain history, it can be improved by introducing avoidance after that history without affecting the active agent's earlier incentives, as their continuation payoffs would remain unchanged. Having established the necessity of having avoidance played infinitely often, we can then utilize Proposition 4 to conclude that bailouts must also occur infinitely often.

5. Automata: endogenous frequency of crises and bailouts

From previous sections, we know that due to the misalignment between the active agent's private incentive and the social objective (represented by $c_1 < \hat{d} < 1$), the optimal arrangement to correct this incentive problem involves bailing out the active agent infinitely often. Two questions arise immediately. The first is implementation. The results from the last section are all qualitative features of the optimal mechanism. What does such a mechanism look like? How can it be implemented? After all, "infinitely often" shirking and bailouts are not actionable plans. The second question is the effectiveness of a mechanism resulting from a constrained efficient PPE if it can be implemented (at least approximately). What are the potential welfare gains it could achieve? How big are the welfare losses relative to the first-best outcome which is not achievable? Furthermore, are these welfare implications sensitive to the model specification? More specifically, when the model parameters vary, such as the avoidance cost (d), the private costs of a crisis to each agent if it is not mitigated (c_1 and c_2), and the effectiveness of the avoidance action in reducing crisis incidence ($\pi_0 - \pi_1$), how would the welfare comparisons change?

In this section, we employ an often-used numerical procedure — finite-state automaton — to address these questions. This procedure explicitly tells us step-by-step how to implement an approximation of a constrained efficient PPE of our infinitely repeated game. It also provides a lower bound on the welfare achievable by the constrained efficient PPE.⁷

⁶ In Section 6.2, we study an extension with perfectly transferable utility and show that the first-best can be achieved when both agents are sufficiently patient.

⁷ An automaton describes a profile of public strategies for the repeated game. If we did not restrict attention to finite-state automata, every public strategy profile could be described by an automaton (Mailath and Samuelson, 2006, pp. 230). However, for computational reasons, we restrict attention to finite-state automata with a fixed upper bound on the number of auxiliary states.

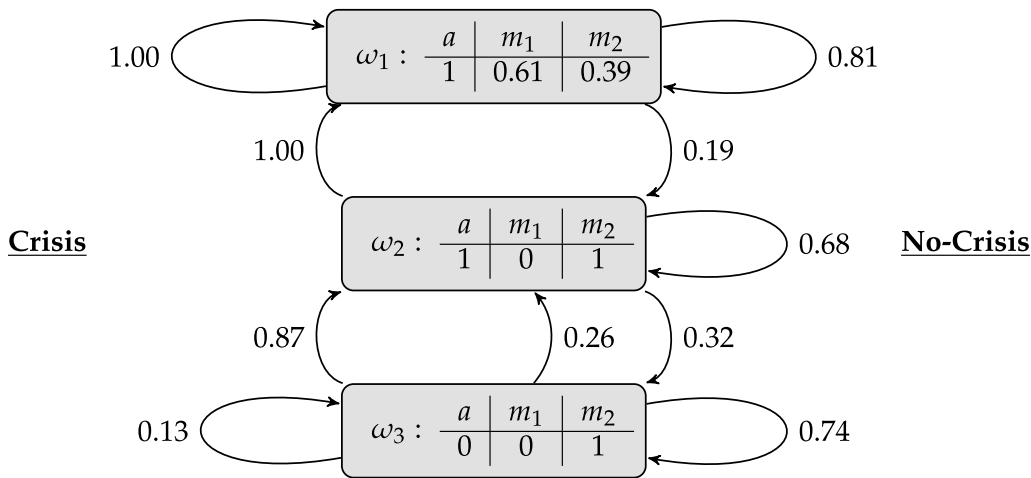


Fig. 2. Automaton.

