

Keister, Todd; Mitkov, Yuliyana

Working Paper

Allocating Losses: Bail-ins, Bailouts and Bank Regulation

ECONtribute Discussion Paper, No. 049

Provided in Cooperation with:

Reinhard Selten Institute (RSI), University of Bonn and University of Cologne

Suggested Citation: Keister, Todd; Mitkov, Yuliyana (2020) : Allocating Losses: Bail-ins, Bailouts and Bank Regulation, ECONtribute Discussion Paper, No. 049, University of Bonn and University of Cologne, Reinhard Selten Institute (RSI), Bonn and Cologne

This Version is available at:

<https://hdl.handle.net/10419/228852>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

ECONtribute

Discussion Paper

Allocating Losses: Bail-ins, Bailouts and Bank Regulation

Todd Keister Yuliyen Mitkov

December 2020

ECONtribute Discussion Paper No. 049

Funding by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany's Excellence Strategy – EXC 2126/1– 390838866 is gratefully acknowledged.

Cluster of Excellence

Allocating Losses: Bail-ins, Bailouts and Bank Regulation*

Todd Keister
Rutgers University
todd.keister@rutgers.edu

Yuliyana Mitkov
University of Bonn
ymitkov@uni-bonn.de

December 18, 2020

Abstract

We study the interaction between a government's bailout policy and banks' willingness to impose losses on (or "bail in") their investors. The government has limited commitment and may choose to bail out banks facing large losses. The anticipation of this bailout undermines a bank's private incentive to impose a bail-in. In the resulting equilibrium, bail-ins are too small and bailouts are too large. Some banks may also face a run by informed investors, creating further distortions and leading to larger bailouts. We show how a regulator with limited information can raise welfare and improve financial stability by imposing a system-wide, mandatory bail-in at the onset of a crisis. In some situations, allowing banks to choose between meeting a minimum bail-in and opting out can raise welfare further.

Keywords: Bank bailouts, moral hazard, financial stability, banking regulation

JEL Codes: E61, G18, G28

*An earlier version of this paper circulated under the title "Bailouts, Bail-ins and Banking Crises." We thank seminar and conference participants at the Bank of Canada, Deutsche Bundesbank, Federal Reserve Banks of Cleveland and Minneapolis, Federal Reserve Board, Hong Kong University of Science and Technology, Nederlandsche Bank, Office of Financial Research, Paris School of Economics, Rutgers University, Swiss National Bank, University of Exeter, University of Mannheim, 2016 Oxford Financial Intermediation Theory (OxFIT) Conference, 2nd Chicago Financial Institutions Conference, Spring 2017 Midwest Macroeconomics Meetings, 2017 Meetings of the Society for Economic Dynamics, 2018 Cornell-Penn State Macroeconomics Workshop, and especially Fabio Castiglionesi, Huberto Ennis, Charles Kahn, Joel Shapiro and Elu von Thadden for helpful comments. Yuliyana Mitkov gratefully acknowledges funding by the German Research Foundation (DFG) through CRC TR 224 (Project C03) and through EXC 2126/1-390838866 under Germany's Excellence Strategy.

1 Introduction

In periods of crisis, banks and other financial institutions suffer losses that are eventually borne in some combination by their own investors and creditors and, possibly, by the public sector in the form of a bailout. How these losses are allocated between private agents and the public sector has important implications for incentives and behavior in normal times as well as for the allocation of resources in society. Following the global financial crisis of 2008 and the subsequent European debt crisis, a broad consensus emerged that too many of the losses associated with these events fell on the public sector, that is, bailouts were too frequent and too large. This perception led policy makers to draft rules requiring financial institutions to impose more losses on (that is, to “bail in”) their investors/creditors in future crises. It remains to be seen how effective these mechanisms will be in practice. Even at a conceptual level, however, it is not well understood how losses *should* be allocated in a crisis, nor what types of bail-in policies are likely to be most effective.

We study the interaction between bail-ins and bailouts with a particular focus on what happens during the early stages of a crisis. Our model builds on the classic framework of Diamond and Dybvig (1983), in which investors facing idiosyncratic liquidity risk pool their resources in banks.¹ Bank assets are risky in our model and, in the event of a crisis, losses are heterogeneous across banks. At the onset of a crisis, some investors have private information about the size of their bank’s loss and can withdraw funds before this information becomes public. Banks have the ability to bail in these investors by paying them less than in normal times. In practice, this bail-in could represent a range of actions that preserve resources within the bank, including lowering dividend payments, restricting withdrawals and/or imposing withdrawal fees. We study banks’ incentives in making these decisions and ask when regulating these actions can improve welfare.

Our model provides a framework for evaluating policies like the reforms to money market mutual funds that were adopted in the U.S. in 2014.² Under the new rules, some funds are permitted to temporarily limit redemptions and impose withdrawal fees – a type of bail-in – during periods of financial stress. A fund is directed to take these actions if doing so is in the best interests of its investors. This policy raises interesting questions: What are the best interests of an institution’s investors in such a situation? Are these rules likely to achieve desirable outcomes? Another example is the debate over whether regulators should

¹While we use the word *bank* throughout the paper, our analysis also applies to intermediation arrangements outside of commercial banks that perform maturity or liquidity transformation. Yorulmazer (2014) discusses several such arrangements that encountered problems during the financial crisis of 2008. See also Chen et al. (2010) for an analysis of open-end mutual funds, Schmidt et al. (2016) on money market mutual funds, and Goldstein et al. (2017) on corporate bond funds.

²See Ennis (2012) for a discussion of the issues involved in reforming money market mutual funds.

restrict dividend payments by banks during the economic crisis caused by the Covid-19 pandemic. The European Central Bank recommended on March 27, 2020, that banks “refrain from making dividend distributions and performing share buy-backs aimed at remunerating shareholders” during this period.³ In the U.S., the Federal Reserve moved on June 25 to prohibit share repurchases and to cap dividend payments by large banks in the U.S. When is it desirable to impose restrictions on the payments banks make to their investors? What types of restrictions are most effective? We develop a model to address these questions.

The efficient allocation of a bank’s losses in our model depends critically on the fiscal capacity of the public sector. If this fiscal capacity is small, a benevolent planner will provide no bailouts and will impose all of a bank’s losses on its investors by bailing them in. When this fiscal capacity is larger, however, the planner will provide bailouts to banks with sufficiently large losses. In other words, the planner will want the public sector to absorb some of the “tail risk” in the economy, which implies that a combination of bail-ins and bailouts is efficient.

In a decentralized equilibrium, banks’ incentive to bail-in their investors depends on what bailout policy they expect. The government chooses whether and how to make bailout payments after banks’ financial conditions become public and some withdrawals have occurred. It cannot commit to a bailout policy in advance; it will choose the bailout payments, if any, as a best response to the situation at hand. We show that the anticipation of being bailed out undermines a bank’s incentive to bail in its investors. Specifically, those banks that receive bailouts in equilibrium choose bail-ins that are smaller than in the planner’s allocation. In addition, total bailout payments are larger in equilibrium than in the planner’s allocation.

The distortion in banks’ incentive to bail in their investors can also lead to runs on some banks. When a bank chooses not to impose a bail-in, it creates a stronger incentive for investors to withdraw early, before the information about the bank’s losses becomes public. We show that, in some cases, withdrawing early becomes a dominant strategy and thus leads to a fundamentals-based run on the bank. Banks recognize this fact in choosing their bail-in policy and can always prevent a run by imposing a sufficiently large bail-in. In some cases, however, investors prefer that their bank not impose any bail-in, even though doing so precipitates a run. In this way, our model identifies a new channel through which bailouts can increase financial fragility: by giving banks and their investors an incentive to delay the recognition of losses, which ends up encouraging early withdrawals.

³See “Recommendation of the European Central Bank of 27 March 2020 on dividend distributions during the COVID-19 pandemic,” *Official Journal of the European Union*, 2020/C 102 I/01. The European Systemic Risk Board made a similar recommendation on May 27. See Acharya et al. (2016) for a discussion of bank dividend payments in the 2007-9 financial crisis and a model in which dividend payments create externalities across banks.

Given these problems, we ask whether a regulator that only has access to public information can improve equilibrium outcomes. The regulator in our model has the ability to restrict the payments banks make to their investors, which can be interpreted as limiting dividend payments, imposing withdrawal fees, or writing down the face value of liabilities. Mandating that all banks bail in their investors by a fixed amount can raise welfare in some situations, but not in others. The effectiveness of this policy depends critically on the distribution of losses across banks; when there is more heterogeneity across banks, imposing a common bail-in is less attractive. We then show how a mandatory *minimum* bail-in policy can often do better. Under this policy, banks must bail in their investors by at least the specified amount, but are allowed to voluntarily impose a larger bail-in. This policy takes advantage of the fact that mandatory bail-ins improve the incentives of banks facing large losses by limiting their ability to shift these losses to the public sector. In particular, it leads banks to internalize more of the costs associated with a run by their investors, which leads some banks to impose a larger bail-in than the mandated minimum as a way of preventing a run. In other words, the threat of a run can provide additional discipline on bank behavior when the bail-in policy is chosen appropriately. We derive the optimal level for the mandatory minimum and provide conditions under which this policy improves welfare.

The regulator can further leverage the disciplining effect of runs by making the minimum bail-in *optional* in the sense that a bank's bail-in must either be zero or at least the minimum level. This policy aims to further separate banks by type, allowing those with no losses to choose zero bail-in while encouraging those with large losses to choose a positive bail-in to prevent a run. We show that, in some cases, the regulator can choose the minimum amount so that this policy yields higher welfare than the optimal mandatory bail-in policy. The optional bail-in policy tends to be more attractive, for example, when there are many banks with no losses and the fiscal capacity of the public sector is moderate.

Overall, our results demonstrate the value of policies that trigger a system-wide bail-in based on aggregate financial or economic conditions. Much of the existing policy discussion has focused on tying bank-specific bail-ins to an idiosyncratic trigger that is observed either publicly or privately by regulators. For example, contingent-convertible bonds (CoCos) can be structured to convert from debt to equity when the book value of a bank's equity falls below some pre-specified level.⁴ In our model, the regulator can directly impose the efficient bail-ins once it observes the status of each bank. There is, however, a period during which the regulator knows a problem exists, but does not yet know how badly each bank is affected. Our results show that a broadly-targeted bail-in mandate during this period can be effective in reducing future bailouts, maintaining financial stability, and raising welfare.

⁴See Flannery (2014) for a detailed discussion of CoCos and a review of the relevant literature.

Related literature. Early discussions of bail-in policy focused on identifying potential advantages and disadvantages compared to other types of resolution for distressed banks and on the practical aspects of implementing a bail-in regime. (See, for example, Dewatripont (2014), Sommer (2014), and Goodhart and Avgouleas (2015).) More recently, a literature has emerged that studies the incentive effects of bail-ins and the resulting policy tradeoffs in formal economic models. Bernard et al. (2017) study a game in which a regulator and banks negotiate over the allocation of losses. They show banks will accept a bail-in in equilibrium if, and only if, the regulator’s cost of funds is large enough that they anticipate no bailout. A similar result appears in our setting: when the fiscal capacity of the government is small enough, no bailouts occur and all banks choose the efficient bail-in. When bailouts distort bail-in incentives, Bernard et al. (2017) focus on how the network structure of interbank linkages affects the credibility of a no-bailout plan, while we focus on how mandatory bail-in policies can improve banks’ incentives and choices.

Walther and White (2019) focuses on how bail-ins affect a bank manager’s incentive to exert an appropriate level of effort. In their setting, the regulator has *more* information than a bank’s creditors about the value of its assets. Bailing in creditors improves the manager’s incentives by increasing her stake, but risks provoking a run if it leads creditors to infer the bank is in bad shape. Our approach is complementary to theirs in that we focus on the distortion in investors’ incentives rather than on an agency problem within the bank. In addition, our assumption that investors have private information captures the fact that some investors and creditors are typically insiders to the bank. Bank runs also play a different role in our analysis. In Walther and White (2019), the threat of a run limits the regulator’s ability to bail in creditors. In our setting, the threat of a run disciplines bank behavior and can *help* the regulator implement better allocations. This insight allows the regulator to design policies that leverage the disciplining effect of runs to give banks facing large losses an incentive to bail in their investors.

