



Guarantees, risk shifting and credit crunches[☆]

Caio Machado^a, Ana Elisa Pereira^{b,*}

^a Instituto de Economía, Pontificia Universidad Católica de Chile, Chile

^b School of Business and Economics, Universidad de los Andes, Chile

ARTICLE INFO

JEL classification:

G21

G28

Keywords:

Credit crunches

Moral hazard

Risk shifting

Guarantees

Global games

ABSTRACT

This paper analyzes the role of guarantees when a policymaker wants to avoid a credit crunch, but at the same time is concerned that its policies may lead to excessive risk taking. Banks face a coordination problem in their investment decisions and are allowed to risk-shift to projects with smaller expected return and higher volatility. Government guarantees can avoid a freeze in funding at the cost of increasing risk-shifting behavior. When fundamentals are sufficiently low, the policymaker prefers not to intervene and a credit crunch happens. The moral hazard problem makes guarantees more powerful, reducing the amount of guarantees needed to avoid a coordination failure.

1. Introduction

Recessions frequently come with credit crunches, understood as scenarios where the funding for productive projects quickly dries up. Such events are often associated with coordination failures. If agents are pessimistic about economic activity, they may refrain from investing and lending, and expectations may become self-fulfilling. As a response, regulators often provide different types of guarantees to cover financial intermediaries' losses in bad times.¹ One possible motivation for such measures is to restore market confidence, enhancing coordination and lending. The overall effects of those policies are unclear, though. As many times argued, the flip side of guarantees is that they might lead to excessive risk taking, distorting financial institutions' investment decisions. Therefore, to better understand the effect of those policies, it seems crucial to assess: (i) how such policies affect the probability of credit crunches directly by incentivizing bank lending; (ii) how they affect banks' asset choices (risk taking); and (iii) how possible changes in banks' risk-taking behavior feed back into the probability of credit crunches.

This paper contributes to the literature on financial fragility by studying the interplay between banks' portfolio choices, credit crunches

and government guarantees. We propose a simple model that captures the feedback between those variables and find that even though banks risk-shift in response to the announcement of guarantees, this policy is still effective in reducing the probability of a credit crunch. However, in bad times, moral hazard distortions eliminate the benefits of avoiding credit freezes, so policymakers should refrain from intervening in such states. In intermediate states, guarantees are desirable as a policy instrument to avoid credit freezes, and in fact the moral hazard distortion makes the policy more effective: a lower level of guarantees suffices to avoid a crisis.

In our model, banks face strategic complementarities in their investment decisions, leaving room for coordination failures. Banks have many projects available, some riskier and with lower expected return than others, leading to the possibility of excessive risk taking. More specifically, a continuum of banks have essentially two investment opportunities: a risk-free asset and a risky project. The riskless asset yields a zero net return (we often refer to this outside option as *not investing*). The expected return of the risky investment opportunity depends positively on two variables: the proportion of banks investing in risky projects and economic fundamentals.

[☆] We thank Luis Araujo, Braz Camargo, Itay Goldstein, Bernardo Guimaraes, Felipe Iachan and seminar participants at SBE (Florianópolis) for helpful comments and suggestions. Machado gratefully acknowledges financial support from the Sao Paulo Research Foundation (Fapesp) through grants 2012/23222-6 and 2013/22873-6, and from ANID/Conicyt through grant Fondecyt Regular 1220615. Pereira gratefully acknowledges financial support from ANID/Conicyt through grant Fondecyt Regular 1240859. A previous version of this paper circulated under the title "Guarantees, Risk Taking and Market Freezes".

* Corresponding author.

E-mail addresses: caio.machado@uc.cl (C. Machado), apereira@uandes.cl (A.E. Pereira).

¹ For instance, the Small Business Administration loan guarantees in the US cover a fraction of banks' losses in case of borrower default. During the 2008 crisis, we have witnessed different policies such as asset purchases programs and bank bailouts aimed at attenuating intermediaries' losses. During the COVID-19 crisis, different government loan guarantees were promised to banks across the globe (e.g., the UK Coronavirus Business Interruption Loan Scheme).

If banks are not sufficiently confident that others will invest, profitable projects do not find the required funding (a credit freeze). Under the assumption that fundamentals are common knowledge, we have the possibility of multiple equilibria if fundamentals are in an intermediate range. In those situations, it is not clear when and why a credit freeze will occur. By assuming that agents have a small amount of private information about fundamentals, we are able to select a unique equilibrium in the model using global games techniques, and hence, the probability of a credit crunch is uniquely determined.²

Importantly, instead of a single type of risky project, we assume that there is a continuum of risky projects with different probabilities of failure. Riskier projects yield a higher return in case of success. The expected return is increasing in the probability of success. Therefore, in the absence of any intervention, agents choose the least risky (highest-quality) project, and this additional dimension of banks' choices on the asset side plays no role. However, with government guarantees, banks' ability to choose among projects opens room for risk-shifting.

As usual in settings with strategic complementarities, there is room for government interventions aiming to eliminate coordination failures. We analyze a policy in which the government commits to cover part of the losses incurred by banks in case a project fails, which directly reduces the probability of credit freezes by affecting banks' incentives to fund projects. This policy has a side effect on the choice of risky projects itself, which affects the probability of a credit crunch indirectly. Moral hazard arises from the fact that, when the government commits to bail out banks in case of project failure, it can lead them to choose worse projects, with smaller expected return (but higher return contingent on success).

In this policy setting, the following trade-off arises: while higher guarantees are still effective to reduce the probability of credit freezes, they induce excessive risk taking. We show that moral hazard may lead the government to allow a credit crunch to happen, even knowing that banks would be better off if they were able to coordinate and invest in sufficiently good risky projects. The reason is that, in low states, the government needs to offer very high guarantees to avoid a credit crunch, as the return of risky projects is not so attractive for banks. But such high guarantees induce too much risk-shifting, eventually making banks invest in very risky projects that pay less than the outside option in expectation, simply because the government is covering their losses in the (likely) event that those projects fail. Anticipating the risk-shifting such high guarantees would cause, the government prefers to let a credit freeze happen.

Interestingly, in states where policymakers do want to intervene, the possibility of moral hazard reduces the level of guarantees needed to avoid a coordination failure. Intuitively, if banks are allowed to risk shift, guarantees have a direct benefit for banks (for any given project, they expect lower losses in case of project failure, for a given level of fundamentals) and an indirect benefit (projects with a higher upside become more attractive from banks' perspective, as part of the losses in case of failure is covered by guarantees). If the government cannot commit to sufficiently high levels of guarantees, in some cases a credit crunch can be avoided with moral hazard with the help of the indirect effect, but that would not be the case if banks could not risk shift.

We also propose and analyze alternative policies that do not suffer from the moral hazard problem, namely, interest rate cuts, capital injections and direct subsidies to investment. Although those policies do not imply moral hazard distortions, they feature other limitations, discussed in the text, which may prevent them from fully eliminating

inefficient credit crunches. Finally, we provide some robustness exercises, namely a model with a more general functional form for banks' returns and a model with different assumptions on the precision of public information, and show that our main insights remain.

Related literature. Government interventions that aim to restore markets functioning have been studied in different contexts. A large part of this literature focuses on how interventions can overcome adverse selection in asset markets (examples include Philippon and Skreta, 2012, Tirole, 2012 and Guerrieri and Shimer, 2014). In Camargo et al. (2016) a trade-off arises from the fact that policies that unfreeze markets harmed by adverse selection may deteriorate the informational content of prices. Ahnert and Kuncel (2024) uncover an additional benefit of government loan guarantees, as in their setting that policy creates a positive externality on liquidity in the market for non-guaranteed loans.

Our model is more related to another branch of the literature, that focuses on how policies can help to mitigate coordination failures when investment and lending decisions present some sort of strategic complementarity. Bebczuk and Goldstein (2011) study policy interventions that aim to mitigate coordination failures in a model of credit freezes similar to Morris and Shin (2004). Our model builds on those papers, but our key contribution on the theoretical front is to study the interplay between banks' risk taking on the asset side, coordination failures and government guarantees. Here, banks' risk taking is endogenous, and all these outcomes are jointly determined. We show that the interaction of risk taking and coordination failures have important implications for the trade-offs policymakers face when implementing guarantees.

Cheng and Milbradt (2012), in a dynamic model built upon He and Xiong (2012), study risk shifting in a context of debt rollover freezes. Similarly to our model, there managers can risk-shift, moving towards a riskier project with smaller expected return. The authors show that debt maturity should be such that it does not lead managers to risk-shift in good times but, at the same time, does not increase too much rollover risk. Their work is related to ours as they consider the moral hazard issues arising from a third party's interventions. However, differently from their paper, we focus on government guarantees that directly protect banks on their asset side and that can be easily implemented through asset purchases that require no information on banks' asset returns. Our simpler and static framework allows us to analytically characterize the optimal policy.

Also, in a setting with coordination failures, Corsetti et al. (2006) contribute to the debate concerning moral hazard distortions due to liquidity provision in international financial crises. They study the role of the IMF as the international lender of last resort during international financial crises, contradicting the conventional view that liquidity support always induces moral hazard. Their results suggest that the availability of liquidity funds can even encourage the government to implement desirable reforms which would otherwise be too costly and risky.

Our paper also relates to recent work on optimal regulations in models of bank runs where policies impact banks' decisions. In Carletti et al. (2023), guarantees affect a bank's endogenous choice of monitoring effort, which impacts its asset returns. In fact, they find that guarantees in general increase banks' monitoring.³ In their model, there are no strategic complementarities in banks' investment decisions, which is crucial in our analysis. There, there is no coordination problem among banks, only among depositors. Allen et al. (2018) study government guarantees and their effects on banks' exposure to runs.

² We follow Morris and Shin (2003) to deal with equilibrium selection. Following the seminal work of Carlsson and Van Damme (1993), many authors have applied what became known as global games techniques to study speculative attacks (Morris and Shin, 1998; Guimaraes and Morris, 2007), bank runs (Goldstein and Pauzner, 2005), optimal subsidies (Sákovics and Steiner, 2012), among others.

³ This happens because they look at long-term guarantees, meaning that they are only due if the bank is not forced to liquidate its assets because of a depositor run. Since better borrower monitoring helps to avoid those runs, banks have higher incentives to incur private costs to monitor as guarantees increase.

Their model focuses on how government guarantees affect the contracts banks offer to creditors (the liability side of banks' balance sheet), while we focus on how guarantees affect the banks' choices of projects (the asset side).⁴ Dávila and Goldstein (2023) analyze optimal deposit insurance (guarantees to depositors) in a Diamond and Dybvig (1983) setting, while we focus on optimal guarantees on the asset side of banks. Kashyap et al. (2024) study optimal regulation of banks' assets and liabilities in the presence of credit and run risk, while leaving aside the role of government guarantees.⁵

The remainder of this paper is organized as follows. In the next section we present the model. Section 3 presents the policy exercise and the main results. Section 4 discusses alternative policies. Section 5 presents a model with more general assumptions and extend the main results to that setting. Section 6 concludes.