A finite-state automaton consists of four components: a (finite) set of auxiliary states, an initial distribution over the states, a transition rule, and a mapping from states to actions. Let Ω be a finite set of auxiliary states. States are mapped into agent 1's avoidance action choice $\alpha_1 : \Omega \rightarrow A$ and agent i 's crisis mitigation contribution $\mu_i : \Omega \times X \rightarrow M$, $i = 1, 2$, where $M = [0, 1]$ is implied by Proposition 3. With this definition, the automaton works as follows. At any state $\omega_t \in \Omega$, the active agent takes the action $\alpha_1(\omega_t)$. After a crisis state ξ_t is realized, when there is a crisis ($\xi_t = 1$), the agent i contributes $\mu_i(\omega_t, 1)$ to mitigate the crisis; otherwise, $\mu_i(\omega_t, 0) = 0$, $i = 1, 2$. Note that in each state ω_t , the two agents play a pure-strategy stage game. No randomization at this point. At the end of each period, given the current state ω_t , observable crisis state ξ_t , the two agents' contribution μ_1 and μ_2 , the next period state ω_{t+1} is drawn randomly according to the transition rule $\eta : \Omega \times X \times M \times M \rightarrow \Delta(\Omega)$. The initial state for period-1 is drawn according to the initial distribution $\eta_0 \in \Delta(\Omega)$.

A N -state automaton can be viewed as an approximation of the original repeated game where the two agents' lifetime expected discounted payoffs of any PPE are represented by convex combinations of the payoffs achievable at the N states of the automaton. The two agents use the publicly observable auxiliary states to coordinate actions in each state. The convex combinations of the payoffs across states are achieved by public randomization, via transition rule η . A resulting PPE of a N -state automaton is often called a correlated equilibrium. The higher the number of states N , the better the approximation to the original game. We look for constrained efficient allocations among all PPEs of an automaton game. One advantage of approximating the original model with finite-state automaton is that the solution of constrained efficient PPE itself is an implementation procedure: following the design of the automaton and the equilibrium strategies, the two agents would achieve lifetime expected discounted payoffs that are hopefully close approximations of the constrained efficient PPE of the original game.

In what follows, we show with an example how a 4-state automaton and an associated PPE work.

5.1. An illustrative example of an equilibrium mechanism

Consider an example with the following set of parameters,

$$\delta = 0.95, \quad \pi^1 = 0.2, \quad \pi^0 = 0.9, \quad d = 0.5, \quad c_1 = 0.6, \quad c_2 = 0.5.$$

In the environment with this set of parameters,

- $c_1, c_2 \in (0, 1)$, and $c_1 + c_2 > 1$, so Assumption 1 is satisfied.
- $c_1 = 0.6 < \hat{d} \approx 0.7143 < 1$, so Assumption 2 is satisfied.

Hence, avoidance and mitigation after a crisis are socially efficient — the first best, which has an expected total cost of 0.7 ($= d + \pi^1 \times 1$). By Proposition 2, the first-best is not obtainable in any PPE. The expected total cost at a static Nash equilibrium is 0.9 ($= \pi^0 \times 1$), which implies a welfare loss (relative to the first-best) of 28.57 percent.

Fig. 2 describes the PPE that minimizes the total expected discounted lifetime cost among all PPEs of a four-state automaton with one of the states being the minimax equilibrium that serves as an off-equilibrium threat. The equilibrium works as follows.

- Agents start in state ω_1 , where the strategy profile is $(a, m_1, m_2) = (1, 0.61, 0.39)$. In this state, agent 1 is supposed to take the avoidance action, but his private cost in the crisis alone does not generate enough incentive to do so since $m_1 = 0.61 < \hat{d} \approx 0.7143$. To induce agent 1 to take the avoidance action, when there is no crisis, the state switches to ω_2 with probability 0.19. State ω_2 is a bailout state: $m_1 = 0 < c_1$. The “reward” of a bailout in the future provides incentives for agent 1 to take the avoidance action now in ω_1 . That is, the probability of going to this bailout state, compensates for the fact that $m_1 = 0.61 < \hat{d}$, aligning the private and social incentives for agent 1 to take the avoidance action.

Table 1
Summary statistics of the PPE.

State ω	Invariant distribution	$u_1(\omega)$	$u_2(\omega)$	$V_1(\omega)$	$V_2(\omega)$	Welfare	Welfare loss (%)
ω_1	0.50	-0.622	-0.078	-0.539	-0.177	-0.716	2.23
ω_2	0.38	-0.500	-0.200	-0.478	-0.250	-0.727	3.87
ω_3	0.12	-0.000	-0.900	-0.420	-0.328	-0.748	6.86
LRA	–	-0.501	-0.222	-0.501	-0.222	-0.724	3.41

Notes: LRA refers to the long-run averages, which correspond to the expected values evaluated using the invariant distribution (the long-run frequencies of the equilibrium in the three on-equilibrium states).