Colliard and Gromb (2018) study the negotiation between a bank’s shareholders and its creditors over how the losses will be distributed. The investors in our model represent both shareholders and creditors; they hold a claim that is a hybrid of debt and equity, as is common in the models based on Diamond and Dybvig (1983). Our focus is on the division of losses between the public sector and these investors as a group, rather than between different categories of investors. Bolton and Oehmke (2019) study the problem of coordinating bail-ins in multinational banking corporations and characterize the incentive problems that arise between national regulators. The public sector in our model is divided between a regulator and a fiscal authority, but both entities have the same objective of maximizing total welfare.

Our paper also contributes to the literature that aims to evaluate financial stability

reforms that have been adopted in recent years. For example, [Cipriani et al. \(2014\)](#) study how the possibility that investors will face future withdrawal restrictions or fees can create a preemptive run on a money market mutual fund.⁵ In their setting, informed investors learn about a bank’s potential losses before the bank itself, which allows a run to occur before the bank can restrict withdrawals. Our results complement theirs by showing that, even if the bank could impose a bail-in quickly enough to prevent a run, doing so would often not be in its investors’ best interests. In this way, both papers cast doubt on the effectiveness of reforms that aim to promote financial stability by allowing intermediaries to restrict withdrawals.

Outline. In the next section, we describe the economic environment and define the concepts of bail-ins and bailouts in this environment. In [Section 3](#), we derive the combination of bail-ins and bailouts that would be chosen by a benevolent planner, which serves as a useful benchmark for what follows. In [Section 4](#), we analyze the decision problems faced by banks, investors and the government and we derive the payoffs of the bail-in game. In [Section 5](#), we document the properties of equilibrium in this game, showing that bail-ins are too small and bailouts are too large, and we derive conditions under which bank runs occur. We introduce regulation in [Section 6](#) and show how system-wide bail-in policies can raise welfare. We offer some concluding remarks in [Section 7](#).

2 The model

We base our analysis on a version of the [Diamond and Dybvig \(1983\)](#) model expanded to include fiscal policy conducted by a government with limited commitment, as in [Keister \(2016\)](#). We introduce private information about the value of banks’ assets into this framework. In this section, we describe the agents, preferences, and technologies that characterize the environment, and we define bail-ins and bailouts within this environment.

2.1 The environment

There are three time periods, labeled $t = 0, 1, 2$. There is a single private consumption good in every period and a public good that can be produced at $t = 1$.

⁵[Engineer \(1989\)](#) was the first to show how threat of a future deposit freeze could create a run in an extended Diamond-Dybvig model. See also [Voellmy \(2019\)](#), which uses a version of Engineer’s model to derive conditions under which withdrawal restrictions and fees can be effective.

Investors. There is a continuum of investors, indexed by $i \in [0, 1]$, in each of a measure one of locations. Investor i in a given location has preferences characterized by

$$u(c_1^i + \omega^i c_2^i) + v(g),$$

where c_t^i denotes her consumption in period $t \in \{1, 2\}$ and g denotes the level of the public good, which is common to all locations. The random variable $\omega^i \in \Omega \equiv \{0, 1\}$ is realized at $t = 1$ and is privately observed by the investor. If $\omega^i = 0$, she is *impatient* and values consumption only in period 1, whereas if $\omega^i = 1$, she is *patient* and values consumption equally in both periods. Each investor will be impatient with a known probability $\pi > 0$, and the fraction of impatient investors in each location will also equal π . The functions u and v are assumed to be smooth, strictly increasing, strictly concave and to satisfy the usual Inada conditions. We assume the coefficient of relative risk aversion for u is constant and strictly greater than unity. Investors are each endowed with one unit of the consumption good at $t = 0$.

Banks. Goods can be stored at a gross return of 1 between $t = 0$ and $t = 1$ and a gross return of $R > 1$ between $t = 1$ and $t = 2$. As in Diamond and Dybvig (1983), the idiosyncratic uncertainty about preference types ω^i creates an incentive for investors to pool resources to insure against liquidity risk, that is, against the possibility of being impatient and needing to consume before the return R is available. There is a banking technology in each location that holds goods and allows investors to withdraw funds at either $t = 1$ or $t = 2$. To simplify the analysis, we begin with the endowment of investors in each location already deposited in their location's banking technology.⁶

Crises. At $t = 0$, before any decisions are made, one of two aggregate states is realized. In one state, which we call *normal times*, the value of all banks' assets remains unchanged. In the other state, which we call a *financial crisis*, banks experience a loss whose size varies across locations. Specifically, a fraction $(1 - \phi)$ of the goods held by a bank in a given location become worthless, leaving the bank with ϕ units of the good per investor. The value of ϕ is an idiosyncratic draw from a distribution F on the interval $\Phi \equiv [\underline{\phi}, 1]$, where $\underline{\phi} \geq 0$. We assume F is continuous, strictly increasing and differentiable on $[\underline{\phi}, 1)$, but may place positive probability on $\phi = 1$. The realized distribution of asset values across locations

⁶That is, we do not study what Peck and Shell (2003) call the *pre-deposit game*, in which investors decide whether or not to pool their resources. Because there is no asymmetry of information between a bank and its investors in our model, this assumption is without loss of generality.

is also given by F , which implies that total losses in the economy can be expressed as

$$\int_{\underline{\phi}}^1 (1 - \phi) dF(\phi).$$

All agents observe the aggregate state and know the distribution F , but the realized value of ϕ in a given location is initially observed only by the investors in that location.

Bank operations. After ϕ has been realized in each location, investors collectively decide how much consumption their bank will give to investors who withdraw at $t = 1$. Each investor then observes her own preference type and decides in which period she will withdraw. Those investors who chose to withdraw in period 1 arrive at their bank one at a time, in a randomly-determined order, and receive the specified payment. Investors who choose to withdraw at $t = 2$ receive an even share of the matured value of their bank's assets. Investors are isolated from each other during the withdrawal process and no trade can occur among them. As in Wallace (1988), this assumption prevents re-trading opportunities from undermining banks' ability to provide liquidity insurance.⁷

Public goods. There is also a technology for converting units of the private good one-for-one into units of the public good in period 1. While any agent can operate this technology, the fact that the set of agents in each location is a negligible fraction of the overall economy implies that there is no private incentive to provide the public good.

Government. There is a benevolent government that acts in two capacities: as a fiscal authority and as a regulator. In both capacities, the government's aim is to maximize the sum of all investors' utilities at all times.

The fiscal authority is endowed with $\tau \geq 0$ units of the good in period 1, where the parameter τ represents the government's *fiscal capacity*. These resources can be used to provide the public good and potentially to make transfers to banks. We call any such transfer a *bailout*. In a financial crisis, after a fraction π of investors have withdrawn in each location, the fiscal authority observes both the realized value of ϕ and the remaining resources in each location and chooses bailout payments.⁸ The fiscal authority is unable to commit to a bailout plan in advance; the bailout payments will be chosen as a best response

⁷See Jacklin (1987), Allen and Gale (2004) and Farhi et al. (2009), among others, for studies of how the presence of markets at $t = 1$ limits the amount of risk-sharing that banks provide to depositors.

⁸The assumption that the fiscal authority observes this information after precisely a fraction π of investors have withdrawn is made for simplicity. It is straightforward to introduce an additional parameter to measure when this information becomes available, as in Keister (2016).

to the situation at hand. All remaining funds are then used to provide the public good.

The regulator is able to restrict the payments that banks make to withdrawing investors in period 1. The anticipation of being bailed out may distort investors' incentives in choosing these payments, which creates the possibility that this type of regulation may raise welfare. However, like the fiscal authority, the regulator has limited information; it observes the realization of ϕ in each location only after a fraction π of investors have withdrawn. This lag implies that regulatory policy cannot be fully state-contingent.

2.2 Allocating losses: Bail-ins and bailouts

Our interest is in studying how the losses that occur during a financial crisis are allocated between bank creditors and the public sector. As a first step, we derive the allocation of private consumption in normal times, which provides the benchmark from which losses will be measured.

A reference allocation. In normal times, the per-capita value of the bank's assets is 1 in all locations. Suppose all patient investors wait until $t = 2$ to withdraw from the bank and there are no bailout payments. Then the efficient allocation of resources in each location gives a common amount c_1 to each investor who withdraws at $t = 1$ and a common amount c_2 to each investor who withdraws at $t = 2$, where these values are chosen to maximize investors' expected utility

$$\pi u(c_1) + (1 - \pi) u(c_2)$$

subject to the feasibility constraint

$$\pi c_1 + (1 - \pi) \frac{c_2}{R} \leq 1. \tag{1}$$

Let (c_1^*, c_2^*) denote the solution to this standard Diamond-Dybvig allocation problem, which satisfies

$$1 < c_1^* < c_2^* < R.$$

We consider c_1^* and c_2^* to represent the *face value* of a bank's liabilities to investors who choose to withdraw in periods 1 and 2, respectively. To be clear: Banks in our model can pay withdrawing investors less than these values in the event of a crisis, and they will do so whenever it is in their investors' best interests. In this sense, the liabilities defined above are not contractually binding, and deviating from these payments does not involve any cost or inefficiency. The role of the reference amounts (c_1^*, c_2^*) is simply to provide a benchmark for measuring what portion of a bank's losses are borne by its own investors.

Bail-ins. In a location where $\phi < 1$ in the crisis state, it will not be feasible for a bank to pay the amounts (c_1^*, c_2^*) to its withdrawing investors. In this case, investors will choose the best feasible allocation of their bank's resources. This allocation will give a common amount of consumption $c_1(\phi)$ to each of the first π investors who withdraw in period 1. If these investors receive less than the reference amount c_1^* , we say they have been *bailed-in*.⁹ It will be convenient to measure the size of the bail-in as the percentage “haircut” from the reference allocation, that is, as the solution $h(\phi)$ to

$$c_1(\phi) = (1 - h(\phi)) c_1^*.$$

In period 2, the bank will divide its matured assets evenly among the remaining investors. Let $\hat{h}(\phi)$ denote the bail-in applied to these investors,¹⁰ which satisfies

$$c_2(\phi) = (1 - \hat{h}(\phi)) c_2^*.$$

If a bank has experienced no loss ($\phi = 1$), it will not bail in its investors; that is, it will set $h(1) = \hat{h}(1) = 0$. If it does have a loss ($\phi < 1$), investors will collectively choose the bail-ins $(h(\phi), \hat{h}(\phi))$ to maximize their expected utility, subject to feasibility constraints and anticipating the actions of the government. We use h and $\hat{h}$ to denote the profile of bail-in decisions across all banks.

Bailouts. After a fraction π of investors have withdrawn from each bank, the fiscal authority observes the value of ϕ of each bank as well as how many resources the bank has left after serving these π withdrawals. It then chooses a bailout payment $b(\phi) \geq 0$ for each bank, with the remaining funds being used to produce the public good. These choices are made with the objective of maximizing the sum of all investors' utilities. Note that the bailout decisions are made after each bank's bail-in has already been chosen (and implemented) for a fraction π of its investors. The fact that the fiscal authority cannot commit to the bailout policy before these withdrawals occur plays a critical role in our analysis.

⁹While some authors apply the term *bail-in* only to losses imposed on certain types of investors (such as long-term debt holders) or in certain situations (such as in resolution), we use the term more broadly to include all losses imposed on a bank's creditors and investors. Our approach aims to capture, in a unified way, a variety of policies and actions observed in reality during financial crises, including restrictions on dividend payments as well as haircuts imposed on depositors, various debt holders, and other creditors.

¹⁰Throughout the analysis, we use h to note the bail-in applied to the first π investors to withdraw and $\hat{h}$ to denote the bail-in applied to all remaining investors, regardless of the period in which these later withdrawals take place.