2. Model

There is a unit-mass continuum of risk-neutral banks, indexed by i . Each bank has one unit of capital to invest. They can invest it in a risk-free asset (bonds), which yields a gross return of 1, or they can provide capital to operating firms with projects with uncertain return. There is a continuum of types of those projects available, indexed by $\pi \in (0, 1]$. We refer to them as risky projects or risky assets. Each risky project has an idiosyncratic probability of success that is increasing in π . In case a project succeeds, it yields a gross return (per dollar invested) given by a function $\tilde{R}(\pi, \theta)$, where $\tilde{R} : (0, 1] \times \mathbb{R} \rightarrow \mathbb{R}$. In case it fails, its gross return is zero. The variable θ is the aggregate state (fundamental) of the economy. We assume the following functional form for $\tilde{R}$:

$$\tilde{R}(\pi, \theta) = \frac{\theta}{\pi} - \frac{\pi}{2} + 1. \quad (1)$$

Note that this function is decreasing in π for positive values of θ . Thus, riskier projects (with lower π) yield higher return in case of success.⁶ Equation (1), while stylized, is a particularly tractable way to capture this. We adopt this functional form to deliver our main insights in the clearest possible way, with closed-form solutions. In Section 5, we extend the model to a more general setting and show that our main insights remain.

The success probability of a project depends positively on the amount of banks investing in risky projects. There are several ways to justify this assumption. If many firms (projects) in the economy get funded, that generates a higher demand and profits for other producers (as in Cooper and John, 1988 and Kiyotaki, 1988). Alternatively, if projects produce goods that are part of a supply chain, the lack of funding for some producers may reduce one's connections and disrupt its demand and supply of inputs (as in Taschereau-Dumouchel, 2025). The important feature here is that there are positive spillovers when projects of operating firms find the required funds. We denote by l the proportion of banks investing in risky projects. Let $p(\pi, l)$ denote the probability of success of project π for a given l . We assume that it is given by

$$p(\pi, l) = \begin{cases} \pi & \text{if } l \geq b, \\ \gamma\pi & \text{if } l < b, \end{cases} \quad (2)$$

⁴ Keister (2016) and Keister and Narasiman (2016) analyze a problem similar to the one in Allen et al. (2018), but differently from the latter, in those papers the probability of runs is exogenously determined by a sunspot.

⁵ In a broader sense, our paper also relates to a large literature studying optimal policies and financial fragility in coordination settings where policies intended to improve coordination often feature unintended consequences due to anticipation effects (recent examples include Schilling, 2023, Matta and Perotti, 2024 and Zhong and Zhou, 2025, to name a few).

⁶ If $\theta < 0$ this is not true, but in that case no one will undertake any risky project, since they could get a higher return by investing in the risk-free asset. Thus, the case where the choice between risky projects is relevant will not include situations where $\theta < 0$.

where $0 < b < 1$ and $0 < \gamma < 1$. When $l < b$, there is a credit freeze along the lines of Bebhuk and Goldstein (2011): banks fail to coordinate on providing enough funding to profitable projects of operating firms, buying safe bonds instead. In that situation, the probability of success of any risky project that happens to get funded is smaller. If we imposed that all projects have $\pi = 1$, our model would essentially be that of Bebhuk and Goldstein (2011).⁷

Letting $R(\pi, \theta)$ be the gross return of project π after all uncertainty realizes, its expectation conditional on some given l can be written as

$$E_l[R(\pi, \theta)] = \begin{cases} E[\theta] - \frac{\pi^2}{2} + \pi & \text{if } l \geq b, \\ \gamma \left(E[\theta] - \frac{\pi^2}{2} + \pi \right) & \text{if } l < b. \end{cases} \quad (3)$$

Note that this function is increasing in π for any given l . Hence, π reflects the quality of risky projects: under no interventions, banks will never invest in projects with $\pi < 1$.⁸ Whether a bank will invest in the best project ($\pi = 1$) or in the safe asset depends on its beliefs about l and the value of θ . If $\theta < 1/2$, then investing in the risk free asset is a dominant action—even if a bank were certain all others would invest in risky projects, its expected return would be lower than 1. If $\theta > 1/\gamma - 1/2$, then investing in the risky project is a dominant choice—even if a bank were certain no other bank would invest in risky projects, the expected return would be greater than 1. Assuming that θ is common knowledge, there are multiple equilibria if $1/2 \leq \theta \leq 1/\gamma - 1/2$.

As usual in the global games literature (see Morris and Shin, 2003 for a survey), if we introduce some private information about θ , multiplicity is ruled out. We assume that θ is drawn from a uniform prior on the real line and that each bank i receives two signals: a public signal $y = \theta + \eta$, where η is a normal random variable with zero mean and variance σ_y^2 ; and a private signal $x_i = \theta + \epsilon_i$, where $(\epsilon_i)_{i \in [0,1]}$ are iid normal random variables with zero mean and variance σ_x^2 . Moreover, ϵ_i and η are independent. Hereafter, we focus on the limiting case where σ_x goes to zero.

The timing of events is the following: (i) Nature draws the fundamentals and banks observe the public and private signals; (ii) Banks choose whether to invest in some risky project $\pi \in (0, 1]$ or in the safe asset.

2.1. Equilibrium without interventions

We say that banks *invest* if they choose to fund a risky project, and *do not invest* if they choose the safe asset. Proposition 1 shows that our setting features an essentially unique equilibrium.⁹ Recall that, in the absence of interventions, whenever banks invest in a risky project, they choose $\pi = 1$.

Proposition 1. *Let θ^* be defined as:*

$$\theta^* = \frac{1}{1 - b(1 - \gamma)} - \frac{1}{2}. \quad (4)$$

In equilibrium, as $\sigma_x \rightarrow 0$, banks play according to the cutoff θ^ : they invest if $x_i > \theta^*$ and do not invest if $x_i < \theta^*$.*

⁷ The technical difference would be that in Bebhuk and Goldstein (2011) fundamentals affect the critical mass of capital needed for the project to succeed (the parameter b in our model), while here they affect the project's return in case of success. See also Section 5.

⁸ It is not critical for our results that expected returns are increasing in π or that the optimal choice of project (without interventions) is $\pi = 1$. In the more general setting of Section 5, neither of those features are necessarily present, and our main insights remain. Also, note that even project $\pi = 1$ is still risky since, if the mass of banks investing is not large enough, success probability falls to $\gamma\pi$ (see Eq. (2)).

⁹ We say the equilibrium is *essentially* unique since, as usual in global games, strategies are not uniquely determined in a zero-measure set of states (for $\theta = \theta^*$ in Eq. (4) below).

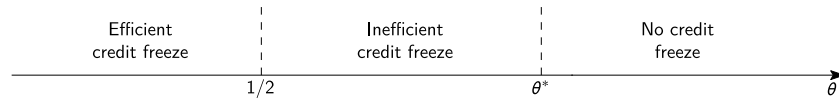


Fig. 1. Credit freezes.

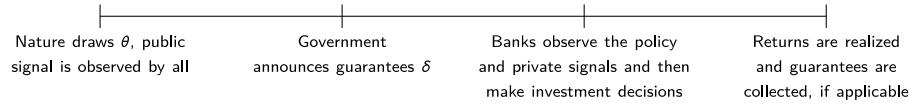


Fig. 2. Timeline.

Proof. The proof follows from Proposition 2.2 in Morris and Shin (2003), which implies that the cutoff θ^* can be computed by imposing that when $\theta = \theta^*$ agents must be indifferent between the (best) risky project and the risk-free asset if they have a uniform belief over l . Therefore, given that agents will choose $\pi = 1$, θ^* must be such that

$$b\gamma\left(\theta^* + \frac{1}{2}\right) + (1-b)\left(\theta^* + \frac{1}{2}\right) = 1,$$

which implies the desired result. $\square$

The result above implies that when $\sigma_x \rightarrow 0$, we have $l = 1 \geq b$ if the realization of θ is higher than θ^* , and $l = 0 < b$ if the realization of θ is smaller than θ^* . In other words, whenever $\theta < \theta^*$, there is a credit freeze.

Note that the cutoff θ^* is between the two dominance region boundaries, i.e., $\theta^* \in (1/2, 1/\gamma - 1/2)$. As illustrated in Fig. 1, we say that if θ is between $1/2$ and θ^* , there are coordination failures (or an inefficient credit freeze), since everyone would be better off if the risky project were undertaken. Agents are failing to coordinate, afraid others might not invest. Therefore, there is room for interventions that aim to mitigate coordination failures. The next section points out that government guarantees can alleviate this problem, but there are some trade-offs involved.

Throughout the remainder of the paper, we adopt the tie-breaking convention that banks invest in a risky project when indifferent between investing and not investing (that is, they invest when observing signals right at the equilibrium cutoff).

3. Government guarantees

We assume in this section that the government commits to pay a part of the loss incurred by a bank in case its project fails. The policy analyzed here is similar to real world bailout policies that aim to rescue banks in case they suffer hard losses. We restrict attention to policies that do not cover more than the capital invested in the project. The policy is the following: after observing the public signal y , the government commits to buy at a price $\delta \in [0, 1)$ any risky asset after the realization of its return. Banks observe the announced policy before choosing how to invest. Note that no information about the banks' assets or returns is needed for the regulator to implement this policy, making it particularly attractive.¹⁰ We summarize the timeline of the model with guarantees in Fig. 2.

Agents will never invest in a project if they believe that $\tilde{R}(\pi, \theta) < \delta$, since they could get $1 > \delta$ with certainty by investing in the risk-free asset. Since banks are infinitely well informed about θ , it will never happen that a bank whose project was successful sells its asset to

the government. Thus, in practice, the government will just pay those guarantees in case some project fails. Gross expected returns become

$$E_l[R(\pi, \theta)] = \begin{cases} E[\theta] - \frac{\pi^2}{2} + \pi + (1-\pi)\delta & \text{if } l \geq b, \\ \gamma\left(E[\theta] - \frac{\pi^2}{2} + \pi\right) + (1-\gamma\pi)\delta & \text{if } l < b. \end{cases} \quad (5)$$

We first study the equilibrium when there is no possibility of moral hazard. To do so, we assume that the only risky project available is the one with $\pi = 1$. Then, we analyze the general case where banks can risk shift.

3.1. Benchmark case: No moral hazard

Assume agents are restricted to invest either in the risk-free asset or in projects with $\pi = 1$. Proposition 2 below is the analogous of Proposition 1 under the government intervention.

Proposition 2. Assume that the government commits to a guarantee $\delta \in [0, 1)$ and banks are not allowed to invest in projects with $\pi \neq 1$. Then, the results in Proposition 1 apply, but the cutoff θ^* is now given by

$$\theta_B^* = \frac{1 - b(1-\gamma)\delta}{1 - b(1-\gamma)} - \frac{1}{2}, \quad (6)$$

which is a strictly decreasing function of δ .