$u_i(\omega)$ denotes agent i 's expected payoff for the period when the state is ω .

$V_i(\omega)$ denotes agent i 's total discounted expected payoff when the state is ω .

The "Welfare loss" is the expected total cost relative to that of the first-best.

- In ω_2 , the strategy profile is $(a, m_1, m_2) = (1, 0, 1)$. Again, agent 1 is supposed to take the avoidance action, but now he has even less incentive to do so since his contribution to mitigation is now zero. This time, to generate sufficient incentive, the equilibrium moves to state ω_3 with probability 0.32 if there is no crisis.
- The state ω_3 has an even stronger form of bailout because the active agent contribution in mitigating crisis is zero, and he takes no avoidance action.⁸
- The fourth state ω_4 , which is not in the figure, has a strategy profile of non-avoidance/no-mitigation (the minimax equilibrium). This state is out of the equilibrium path and works as a punishment state in case of a detectable deviation by either agent.

This automaton PPE is not on the Pareto frontier, but it provides a lower bound on what can be achieved by a constrained efficient allocation. Table 1 provides summary statistics of the equilibrium. In state ω_1 , the expected discounted total cost is 0.716, which is only 2.23 percent greater than the first-best one, 0.7. On average, crisis occurs 28.4 percent of the time, compared to 20 percent at the first-best. When a crisis happens, a bailout occurs 70.3 percent of the time. The striking result is that even though the passive agent bailouts the active agent over 70 percent of the crises, the expected present value of his cost is only 0.17. For comparison, in the best equilibrium for the passive agent with no bailouts, the expected present value of his cost is 0.36. That is, by optimally choosing a bailout policy, the passive agent can reduce his cost with crises by half.

5.2. Comparative statics

The equilibrium displayed in Fig. 2 illustrates how bailouts can be used to induce avoidance in equilibrium. In this subsection, we study how properties of this equilibrium change with key parameters of the model: the avoidance cost d , the private costs of nonmitigated crisis (c_1, c_2), and the probabilities of crises (π^0, π^1). Throughout this analysis, Assumptions 1 and 2 are maintained. For each set of parameters, we found the PPE that minimizes the total discounted long-run cost among six-state automata.⁹ Then, we compare the implied long-run probabilities of avoidance, crisis, and bailouts, as well as the long-run average cost of avoidance and the agents' mitigation payments.

The impact of changes in the avoidance cost

Table 2 displays features of the equilibrium outcome for different avoidance cost d , while keeping other parameters at $\delta = 0.9$, $\pi^1 = 0.2$, $\pi^0 = 0.9$, $c_1 = 0.6$ and $c_2 = 0.5$. As one could expect, when the avoidance cost increases, avoidance action is taken less frequently, therefore, crisis happens more often. The average avoidance cost (column 5) is non-monotone, reflecting the fact that the more costly avoidance action is taken less often and at a slower reduction pace. The incidence of bailouts (column 4) is non-monotone, similar to agent 1's mitigation cost (column 6). These changes reflect the structure of the equilibrium. Bailouts are the mechanism where agent 2 compensates agent 1 for bearing the avoidance cost alone. As d increases, the compensation needed to generate incentives for the avoidance action also increases.

The impact of changes in the private costs of a crisis

Table 3 displays features of the equilibrium outcome for different values of (c_1, c_2) while keeping other parameters of the model at $\delta = 0.9$, $\pi^1 = 0.2$, $\pi^0 = 0.9$, and $d = 0.5$. When c_1 and/or c_2 increase, both agents' minimax payoff decreases. The impact of c_1 on the equilibrium outcome is substantial. Increasing c_1 from 0.55 to 0.65 reduces the long-run expected total cost from about 106.4 percent to 101.8 percent of the first-best; the incidence of crisis is reduced by approximately one-third (from 0.36 to 0.24); and bailouts are reduced to 60 percent from 89 percent. The reason for these changes is that higher c_1 improves the alignment of agent 1 private cost of a crisis, c_1 , with the social cost of a crisis, the total mitigation cost of 1. Agent 2 private cost of crisis, c_2 , has little effect on the equilibrium outcome since he is a passive agent and has no private information that impedes equilibrium efficiency.