Feasibility. The bank in a location with realized asset value ϕ (hereafter, “bank ϕ ”) will have a total of $\phi + b(\phi)$ units of the good in period 1. Suppose for the moment that patient investors wait until period 2 to withdraw, so that only a fraction π of investors withdraw at $t = 1$. Then bank ϕ must choose its bail-ins $(h(\phi), \hat{h}(\phi))$ to satisfy the feasibility constraint

$$\pi(1 - h(\phi))c_1^* + (1 - \pi)\left(1 - \hat{h}(\phi)\right)\frac{c_2^*}{R} \leq \phi + b(\phi). \quad (2)$$

Using equation (1), we can rewrite this constraint as

$$h(\phi)\pi c_1^* + \hat{h}(\phi)(1 - \pi)\frac{c_2^*}{R} + b(\phi) \geq 1 - \phi. \quad (3)$$

This expression shows that the losses in bank ϕ are divided between bail-ins and bailouts. The first two terms of the left-hand side measure the period-1 value of the bank’s bail-ins: an amount πc_1^* of the bank’s liabilities is bailed in at rate $h(\phi)$ at $t = 1$, while the amount $(1 - \pi)c_2^*$ of liabilities that will be bailed in at rate $\hat{h}(\phi)$ at $t = 2$ is discounted by the return R . Feasibility requires that the sum of these bail-ins plus the bailout payment $b(\phi)$ be enough to cover the bank’s loss, $1 - \phi$.

Bank runs and resolution. Throughout our analysis, we assume that patient investors choose to wait until $t = 2$ unless withdrawing early is a strictly dominant strategy in the withdrawal game for their bank. In other words, we do not focus on the type of self-fulfilling bank runs studied by Diamond and Dybvig (1983) and many others. There may, however, be situations in which patient investors receive strictly more from their bank by withdrawing early regardless of the actions of others. In such cases, a bank run is inevitable.¹¹

If investors continue to arrive at a bank in period 1 after π withdrawals have been made, the bank is placed into a *resolution* process. We assume that, as part of this process, the run on the bank stops, meaning that only the remaining impatient investors withdraw in period 1 and all remaining patient investors withdraw in period 2. In addition, the fiscal authority observes the fraction of the remaining investors who are impatient and can condition the bailout payment (if any) on this information. When a bank is in resolution, the regulator dictates the bail-ins applied to all remaining investors, which implies that the bank’s remaining resources will be allocated efficiently among its remaining investors.

¹¹In focusing on bank runs that are driven by the “fundamentals” of the withdrawal game, we follow Chari and Jagannathan (1988) and Allen and Gale (1998), among others. As we show below, however, multiple equilibria may still arise in our model, driven by the bail-in decisions of banks, which shape the fundamentals of the withdrawal game.

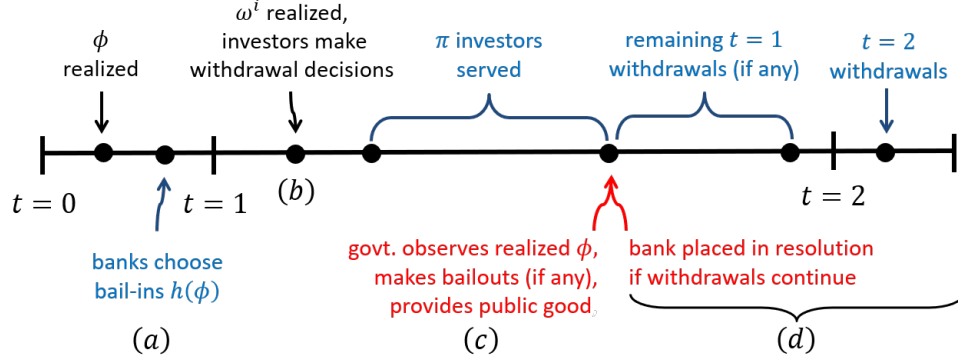


Figure 1: Timeline

2.3 Timeline

The sequence of events in the crisis state is depicted in Figure 1, where items in black represent moves by nature and individual investors, items in blue represent the actions of banks, and items in red correspond to the actions of the public sector. In period 0, investors observe the realized loss in their location and collectively choose the initial bail-in $h(\phi)$. Moving to period 1, these choices determine the payoffs of a Diamond-Dybvig-style withdrawal game in which each investor observes her own preference type ω^i and decides whether to withdraw in period 1 or in period 2. Banks pay an amount $(1 - h(\phi))c_1^*$ to withdrawing investors as they arrive. After a measure π of investors have withdrawn, the fiscal authority observes the type and remaining resources of each bank and chooses bailout payments $b(\phi)$. After bailout payments have been made, all remaining government funds are used to provide the public good. If withdrawals at a bank continue past π , the bank is placed into resolution and the regulator chooses the allocation of the bank's resources between the remaining investors. Otherwise, all remaining investors are given an even share of the banks' assets in period 2.

2.4 Discussion

Aggregate uncertainty. While our model has two aggregate states, our analysis focuses on decisions made and actions taken in the crisis state. The only role of the good aggregate state in our model is to establish what investors would receive from their bank in normal times. Because banks are free to adjust payments to investors based on the realization of ϕ , the ex ante probabilities of the two aggregate states have no effect on the analysis. For interpretation purposes, we think of the probability of the crisis state as being small, but the precise value is irrelevant. What matters for our analysis is the distribution of losses across

banks conditional on the crisis state, which is given by the function F .

We do assume that the probabilities of the two aggregate states are exogenous in our model, as is the distribution F . In this sense, our focus is on how the losses in a crisis are allocated and not on the determinants of a crisis or *ex ante* moral hazard issues.¹² Our approach seems particularly relevant for studying the effects of unexpected economic shocks originating outside of the financial sector, such as a pandemic. However, the effects we highlight in our analysis will be present any time there are significant losses in the banking system, regardless of the underlying cause.

Limited risk sharing. In the crisis state, banks face idiosyncratic risk about the value of their assets. In practice, interbank insurance arrangements can be used to share this type of risk across institutions. Our approach implicitly assumes that these arrangements are imperfect. A sizeable literature has studied how information and incentive problems can limit the effectiveness of interbank markets and other types of insurance arrangements.¹³ We think of whatever arrangements exist as being embedded in the function F , so that this probability distribution of ϕ in the crisis state represents the *uninsured* component of bank-specific risk.

The restriction that bailout payments be non-negative prevents the government from using fiscal policy to replace these missing insurance markets. If the government were allowed to tax banks with small losses and transfer the proceeds to banks with larger losses, it could effectively insure banks against all idiosyncratic risk. In practice, one would expect this type of arrangement to face incentive problems at least as severe as those in private insurance markets. Our assumptions ensure the role of bailouts in the model is to transfer risk between the private and public sectors, rather than to replace missing private insurance markets.

Fiscal capacity. The parameter τ measures the government's access to resources in the event of a crisis. In our model, the government is simply endowed with these resources at $t = 1$. One can interpret these goods as coming from tax revenue raised before our model begins (and stored until $t = 1$) or from taxing activities outside of the scope of the model. More generally, one can think of the government's fiscal capacity as including funds that could be raised by issuing new debt or by selling public assets. The key point for our analysis is that, whatever the source of these funds, using them to bail out banks is costly.

¹²A large literature has studied how government guarantees, both explicit and implicit, affect the riskiness of banks' assets and, therefore, the probability of a crisis state. Kareken and Wallace (1978) is one classic reference. More recently, Acharya and Yorulmazer (2007) study how the anticipation of intervention affects the correlation of banks' asset choices and, therefore, the distribution F of losses across banks in a crisis.

¹³See, for example, Bhattacharya and Gale (1987), Flannery (1996), Freixas and Jorge (2008), Heider et al. (2015) and Castiglionesi and Navarro (2020), among many others.

This cost is captured in our model by a decrease in the level of the public good g .¹⁴ We assume the government’s fiscal capacity is small enough that

$$v' \left(\tau - \int_{\underline{\phi}}^1 (1 - \phi) dF(\phi) \right) > u'(c_1^*) \quad (4)$$

holds. This condition guarantees that it is not efficient for the public sector to absorb all of banks’ losses; at least some losses will be imposed on investors in the form of a bail-in.

Informed investors. We assume that a bank’s investors are informed about the value of the bank’s assets at the beginning of period 1. Because of the sequential service constraint, however, only a fraction π of these investors can act on this information before the fiscal authority and regulator intervene. In effect, the private information is only relevant for this group of investors, which we interpret as at least partly representing insiders to the bank. Several recent studies have highlighted the importance of withdrawals by bank insiders in the period before regulatory actions and/or bank failure occur. [Acharya et al. \(2011\)](#), for example, discuss how dividend payments in 2007-9 effectively transferred banks’ losses onto the public sector. [Henderson et al. \(2015\)](#) study smaller community banks in the U.S. during the same period and document a “run on capital by insiders” in the form of higher dividend payouts by troubled banks. [Iyer et al. \(2016\)](#) study depositor-level data from a bank in India that suffered a run and document withdrawals by bank staff in the period before losses on the banks’ loan portfolio became public information. Assuming that all investors are informed and, hence, face the same decision problem helps simplify the presentation of our model. The important point, however, is that at least some investors have private information and are able to act on this information before any intervention by the public sector.

Delayed intervention. We assume the government observes bank-specific information with a delay, reflecting the fact that, at least in the early stages of a crisis, banks are likely to have more information about their own situation than is available to regulators. In direct terms, this assumption aims to capture the time required to carry out detailed examinations and to verify the information that forms the basis for supervisory action. More broadly, the delay in the model can also represent a variety of practical and political concerns that make policy makers slow to react to an incipient crisis. For example, [Kroszner and Strahan \(1996\)](#) argue that the Federal Savings and Loan Insurance Corporation (FSLIC) faced a severe shortage of funds for resolving insolvent thrift institutions in the 1980s. This lack of funds led FSLIC to practice regulatory forbearance and to delay its intervention in insolvent thrifts, anticipating

¹⁴Other papers using this approach include [Keister \(2016\)](#), [Allen et al. \(2018\)](#), [Mitkov \(2019\)](#), and [Li \(2020\)](#).

that the federal government would eventually supply additional resources. During this delay, insolvent thrift institutions could continue paying high dividends to their investors as a way of maximizing the private value of the government guarantee of their liabilities, at the taxpayers' eventual expense. In a similar spirit, [Brown and Dinc \(2005\)](#) provide evidence that the timing of a government's intervention to resolve a failing financial institution depends on the electoral cycle. By examining episodes from 21 major emerging market economies in the 1990s, they find that interventions that would impose large costs on taxpayers and/or would more fully reveal the extent of the crisis were significantly less likely to occur before elections. (See also [Rogoff and Sibert, 1988.](#)) Whatever the underlying cause of the delayed policy response, the key point for our analysis is that some investors are able to withdraw from a bank facing losses before decisions about bailouts and bank resolution are made.

Resolution. We assume a bank is placed in resolution as soon as it becomes apparent to regulators that a run on the bank is underway. Once in resolution, the bank's available resources – including any bailout it receives – are allocated efficiently among its remaining investors. There are a variety of ways to implement this type of resolution process, all of which would lead to the same outcome in our model. One could, for example, think of a court system intervening to verify investors' preference types, as discussed in [Ennis and Keister \(2009a\)](#). Alternatively, one could allow investors to write a “living will” that specifies how their bank will be operated following a run and intervention. Regardless of the specific mechanism, the run on the bank will necessarily end once it is in resolution because there is no further uncertainty and the payments offered to investors are adjusted accordingly.¹⁵ Moreover, because there are no further bailouts, investors' incentives will no longer be distorted from the regulator's point of view and the allocation of the bank's resources will be the same regardless of who chooses the payments. Our approach here of having the regulator dictate all remaining payments serves only to simplify the notation.

3 A planner's problem

In this section, we derive the combination of bail-ins and bailouts that would be chosen by a benevolent planner in the crisis state. This planner controls the operations of all banks, the actions of the fiscal authority, and investors' withdrawal decisions. It observes all of the information available to banks and investors, including each investor's preference type. The

¹⁵[Ennis and Keister \(2010\)](#) show how a bank run may continue after the initial policy reaction if there is still some uncertainty about the state of nature. In their setting, a run may occur in *waves* as the policy maker gradually learns the state.

planner faces the same restrictions as agents in the environment, including the inability to directly redistribute resources across locations. How would this planner allocate resources if its objective is to maximize the sum of all investors' expected utilities?