Proof. Note that banks only collect guarantees if they invest in the risky project, $l < b$ and the project happens to fail. The latter happens with probability $1 - \gamma$ conditional on $l < b$. Following the same arguments in the proof of Proposition 1 and using (5), θ_B^* must then satisfy

$$b\left[\gamma\left(\theta_B^* + \frac{1}{2}\right) + (1-\gamma)\delta\right] + (1-b)\left(\theta_B^* + \frac{1}{2}\right) = 1,$$

which yields the desired result. $\square$

3.2. General case

We now turn to the case where banks optimally choose in which risky project to invest, given the government guarantee. If $\delta > 0$, a bank investing in the risky asset will no longer choose the higher expected return project, $\pi = 1$. Since the government is covering some losses, banks have incentives to take excessive risk. The next lemma shows the choice of project that maximizes a bank's expected returns conditional on it investing in a risky asset.

Lemma 1. Given an announced guarantee δ , whenever banks invest in a risky project, their project choice is

$$\pi^* = 1 - \delta. \quad (7)$$

Proof. Suppose a bank chooses to invest in a risky project. Given some belief about l , the bank's expected return is given by (5), and it is strictly concave on π . Taking the first order condition on expected returns, conditional on both $l \geq b$ and $l < b$, yield $\pi^* = 1 - \delta$. $\square$

¹⁰ An alternative policy would be for the government to provide a subsidy to banks funding risky projects, but that would require the assumption that the government observes banks' investment choices. We focus on another kind of policy that does not require such detailed information and is widely used by policy makers. See Section 4 for a more detailed discussion of other policies.

Lemma 1 shows that the more guarantees the government commits to, the worse are the projects chosen by banks, who turn to projects with lower expected value, but higher payoff in case of success.

Proposition 3 characterizes banks' decisions between investing in project π^* or investing in the safe asset.

Proposition 3. Assume the government commits to guarantees $\delta \in [0, 1)$. Then, the results in [Proposition 1](#) apply, but the cutoff θ^* is now given by

$$\theta_G^* = \frac{1 - b(1 - \gamma)\delta}{1 - b(1 - \gamma)} - \frac{1}{2} - \frac{\delta^2}{2}, \quad (8)$$

which is a decreasing function of δ .

Proof. Following the same arguments in the proof of [Proposition 1](#) and using (5), θ_G^* must satisfy

$$b \left[\gamma \left(\theta_G^* - \frac{\pi^{*2}}{2} + \pi^* \right) + (1 - \gamma\pi^*)\delta \right] + (1 - b) \left[\theta_G^* - \frac{\pi^{*2}}{2} + \pi^* + (1 - \pi^*)\delta \right] = 1,$$

with π^* given by (7), which yields the desired result. $\square$

Note that the cutoff in the model with moral hazard, θ_G^* in (8), is very similar to the one of the benchmark case, θ_B^* in (6), except for the term $-\frac{\delta^2}{2}$. When agents are allowed to risk-shift, the effect of the intervention over θ^* is larger. The direct effect of the intervention is captured by the term $\frac{-b(1-\gamma)\delta}{1-b(1-\gamma)}$ in both thresholds: with guarantees, banks' expected payoffs of investing are higher, since the government is covering part of the losses in case a project fails; hence, banks are willing to invest for lower fundamentals—the cutoff moves to the left as δ increases. The indirect effect, only present in the model with moral hazard, is captured by the term $-\frac{\delta^2}{2}$: Banks now have an extra benefit due to the chance to “gamble with government money”—if the project succeeds, the payoff is higher for projects with lower π , and the downside in case of failure is protected by guarantees. This also increases the expected payoff of investing in risky projects, moving the cutoff further to the left.

We summarize the result that moral hazard makes guarantees more powerful in reducing the range of states for which a credit crunch occurs in the following corollary.

Corollary 1. For any $\delta \in (0, 1)$, a marginal increase in δ has a stronger effect on the cutoff θ_G^* than on the benchmark cutoff θ_B^* , that is,

$$\frac{\partial \theta_G^*}{\partial \delta} < \frac{\partial \theta_B^*}{\partial \delta} < 0.$$

Proof. See [Appendix A.1](#). $\square$

3.3. Optimal guarantees

A strategy for the government is a map from its signal y to a level of guarantees $\delta \in [0, 1)$. Here, for tractability, we also take the limit as the noise in the public signal vanishes ($\sigma_y \rightarrow 0$), after taking the limit of the private noise ($\sigma_x \rightarrow 0$) as we have done so far. As usual in global games, the order of that limit here possibly matters for some realizations of the fundamental, since to ensure equilibrium uniqueness one needs private information to be sufficiently more precise than public information, as shown by [Morris and Shin \(2003\)](#).¹¹ [Appendix B](#) discusses the case where σ_y does not vanish, yielding similar insights. Since the government observes an infinitely precise signal regarding θ , it assigns probability one that $y = \theta$ in the limit. Thus, to simplify the exposition we sometimes replace the public signal y by the actual value of θ (one can think that the planner actually knows the fundamental).

In what follows, we characterize optimal guarantees in the benchmark without risk shifting and in the main model with moral hazard. In both cases, the government objective is to maximize the expected total wealth in the economy, that is, the sum of banks' expected returns. We often refer to the government objective as welfare. Also, to ease the exposition, we adopt the tie-breaking convention than when indifferent between $\delta = 0$ or some $\delta > 0$, the government chooses the former.

3.3.1. Optimal guarantees in the benchmark without moral hazard

Suppose it is possible to prohibit banks from investing in projects with $\pi \neq 1$ (our benchmark case of [Section 3.1](#)). Given the assumption that signals are infinitely precise, using (3), $\pi = 1$ and [Proposition 2](#), the government chooses $\delta \in [0, 1)$ to maximize

$$\mathcal{W}_B = \begin{cases} 1 & \text{if } y < \theta_B^*, \\ y + \frac{1}{2} & \text{if } y \geq \theta_B^*, \end{cases} \quad (9)$$

where θ_B^* is given by (6). Since $\sigma_y \rightarrow 0$, we could replace y by the actual value of θ in the expression above. Note that for $\theta \leq 1/2$, the risk-free asset yields (weakly) higher returns, so in that area it is optimal to set $\delta = 0$. Similarly, if $\theta \geq \theta^* = \theta_B^*|_{\delta=0}$, all banks invest and earn more than the risk-free return even without guarantees, so $\delta = 0$ is also optimal. The next proposition presents the solution to the government problem in the coordination failure area.

Proposition 4. Suppose $\theta \in (1/2, \theta^*)$. In the benchmark model with $\pi = 1$, the optimal policy is to always offer enough guarantees to avoid a coordination failure, and the lowest optimal guarantee is

$$\underline{\delta}(\theta) = 1 - \frac{[1 - b(1 - \gamma)]}{b(1 - \gamma)} \left(\theta - \frac{1}{2} \right).$$

Proof. See [Appendix A.2](#). $\square$

If it were possible to enforce the project choice $\pi = 1$ whenever banks invest, the government would always be willing to avoid a coordination failure. Hence, whenever $\theta \in (1/2, \theta^*)$, it is optimal to offer the amount of guarantees needed to lower θ_B^* until it reaches θ , thus triggering investment in the risky project.¹²

3.3.2. Optimal guarantees with moral hazard

Now, in the model with multiple projects with different levels of risk, the government faces a trade-off: a higher δ induces more moral hazard, in the sense that banks move towards riskier projects with smaller expected return (lower quality projects). But also, a higher δ reduces the values of θ for which coordination failures occur. Thus, in principle, even if θ is in the coordination failure interval, the government might not want to give guarantees to avoid the credit crunch, since banks could potentially gamble with taxpayer funds by choosing projects with smaller expected return than the risk-free asset.

Using [Lemma 1](#) and [Proposition 3](#), the government chooses $\delta \in [0, 1)$ to maximize the expected total wealth, which in this case can be written as

$$\mathcal{W} = \begin{cases} 1 & \text{if } y < \theta_G^*, \\ y - \frac{(1-\delta)^2}{2} + (1 - \delta) & \text{if } y \geq \theta_G^*, \end{cases} \quad (10)$$

where θ_G^* is given by (8). The second line is just the expected return of projects when there is no coordination failure, already considering that agents will choose $\pi^* = 1 - \delta$. Again, since $\sigma_y \rightarrow 0$, we can replace y by θ above. [Proposition 5](#) pins down the government decision for every realization of θ .

¹¹ For a given $\sigma_x > 0$ the equilibrium set of the coordination game may not be unique as $\sigma_y \rightarrow 0$, but for any $\sigma_y > 0$ the equilibrium is unique as $\sigma_x \rightarrow 0$, as shown in [Proposition 2.2](#) in [Morris and Shin \(2003\)](#).

¹² The government is indifferent between offering that minimum guarantee or any higher guarantees that makes $\theta_B^* < \theta$, since those also avoid the coordination failure and yield the maximum possible level of welfare.

Proposition 5. *Offering positive guarantees improves welfare if, and only if, θ is in a non-empty interval (θ°, θ^*) , where $\theta^\circ > 1/2$. In that case, the government commits to the smallest value of δ that implements the threshold $\theta_G^* = \theta$, which is given by*

$$\delta^*(\theta) = \sqrt{\frac{1}{[1 - b(1 - \gamma)]^2} - 2\theta} - \frac{b(1 - \gamma)}{1 - b(1 - \gamma)}. \quad (11)$$

Proof. See Appendix A.3. $\square$

The results of Proposition 5 are illustrated in Fig. 3. It shows that the optimal policy never avoids a credit crunch in some neighborhood of $1/2$, in contrast to the case with no moral hazard. There is a range of fundamentals, $\theta \in (1/2, \theta^\circ]$, for which the government deliberately chooses not to intervene despite knowing that a credit crunch will happen in the absence of intervention. Note that, by promising enough guarantees, the government could avoid an inefficient credit freeze for all $\theta \in (1/2, \theta^*)$, since θ_G^* converges to zero as δ approaches one, but for $\theta < \theta^\circ$ it prefers not to do so. The intuition for the optimality of not intervening when $\theta \in (1/2, \theta^\circ]$ stems from the main trade-off facing the authority. On one hand, guarantees increase the expected return of any risky project for banks, which incentivizes them to fund those projects and alleviates the coordination issue. On the other hand, guarantees distort the choice of banks, who move towards riskier projects (but with higher return in case of success) when their downside is protected by government guarantees. In light of this trade-off, the government optimally chooses the optimal guarantee.

Recall that to the left of $\theta = 1/2$, banks have a dominant strategy to invest in the safe asset. For $\theta \in (1/2, \theta^*)$, banks are better off if they coordinate on investing in (sufficiently good) risky projects, but in the absence of guarantees, they fail to do so. If fundamentals are relatively low, $\theta \in (1/2, \theta^\circ)$, large guarantees would be needed to overcome the coordination problem and induce investment in risky projects. But that level of guarantees required to solve the coordination problem generates too much moral hazard: Banks would distort their project choice too much, reducing aggregate welfare below its no-intervention level. The authority then gives up avoiding the coordination failure for $\theta \in (1/2, \theta^\circ)$, since the welfare gains of preventing the credit crunch would not compensate the welfare loss induced by the risk-shifting behavior of banks.