⁸ The probability of moving to a state preferred by agent 1 is always higher when there is no crisis. Hence, the automata are reminiscent of the revision strategies used in Rubinstein and Yaari (1983) and Radner (1985).

⁹ We used in Section 5.1 a four-state automaton since it helps to illustrate the equilibrium dynamics. We now move to a six-state automaton for better computational precision. We found no significant increase in welfare from further increasing the number of states.

Table 2
The impact of changes in the avoidance cost.

d	$P(a = 1)$	$P(\xi = 1)$	$P(m_1 < c_1)$	$E(d)$	$E(m_1)$	$E(m_2)$	Expected total cost (%)
0.45	0.9569	0.2302	0.4647	0.4306	0.0784	0.1518	101.66
0.50	0.8571	0.3000	0.7937	0.4285	0.0340	0.2660	104.08
0.55	0.7598	0.3682	0.8942	0.4179	0.0182	0.3500	104.81
0.60	0.7115	0.4019	0.9301	0.4269	0.0143	0.3877	103.61
0.65	0.6851	0.4204	0.8727	0.4453	0.0227	0.3977	101.85

Note: The probabilities and expectations are evaluated using the implied invariant distribution.
The “Expected total cost” is expressed as a percentage of the first-best.

Table 3
The impact of changes in the private costs of a crisis.

c_1	c_2	$P(a = 1)$	$P(\xi = 1)$	$P(m_1 < c_1)$	$E(d)$	$E(m_1)$	$E(m_2)$	Expected total cost (%)
0.55	0.5	0.7763	0.3566	0.8870	0.3882	0.0177	0.3388	106.39
	0.7	0.7729	0.3590	0.8875	0.3864	0.0176	0.3414	106.49
0.60	0.5	0.8571	0.3000	0.7937	0.4285	0.0340	0.2660	104.08
	0.7	0.8529	0.3029	0.7119	0.4265	0.0391	0.2638	104.20
0.65	0.5	0.9363	0.2446	0.6136	0.4681	0.0655	0.1791	101.82
	0.7	0.9349	0.2456	0.5420	0.4675	0.0722	0.1733	101.86

Note: The probabilities and expectations are evaluated using the implied invariant distribution.
The “Expected total cost” is expressed as a percentage of the first-best.

Table 4
Total cost above the first-best (%).^a

$\pi^0 \backslash \pi^1$	0.65	0.70	0.75	0.80	0.85	0.90	0.95
0.10	3.25	5.71	8.49	9.05	8.00	2.94	—
0.15	—	2.70	4.65	6.46	6.59	4.41	1.28
0.20	—	—	2.20	3.68	4.67	4.54	2.74
0.25	—	—	—	1.75	2.89	3.30	2.89
0.30	—	—	—	—	1.38	2.19	2.26

^a Long-run average as percentage of the first-best.

The impact of changes in crisis probabilities

Table 4 displays the long-run expected total cost above that of the first-best for different combinations of (π^0, π^1) , while other parameters are set to $\delta = 0.95$, $d = 0.5$, $c_1 = 0.6$, and $c_2 = 0.5$. The cells with the symbol “—” represent the cases where the parameters do not satisfy Assumption 2. The effect of (π^0, π^1) on the welfare cost is not uniform. Combinations of (π^0, π^1) , with $\hat{d} (= d\pi^0 - \pi^1)$ either closer to c_1 or 1, lead to lower cost. This means that sometimes decreasing π^0 reduces the welfare cost, while other times increasing π^0 reduces the welfare cost. The same is true for π^1 .

On the other hand, for combinations of (π^0, π^1) with the same $\hat{d}$ (that is, $\pi^0 - \pi^1$ constant), higher π^0 and π^1 always lead to a lower welfare cost. The interpretation of these results is more subtle. One might think that higher π^1 means that avoidance is less effective in preventing crisis, which could imply a higher cost, but this is not true. The correct measure is $\pi^0 - \pi^1$ — the reduction of crisis probability by taking the avoidance action. With $\pi^0 - \pi^1$ held constant (corresponding to cells along diagonal and off-diagonals), the only impact is increasing π^0 . Higher π^0 implies that the minimax utility of agent 1 is lower, hence, it is easier to generate incentives for avoidance.