The planner will clearly direct all impatient investors to withdraw at $t = 1$ since they only value consumption in that period, and will direct all patient investors to withdraw at $t = 2$, after the bank's assets have earned the return $R > 1$. It will treat banks with the same value of ϕ symmetrically in choosing the bail-ins and bailouts; we denote these choices by $\{h(\phi), \hat{h}(\phi), b(\phi)\}$ for $\phi \in \Phi$. The planner will choose these functions to maximize the sum of expected utilities across all investors

$$\int_{\underline{\phi}}^1 \left(\pi u((1 - h(\phi))c_1^*) + (1 - \pi)u((1 - \hat{h}(\phi))c_2^*) \right) dF(\phi) \\ + v \left(\tau - \int_{\underline{\phi}}^1 b(\phi) dF(\phi) \right)$$

subject to the feasibility constraint (3) and the non-negativity restrictions

$$h(\phi) \in [0, 1], \quad \hat{h}(\phi) \in [0, 1], \quad \text{and} \quad b(\phi) \geq 0$$

for all ϕ . Our first result shows that the planner will choose to impose the same level of bail-in on all investors within a bank.

Proposition 1. The efficient plan satisfies $h^*(\phi) = \hat{h}^*(\phi)$ for each $\phi \in \Phi$. That is, all investors within a bank face the same bail-in.

Proofs of all propositions are provided in the appendix. The result in Proposition 1 relies on the assumption that the coefficient of relative risk aversion for the utility function u is constant, which implies that investors' expected-utility preferences over private consumption are homothetic. As a result, the efficient levels of consumption for impatient and patient investors scale down in proportion when a bank experiences losses. In the remainder of this section, we use $h(\phi)$ to denote the bail-in applied by the planner to all investors in bank ϕ .

Using the result in Proposition 1 together with the resource constraint for the reference allocation in equation (1), we can rewrite the feasibility constraint (3) for allocating the losses in each bank in a particularly simple form:

$$h(\phi) + b(\phi) \geq 1 - \phi \quad \text{for all } \phi \in \Phi. \quad (5)$$

This constraint highlights how the loss $(1 - \phi)$ must be covered by a combination of bail-ins $h(\phi)$ of the bank's investors and bailouts $b(\phi)$ by the public sector. We can also simplify

the objective function in the planner's problem by defining

$$U(1 - h(\phi)) \equiv \pi u((1 - h(\phi))c_1^*) + (1 - \pi)u((1 - h(\phi))c_2^*) \quad (6)$$

and using equation (5), which will hold with equality at an optimum, to replace $b(\phi)$. Then the planner will choose a bail-in function h to solve

$$\max_{\{h\}} \left\{ \int_{\underline{\phi}}^1 U(1 - h(\phi)) dF(\phi) + v \left(\tau - \int_{\underline{\phi}}^1 (1 - \phi - h(\phi)) dF(\phi) \right) \right\} \quad (7)$$

subject to

$$0 \leq h(\phi) \leq 1 - \phi \quad \text{for all } \phi \in \Phi.$$

We call the solution to this problem the *efficient plan* and denote it (h^*, b^*) . Our main result in this section characterizes this plan.

Proposition 2. There exists $\phi^* \in \Phi$ such that the efficient plan (h^*, b^*) has the form:

$$h^*(\phi) = \begin{cases} 1 - \phi \\ 1 - \phi^* \end{cases} \quad \text{and} \quad b^*(\phi) = \begin{cases} 0 \\ \phi^* - \phi \end{cases} \quad \text{as } \phi \begin{cases} \geq \\ < \end{cases} \phi^*.$$

The efficient plan is characterized by a maximum bail-in $1 - \phi^*$. For banks with a loss smaller than the maximum bail-in, the efficient bail-in $h^*(\phi)$ equals the total loss $1 - \phi$ and no bailout is received. For banks with a loss greater than $1 - \phi^*$, the maximum bail-in is applied and the remaining loss, $\phi^* - \phi$, is covered by a bailout. In this way, the planner uses public resources to provide insurance against large losses, but not against smaller losses.

Panel (a) of Figure 2 illustrates this result.¹⁶ The graph shows, for each value of ϕ , how the loss $1 - \phi$ is divided between a bail-in of the first π investors to withdraw (bottom region, light-blue), a bail-in of the remaining $(1 - \pi)$ investors (middle region, darker blue), and a bailout (top region, red). Note that the relative sizes of the first two regions are the same for all ϕ , in line with Proposition 1. For banks with a loss smaller than $1 - \phi^*$, these bail-ins cover the entire loss and there is no bailout. For banks with a loss larger than $1 - \phi^*$, the bail-ins take their maximum value and the bailout covers the remaining loss.

The cutoff value ϕ^* depends critically on the amount of resources τ available to the government, as illustrated in panel (b) of the figure. If τ is sufficiently small, the cutoff equals the lower bound $\underline{\phi}$, meaning there are no bailouts and the bail-in applied at each

¹⁶All of the numerical examples in the paper use the functional forms $u(x) = v(x) = x^{1-\gamma}/(1-\gamma)$ with $\gamma = 2$ and parameter values $\pi = 1/2$ and $R = 3$. The distribution F places measure $1/2$ on $\phi = 1$ and the other half of banks are uniformly distributed on $\Phi = [1/4, 1]$. Panel (a) of Figure 2 uses $\tau = 0.9$.

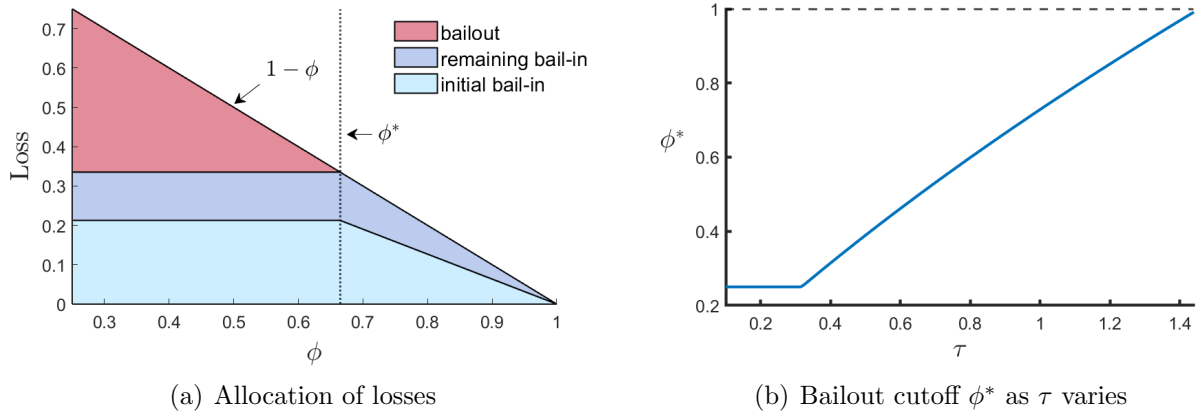


Figure 2: The planner's allocation

bank covers the total loss. For larger values of τ , the planner uses a combination of bail-ins for all banks and bailouts for some banks, which corresponds to the case depicted in panel (a). The cutoff ϕ^* is increasing in τ in this region, meaning that when the fiscal capacity of the government is larger, the planner shifts more of the losses to the public sector and provides bailouts to a larger number of banks. As τ approaches the upper bound in equation (4), ϕ^* approaches 1 and all banks with losses are bailed out. Our final result of this section shows that the increasing relationship depicted in the figure holds in general.

Proposition 3. The efficient bailout cutoff ϕ^* is strictly increasing in τ whenever $\phi^* \in (\underline{\phi}, 1)$.

4 The bail-in game

We now return to the decentralized economy, where investors' preference types ω^i are private information and the bail-in decisions are made separately by investors in each location. In this section, we formulate the game played at $t = 0$ when the initial bail-ins are chosen. We derive the payoffs in this game by working backward through the timeline in Figure 1, focusing on the decision points labeled (a) – (d).

4.1 Remaining withdrawals

At point (d) in the timeline, a fraction π of the investors in each bank have already withdrawn and consumed. Let $h(\phi)$ denote the bail-in that has been imposed on these initial withdrawals from bank ϕ , and let $h : \Phi \rightarrow [0, 1]$ denote the profile of initial bail-ins across all banks. Let $b : \Phi \rightarrow \mathbb{R}_+$ denote the bailout payment that has been made to each bank by the fiscal authority at point (c) in the timeline. Then the amount of resources bank ϕ has

available for each of its $(1 - \pi)$ remaining investors at point (d) is given by

$$\psi(\phi) \equiv \frac{\phi - \pi(1 - h(\phi))c_1^* + b(\phi)}{1 - \pi}. \quad (8)$$

The composition of these remaining investors between patient and impatient types depends on whether or not the bank experienced a run at point (b) . Let $y : \Phi \rightarrow \{1, 2\}$ denote the withdrawal behavior of patient investors in each bank.¹⁷ Specifically, $y(\phi) = 2$ represents a situation where all patient investors in bank ϕ chose to withdraw in period 2. In this case, the first π withdrawals were made by impatient investors and, therefore, all of the bank's remaining investors at point (d) are patient. Each of these investors will receive $R\psi(\phi)$ at $t = 2$. If $y(\phi) = 1$, the bank experienced a run at point (c) and its remaining investors are a mixture of patient and impatient types. In this case, the bank is placed into resolution and the regulator will choose all remaining payments to maximize the sum of these remaining investors' utilities. It is straightforward to show that the solution to this allocation problem gives $\psi(\phi)c_1^*$ to each remaining impatient investor at $t = 1$ and $\psi(\phi)c_2^*$ to each remaining patient investor at $t = 2$.¹⁸ Letting $V(\psi, y)$ denote the average utility of a bank's remaining investors when its resources are allocated in this way, we then have

$$V(\psi(\phi); y(\phi)) \equiv \left\{ \begin{array}{c} u(R\psi(\phi)) \\ \pi u(\psi(\phi)c_1^*) + (1 - \pi)u(\psi(\phi)c_2^*) \end{array} \right\} \quad \text{as} \quad \left\{ \begin{array}{c} y(\phi) = 2 \\ y(\phi) = 1 \end{array} \right\}. \quad (9)$$

Given any profiles of initial bail-ins h , withdrawal behavior y , and bailouts b , equations (8) and (9) show how the resources remaining in each bank at point (d) in the timeline will be allocated among that bank's remaining investors.

4.2 Bailouts

Next, we consider point (c) in the timeline, where the fiscal authority makes bailout decisions and provides the public good. The fiscal authority knows that the remaining resources in each bank will be utilized as derived above. For a given pair (h, y) , the fiscal authority will choose the bailout policy b to maximize

¹⁷Note that we restrict attention to symmetric outcomes of the withdrawal game, in which all patient investors in a given bank take the same action. Given our focus on fundamentals-based bank runs, this approach is without loss of generality.

¹⁸This result can be viewed as an application of Proposition 1 to the regulator's allocation problem for a bank that has experienced a run. Recall that the proposition shows how the efficient allocation of resources within a bank scales the consumption of impatient and patient investors in proportion to the reference allocation (c_1^*, c_2^*) .

$$W(h, y, b) \equiv \int_{\underline{\phi}}^1 (1 - \pi) V(\psi(\phi); y(\phi)) dF(\phi) + v(g) \quad (10)$$

subject to

$$0 \leq b(\phi) \leq \bar{b}(\phi, h(\phi), y(\phi)) \quad \text{for all } \phi,$$

where the remaining resources $\psi(\phi)$ of each bank depend on the bailout policy $b(\phi)$ as shown in equation (8) and the level of public good is given by

$$g = \tau - \int_{\underline{\phi}}^1 b(\phi) dF(\phi). \quad (11)$$

The objective function in equation (10) illustrates the trade-off faced by the fiscal authority: bailouts raise the private consumption of investors through $\psi(\phi)$, but decrease the provision of the public good g for all agents.¹⁹ The upper bound $\bar{b}$ on the bailout for bank ϕ ensures that no investors receive more consumption than in the reference allocation (c_1^*, c_2^*) . This bound can be shown to equal

$$\bar{b}(\phi, h(\phi), y(\phi)) \equiv 1 - \phi - h(\phi) \pi c_1^* + \mathbb{I}_{(y(\phi)=1)} \pi (c_1^* - 1).$$

Intuitively, the bailout payment cannot be larger than the bank's loss $1 - \phi$, minus the bail-in already imposed on the π investors who have withdrawn, plus an adjustment to reflect the misallocation of resources that occurs if the bank has experienced a run.²⁰

The first-order conditions for this problem require that, for all ϕ , we have either

$$V_1(\psi(\phi), y(\phi)) \leq v'(g) \quad \text{and} \quad b(\phi) [v'(g) - V_1(\psi(\phi), y(\phi))] = 0 \quad (12)$$

or

$$V_1(\psi(\phi), y(\phi)) > v'(g) \quad \text{and} \quad b(\phi) = \bar{b}(\phi, h(\phi), y(\phi)), \quad (13)$$

where the derivative V_1 represents the marginal value of resources in bank ϕ . If the choice of $b(\phi)$ is interior, it must be the case that V_1 is equal to the marginal value of the public good. If the marginal utility of a bank's investors is lower than the marginal value of public consumption, that bank receives no bailout. In the opposite case, the bank receives the maximum bailout and its remaining investors are not bailed-in. Note that the marginal value of public consumption depends on the bailout payments made to all banks, which

¹⁹See Mitkov (2019) for a model with ex ante heterogeneity across investors where bailout policy is shaped by distributional considerations in addition to these concerns.