For higher fundamentals, $\theta \in (\theta^\circ, \theta^*)$, the level of guarantees needed to avoid the coordination failure is not so large, and consequently the moral hazard distortion is not so severe. Although banks move towards riskier projects in equilibrium, the welfare gains of avoiding the credit crunch still compensate. The optimal policy in that range implements the minimum amount of guarantees that ensures coordination in funding risky projects, given by (11).

The next corollary compares the guarantees needed to avoid a coordination failure with and without moral hazard, in the range where positive guarantees are optimal in both settings.

Corollary 2. *Consider $\theta \in (\theta^\circ, \theta^*)$. The level of guarantees needed to avoid a coordination failure are lower in the presence of moral hazard, i.e., $\delta^*(\theta) < \underline{\delta}(\theta)$.*

Proof. See Appendix A.4. $\square$

Fig. 3 also illustrates that result, showing the minimal level of guarantees $\underline{\delta}(\theta)$ to which the government commits in the benchmark case without moral hazard together with the optimal guarantee with moral hazard, $\delta^*(\theta)$. Without moral hazard, higher guarantees are needed to prevent a credit crunch, since there is only the direct effect of guarantees discussed in Section 3.2. The possibility of “betting for resurrection” is an additional incentive for banks to finance risky projects, so it is possible to induce coordination and avoid a credit crunch with a lower level of guarantees in that case. Thus, an additional insight of our results is that, if the government is in financial distress and cannot

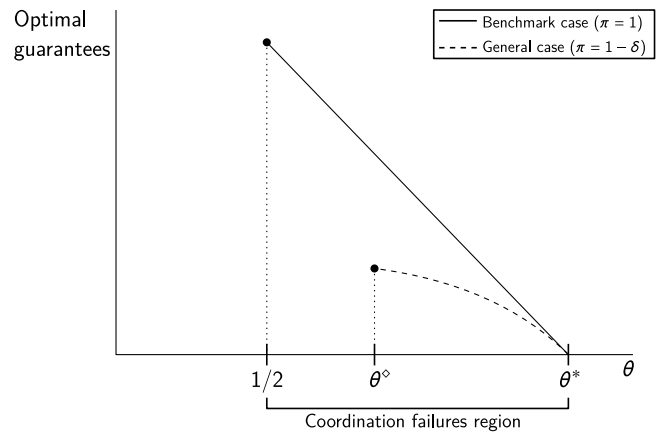


Fig. 3. Optimal level of guarantees.

credibly commit to a high level of guarantees, it might actually benefit from the risk-shifting behavior of financial institutions.¹³

Naturally, when the government has enough funds, moral hazard is detrimental to welfare for $\theta \in (\theta^\circ, \theta^*)$, as total expected wealth is lower with risk-shifting under the optimal policy—subtracting (9) and (10), the welfare loss due to risk-shifting is $\delta^*(\theta)^2/2$. However, an implication of Corollary 2 is that, if the government is unable to commit to guarantees as high as $\underline{\delta}(\theta)$, but is able to offer $\delta^*(\theta)$, inducing coordination is only possible in a world with risk shifting, and hence, in such circumstance, moral hazard improves welfare.

4. Alternative policies

We now analyze some additional policies with the potential to overcome the moral hazard problem implied by guarantees, discussing their benefits and limitations.

4.1. Interest rate cuts

Suppose that banks' outside option is a (safe) government bond that pays $1 + r$, rather than 1 as previously assumed. Our framework allows us to compare guarantees with a policy that lowers the interest rate on government bonds. Note that changing r does not affect banks' choice of risky project: the expected return of risky projects in (3) is independent of r , and is strictly increasing in π for any l . Hence, regardless of r , and in the absence of other interventions, if a bank invests in a risky project, it chooses $\pi = 1$. Following the same procedure as in the proof of Proposition 1, the cutoff played by banks must satisfy the indifference condition:

$$b\gamma \left(\theta_r^* + \frac{1}{2} \right) + (1 - b) \left(\theta_r^* + \frac{1}{2} \right) = 1 + r, \quad (12)$$

and is thus given by

$$\theta_r^* = \frac{1 + r}{1 - b(1 - \gamma)} - \frac{1}{2}.$$

This cutoff is increasing in r , implying that interest rate cuts shift the cutoff to the left, reducing the range of fundamentals for which a credit crunch occurs. As in Bebchuk and Goldstein (2011), lower rates reduce the probability of a credit crunch, but if there is a lower bound on

¹³ Fiscal issues were very relevant to some countries in the Eurozone during the 2008 financial crisis. See Allen et al. (2015).

interest rates ($r = 0$), interest rate cuts alone cannot eliminate credit crunches, leaving room for other policy tools. Beyond this parallel, our setting brings up a new insight: when feasible, interest rate reductions have an advantage over guarantees, as they avoid the side effect of distorting banks' risk-taking incentives.

4.2. Capital injections

Capital injections are another common policy response to crises that lead to capital losses in the banking sector. Following [Bebchuk and Goldstein \(2011\)](#), we can incorporate capital losses and capital injections in our model as follows: Suppose banks suffer a capital loss $L \in (0, 1)$. Such losses could arise, for instance, if banks themselves suffer a run from depositors. If the government covers a fraction $\alpha \in [0, 1]$ of those losses by injecting new capital into all banks, each bank will have $1 - (1 - \alpha)L$ to invest, instead of the original endowment of 1. To preserve the interpretation of b as the mass of funds that must be allocated to risky projects to ensure a higher success rate, we replace the probability of success in (2) and expected returns in (3) by

$$p(\pi, l) = \begin{cases} \pi & \text{if } l[1 - (1 - \alpha)L] \geq b, \\ \gamma\pi & \text{if } l[1 - (1 - \alpha)L] < b, \end{cases}$$

and

$$E_l[R(\pi, \theta)] = \begin{cases} E[\theta] - \frac{\pi^2}{2} + \pi & \text{if } l[1 - (1 - \alpha)L] \geq b, \\ \gamma \left(E[\theta] - \frac{\pi^2}{2} + \pi \right) & \text{if } l[1 - (1 - \alpha)L] < b, \end{cases} \quad (13)$$

respectively. Here, although α enters the expression for expected returns, it does not alter banks' project choice: again, expected returns are increasing in π —regardless of α , L and l —so banks optimally choose $\pi = 1$. Hence, capital injections do not induce moral hazard. The next proposition presents the equilibrium for a given loss L and capital injection α .

Proposition 6. *When banks suffer a capital loss $L \in (0, 1)$ and the government covers a fraction $\alpha \in [0, 1]$ of those losses, the results in [Proposition 1](#) apply, but with the equilibrium cutoff given by*

$$\theta_\alpha^* = \begin{cases} \frac{1}{\gamma} - \frac{1}{2} & \text{if } 1 - (1 - \alpha)L < b, \\ \frac{1}{1 - b(1 - \gamma)} - \frac{1}{2} & \text{if } 1 - (1 - \alpha)L > b. \end{cases} \quad (14)$$

Proof. See [Appendix A.5](#). $\square$

Larger capital injections α either have no effect—if capital losses are too large or capital injections are insufficient, $1 - (1 - \alpha)L < b$ —or they reduce the range of fundamentals where a credit crunch occurs ($\frac{d\theta_\alpha^*}{d\alpha} < 0$ for $1 - (1 - \alpha)L > b$).

The insights from this analysis are similar to the case of interest rate cuts. Large enough capital injections always reduce the range of fundamentals for which a credit crunch occurs (as α approaches 1, $1 - (1 - \alpha)L > b$ and thus $\frac{d\theta_\alpha^*}{d\alpha} < 0$). Such policy may also be limited: if capital injections exceeding the capital losses of banks are fiscally or politically infeasible (so $\alpha \leq 1$), coordination failures cannot be fully eliminated—note that for $\alpha = 1$, we are back to our original setting, which features coordination failures. Still, to the extent that capital injections are enough to avoid a credit crunch, they have an advantage over guarantees as they do not distort banks' risk-taking decisions.

An additional point, not related to the capital injection itself, is that capital losses by banks—for instance, due to a depositor run, which we can interpret in reduced form as an increase in L —have the potential to generate a credit crunch, as the cutoff θ_α^* is strictly decreasing in L if $1 - (1 - \alpha)L > b$.

4.3. Direct subsidies to investment

Suppose the government pays a subsidy s to a bank that funds any risky project instead of purchasing the safe asset. Mathematically, this policy adds a constant s to each line of the expression for expected returns in (3), with no implication for banks' project choice, and hence whenever a bank invests, it chooses $\pi = 1$. The indifference condition pinning down the equilibrium threshold becomes

$$b\gamma \left(\theta_s^* + \frac{1}{2} \right) + (1 - b) \left(\theta_s^* + \frac{1}{2} \right) + s = 1, \quad (15)$$

which yields

$$\theta_s^* = \frac{1 - s}{1 - b(1 - \gamma)} - \frac{1}{2}.$$

Comparing Eq. (15) with (12), one can see that offering a positive s has the same effect on the threshold as imposing a negative interest rate r . In fact, large enough subsidies to investment could eliminate coordination failures altogether—which occurs when θ_s^* reaches the lower dominance region boundary, $1/2$.

There are, however, a few caveats with such a policy. First, it requires the government to be able to observe banks' investment decisions. Otherwise, banks would claim the subsidies and continue to invest in the risky project only when to the right of the original threshold of [Proposition 1](#), leaving the probability of a credit freeze unchanged.¹⁴ Hence, when compared to guarantees, direct subsidies might be more difficult to implement. Guarantees can operate like an asset purchase program, and banks have a natural incentive to turn to the government to offload failed projects, so its implementation does not rely on the government having knowledge about banks' investments and returns, as discussed in [Section 3](#).

Second, direct subsidies involve transfers in every state of the world where investment takes place, regardless of project success or failure, which could pose an extra fiscal burden for the government.

Third, and relatedly, if the government has a tight budget, it may not be able to avoid a coordination failure using subsidies, but precisely due to the extra incentive that the possibility of risk-shifting creates for investment (see [Corollaries 1](#) and [2](#)), it may still be able to avoid the credit crunch with guarantees.

To illustrate this last point, [Fig. 4](#) shows the fiscal cost of the (minimum) subsidy policy and guarantee policy that avoid a coordination failure in a numerical example. Recall that under the optimal guarantee $\delta^*(\theta)$, for $\theta \in (\theta^\circ, \theta^*)$, there is no credit crunch, and banks invest in project $\pi^* = 1 - \delta^*(\theta)$. If projects are independent, the fraction of projects that fail in equilibrium when the realized fundamental is $\theta \in (\theta^\circ, \theta^*)$ is $1 - \pi^* = \delta^*(\theta)$, and hence the fiscal cost of the policy is $[\delta^*(\theta)]^2$. For subsidies, the minimum spending necessary to avoid the credit crunch is the value of s satisfying $\theta = \theta_s^*$, and a credit crunch may be avoided in the whole range $(1/2, \theta^*)$ if budget allows.