The impact of changes in agents’ discount factor

Table 5 displays features of the equilibrium outcome for different values of the discount factor δ , while other parameters are set to $\pi^1 = 0.2$, $\pi^0 = 0.9$, $c_1 = 0.6$, $c_2 = 0.5$, and $d = 0.5$. As one could expect, lower δ is associated with lower welfare. Increasing δ from 0.6 to 0.9 decreases the total cost by about 4 percent of the first-best. This pattern reflects that a higher δ implies a higher future payoff as well as punishment, so agents are more willing to cooperate.

6. Alternative mechanisms

We have shown that the first-best cannot be achieved as a PPE of the repeated game, and that whenever avoidance is possible in equilibrium, every constrained efficient PPE involves bailouts infinitely often. In this section, we consider two alternative

Table 5
The impact of changes in agents' discount factor.

δ	$P(a = 1)$	$P(\xi = 1)$	$P(m_1 < c_1)$	$E(d)$	$E(m_1)$	$E(m_2)$	Expected total cost (%)
0.6	0.7111	0.4022	0.5978	0.3555	0.1063	0.2960	108.25
0.7	0.7630	0.3659	0.8345	0.3815	0.0636	0.3023	106.77
0.8	0.7979	0.3415	0.7491	0.3989	0.0444	0.2970	105.78
0.9	0.8571	0.3000	0.7937	0.4285	0.0340	0.2660	104.08

Note: The probabilities and expectations are evaluated using the implied invariant distribution.

The "Expected total cost" is expressed as a percentage of the first-best.

mechanisms that can help to improve welfare.¹⁰ We analyze one model where the avoidance action is perfectly observed, and one where the passive agent can directly subsidize the active agent. In both cases, some forms of compensation (either bailouts or direct transfers) are still necessary for the active agent to take the avoidance action. However, unlike our benchmark model, these alternative specifications admit PPE that achieves the first-best when agents are patient enough.

6.1. The avoidance action is observable

In the benchmark model, we assumed that the avoidance action of the active agent is private. The passive agent could only make imperfect inferences about it via the realization of crises. Now, consider the alternative specification where a is perfectly observable to both agents. This allows agents to use strategy profiles conditional on a , in particular, bailing out the active agent if and only if he takes the avoidance action. This additional possibility does not change the fact that bailouts are necessary for avoidance.

Proposition 6. *In any PPE of the game with observable actions, if the avoidance action is taken with a positive probability, bailouts occur with positive probability.*

To illustrate the difference from the unobservable-action case, consider the following simple strategy profile. Along the equilibrium path, the active agent always takes the avoidance action and crises are always mitigated.

$$(\alpha_t(h_t), \mu_{1t}(h_t), \mu_{2t}(h_t)) = (1, m_1^*, 1 - m_1^*),$$

for some fixed constant $m_1^* > 0$, which is specified in Appendix B.7. After any deviation, the active agent chooses $a = 0$ forever after, and both agents never again make positive mitigation contributions. We show in the appendix that, if the discount factor is sufficiently high, then a PPE with such a grim trigger strategy exists. Along this equilibrium path, avoidance action is taken every period and mitigation always happens if there is a crisis. Hence this strategy profile implements the first-best.

Proposition 7. *There exists $\bar{\delta}' \in (0, 1)$ such that, if $\delta > \bar{\delta}'$, the game with observable actions admits a PPE where the active agent takes the avoidance action at every period and after every history.*

6.2. Monetary transfers

The previous analysis depends crucially on the assumption of nontransferable utility. In particular, if the active agent takes the avoidance action, he has to pay the cost d by himself. Moreover, both agents' contributions to mitigation can only be used to clean up crises. Suppose we relax this assumption by allowing the passive agent to directly transfer resources to the active agent for consumption. More precisely, suppose that at any date t and after any history h_t , agent 2 can make a transfer $\beta_{t1}(h_t) \geq 0$ to agent 1 if there is a crisis and a transfer $\beta_{t0}(h_t) \geq 0$ otherwise. These transfers enter the stage-game payoffs as an additive term. That is, the stage-game payoffs for the active (passive) agent in the game with transfers are exactly those from the game without transfers plus (minus) whatever transfer he receives (makes).