²⁰We show below that equation (4) guarantees this upper bound is never binding in equilibrium. It may, however, bind when solving the fiscal authority's choice problem for an arbitrary profile h , as we do here.

implies that these first-order conditions must be solved jointly for all ϕ .

One way to understand the implications of these first-order conditions is to look at the level of bail-ins faced by each bank's remaining investors. Define the *remaining bail-in* function $\hat{h} : \Phi \rightarrow [0, 1]$ by

$$\hat{h}(\phi) = \begin{cases} 1 - \frac{R}{c_2^*} \psi(\phi) \\ 1 - \psi(\phi) \end{cases} \quad \text{as} \quad y(\phi) = \begin{cases} 2 \\ 1 \end{cases}, \quad (14)$$

where $\psi(\phi)$ depends on the bailout policy as shown in equation (8). We then have the following result.

Proposition 4. For any (h, y) , there exists a unique $\hat{h}_{max} \in [0, 1]$ such that the fiscal authority's bailout policy b will imply

$$\hat{h}(\phi) \leq \hat{h}_{max} \text{ for all } \phi, \text{ with equality if } b(\phi) > 0.$$

This result shows that, after the first π withdrawals have occurred, the allocation of the remaining losses in the banking system takes the same general form as the solution to the planner's problem characterized in Proposition 2. In particular, the investors in any bank that receives a bailout will experience the same bail-in. Investors in a bank that is not bailed out experience a smaller bail-in that fully covers the bank's remaining losses.

Looking ahead, Proposition 4 also illustrates how the equilibrium bailout policy will tend to distort banks' choice of initial bail-in $h(\phi)$. Among all banks receiving a bailout, this result implies that the fiscal authority will give a larger payment to those banks that have fewer remaining resources. In other words, a bank can increase the bailout it receives by imposing a smaller bail-in on its first π investors to withdraw. This incentive distortion creates a wedge between the equilibrium outcome in our model and the solution to the planner's problem in Section 3. Before discussing this wedge in detail, however, we need to analyze investors' withdrawal choices.

4.3 Withdrawal choices

At point (b) in the timeline, investors choose when to withdraw from their bank. Impatient investors only value consumption in period 1 and, therefore, will always choose to withdraw early. A patient investor will choose to withdraw in whichever period she would receive a higher payment from her bank. She anticipates that the bailout payments at point (c) and the subsequent bail-in of remaining investors at point (d) will be as described above. Moreover, if her bank experiences a run and she is not among the first π investors to withdraw, she

knows the bank will be placed into resolution and she will receive the payment chosen by the regulator in period 2. Using the function $\hat{h}$ defined in equation (14), we can say that waiting to withdraw is a best response for a patient investor in bank ϕ if and only if

$$(1 - h(\phi))c_1^* \leq (1 - \hat{h}(\phi))c_2^*. \quad (15)$$

Recall that the bail-in of remaining investors $\hat{h}(\phi)$ depends on the profile of withdrawal behavior y both directly, as shown in equation (14), and indirectly through its effect on the bailout policy and thus on the bank's remaining resources in equation (8). For this reason, the optimal choice for an individual patient investor may depend on the choices of all other patient investors in the economy. We derive the optimal decision rule for a patient investor in two steps, first looking at a single bank in isolation and then considering all banks together.

The withdrawal game within a bank. To begin, we focus on the withdrawal game played by the patient investors within a single bank, holding the actions of investors in all other banks fixed. Suppose the bank is receiving a bailout, that is, $b(\phi) > 0$. Since the size of the bailout payment to a single bank has a negligible effect on the government's finances, the level of the public good g in the first-order conditions (12) – (13) is independent of the initial bail-in and withdrawal behavior in bank ϕ . This fact, together with equations (9) and (14), implies that the bail-in $\hat{h}(\phi)$ of the bank's remaining investors will be the same regardless of how many investors attempt to withdraw early. Intuitively, if the bank experiences a run, the fiscal authority will choose to increase the bailout payment precisely so that the consumption level of the bank's remaining patient investors remains unchanged. When an individual patient investor is evaluating whether condition (15) holds, therefore, none of the terms in the expression depend on the withdrawal behavior of the other investors in her bank. It follows that, apart from the knife-edge case in which she is indifferent between withdrawing early and waiting, the withdrawal game within any individual bank with $b(\phi) > 0$ has a *unique equilibrium*. The strategic complementarity that usually generates multiple equilibria in the Diamond-Dybvig framework is eliminated here by the fiscal authority's choice of bailout policy.

For a bank that is not receiving a bailout, in contrast, the standard strategic complementarity is present and the withdrawal game within the bank may have multiple equilibria. In particular, early withdrawals by other patient investors in the bank increase the bail-in $\hat{h}(\phi)$ of the remaining investors and thus increase the incentive to withdraw early. Given that this type of bank run has been extensively studied, we do not focus on it here. Instead, we assume that patient investors withdraw early only if doing so is a strictly dominant strategy

of the withdrawal game within their bank.

To operationalize this assumption, we define $\hat{h}_{NB}(\phi)$ to be the bail-in that would be imposed on bank ϕ 's remaining investors if (i) all of its patient investors withdraw at $t = 2$ and (ii) it receives no bailout. Using equations (8) and (14) with $b(\phi) = 0$, we have

$$\hat{h}_{NB}(\phi) = 1 - \left(\frac{R}{c_2^*} \right) \frac{\phi - \pi(1 - h(\phi))c_1^*}{1 - \pi}. \quad (16)$$

Using Proposition 4, we can determine whether or not bank ϕ will receive a bailout when it does not experience a run by comparing $\hat{h}_{NB}(\phi)$ with $\hat{h}_{max}$, which depends on the average choices of all banks but not on the individual choice of bank ϕ . If $\hat{h}_{NB}(\phi) > \hat{h}_{max}$ holds, bank ϕ will be bailed out and equilibrium play in the withdrawal game within the bank will depend on the choice of initial bail-in $h(\phi)$ according to

$$y(\phi) = \left\{ \begin{array}{c} 2 \\ 1 \end{array} \right\} \quad \text{as} \quad \left(1 - \hat{h}_{max}\right) c_2^* \left\{ \begin{array}{c} \geq \\ < \end{array} \right\} (1 - h(\phi)) c_1^*. \quad (17)$$

If $\hat{h}_{NB}(\phi) \leq \hat{h}_{max}$ holds, bank ϕ will not receive a bailout and our assumption that patient investors only run if doing so is a strictly dominant strategy can be written as

$$y(\phi) = \left\{ \begin{array}{c} 2 \\ 1 \end{array} \right\} \quad \text{as} \quad \left(1 - \hat{h}_{NB}(\phi)\right) c_2^* \left\{ \begin{array}{c} \geq \\ < \end{array} \right\} (1 - h(\phi)) c_1^*. \quad (18)$$

We refer to (17) and (18) together as the *equilibrium withdrawal behavior* of investors within bank ϕ . When a bank is choosing its initial bail-in $h(\phi)$, it recognizes that these conditions determine whether or not it will experience a run.

The overall withdrawal game. We now turn to the overall withdrawal game, in which all investors in all banks simultaneously make their withdrawal decisions. We define an equilibrium of this game as follows.

Definition 1. Given a profile of initial bail-ins h , an *equilibrium of the overall withdrawal game* is a profile $y^e : \Phi \rightarrow \{1, 2\}$ such that:

- (i) given $\hat{h}_{max}$, $y^e(\phi)$ satisfies (17) or (18), as appropriate, for each ϕ , and
- (ii) given y^e , $\hat{h}_{max}$ is determined by the bailout policy characterized in (12) and (13).

While our assumptions above assign a unique equilibrium to the withdrawal game within each bank, there still may be multiple equilibria of the overall withdrawal game. This multiplicity arises because the outcomes of the withdrawal games within other banks affect

the government's fiscal position and, through the bailout policy, affect the value of the maximum bail-in $\hat{h}_{max}$. A change in $\hat{h}_{max}$, in turn, may affect the equilibrium withdrawal behavior of investors in bank ϕ .²¹ Given our focus on fundamentals-based runs, we select the equilibrium that has a run occurring at the smallest measure of banks when multiple equilibria exist, which we call the *minimal-run equilibrium* associated with a given profile of initial bail-ins h . Our next result shows that there is a unique minimal-run equilibrium associated with each h . An algorithm for computing the minimal-run equilibrium is provided as part of the proof of this result in the appendix.

Proposition 5. For each profile of initial bail-ins h , there is unique minimal-run equilibrium of the overall withdrawal game.

To summarize the analysis so far in this section, it is helpful to look back at the timeline in Figure 1. Given any profile of initial bail-ins h chosen by banks at point (a) of the timeline, the results above uniquely determine investors' withdrawal behavior at point (b), the bailout payments made by the government at point (c), and the allocation of the remaining resources in each bank at point (d). In other words, we now have a mapping from banks' choice of initial bail-ins h to the entire allocation of resources in the economy. This mapping determines the payoffs of the *bail-in game*.

4.4 Equilibrium of the bail-in game

At point (a) in the timeline, each bank chooses its initial bail-in $h(\phi)$, taking the choices of all other banks as given. Let $\mathcal{W}$ denote the expected utility of investors in an individual bank as a function of the full profile of initial bail-ins and the bank's resources, that is,

$$\mathcal{W}(h(\phi), h_{-\phi}; \phi) \equiv \pi u((1 - h(\phi))c_1^*) + (1 - \pi)V(\psi(\phi); y^e(\phi)), \quad (19)$$

where $h_{-\phi}$ represents the initial bail-ins at all banks except bank ϕ . Recall that $\psi(\phi)$ depends on the bailout $b(\phi)$, which in turn depends on both $h(\phi)$ and $h_{-\phi}$ as determined by equations (12) – (13). The withdrawal behavior $y^e(\phi)$ also depends on both $h(\phi)$ and $h_{-\phi}$ as determined by equations (17) – (18) together with the selection of the minimum-run equilibrium of the overall withdrawal game.

²¹In other words, there is a complementarity in the withdrawal decisions of investors *across* banks that operates through the public sector's budget constraint. This complementarity is also present in Mitkov (2019), where early withdrawals at other banks can undermine the government's *ex post* willingness to fully insure deposits at an individual bank, giving its investors a stronger incentive to withdraw early. A similar cross-bank complementarity arises in Goldstein et al. (2020), where early withdrawals at other banks can drive down asset prices, which strengthens the strategic complementarity in the withdrawal decisions of investors within a given bank.

The first term on the right-hand side of equation (19) is decreasing in $h(\phi)$, while the second term is increasing in $h(\phi)$ if we hold $b(\phi)$ and $y^e(\phi)$ fixed. These direct effects of the initial bail-in are straightforward: imposing losses on the first π investors to withdraw makes these investors worse off, but leaves more resources for the bank's remaining investors. However, the bank's choice of $h(\phi)$ may also affect both the bailout it receives $b(\phi)$ and the withdrawal behavior of its investors $y^e(\phi)$. Moreover, these relationships depend on the government's fiscal position, which in turn depends on the initial bail-ins chosen by other banks. These payoff spillovers imply that each bank's desired initial bail-in will, in general, depend on the initial bail-ins chosen by other banks.

Given these payoffs, the definition of equilibrium for the bail-in game is straightforward.