In the numerical example of [Fig. 4](#), the budget available for policy is 0.1 (10% of banks' capital). Note that, if fundamentals are sufficiently high (above B in the figure), the government can afford either policy, and abstracting from the observability issues mentioned, the best alternative for the government is to give subsidies to avoid the credit crunch, since it does not distort project choice. If, instead, fundamentals are in an intermediate range (between A and B in the figure), the government budget is insufficient to cover the subsidies needed to avoid the coordination failure. In that case, the best feasible policy is to give out guarantees: precisely because of the moral hazard distortion, that instrument is more powerful in incentivizing investment ([Corollary 1](#)), and it is cheaper to avoid the credit crunch through guarantees.

¹⁴ If investment is unverifiable, banks have incentives to claim the subsidies regardless of whether they invest or not, so the indifference condition in (15) would have $1 + s$ on the right-hand side as well, and the cutoff would be given by the original θ^* in (4).

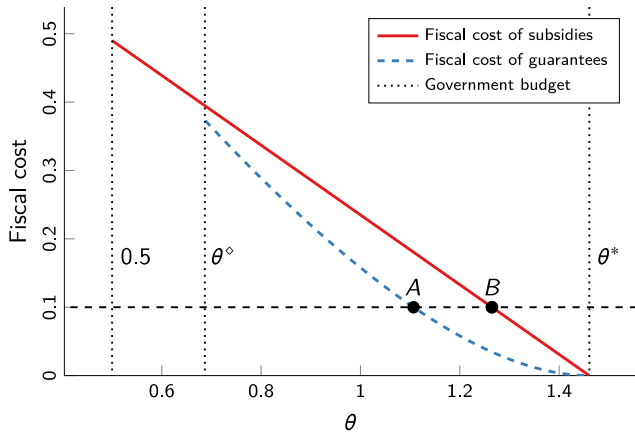


Fig. 4. Fiscal cost of subsidies and guarantees.

Notes. Parameters are $b = 0.7$ and $\gamma = 0.3$, which implies that $\theta^* = 1.461$ and $\theta^\circ = 0.687$.

Although banks risk shift, avoiding the coordination failure through guarantees still pays off since $\theta \in (\theta^\circ, \theta^*)$. For lower fundamentals (below A), with a tight budget, neither alternative is feasible, and a coordination failure occurs.

5. Extension: A more general setting

In this section, we extend the analysis of Sections 2 and 3, allowing for a more general functional form for returns. The only difference between the model presented here and the model of Sections 2 and 3 are the assumptions about the functions that determine returns and the probability of success of a project ($\tilde{R}(\cdot)$ and $p(\cdot)$). All other assumptions and notation remain the same.

We now assume that in case of success, each project $\pi \in [0, 1]$ delivers a gross return $\tilde{R}(\pi) > 1$. The function $\tilde{R} : [0, 1] \rightarrow \mathbb{R}$ is assumed to be continuously differentiable and to have a bounded derivative in its domain. In case of failure, projects still yield a gross return of zero. The probability of success $p(\cdot)$ now depends on π , l and on the fundamental θ :

$$p(\pi, l, \theta) = \begin{cases} \pi & \text{if } l \geq b(\theta), \\ 0 & \text{if } l < b(\theta), \end{cases}$$

where $b(\theta)$ is a strictly decreasing and differentiable function. Note that in the previous model $b(\theta)$ was constant, and the fundamental θ appeared directly in the final return of the firm (conditional on success). Here, it is more tractable to assume the fundamental affects the critical mass needed to avoid a coordination failure b , as in [Bebchuk and Goldstein \(2011\)](#), which is a common assumption in global games (see [Morris and Shin, 2004](#) and [Sákovics and Steiner, 2012](#), for other examples). Also, we assume here that the project fails with certainty in case of a coordination failure ($l < b(\theta)$), which is useful for technical reasons, as it helps ensure the existence of a lower dominance region for any level of guarantees $\delta \in (0, 1)$.

Let $\hat{R}(\pi) = \pi \tilde{R}(\pi)$ represent the expected (gross) return of the project without guarantees and assuming no credit freeze occurs ($l \geq b(\theta)$). We assume that $\hat{R}(\pi)$ is strictly concave and that there exists $\pi_E \in (0, 1]$ such that $\frac{d\hat{R}(\pi_E)}{d\pi} = 0$. These assumptions ensure the optimal choice of project is always characterized by a first-order condition, without guarantees. Note that the project π_E represents the most efficient project, as it has a higher expected return than any other project conditional on $l \geq b(\theta)$ (for $l < b(\theta)$, all projects pay zero).

To ensure the existence of dominance regions, which are needed to use global games techniques, we impose $\lim_{\theta \rightarrow -\infty} b(\theta) > 1$ and $\lim_{\theta \rightarrow \infty} b(\theta) < 0$. For future reference, we define $\underline{\theta}$ as the solution to

$b(\underline{\theta}) = 1$ and $\bar{\theta}$ as the solution to $b(\bar{\theta}) = 0$. Note that those bounds determine the dominance regions without guarantees: If banks know $\theta > \bar{\theta}$, then it is optimal to invest in some risky project under any belief about the actions of others. Conversely, if banks know that $\theta < \underline{\theta}$, then it is always optimal not to invest in risky projects.

Proposition 7 below generalizes the results in [Propositions 1–3](#), [Lemma 1](#) and [Corollary 1](#).

Proposition 7. As $\sigma_x \rightarrow 0$:

1. Without guarantees, banks play according to a cutoff given by θ^* , where θ^* is the unique solution to

$$b(\theta^*) = 1 - \frac{1}{\pi_E \tilde{R}(\pi_E)}.$$

2. If the government commits to guarantees $\delta \in [0, 1)$ and banks are not allowed to choose projects with $\pi \neq \pi_E$, then the equilibrium cutoff is given by the unique θ_B^* that solves

$$b(\theta_B^*) = 1 - \frac{1 - \delta}{\pi_E (\tilde{R}(\pi_E) - \delta)}. \quad (16)$$

Moreover, $\frac{d\theta_B^*}{d\delta} < 0$.

3. Given an announced guarantee $\delta \in [0, 1)$, whenever banks invest in a risky project, their project choice is given by the unique π^* that solves

$$\pi^* \frac{d\tilde{R}(\pi^*)}{d\pi} + \tilde{R}(\pi^*) = \delta. \quad (17)$$

Moreover, $\frac{d\pi^*}{d\delta} < 0$.

4. When the government commits to guarantees $\delta \in [0, 1)$, the equilibrium cutoff is given by the unique θ_G^* that solves

$$b(\theta_G^*) = 1 - \frac{1 - \delta}{\pi^* (\tilde{R}(\pi^*) - \delta)}. \quad (18)$$

where π^* is such that (17) holds. Moreover, $\frac{d\theta_G^*}{d\delta} < 0$.

5. For any $\delta \in (0, 1)$, $\theta_G^* < \theta_B^*$.

Proof. See [Appendix A.6](#). $\square$

The proposition above shows that the predictions of the main model regarding the equilibrium for a given level of guarantees extend to the model of this section. In particular, it implies: (i) Banks choose inefficient projects with higher risk when guarantees are introduced, that is, $\pi^* < \pi_E$ for $\delta > 0$; (ii) The more guarantees the government provides, the larger the range of states for which banks invest in risky projects, since θ_G^* decreases with δ ; (iii) Guarantees are more effective in avoiding credit crunches when banks can risk-shift, since $\theta_G^* < \theta_B^*$ for any $\delta \in (0, 1)$.

Note that, without guarantees, inefficient credit freezes then happen when $\theta \in (\underline{\theta}, \theta^*)$, as in this region all banks would be better off if they could all agree to invest in any risky project with gross expected return $\pi \tilde{R}(\pi) > 1$.

We now characterize the optimal guarantees, which will necessarily have to balance the trade-off between inducing coordination and limiting moral hazard.

5.1. Optimal guarantees

The problem of the government is the same as before: After observing its (infinitely precise) public signal, the government chooses the level of guarantees that maximizes the expected gross returns of banks, that is, the total wealth generated in the economy, taking as given that, for any level of guarantees announced, banks will play according to the cutoff (18), choosing among risky projects according to (17). Since

returns have changed, the government objective function now needs to be written in its general form:

$$\mathcal{W} = \begin{cases} 1 & \text{if } y < \theta_G^*, \\ \pi^* \bar{R}(\pi^*) & \text{if } y \geq \theta_G^*, \end{cases}$$

where π^* and θ_G^* are given by (17) and (18). Hence, the government chooses $\delta \in [0, 1)$ to maximize $\mathcal{W}$ subject to (17) and (18).

The next proposition characterizes the optimal guarantee policy. As before, we assume that if the government is indifferent between $\delta = 0$ or some $\delta > 0$, it chooses the former. Also, it is useful to define the following property on $\bar{R}(\pi)$ —which may or may not hold—for future reference:

$$\pi \bar{R}(\pi) = 1 \implies \pi \frac{d\bar{R}(\pi)}{d\pi} + \bar{R}(\pi) < 1. \quad (\text{P1})$$

In (P1), the left-hand side of the inequality is the derivative of expected gross returns in the absence of coordination failures ($\pi \bar{R}(\pi)$) with respect to π . Hence, in words, (P1) holds when $\pi \bar{R}(\pi)$ does not increase much in π for the project that has the same return as the risk-free asset.

Proposition 8. Suppose property (P1) holds. Then, there exists $\theta^\circ \in (\underline{\theta}, \theta^*)$ such that the government pays positive guarantees if, and only, $\theta \in (\theta^\circ, \theta^*)$. If (P1) does not hold, then the government pays positive guarantees if, and only if, $\theta \in (\underline{\theta}, \theta^*)$. In both cases, in the range of states where guarantees are paid, guarantees are a strictly decreasing function of θ .

Proof. See Appendix A.7. $\square$

Therefore, if (P1) holds, in some cases, the government allows inefficient credit freezes to happen, even though it has available a policy instrument that could induce investment in risky projects. The intuition is similar to before: Avoiding those credit freezes would imply too much risk shifting, pushing banks to project with expected return $\pi \bar{R}(\pi)$ lower than the return of the risk-free asset.

However, if (P1) does not hold, then whenever there is an inefficient credit freeze, the government intervenes, even if that implies some risk shifting on banks' choices of projects. The reason is that, in that case, risk-shifting is never strong enough to push banks into projects with expected return $\pi \bar{R}(\pi)$ below one.

Property (P1) is a bit intricate and technical, so it is useful to consider the following examples to fix ideas: (i) $\bar{R}(\pi) = 2.688 - 1.68\pi$ and (ii) $\bar{R}(\pi) = 3.52 - 2.2\pi$. Both functions satisfy all the assumptions we have made, but in case (i) property (P1) holds and in case (ii) it does not. Fig. 5 shows the expected (gross) returns $\pi \bar{R}(\pi)$ under both assumptions.