A version of Proposition 4 continues to hold in this modified model. For the active agent to be willing to take the avoidance action, he must expect some form of compensation. The only difference is that the passive agent now has new forms of compensation available. He can still compensate the active agent by bailout — contributing sufficient resources to mitigation so that the cost to the active agent is less than c_1 in case of a crisis. Additionally, agent 2 can transfer resources to agent 1 in periods where there are no crises. Any equilibrium with avoidance must involve at least one of these forms of compensation.

Proposition 8. *In any PPE of the game with transfers where the active agent takes the avoidance action with positive probability, agent 2 compensates agent 1 by having either $\beta_{t0}(h_t) > 0$ or $\mu_{1t}(h_t) - \beta_{t1}(h_t) < -c_1$, or both with positive probability.*

¹⁰ Another interesting extension is one where the passive agent can impose punishments to the active agent. If punishments are a sunk cost, then it does not implement the first best. If such punishments take the form of transfers from the active agent to the passive agent, like a fee paid in case of a crisis, this exercise is isomorphic to the one with monetary transfers. We omit this extension for this reason.

To illustrate the difference from the nontransferable-utility case, consider the following simple strategy profile for the game with transfers.

$$(\alpha_t(h_t), \mu_{1t}(h_t), \mu_{2t}(h_t)) = (1, m_1^*, 1 - m_1^*),$$

and

$$(\beta_{t0}(h_t), \beta_{1t}(h_t)) = (b^*, 0),$$

for all t and every h_t along the equilibrium path, where $m_1^* \in (0, 1)$ and $b^* > 0$ are fixed constants specified in Appendix B.8. That is, agent 1 always takes the avoidance action and contributes m_1^* when there is a crisis, and agent 2 compensates agent 1 with b^* units of consumption when there is no crisis. The transfer b^* can be viewed as a subsidy to agent 1 from agent 2 in no-crisis time. In case of a detectable deviation, the agents switch to play the one-shot Nash equilibrium with no avoidance and no mitigation forever. We show in the appendix that, if the discount factor is high enough, this strategy profile constitutes a PPE of the game with transfers. Hence, when the agents are patient enough, the first-best is attainable in equilibrium.

Proposition 9. *There exists $\bar{\delta}'' \in (0, 1)$ such that, if $\delta > \bar{\delta}''$, the game with transfers admits a PPE where the active agent takes the avoidance action at every period and after every history.*

This subsidy scheme is simple theoretically but may not be easy to implement in reality. For example, it might be difficult to justify paying Greece's government every period — subsidy in normal times and mitigation in crisis times to the public.

7. Conclusion

Our study of a liability-sharing problem between two asymmetric parties in an infinitely repeated game has highlighted some key points. First, the presence of shirking and bailouts may be necessary to achieve constrained efficient outcomes within a social arrangement (e.g., the European Monetary Union). Insisting on eliminating these elements may not be realistic due to high incentive costs. Second, stochastic shirking and bailout can be essential features of an approximately efficient outcome. It does so by adopting a differential approach towards bailouts, favoring institutions with a clean track record and exercising caution when considering bailouts for those with a history of crises. Coordination between the active and passive players can be accomplished using n -state automata and correlated equilibrium, with the levels of mitigation contribution and the transition probabilities serving as fine-tuning tools for incentive provision. Our numerical simulations demonstrate that the equilibrium of the n -state automata can achieve high levels of welfare relative to the first-best. Last, even if one relaxes the assumption of nontransferable utility, the payment from the passive player to the active player simply shifts from ex-post to ex-ante but does not disappear.