Definition 2. An *equilibrium of the bail-in game* is a profile of strategies $h^e : \Phi \rightarrow [0, 1]$ such that, for all $\phi \in \Phi$, we have

$$\mathcal{W}(h^e(\phi), h_{-\phi}^e; \phi) \geq \mathcal{W}(h, h_{-\phi}^e; \phi) \quad \text{for all } h \in [0, 1].$$

Our next result shows that an equilibrium of the bail-in game always exists in pure strategies.

Proposition 6. There exists a pure strategy equilibrium of the bail-in game.

The remainder of the paper studies the properties and policy implications of equilibrium in the bail-in game. In the next section, we derive conditions under which the equilibrium allocation is inefficient and under which bank runs occur in equilibrium. In Section 6, we study how regulation can improve equilibrium outcomes.

5 Properties of equilibrium

We begin the analysis of equilibrium play in the bail-in game by discussing the incentives each bank faces when choosing its initial bail-in.

5.1 Incentives to bail in

Consider first the decision problem facing a bank that anticipates receiving no bailout, regardless of its choice of initial bail-in $h(\phi)$. In this case, the bail-ins chosen by other banks, $h_{-\phi}$, have no effect on the bank's payoffs. It is straightforward to show that this bank's optimal choice is $h(\phi) = 1 - \phi$. Substituting this choice into equations (8) and (14) when $b(\phi) = 0$ yields $\hat{h}(\phi) = 1 - \phi$ as well, meaning that the bank's losses are shared evenly by all of its investors. Notice that a bank in this situation will choose bail-ins that match those

chosen by the planner for a bank that receives no bailout (see Propositions 1 and 2). Given this choice, patient investors in this bank will have no incentive to withdraw early and the withdrawal behavior assigned by equation (18) is $y^e(\phi) = 2$.

For a bank that does anticipate receiving a bailout, Proposition 4 shows that the bail-in experienced by the bank's remaining investors will equal $\hat{h}_{max}$, which depends on aggregate conditions but not on the bank's own choice of initial bail-in $h(\phi)$. If we hold the withdrawal decisions of the bank's investors fixed, therefore, it will want to set the lowest bail-in possible, $h(\phi) = 0$. In this way, bailouts distort the incentives of banks in the bail-in game. Why impose any loss on the bank's first π investors if doing so reduces the bailout the bank will receive dollar-for-dollar?

There is one caveat to this logic: in some situations, setting $h(\phi) = 0$ will lead to a run on the bank. In these cases, the bank can prevent the run by setting its initial bail-in appropriately. For a given value of $\hat{h}_{max}$, the withdrawal behavior specified in equation (17) shows that investors will not run on a bank that is being bailed out as long as

$$(1 - \hat{h}_{max})c_2^* \geq (1 - h(\phi))c_1^*.$$

Let $\underline{h}$ denote the smallest initial bail-in $h(\phi)$ that will prevent a run, that is,

$$\underline{h} \equiv \max \left\{ 1 - \left(1 - \hat{h}_{max} \right) \frac{c_2^*}{c_1^*}, 0 \right\}. \quad (20)$$

Note that $c_2^* > c_1^*$ implies $\underline{h} \leq \hat{h}_{max}$, with strict inequality whenever $\hat{h}_{max} > 0$. In other words, the initial bail-in required to prevent a run is always smaller than the bail-in that will apply to the bank's remaining investors.

If $\underline{h} = 0$, meaning the bank will not experience a run even if the initial bail-in is zero, setting $h(\phi) = 0$ is clearly its best response. If $\underline{h}$ is positive, the bank must choose between (i) imposing no initial bail-in but suffering a run and (ii) imposing the bail-in $\underline{h}$ on its first π investors to withdraw. Which of these two choices is optimal depends on the value of $\hat{h}_{max}$, which is determined in equilibrium. Specifically, preventing a run will be optimal for bank ϕ if the expected utility of its investors when the initial bail-in is zero and a run occurs is smaller than the expected utility associated with an initial bail-in of $\underline{h}$ and no run, that is,

$$\begin{aligned} & \pi u(c_1^*) + (1 - \pi) \left[\pi u \left(\left(1 - \hat{h}_{max} \right) c_1^* \right) + (1 - \pi) u \left(\left(1 - \hat{h}_{max} \right) c_2^* \right) \right] \\ & \leq \pi u \left((1 - \underline{h}) c_1^* \right) + (1 - \pi) u \left(\left(1 - \hat{h}_{max} \right) c_2^* \right). \end{aligned}$$

Rearranging terms, we can write this condition as

$$\pi [u(c_1^*) - u((1 - \underline{h})c_1^*)] \leq \pi(1 - \pi) \left[u\left(\left(1 - \hat{h}_{max}\right)c_2^*\right) - u\left(\left(1 - \hat{h}_{max}\right)c_1^*\right) \right]. \quad (21)$$

The left-hand side of equation (21) measures the gain from setting the initial haircut to zero: the first π investors to withdraw from the bank receive c_1^* rather than only a fraction of this amount. The right-hand side of the equation measures the cost of a run: because some of the first π withdrawals will be made by patient investors, a fraction π of the remaining $(1 - \pi)$ investors will be impatient and need to consume in period 1. Because $c_2^* > c_1^*$, the total consumption of the bank's investors will then be lower than if there had been no run and all of the remaining investors were patient.

Notice that none of the terms in equation (21) depend on ϕ , which implies that all banks receiving a bailout will choose the same initial bail-in. We use $\tilde{h} \in \{0, \underline{h}\}$ to denote this common choice. For a given value of $\hat{h}_{max}$, equation (21) indicates that these banks will choose to prevent a run whenever the necessary bail-in $\underline{h}$ is sufficiently small.

5.2 Inefficiency of equilibrium

Our next result builds on the discussion above to derive the equilibrium allocation of losses between bail-ins and bailouts. It shows that any equilibrium is characterized by a cutoff ϕ^e such that only those banks whose realized ϕ falls below the cutoff are bailed out.

Proposition 7. In any equilibrium of the bail-in game, there exists $\phi^e \in \Phi$ such that

$$h^e(\phi) = \begin{cases} 1 - \phi \\ \tilde{h} \in \{0, \underline{h}\} \end{cases} \quad \text{and} \quad \begin{cases} b^e(\phi) = 0 \\ b^e(\phi) > 0 \end{cases} \quad \text{as} \quad \phi \begin{cases} > \\ < \end{cases} \phi^e.$$

The value of ϕ^e satisfies the following condition

$$\mathcal{W}(\tilde{h}, h_{-\phi}^e; \phi) \begin{cases} > \\ = \\ < \end{cases} U(1 - \phi) \quad \text{as} \quad \phi \begin{cases} < \\ = \\ > \end{cases} \phi^e.$$

To understand how the equilibrium bailout cutoff ϕ^e is determined, consider first a bank whose realized ϕ is so small that it would receive a bailout even if it set the “full” initial bail-in $h(\phi) = 1 - \phi$. Such a bank clearly faces the incentive distortion described above and will choose an initial bail-in of $\tilde{h} \in \{0, \underline{h}\}$ according to equation (21). At the other extreme, a bank with ϕ sufficiently close to 1 may not receive a bailout even if it set the smallest

possible bail-in, $h(\phi) = 0$. In this situation, the bank's best response is to spread the loss evenly among its investors by setting $h(\phi) = 1 - \phi$.

In between these two extremes, there is a range of values of ϕ for which the bank would be bailed out if it set $h(\phi) = \tilde{h}$, but not if it chose the initial bail-in $h(\phi) = 1 - \phi$. A bank in this range recognizes that if it takes the best action conditional on not being bailed out – allocating its loss evenly across its investors – it will not receive a bailout. If it instead chooses the smaller initial bail-in $\tilde{h}$, it will be bailed out. In choosing its initial bail-in, therefore, this bank is effectively *choosing* whether or not it will receive a bailout. The benefit of receiving a bailout is obvious: it raises the total consumption of the bank's investors. The cost is that, in order to qualify for the bailout, the bank must set its initial bail-in in a way that inefficiently allocates this consumption across its investors. The bank with $\phi = \phi^e$ is exactly indifferent between these two options in equilibrium.

This indifference implies that the following ordering must hold in equilibrium:

$$\underline{h} < 1 - \phi^e < \hat{h}_{max}. \quad (22)$$

If bank ϕ^e chooses to set $h(\phi) = \underline{h}$ instead of $h(\phi) = 1 - \phi^e$, its first π investors will experience a smaller bail-in, but its remaining $1 - \pi$ investors will experience the larger bail-in $\hat{h}_{max}$. In other words, qualifying for a bailout requires the bank to shift the burden of the losses away from its first π investors and onto the remaining fraction $1 - \pi$.

Figure 3 illustrates the equilibrium allocation of losses for two different values of the government's fiscal capacity τ . In panel (a), which is based on $\tau = 0.9$, all banks receiving a bailout set their initial bail-ins to zero. As a result, there is no light-blue region in the figure for banks with $\phi < \phi^e$; all of their losses fall on the remaining $1 - \pi$ investors and the public sector. Panel (b) of the figure is based on a smaller fiscal capacity, $\tau = 0.6$, which generates a lower bailout cutoff ϕ^e and a larger maximum bail-in $\hat{h}_{max}$. Banks that receive a bailout in this example would experience a run if they set their initial bail-in to zero. Instead, they choose to set their initial bail-in to $\underline{h} > 0$, as indicated by the presence of the light-blue region in the figure for $\phi < \phi^e$. Notice that both panels in the figure illustrate the ordering in equation (22): the initial bail-in for all banks below the cutoff ϕ^e is smaller than for those banks just above the cutoff, while the bail-in for their remaining investors is larger.

Comparing panel (a) in Figures 2 and 3, which are based on the same parameter values, it is apparent for this example that bail-ins cover smaller portion of the losses in equilibrium than in the planner's allocation and that bailouts cover a larger portion. Our next result shows that these properties are general.

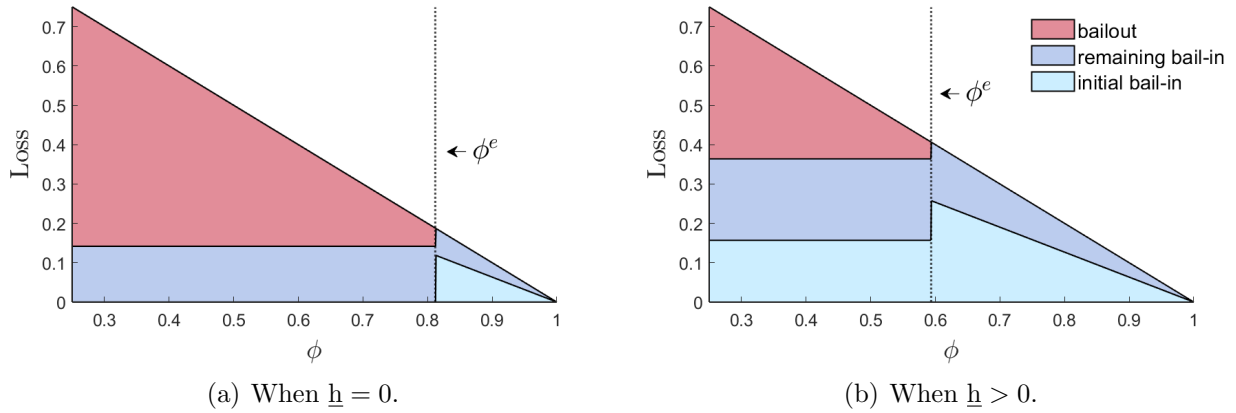


Figure 3: Equilibrium allocation of losses

Proposition 8. In any equilibrium of the bail-in game, $h^e(\phi) < h^*(\phi)$ holds for all ϕ with $b^e(\phi) > 0$. In addition, if $b^e(\phi) > 0$ holds for some ϕ , we have

$$\int_{\underline{\phi}}^1 b^e(\phi) dF(\phi) > \int_{\underline{\phi}}^1 b^*(\phi) dF(\phi).$$

Whenever bailouts occur in equilibrium, the initial bail-ins chosen by those banks being bailed out are smaller than in the planner's allocation. In this sense, bailouts *undermine* banks' incentive to choose the socially-efficient bail-in. In addition, the fact that the initial bail-ins are too small at some banks implies that the total bailout expenditure is larger in equilibrium than in the planner's solution. The equilibrium allocation is efficient only if the fiscal capacity of the government is small enough that no bailouts occur in equilibrium.