Note that in both cases, the project that maximizes the expected returns are the same, $\pi_E = 0.8$. However, returns are higher for any $\pi > 0$ in the second example. Hence, in the first example, there is much less margin to induce risk-shifting with guarantees without turning the project socially undesirable, as a relatively small fall in returns can make risky projects pay less than the outside option (which pays one). This fact would not prevent the government from always avoiding an inefficient credit freeze if in the setting of example (i), the choice of project was not much sensitive to guarantees—perhaps then even almost full guarantees ($\delta \approx 1$) would not be enough to make banks risk-shift to projects dominated by the outside option. However, since example (i) satisfies (P1), for large enough guarantees we will reach the point where $\pi \bar{R}(\pi)$ crosses one, since the optimality condition in the choice of projects for banks is $\frac{d\pi \bar{R}(\pi)}{d\pi} = \pi \frac{d\bar{R}(\pi)}{d\pi} + \bar{R}(\pi) = \delta$.

To see this more clearly, the figure shows the point where the slope of each curve is 0.99. That point corresponds to the banks' choice of project when guarantees are set to $\delta = 0.99$. As one can see, in the first example, that guarantee yields expected gross returns below one, and hence the government would not be willing to offer such a high guarantee to avoid a credit freeze. In the second example, even with almost full guarantees, banks do not risk-shift enough for $\pi \bar{R}(\pi)$ to go below one, and the government is always willing to give any guarantee needed to avoid an inefficient credit freeze.

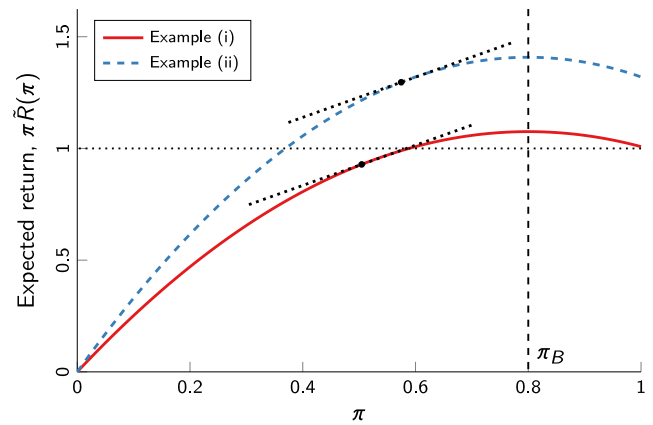


Fig. 5. Expected returns under different assumptions about $\bar{R}(\pi)$.

Notes. Example (i) refers to $\bar{R}(\pi) = 2.688 - 1.68\pi$ and Example (ii) refers to $\bar{R}(\pi) = 3.52 - 2.2\pi$. The points on the depicted tangent lines refer to the point where the derivative of each curve equals 0.99, which denote the projects chosen by banks when guarantees are equal to $\delta = 0.99$.

6. Final remarks

This paper studies the problem of a regulator that would like to avoid a credit crunch, but is concerned with the effects of the opportunistic behavior that its policies may induce. Banks face a coordination problem in their investment decisions and may not fund profitable projects if they are not confident other financial intermediaries in the economy will do the same. Government guarantees that protect banks from bad outcomes may enhance banks' ability to coordinate investment decisions, but at the same time, they worsen the quality of projects that get funded by inducing excessive risk taking. In intermediate states, guarantees are still desirable, despite their moral hazard costs. In low states, the regulator prefers not to intervene since the moral hazard distortions would be too large, and a credit freeze happens. Whenever guarantees are implemented, the amount needed to avoid a credit freeze is smaller in the presence of moral hazard.

We conclude with some conjectures about how our mechanism would withstand in a model where banks are subject to runs on the liability side. As shown in the seminal contribution of Calomiris and Kahn (1991), demandable debt can act as a disciplining device for banks, and hence one could wonder whether depositor runs could help to mitigate the risk-shifting problem in our setting. We conjecture depositor runs would not change our main results for the following reason. In our model, banks risk-shift after guarantees are promised not because bank managers are reaping some private benefit that does not accrue to the returns of the bank's assets. They are doing so precisely because that is the way to maximize bank returns, including the guarantees possibly paid by the government. Therefore, creditors of a bank, whose payoff weakly increases in the bank's final cash flow, would not be more likely to run on banks that risk-shift because of guarantees, hence not disciplining banks' asset choices.

CRedit authorship contribution statement

Caio Machado: Writing – review & editing, Writing – original draft, Visualization, Validation, Supervision, Software, Resources, Project administration, Methodology, Investigation, Funding acquisition, Formal analysis, Conceptualization. **Ana Elisa Pereira:** Writing – review & editing, Writing – original draft, Visualization, Validation, Supervision, Software, Resources, Project administration, Methodology, Investigation, Funding acquisition, Formal analysis, Conceptualization.

Appendix A. Proofs

A.1. Proof of Corollary 1

Using (6) and (8), for $\delta \in (0, 1)$, we have that

$$\frac{d\theta_G^*}{d\delta} = -\frac{b(1-\gamma)}{1-b(1-\gamma)} - \delta < \frac{d\theta_B^*}{d\delta} = -\frac{b(1-\gamma)}{1-b(1-\gamma)} < 0. \quad \square$$

A.2. Proof of Proposition 4

Fix $\theta \in (1/2, \theta^*)$. Recall from (6) that θ_B^* is strictly decreasing in δ and that $\theta_B^*|_{\delta=0} = \theta^*$. From (9), expected welfare for $\theta_B^* \leq \theta$ is larger than one and the expected welfare for $\theta_B^* > \theta$ equals one. Hence, it is optimal to set δ such that $\theta \leq \theta_B^*$, and the minimum optimal level of guarantees is the one that satisfies $\theta = \theta_B^*$. Using (6) and solving $\theta = \theta_B^*$ for δ yields the desired expression for $\underline{\delta}(\theta)$. Finally, note that $\underline{\delta}(\theta) \in (0, 1)$, since $\theta_B^*|_{\delta=0} = \theta^*$, $\theta_B^*|_{\delta=1} = 1/2$ and θ_B^* is strictly decreasing in δ . $\square$

A.3. Proof of Proposition 5

With some abuse of notation, in this proof we write θ_G^* as a function of δ when convenient. First note that if $\theta \notin (1/2, \theta^*)$ it is optimal to give no guarantees: if $\theta \geq \theta^*$ the government knows that agents will invest with probability one, and thus there is no reason to give guarantees and make them choose worse projects. If $\theta \leq 1/2$ the government knows that the return of any risky project is weakly smaller than one, so inducing investment through guarantees cannot increase welfare. Also, for any given θ , we can show that the government will never give a guarantee level greater than $\delta_{\max}(\theta) \equiv \sqrt{2\theta - 1}$. Otherwise, banks would invest in projects with expected return smaller than 1. Therefore, for a given $\theta \in (1/2, \theta^*)$, the government problem is to verify whether it can achieve an equilibrium cutoff equal to θ by giving guarantees smaller than $\delta_{\max}(\theta)$. Whenever $\theta \in (1/2, \theta^*)$, the government is willing to intervene and eliminate coordination failures if there is a $\delta < \delta_{\max}(\theta)$ such that $\theta_G^*(\delta) \leq \theta$. Since $\theta_G^*(\delta)$ is a decreasing function of δ , if $\theta_G^*(\delta) \leq \theta$ for some $\delta \in [0, \delta_{\max}(\theta)]$, then $\theta_G^*(\delta_{\max}(\theta)) \leq \theta$. Thus, given $\theta \in (1/2, \theta^*)$, the government gives guarantees if, and only if,

$$\theta_G^*(\delta_{\max}(\theta)) = \frac{1-b(1-\gamma)\sqrt{2\theta-1}}{1-b(1-\gamma)} - \frac{1}{2} - \frac{2\theta-1}{2} < \theta,$$

which can be expressed as

$$\frac{1-b(1-\gamma)\sqrt{2\theta-1}}{1-b(1-\gamma)} < 2\theta. \quad (19)$$

When $\theta = 1/2$, the inequality above is violated, since $\frac{1}{1-b(1-\gamma)} > 1$. By continuity, we know that there is a non-empty interval $(1/2, \theta^*)$ in which condition (19) does not hold. The left-hand side is a strictly decreasing function of θ and the right-hand side is strictly increasing in θ . We need to show that the left- and the right-hand side are equal at some point $\theta^* < \theta^*$. To see if this is true, it suffices to check if (19) holds when $\theta = \theta^*$, that is, we want to verify whether

$$\frac{1-b(1-\gamma)\sqrt{2\theta^*-1}}{1-b(1-\gamma)} < 2\theta^*.$$

Substituting the expression for θ^* (with no guarantees) in the right-hand side, we get

$$-\sqrt{2\theta^*-1} < 1,$$

which is satisfied since $\theta^* > 1/2$. That is, if $\theta = \theta^*$, condition (19) holds, meaning that there is always an interval (θ^*, θ^*) in which the government is willing to intervene. θ^* is the value that equates the left- and right-hand side of (19). Finally, the optimal value of guarantees, $\delta^*(\theta)$ in Eq. (11), is obtained by simply imposing $\theta_G^* = \theta$ in the threshold equation given by (8) and solving for δ (ignoring the negative root). $\square$

A.4. Proof of Corollary 2

Fix $\theta \in (\theta^*, \theta^*)$. Recall that $\underline{\delta}(\theta)$ and $\delta^*(\theta)$ are the values of δ that solve $\theta_B^* = \theta$ and $\theta_G^* = \theta$, respectively, where those thresholds are given by (6) and (8) (see Propositions 4 and 5). For this proof, with some abuse of notation, it is useful to write θ_B^* and θ_G^* as functions of δ . Note that $\theta_B^*(0) = \theta_G^*(0) = \theta^*$. We then have that, by the fundamental theorem of calculus,

$$\int_0^{\underline{\delta}(\theta)} \frac{d\theta_B^*}{d\delta} d\delta = \theta_B^*(\underline{\delta}(\theta)) - \theta_B^*(0) = \theta - \theta^*$$

and

$$\int_0^{\delta^*(\theta)} \frac{d\theta_G^*}{d\delta} d\delta = \theta_G^*(\delta^*(\theta)) - \theta_G^*(0) = \theta - \theta^*,$$

which we can rewrite as

$$\theta^* - \theta = - \int_0^{\delta^*(\theta)} \frac{d\theta_G^*}{d\delta} d\delta = - \int_0^{\underline{\delta}(\theta)} \frac{d\theta_B^*}{d\delta} d\delta.$$

Since, by Corollary 1, $\frac{d\theta_G^*}{d\delta} < \frac{d\theta_B^*}{d\delta} < 0$ for all $\delta \in (0, 1)$, it must be that $\delta^*(\theta) < \underline{\delta}(\theta)$. $\square$

A.5. Proof of Proposition 6

Note that changes in α are isomorphic to changes in parameter b in the model of Section 2: payoffs in the model with capital losses are the same as in the original model if we replace b by $\hat{b} \equiv b[1 - (1-\alpha)L]^{-1}$, i.e., the fraction of banks that must invest to ensure a higher probability of project success is now $\hat{b}$. If $1 - (1-\alpha)L < b$, $\hat{b} > 1$, and then $l < \hat{b}$ for all $l \in [0, 1]$, so $E_l[R(\pi, \theta)]$ is given by the second line in (13) with probability one. Hence, for $\hat{b} > 1$, using (13) and $\pi = 1$, banks invest if, and only if, $\gamma(\theta + 1/2) \geq 1$, that is, they invest only when fundamentals lie in the upper dominance region, $\theta \geq 1/\gamma - 1/2$.