While we acknowledge the lack of empirical evidence on the use of such an arrangement, we believe there are hints that it is at play in reality. Policymakers are generally more inclined to be lenient towards institutions with a good track record. A rigorous empirical exercise to test this hypothesis could be a valuable avenue for future research. Our model is also highly schematic and lacks realistic features such as different maturities of debt instruments, sovereign default, and other fiscal policies. This is by design. The simplified model economy allows us to effectively illustrate the role of stochastic bailouts in repeated games. Future research can further enhance our understanding of the use of bailouts in generating incentives by introducing these more realistic features.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Appendix A and B. Recursive analysis of the set of perfect public equilibrium and proofs

Supplementary material related to this article can be found online at <https://doi.org/10.1016/j.euroecorev.2025.104999>. Appendix A shows that Perfect Public Equilibria can be characterized recursively. Appendix B contains all proofs.

Data availability

No data was used for the research described in the article.

References

- Abreu, Dilip, Pearce, David, Stacchetti, Ennio, 1986. Optimal cartel equilibria with imperfect monitoring. *J. Econom. Theory* 39 (1), 251–269.
- Abreu, Dilip, Pearce, David, Stacchetti, Ennio, 1990. Toward a theory of discounted repeated games with imperfect monitoring. *Econometrica* 58 (5), 1041–1063.
- Brunnermeier, Markus K., Garicano, Luis, Lane, Philip R., Pagano, Marco, Reis, Ricardo, Santos, Tano, Thesmar, David, Van Nieuwerburgh, Stijn, Vayanos, Dimitri, 2016. The sovereign-bank diabolic loop and ESBies. *Am. Econ. Rev.* 106 (5), 508–512.
- Chari, V.V., Kehoe, Patrick J., 2016. Bailouts, time inconsistency, and optimal regulation: A macroeconomic view. *Am. Econ. Rev.* 106 (9), 2458–2493.
- Diamond, Douglas W., Dybvig, Philip H., 1983. Bank runs, deposit insurance, and liquidity. *J. Political Econ.* 91 (3), 401–419.
- Ennis, Huberto M., Keister, Todd, 2009. Bank runs and institutions: The perils of intervention. *Am. Econ. Rev.* 99 (4), 1588–1607.
- Farhi, Emmanuel, Tirole, Jean, 2012. Collective moral hazard, maturity mismatch, and systemic bailouts. *Am. Econ. Rev.* 102 (1), 60–93.
- Farhi, Emmanuel, Tirole, Jean, 2018. Deadly embrace: Sovereign and financial balance sheets doom loops. *Rev. Econ. Stud.* 85 (3), 1781–1823.
- Green, Edward J., 2010. Bailouts. *FRB Richmond Econ. Q.* 96 (1), 11–32.
- Green, Edward J., Porter, Robert H., 1984. Noncooperative collusion under imperfect price information. *Econometrica* 52 (1), 87–100.
- Keister, Todd, 2016. Bailouts and financial fragility. *Rev. Econ. Stud.* 83 (2), 704–736.
- Mailath, George J., Samuelson, Larry, 2006. *Repeated Games and Reputations: Long-Run Relationships*. Oxford University Press.
- Radner, Roy, 1985. Repeated principal-agent games with discounting. *Econometrica* 53 (5), 1173–1198.
- Rubinstein, Ariel, 1979. An optimal conviction policy for offenses that may have been committed by accident. In: *Applied Game Theory: Proceedings of a Conference At the Institute for Advanced Studies, Vienna, June 13–16, 1978*. Springer, pp. 406–413.
- Rubinstein, Ariel, Yaari, Menahem E., 1983. Repeated insurance contracts and moral hazard. *J. Econom. Theory* 30 (1), 74–97.
- Sannikov, Yuliy, 2008. A continuous-time version of the principal-agent problem. *Rev. Econ. Stud.* 75 (3), 957–984.
- Schneider, Martin, Tornell, Aaron, 2004. Balance sheet effects, bailout guarantees and financial crises. *Rev. Econ. Stud.* 71 (3), 883–913.
- Spear, Stephen E., Srivastava, Sanjay, 1987. On repeated moral hazard with discounting. *Rev. Econ. Stud.* 54 (4), 599–617.
- Watson, Joel, 2021. Theoretical foundations of relational incentive contracts. *Annu. Rev. Econ.* 13 (1), 631–659.
- Williams, Noah, 2015. A solvable continuous time dynamic principal-agent model. *J. Econom. Theory* 159, 989–1015.