5.3 Bank runs and multiplicity

In the planner's allocation, investors withdrawing at $t = 2$ always receive more than investors withdrawing from the same bank at $t = 1$. In equilibrium, however, this relationship may not always hold. Proposition 8 shows that, for banks that are bailed out in equilibrium, the initial bail-in is smaller than in the planner's allocation. In addition, the ordering in equation (22) implies that the equilibrium bail-ins applied at $t = 2$ in banks receiving bailouts are larger than in the planner's allocation. In some economies, these distortions are large enough that investors who withdraw early receive more from their bank than investors who wait until $t = 2$. These banks experience a run in which all investors attempt to withdraw at $t = 1$.

Additional losses. A bank run creates a misallocation of resources because some patient investors are served in period 1, before the bank's investment has matured. As a result, the total loss that must be allocated between bail-ins and bailouts in such a bank is larger than $1 - \phi$. For a bank experiencing a run, the feasibility constraint (2) becomes

$$\pi (1 - h(\phi)) c_1^* + (1 - \pi) (1 - \hat{h}(\phi)) \left(\pi c_1^* + (1 - \pi) \frac{c_2^*}{R} \right) \leq \phi + b(\phi).$$

As the second term on the left-hand side of this constraint indicates, a fraction π of the $1 - \pi$ investors who remain in the bank when it is placed in resolution are impatient and will withdraw at $t = 1$. The other fraction $(1 - \pi)$ are patient and will withdraw at $t = 2$. The resolution process will bail in all of these investors at a common rate $\hat{h}(\phi)$. Using the resource constraint in equation (1) with equality, we can rewrite this feasibility constraint as

$$h(\phi) \pi c_1^* + \hat{h}(\phi) (1 - \pi) + b(\phi) \geq 1 - \phi + \pi (1 - \pi) \left(c_1^* - \frac{c_2^*}{R} \right). \quad (23)$$

The left-hand side of this expression is the sum of bail-ins and bailouts, as in equation (2). The right-hand side is the loss on the bank's assets *plus* the cost of the misallocation created by the run, in which an additional measure $\pi (1 - \pi)$ of investors are served at $t = 1$ rather than at $t = 2$. That fact that the reference allocation satisfies $c_1^* > 1$ and $c_2^* < R$ implies that this misallocation cost is always strictly positive.

Figure 4 depicts the allocation of losses for an economy in which a bank run occurs in equilibrium. For banks that are not bailed out ($\phi > \phi^e$), the bail-ins cover the loss on the bank's assets, as before. For banks that are bailed out, however, the sum of bail-ins and bailouts exceeds the loss on the bank's assets, $1 - \phi$, because of the misallocation. The additional area at the top of the red region corresponds to the final term in equation (23) for these banks.

Multiple equilibria. For some parameter values, the bail-in game will have multiple equilibria: one in which a run occurs on those banks being bailed out and another in which no run occurs. This multiplicity arises because of a strategic complementarity in the bail-in choices of banks receiving a bailout. Consider the decision problem of a bank that anticipates being bailed out, and suppose the bank's investors will run if it sets its initial bail-in to zero. Then the bank will set $h(\phi) = \underline{h}$ and avoid the run if condition (21) holds and will set $h(\phi) = 0$ if the inequality is reversed. Note that the right hand side of condition (21) depends on $\hat{h}_{max}$, the bail-in of the bank's remaining investors, which in turn depends on the size of the bailout it receives.

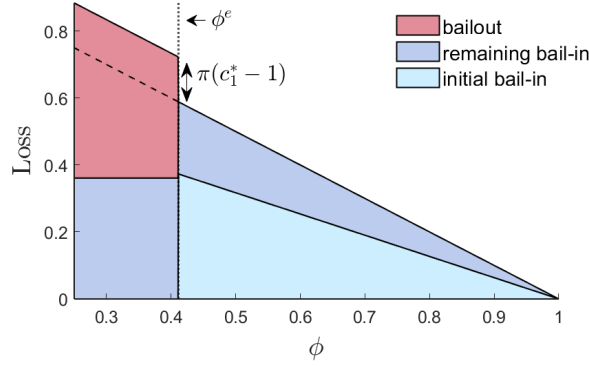


Figure 4: Allocation of losses with a run on some banks

When other banks choose an initial bail in of zero, they will be in worse condition after the first π withdrawals for two reasons: no loss was imposed on their first π investors to withdraw and the run has increased their total losses as shown in equation (23). This fact will lead the fiscal authority to provide larger bailouts to these banks. These larger bailouts, in turn, worsen the fiscal position for the government and lead to larger bail-ins $\hat{h}_{max}$ for the remaining investors in all banks receiving bailouts. When $\hat{h}_{max}$ is larger, the right hand side of equation (21) is smaller, making it more likely that the condition is violated and an individual bank will find it optimal to set $h(\phi) = 0$. In other words, when other banks choose not to bail in their first π investors, it becomes more attractive for an individual bank to take the same action, even though doing so causes a run on the bank. Conversely, when other banks choose an initial bail-in to prevent a run, preventing a run becomes more attractive to an individual bank.

Our next result shows that this strategic complementarity is the only source of multiplicity of equilibrium in the bail-in game.

Proposition 9. Either equilibrium in the bail-in game is unique or there are exactly two pure-strategy equilibria, one in which no bank runs occur and one in which a run occurs on all banks that are bailed out.

It bears emphasizing that this source of strategic complementarity is fundamentally different from the usual complementarity in withdrawal decisions that arises in models in the Diamond-Dybvig tradition. As described above, we remove that source of multiplicity by assuming that investors withdraw early only if doing so is a strictly dominant strategy of the withdrawal game within their bank. We also select the (unique) minimal-run equilibrium in the overall withdrawal game across banks. The multiplicity that arises in our model comes from a strategic complementarity in how banks allocate their resources, recognizing that the allocation in one bank affects others through the government's bailout policy. For the

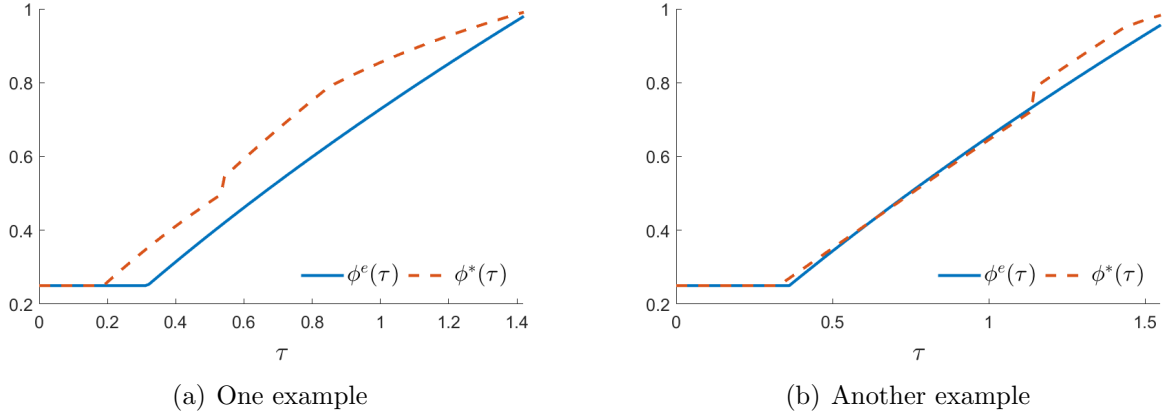


Figure 5: The bailout cutoff

remainder of the analysis, when multiple equilibria exist in the withdrawal game, we select the equilibrium in which no runs occur on any banks.

5.4 The equilibrium bailout cutoff

As suggested by the two panels of Figure 3, the equilibrium cutoff ϕ^e is increasing in the amount of resources τ available to the government. If τ is sufficiently small, ϕ^e will equal $\underline{\phi}$ and the bail-in covers the total loss in all banks. When τ is larger, the equilibrium bailout cutoff is interior, which implies that some banks' losses are divided between the public sector and the private sector. Our next proposition formalizes these results.

Proposition 10. The equilibrium bailout cutoff ϕ^e is increasing in τ , and is strictly increasing whenever $\phi^e \in (\underline{\phi}, 1)$.

This result is illustrated in Figure 5, which plots the equilibrium cutoff ϕ^e as a function of τ for two different sets of parameters.²² In both cases, the cutoff is strictly increasing in τ whenever it lies above $\underline{\phi}$.

The figure also shows that the equilibrium bailout cutoff ϕ^e can be either larger or smaller than the planner's cutoff ϕ^* . To understand why, consider the bank that lies exactly at the planner's bailout cutoff. In the planner's allocation, the bail-in applied to this bank's investors is the same as that applied in all banks that receive a bailout: $1 - \phi^*$. (See Proposition 2.) Suppose first that we hold the bail-in of the bank's remaining $1 - \pi$ investors fixed at $1 - \phi^*$ and allow the bank to choose its initial bail-in. In this situation, the bank would clearly choose a smaller bail-in (either zero or $\underline{h}$, according to condition (21)) since

²²Panel (a) uses the same parameter values as Figure 2, while in panel (b) the coefficient of relative risk aversion for the functions u and v is increased from 2 to 10.

the extra resources paid out to the first π investors are fully replaced by a bailout. The same logic applies to banks with realizations slightly above ϕ^* , which implies that these banks would also choose smaller initial bail-ins and receive bailouts. In this way, the moral hazard problem in the marginal bank's choice of initial bail-in tends to push the equilibrium bailout cutoff ϕ^e higher relative to the planner's cutoff ϕ^* .

When we take into account the change in behavior of other banks, however, a countervailing effect arises. In equilibrium, all banks receiving a bailout set their initial bail-ins smaller than in the planner's allocation. These actions will lead the fiscal authority to spend more resources on bailout payments, which decreases the level of the public good and raises the marginal value of funds in the public sector. As a result, the bail-in $\hat{h}_{max}$ applied to the remaining investors in banks that are bailed out increases. Returning to the decision problem of bank ϕ^* , this increase in $\hat{h}_{max}$ makes it *less* attractive for the bank to lower its initial bail-in, since the bailout it receives will be less generous. This same logic applies to banks with realizations slightly below ϕ^* : choosing $h(\phi) = 1 - \phi$ and not receiving a bailout becomes more attractive. In this way, the moral hazard problem at those banks that have large losses and are bailed out in equilibrium tends to discourage the marginal bank from seeking a bailout and thereby push the equilibrium bailout cutoff ϕ^e *lower* relative to the planner's cutoff ϕ^* .

In panel (a) of Figure 5, the equilibrium cutoff is higher than the planner's cutoff for all values of τ , meaning that more banks receive bailouts in equilibrium than in the planner's allocation. In panel (b), in contrast, the countervailing effect dominates for some values of τ and fewer banks receive bailouts than in the planner's allocation. Recall that, even in these cases, total bailout payments are larger in equilibrium than in the planner's allocation, as established in Proposition 8.

It is also worth noting that the equilibrium bailout cutoff increases discontinuously in both panels, at about $\tau = 0.5$ in panel (a) and at about $\tau = 1.2$ in panel (b). These jumps reflect a shift in the initial bail-in chosen by banks being bailed out from zero to $\underline{h}$. This shift improves the fiscal position of the government, which leads to banks with even higher realizations of ϕ receiving bailouts.

5.5 Discussion

Source of fragility. In choosing their initial bail-in $h(\phi)$, banks know whether or not each choice will lead to a run by their investors. A run leads to a misallocation of the bank's resources and lowers the welfare of its investors in much the same way as the existing literature. Moreover, the bank always has the ability to prevent a run by choosing a larger initial

bail-in. However, by putting the bank in better financial condition, this bail-in may reduce the bailout payment the bank receives. In some cases, the bank would choose to tolerate a run as a side effect of the plan that maximizes the total amount of payments it can make to its investors, including those financed by a bailout. This approach is in sharp contrast to the existing literature, which assumes that banks are initially unaware that a run is underway and/or cannot take any action to prevent the run.²³ The source of financial fragility in our model is novel: the anticipation of being bailed out can make bail-ins unattractive, even when a bail-in could prevent a run.