Consider now $\hat{b} < 1$. In this case, the model with b replaced by $\hat{b}$ satisfies all the assumptions of the main model, and the equilibrium is then identical to that of Proposition 1, but with b replaced by $\hat{b}$, which implies the desired expression for the second line of θ_α^* in (14). $\square$

A.6. Proof of Proposition 7

Before proving each statement separately, it is useful to establish that for any level of guarantees $\delta \in [0, 1)$ there are dominance regions. First, note that even if a bank believes no other bank will invest ($l = 0$), for sufficiently high fundamentals θ , projects succeed with probability π , since $\lim_{\theta \rightarrow \infty} b(\theta) < 0$. Hence, that bank has risky projects available that dominate the risk-free asset, since $\bar{R}(\pi) > 1$ for all $\pi \in [0, 1]$ ensures that $\pi \bar{R}(\pi) > 1$ in a neighborhood of $\pi = 1$. This proves the existence of the upper dominance region. Now, suppose a bank believes $l = 1$ with certainty. Even then, for sufficiently low fundamentals, the success probability of the risky project is zero since $\lim_{\theta \rightarrow -\infty} b(\theta) > 1$. In that case, if it invests in the risky project it gets $\delta < 1$, which is dominated by the outside option that pays one, and hence there exists a lower dominance region. Having established those regions, our model satisfies all the assumptions in Proposition 2.2 in Morris and Shin (2003).

First statement. Note that without guarantees, conditional on investing in a risky project, it is optimal for banks to invest in project π_E as it maximizes both the return in case of a credit freeze ($l < b(\theta)$) and in the complementary case. Following Proposition 2.2 in Morris and Shin (2003), the equilibrium cutoff θ_B^* is computed by finding the indifference condition of an agent that believes $\theta = \theta_B^*$ and that $l \sim U[0, 1]$. This implies the following condition:

$$(1 - b(\theta^*)) \pi_E \bar{R}(\pi_E) = 1,$$

which after solving for $b(\theta^*)$ yields the desired result.

Second statement. Following the same argument as in the first statement, the equilibrium condition is now given by

$$b(\theta_B^*)\delta + (1 - b(\theta_B^*))[\pi_E \tilde{R}(\pi_E) + (1 - \pi_E)\delta] = 1,$$

which after solving for $b(\theta_B^*)$ yields the desired expression. Differentiating with respect to δ we have:

$$\frac{db(\theta_B^*)}{d\delta} = \frac{\tilde{R}(\pi_E) - 1}{\pi_E (\tilde{R}(\pi_E) - \delta)^2},$$

which is strictly positive since $\tilde{R}(\pi_E) > 1$, implying that $\frac{d\theta_B^*}{d\delta} < 0$ as $b(\theta)$ is strictly decreasing.

Third statement. Suppose a bank decides to invest in a risky project and let q denote the probability it assigns to $l \geq b(\theta)$. Note that if the bank decides to invest in a risky project, we must have $q > 0$, otherwise the bank would choose the safe asset, since it would pay more than any guarantee $\delta \in [0, 1]$. Then, the bank chooses π to maximize its expected return given by

$$q[\pi \tilde{R}(\pi) + (1 - \pi)\delta] + (1 - q)\delta. \quad (20)$$

Note that the expression above is strictly concave in π given that $\tilde{R}(\pi) = \pi \tilde{R}(\pi)$ is strictly concave in π . Taking derivatives of the expression above with respect to π and equating it to zero, we can write the first-order condition as

$$\underbrace{\pi \frac{d\tilde{R}(\pi)}{d\pi}}_{\frac{d\tilde{R}(\pi)}{d\pi}} + \tilde{R}(\pi) = \delta. \quad (21)$$

This condition is sufficient for a strict optimal by strict concavity of $\tilde{R}(\pi)$. Strict concavity of $\tilde{R}(\pi)$ also implies that the left-hand side above is strictly decreasing in π . It is also continuous, since $\tilde{R}(\pi)$ is continuously differentiable. Also, note that $\frac{d\tilde{R}(0)}{d\pi} = \tilde{R}(0) > 1$, where the equality follows from the fact that $\tilde{R}(\pi)$ has bounded derivative. Therefore, $\frac{d\tilde{R}(0)}{d\pi} > 1$ together with $\frac{d\tilde{R}(\pi_E)}{d\pi} = 0$ imply that there exists a $\pi = \pi^* \in (0, \pi_E)$ that satisfies (21) for any $\delta \in [0, 1]$. The fact that $\frac{d\pi^*}{d\delta} < 0$ follows immediately from the fact that $\frac{d\tilde{R}(\pi)}{d\pi}$ is strictly decreasing in π .

Fourth statement. The proof to find the expression for the equilibrium cutoff θ_G^* is identical to that of the second statement, replacing π^B by the optimal π^* that satisfies (17).

Taking derivatives of $b(\theta_G^*)$ with respect to δ we get

$$\frac{db(\theta_G^*)}{d\delta} = \frac{(1 - \delta) \frac{d\pi^*}{d\delta} \left[\pi^* \frac{d\tilde{R}(\pi^*)}{d\pi} + \tilde{R}(\pi^*) - \delta \right] + \pi^* [\tilde{R}(\pi^*) - 1]}{[\pi^* (\tilde{R}(\pi^*) - \delta)]^2},$$

which, using (17), simplifies to

$$\frac{db(\theta_G^*)}{d\delta} = \frac{\tilde{R}(\pi^*) - 1}{\pi^* (\tilde{R}(\pi^*) - \delta)^2}.$$

This is strictly larger than zero, since $\tilde{R}(\pi^*) > 1$ and since $\frac{d\tilde{R}(0)}{d\pi} > 1$ ensures $\pi^* > 0$. Since $b(\theta)$ is strictly decreasing, we have $\frac{d\theta_G^*}{d\delta} < 0$.

Fifth statement. Since $b(\cdot)$ is strictly increasing, we need to show that $b(\theta_G^*) > b(\theta_B^*)$, for a given $\delta \in (0, 1)$. Using (16) and (18) and simplifying, $b(\theta_G^*) > b(\theta_B^*)$ boils down to

$$\pi^* (\tilde{R}(\pi^*) - \delta) > \pi_E (\tilde{R}(\pi_E) - \delta).$$

Using (20), note that π^* is the unique value of π that maximizes $\pi (\tilde{R}(\pi) - \delta)$ (as all the other terms in (20) are independent of π). Since $\pi_E \neq \pi^*$ for any $\delta \in (0, 1)$, the inequality above then holds. $\square$

A.7. Proof of Proposition 8

First, note that if $\theta < \underline{\theta}$, risky projects fail with certainty, regardless of banks' decisions, in which case they yield zero gross return. Hence, in that case, the government prefers not to give guarantees ($\delta = 0$), and since $\theta^* > \underline{\theta}$, $l = 0$. If $\theta \geq \theta^*$, then even if $\delta = 0$, $l = 1$. With $\delta = 0$, $\pi^* = \pi_E$, which is the only value of π that maximizes $\pi \tilde{R}(\pi)$, since $\pi \tilde{R}(\pi)$ is strictly concave. Moreover, the project $\pi = \pi_E$ has expected gross returns above one (otherwise, it would yield returns lower than the project $\pi = 1$, which yields an expected gross return above one, contradicting that $\pi = \pi^*$ maximizes $\pi \tilde{R}(\pi)$). Hence, for $\theta \geq \theta^*$, the government also chooses $\delta = 0$, as it leads to the best possible outcome. It remains to be determined the optimal δ in the range $\theta \in (\underline{\theta}, \theta^*)$, and therefore we assume $\theta \in (\underline{\theta}, \theta^*)$ in the remaining of this proof.

Note that as $\delta \rightarrow 1$, $\theta_G^* \rightarrow \underline{\theta}$, and $\theta_G^* = \theta^*$ when $\delta = 0$. Hence, when $\theta \in (\underline{\theta}, \theta^*)$, the government has the option to avoid a credit freeze by giving enough guarantees, but we need to determine when it is optimal to do so. Denote by $\delta^\dagger(\theta)$, for $\theta \in (\underline{\theta}, \theta^*)$, the guarantee required to make $\theta_G^* = \theta$, which is continuous and strictly decreasing in θ given that $\frac{d\theta_G^*}{d\delta} < 0$. Note that $\lim_{\theta \rightarrow \underline{\theta}} \delta^\dagger(\theta) = 1$, since $\lim_{\delta \rightarrow 1} \theta_G^* = \underline{\theta}$, and $\delta^\dagger(\theta^*) = 0$, since θ_G^* and θ^* are equal without guarantees.

We now argue that it is never optimal to set a guarantee above $\delta^\dagger(\theta)$. To see that, first note that $\pi \tilde{R}(\pi)$ is strictly increasing for $\pi < \pi_E$. This comes from strict concavity of $\pi \tilde{R}(\pi)$ and the fact that $\pi = \pi_E$ is a critical point of $\pi \tilde{R}(\pi)$. Since $\frac{d\pi^*}{d\delta} < 0$, conditional on avoiding a credit freeze, the government then always strictly prefers the lowest level of guarantees that achieves that goal, which is $\delta^\dagger(\theta)$. Therefore, the government problem boils down to choosing $\delta = \delta^\dagger(\theta)$ or $\delta = 0$.

Let $\underline{\pi} \in (0, 1)$ be the lowest solution to $\pi \tilde{R}(\underline{\pi}) = 1$. Such value exists, since $\pi \tilde{R}(\pi)|_{\pi=0} = 0$, as $\tilde{R}(0)$ is a real number, and since $\pi \tilde{R}(\pi)|_{\pi=1} = \tilde{R}(1) > 1$. Also note that $\underline{\pi} < \pi_E$. This is so because $\pi_E \tilde{R}(\pi_E) \geq \tilde{R}(1) > 1$ and because $\pi \tilde{R}(\pi)$ increases continuously from $\pi = 0$ to $\pi = \pi_E$. Let δ_{max} be defined as

$$\delta_{max} = \underline{\pi} \frac{d\tilde{R}(\underline{\pi})}{d\pi} + \tilde{R}(\underline{\pi}).$$

That is, if δ were set to δ_{max} , banks would choose $\pi = \underline{\pi}$ (see Eq. (17)), for any $\delta > \delta_{max}$, banks choose projects with $\pi \tilde{R}(\pi) < 1$ (since $\frac{d\pi^*}{d\delta} < 0$), and the opposite is true for $\delta < \delta_{max}$. Therefore, if $\delta^\dagger(\theta) < \delta_{max}$, the government chooses $\delta = \delta^\dagger(\theta)$. If $\delta^\dagger(\theta) \geq \delta_{max}$, the government chooses $\delta = 0$ (here we are using the tie-breaking convention that the government chooses $\delta = 0$ when indifferent).

Suppose first (P1) does not hold. Then $\delta_{max} \geq 1$, and hence $\delta^\dagger(\theta) < \delta_{max}$ for all $\theta \in (\underline{\theta}, \theta^*)$. Therefore, in that case, the government chooses $\delta = \delta^\dagger(\theta)$.

Suppose now (P1) holds. Then, $\delta_{max} \in (0, 1)$. Hence, $\lim_{\theta \rightarrow \underline{\theta}} \delta^\dagger(\theta) = 1 > \delta_{max}$ and $\delta^\dagger(\theta^*) = 0 < \delta_{max}$. By continuity and monotonicity, there is a unique $\theta^\circ \in (\underline{\theta}, \theta^*)$ such that $\delta^\dagger(\theta^\circ) = \delta_{max}$. Hence, the government offers guarantees equal to $\delta^\dagger(\theta)$ if $\theta > \theta^\circ$ and does not offer any guarantees if $\theta \leq \theta^\circ$. $\square$

Appendix B. Optimal guarantees with a finite precision of the public signal

Suppose that, before announcing a guarantee policy, the government observes a public signal $y = \theta + \eta$, with $\eta \sim \mathcal{N}(0, \sigma_y^2)$, for some finite σ_y . We continue to assume that banks observe an infinitely precise private signal before deciding on investment, and this ensures uniqueness of equilibrium. Let $\Phi(\cdot)$ and $\phi(\cdot)$ denote the standard normal distribution and density functions, respectively. In this setting, the best guarantee policy for the government is the value of $\delta \in [0, 1]$ that maximizes the expected total wealth

$$\int_{-\infty}^{\theta_G^*} \frac{1}{\sigma_y} \phi\left(\frac{\theta - y}{\sigma_y}\right) d\theta + \int_{\theta_G^*}^{\infty} \left[\theta - \frac{\pi^*(\delta)^2}{2} + \pi^*(\delta) \right] \frac{1}{\sigma_y} \phi\left(\frac{\theta - y}{\sigma_y}\right) d\theta, \quad (22)$$

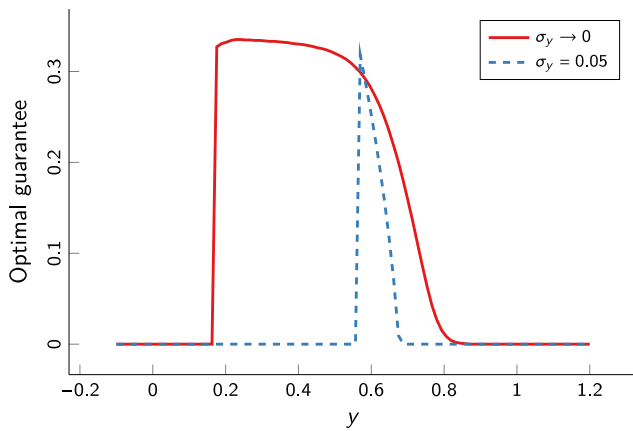


Fig. 6. Fiscal cost of subsidies and guarantees.

Notes. Parameters are $b = 0.5$ and $\gamma = 0.7$, which implies that $\theta^* = 0.676$ and $\theta^\circ = 0.558$.

subject to $\pi^*(\delta) = 1 - \delta$ and θ_G^* given by (8). The objective in (22) simplifies to

$$\Phi\left(\frac{\theta_G^* - y}{\sigma_y}\right) + \left[1 - \Phi\left(\frac{\theta_G^* - y}{\sigma_y}\right)\right] \left[1 - \delta - \frac{(1 - \delta)^2}{2} + y\right] + \sigma_y \phi\left(\frac{\theta_G^* - y}{\sigma_y}\right).$$

Given that the problem becomes a lot less tractable, we illustrate how our results change when we move away from the limiting case of $\sigma_y \rightarrow 0$ with a numerical example, which is depicted in Fig. 6.

For very high signals y , the optimal guarantee is zero. Intuitively, after observing a very high y , the government is sufficiently confident that fundamentals are in fact good and that coordination is likely. Hence, offering a positive δ and distorting banks' risk taking to avoid the low-probability event of a credit crunch is not worth it. As signal y lowers, the probability assigned to the actual fundamental falling in the coordination failure area increases, and the government optimally gives positive guarantees for signals in an intermediate range.

Given the higher uncertainty, positive guarantees are offered in a broader range of states than in the case with infinitely precise signals, $\sigma_y \rightarrow 0$, as can be seen by comparing the two curves in Fig. 6. Even if the public signal is slightly outside the coordination failure region ($1/2, \theta^*$), with positive σ_y , the government still assigns some significant probability of the fundamental falling in that region, and it is optimal to promise some guarantees. Importantly, the main message that for low (expected) fundamentals the government gives up offering guarantees to reduce the occurrence of credit crunches remains.

For intuition, consider a very low y . From the government's perspective, there is a high probability that the true fundamental is so low that no level of guarantees can induce investment. In the low-probability event that fundamentals are not so low and instead fall within the coordination failure area, $\theta \in (0.5, \theta^*)$, lower fundamentals are still more likely within that range, so for guarantees to work they must be relatively high, inducing too much risk-shifting—if fundamentals fall in the range $(0.5, \theta^\circ)$, no guarantees is better than the guarantee that avoids the credit crunch. If fundamentals happen to fall in the range (θ°, θ^*) , we know that there are positive guarantees that dominate zero guarantees (Proposition 5). However, in the setting with $\sigma_y > 0$, committing to a positive δ also distorts project choice for all realizations of $\theta \geq \theta^*$, with a negative impact on welfare. Numerically, this negative

effect dominates, and as y decreases, eventually the optimal guarantee falls back to zero.

Data availability

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

References

- Ahnert, T., Küncl, M., 2024. Government loan guarantees, market liquidity, and lending standards. *Manag. Sci.* 70 (7), 4502–4532.
- Allen, F., Carletti, E., Goldstein, I., Leonello, A., 2015. Moral hazard and government guarantees in the banking industry. *J. Financ. Regul.* 1 (1), 30–50.
- Allen, F., Carletti, E., Goldstein, I., Leonello, A., 2018. Government guarantees and financial stability. *J. Econom. Theory* 177, 518–557.
- Bebchuk, L., Goldstein, I., 2011. Self-fulfilling credit market freezes. *Rev. Financ. Stud.* 24 (11), 3519–3555.
- Calomiris, C.W., Kahn, C.M., 1991. The role of demandable debt in structuring optimal banking arrangements. *Am. Econ. Rev.* 497–513.
- Camargo, B., Kim, K., Lester, B., 2016. Information spillovers, gains from trade, and interventions in frozen markets. *Rev. Financ. Stud.* 29 (5), 1291–1329.
- Carletti, E., Leonello, A., Marquez, R., 2023. Loan guarantees, bank underwriting policies and financial stability. *J. Financ. Econ.* 149 (2), 260–295.
- Carlsson, H., Van Damme, E., 1993. Global games and equilibrium selection. *Econometrica* 61 (5), 989–1018.
- Cheng, I.H., Milbradt, K., 2012. The hazards of debt: Rollover freezes, incentives, and bailouts. *Rev. Financ. Stud.* 25 (4), 1070–1110.
- Cooper, R., John, A., 1988. Coordinating coordination failures in Keynesian models. *Q. J. Econ.* 103 (3), 441–463.
- Corsetti, G., Guimaraes, B., Roubini, N., 2006. International lending of last resort and moral hazard: A model of IMF's catalytic finance. *J. Monet. Econ.* 53 (3), 441–471.
- Dávila, E., Goldstein, I., 2023. Optimal deposit insurance. *J. Political Econ.* 131 (7), 1676–1730.
- Diamond, D., Dybvig, P., 1983. Bank runs, deposit insurance, and liquidity. *J. Political Econ.* 91 (3), 401–419.
- Goldstein, I., Pauzner, A., 2005. Demand–deposit contracts and the probability of bank runs. *J. Financ.* 60 (3), 1293–1327.
- Guerrieri, V., Shimer, R., 2014. Dynamic adverse selection: A theory of illiquidity, fire sales, and flight to quality. *Am. Econ. Rev.* 104 (7), 1875–1908.
- Guimaraes, B., Morris, S., 2007. Risk and wealth in a model of self-fulfilling currency attacks. *J. Monet. Econ.* 54 (8), 2205–2230.
- He, Z., Xiong, W., 2012. Dynamic debt runs. *Rev. Financ. Stud.* 25 (6), 1799–1843.
- Kashyap, A.K., Tsomocos, D.P., Vardoulakis, A.P., 2024. Optimal bank regulation in the presence of credit and run risk. *J. Political Econ.* 132 (3), 772–823.
- Keister, T., 2016. Bailouts and financial fragility. *Rev. Econ. Stud.* 83 (2), 704–736.
- Keister, T., Narasiman, V., 2016. Expectations vs. fundamentals-driven bank runs: When should bailouts be permitted? *Rev. Econ. Dyn.* 21, 89–104.
- Kiyotaki, N., 1988. Multiple expectational equilibria under monopolistic competition. *Q. J. Econ.* 103 (4), 695–713.
- Matta, R., Perotti, E., 2024. Pay, stay, or delay? how to settle a run. *Rev. Financ. Stud.* 37 (4), 1368–1407.
- Morris, S., Shin, H.S., 1998. Unique equilibrium in a model of self-fulfilling currency attacks. *Am. Econ. Rev.* 587–597.
- Morris, S., Shin, H.S., 2003. Global games: Theory and applications. In: Dewatripont, M., Hansen, L., Turnovsky, S. (Eds.), *Advances in Economics and Econometrics: Theory and Applications, Eighth World Congress*, vol. 1, Cambridge University Press.
- Morris, S., Shin, H.S., 2004. Coordination risk and the price of debt. *Eur. Econ. Rev.* 48 (1), 133–153.
- Philippon, T., Skreta, V., 2012. Optimal interventions in markets with adverse selection. *Am. Econ. Rev.* 102 (1), 1–28.
- Sákovics, J., Steiner, J., 2012. Who matters in coordination problems? *Am. Econ. Rev.* 102 (7), 3439–3461.
- Schilling, L.M., 2023. Optimal forbearance of bank resolution. *J. Financ.* 78 (6), 3621–3675.
- Taschereau-Dumouchel, M., 2025. Cascades and fluctuations in an economy with an endogenous production network. *Rev. Econ. Stud.* rda036.
- Tirole, J., 2012. Overcoming adverse selection: How public intervention can restore market functioning. *Am. Econ. Rev.* 102 (1), 29–59.
- Zhong, H., Zhou, Z., 2025. Dynamic coordination and bankruptcy regulations. *Rev. Financ. Stud.* hhaf039.