Disciplining effect of runs. In other cases, such as the one depicted in panel (b) of Figure 3, banks do find it optimal to impose the bail-in $\underline{h}$ to prevent a run. In these situations, the threat of a run is *disciplining* the behavior of banks that will be bailed out in the sense of moving their choice of initial bail-in closer to what the planner would choose. This disciplining role of bank runs is similar in spirit to Calomiris and Kahn (1991) and Diamond and Rajan (2001), where depositors design a fragile banking contract that will lead to a run if the banker tries to misappropriate funds. In our setting, fragility is not a design choice; it emerges naturally from the incentives created by the bailout policy chosen by a government with limited commitment. Inspired by this earlier literature, however, one might ask whether the regulator can use this disciplining effect of runs to design regulations that improve on the equilibrium allocation. We address this question in the next section.

6 Regulation

Proposition 8 demonstrates that the allocation of resources resulting from equilibrium play in the bail-in game is inefficient: the initial bail-ins are too small in the aggregate and the bailouts are too large. These distortions can, in addition, lead to runs on those banks that are bailed out, as illustrated in Figure 4. These results naturally raise the question of whether regulation can improve equilibrium outcomes. The regulator in our model can restrict the payments made by banks to the first π investors who withdraw. If the regulator could observe the realized ϕ in each bank, it could mandate that each bank set the same initial bail-in as the planner, $h^*(\phi)$. It is straightforward to show that the subsequent actions – the withdrawal decisions of investors, the bailouts chosen by the fiscal authority, and the remaining bail-ins

²³Diamond and Dybvig (1983), Cooper and Ross (1998), Goldstein and Pauzner (2005) and many others assume the bank must pay the promised amount to depositors at $t=1$ until it has run out of resources; no bail-in is allowed. Peck and Shell (2003), Ennis and Keister (2009b), and Sultanum (2014) and others allow the bank to freely bail-in investors, but assume withdrawal decisions can be conditioned on an extrinsic, sunspot variable that is not observed by the bank.

$\hat{h}$ – would all match the planner’s allocation as well. However, the regulator faces the same information friction as the rest of the public sector: it observes bank-specific information with a delay. As the first fraction π of investors withdraw, the regulator only knows the aggregate state and the distribution of losses across banks.

What should the regulator do? In this section, we study three different ways the regulator might restrict banks’ choice of initial bail-in and discuss the implications of each policy for equilibrium behavior and welfare.

6.1 System-wide mandatory bail-ins

The simplest option is for the regulator to mandate a system-wide bail-in, that is, a common initial bail-in to be applied at all banks. We denote the level of this bail-in by h_{req} . Under this policy, each bank’s choice set in the bail-in game is a single element, $h(\phi) = h_{req}$, so equilibrium in the game is trivial. The regulator chooses h_{req} to maximize the sum of utilities of all investors in the economy.

Looking back at the two panels in Figure 3, this policy will make the height of the light blue region, which measures the losses imposed on the first π investors to withdraw, uniform across all banks. The primary benefit of this policy is that it can be used to increase the initial bail-ins at those banks that are bailed out in equilibrium, moving them closer to the planner’s allocation. One cost of the policy is that it requires an initial bail-in that is larger than appropriate at banks with small or no losses. These points highlight the fundamental tradeoff facing a regulator with limited information: mandatory bail-ins improve the allocation of resources in some banks, but are counterproductive for others. The optimal policy of this type will tend to raise welfare when the distribution F has relatively few banks with small or no losses and when the fiscal capacity of the government is large, so that many banks are bailed out in the equilibrium with no regulation. In such cases, the benefit of increasing the initial bail-in at banks whose incentives are distorted by the anticipation of being bailed out is large, while the aggregate cost of imposing a bail-in on banks whose incentives are not distorted will be small because relatively few banks are in this situation.

There is another, more subtle cost of a system-wide bail in policy: the mandated bail-in may be *smaller* than what some banks choose in the absence of regulation. Looking again at Figure 3, consider a bank that is just above the equilibrium bailout cutoff ϕ^e when there is no regulation. Because it is not being bailed out, this bank chooses $h^e(\phi) = 1 - \phi$ to allocate its loss evenly across its investors. The best mandatory bail-in h_{req} may require that this bank *decrease* its initial bail-in, which implies shifting some of the loss away from the

first π investors to withdraw and onto the remaining $1 - \pi$ investors. This shift not only distorts the allocation of resources in the bank, it can also push the bank into the bailout region and, in some cases, may even provoke a run on the bank. In such a situation, both the regulator and the investors in these banks would prefer that the bail-in be set higher than the mandated value. This logic points to a potentially better policy option: allowing banks, at their discretion, to set a larger initial bail-in than the one specified by the regulator. We call this type of policy a *mandatory minimum bail-in*.

6.2 Mandatory minimum bail-ins

Suppose now that the regulator chooses a minimum bail-in requirement, denoted h_{min} , but allows banks to set a larger initial bail-in if doing so raises the expected utility of their investors. Specifically, banks' choice of initial bail-in must now satisfy

$$h(\phi) \geq h_{min} \quad \text{for all } \phi \in \Phi.$$

Under this policy, banks once again have a non-trivial choice set in the bail-in game, and we study equilibrium in this game as specified in Definition 2, but with the choice set for each bank adjusted to $h \in [h_{min}, 1]$.

Improving the allocation of resources. The regulator will choose h_{min} to maximize the expected utility of all investors in the economy; let h_{min}^* denote the optimal policy of this type. Note that the economy without regulation studied in Section 5 corresponds to the special case of this policy where $h_{min} = 0$. A mandatory minimum bail-in policy will be useful, therefore, if and only if the optimal policy satisfies $h_{min}^* > 0$. Our next proposition provides a sufficient condition for this result to obtain.

Proposition 11. If the equilibrium of the economy without regulation has $h^e(\phi) = 0$ for those ϕ with $b^e(\phi) > 0$, then $h_{min}^* > 0$.

This result applies when, in the equilibrium with no regulation, banks that are bailed out choose an initial bail-in of zero, as in panel (a) of Figure 3. In setting $h_{min} > 0$, the regulator faces the tradeoff described above. The benefit comes from decreasing the consumption of the first π investors in banks that are bailed out and increasing the consumption of both their remaining $1 - \pi$ investors and of the public good. The cost comes from distorting the allocation of resources in banks with small or no losses, where the efficient bail-in would be close or equal to zero. However, since the allocation of resources within the latter banks

was efficient without any regulation, increasing h_{min} above zero initially has only a second-order effect on the utility from private consumption in these banks. The welfare gain from imposing a bail-in at those banks that are bailed out, in contrast, is first order. For this reason, welfare is initially increasing in h_{min} and the optimal policy will involve a positive minimum bail-in. Note that allowing banks to choose a bail-in above the specified minimum is crucial for this result to obtain. If banks with moderate losses were required to *decrease* their bail-in to the mandated level, as with the previous policy considered, the welfare cost of the resulting misallocation would be first order and there would be no guarantee that setting $h_{min} > 0$ would be desirable.

Panel (a) of Figure 6 presents the allocation of losses under the optimal minimum bail-in policy for the same parameter values as in panel (a) of Figure 3. Comparing the two panels shows how the policy creates a positive initial bail-in for all banks, including those that are bailed out. This change makes both the bail-in of the remaining investors and the bailout at these banks smaller. In addition, the policy shifts the equilibrium bailout cutoff ϕ^e to the left, meaning that fewer banks are bailed out. This shift occurs because the minimum bail-in reduces the incentive for a bank at the margin to distort its bail-in choice in order to qualify for a bailout. As a result, more banks choose to set $h(\phi) = 1 - \phi$ and not be bailed out. Finally, note that the banks just above the bailout cutoff ϕ^e in panel (a) of Figure 6 are choosing to set their initial bail-in higher than the minimum value, which demonstrates that providing this option does indeed raise welfare in this example.

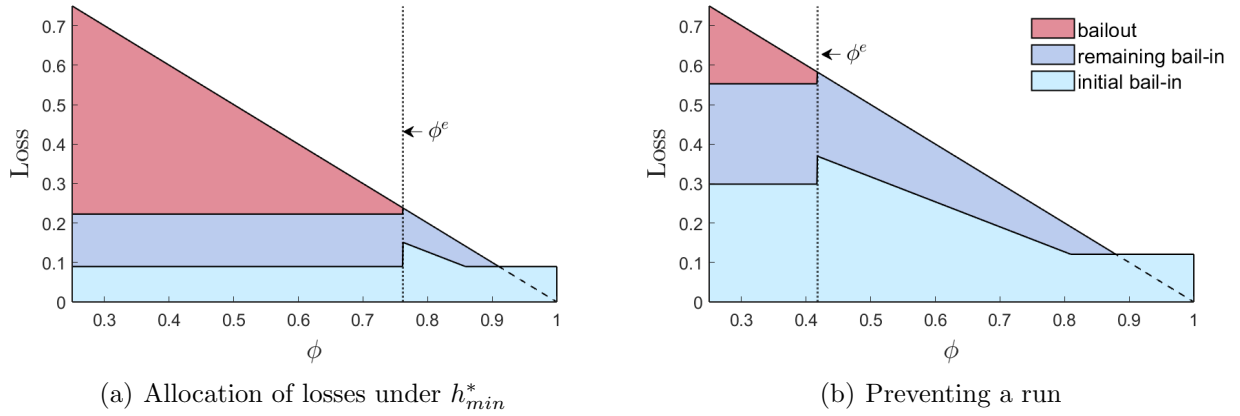


Figure 6: Mandatory minimum bail-in policy

Preventing runs. In some cases, a mandatory minimum bail-in can also enhance financial stability. Consider, for example, the economy depicted in Figure 4 above, where investors run on all banks below the bailout cutoff in the equilibrium with no regulation. Because these

banks are setting $h(\phi) = 0$, Proposition 11 implies that a minimum bail-in policy can raise welfare. Panel (b) of Figure 6 shows that, in addition, the optimal minimum bail-in policy eliminates the bank runs in this economy. The intuition for this result is straightforward: imposing a bail-in on the early withdrawals from these banks makes it less attractive for a patient investor to withdraw early. In addition, by improving the fiscal position of the government, the policy decreases the bail-in applied to the bank's remaining investors and thereby makes waiting to withdraw more attractive. If the minimum bail-in is set high enough, no bank runs will occur in the resulting equilibrium.

Panel (b) of Figure 6 also illustrates a more subtle result: banks with the highest values of ϕ (above about 0.8) are choosing the minimum bail-in, but all other banks are choosing a *larger* bail-in, including those banks that are being bailed out. Why would a bank that anticipates being bailed out set a larger initial bail-in than required? The answer is that these banks would experience a run if they choose h_{min} and are instead choosing to set the smallest bail-in that will prevent a run, $\underline{h}$ as defined in equation (20). When there was no regulation, the benefit to these banks of choosing an initial bail-in of zero outweighed the cost of suffering a run. The minimum bail-in policy changes this calculation. When the first π investors must be bailed in by at least h_{min}^* , these banks find it optimal to bail them in a bit more and avoid the costs associated with a run. In other words, these banks' choices are being disciplined partially by the regulation and partially by the threat of a run. The regulator can use this fact to its advantage when choosing h_{min} , which helps minimize the distortion of the allocation in banks with small or no losses.

This example illustrates that care must be taken when assessing the observed effects of a minimum bail-in policy. Under the optimal policy in panel (b) of Figure 6, the minimum bail-in is binding only at the banks with the smallest losses. One might be tempted to conclude that the policy is ineffective: it is distorting the allocation of resources in the strongest banks, but appears not to be affecting the choices of weaker banks. This conclusion is incorrect, of course; absent the policy, the weakest banks would set their initial bail-ins to zero and bank runs would occur, as shown in Figure 4.

6.3 Optional minimum bail-ins

One clear cost of the policies discussed so far is that they inevitably distort the allocation of resources in banks with small or no losses. If there are many such banks, imposing a mandatory minimum bail-in may be undesirable. However, the discussion above of how regulation interacts with the disciplining effect of bank runs points to another policy option,

which we call an *optional minimum bail-in*. Banks' choice set under this policy is given by

$$h(\phi) \geq h_{min} \quad \text{or} \quad h(\phi) = 0 \quad \text{for all } \phi \in \Phi.$$

In other words, banks can choose an initial bail-in equal to h_{min} or larger, or they can choose zero. The restriction imposed by the regulation is that the initial bail-in cannot lie in between zero and h_{min} .

This type of policy may be desirable when, in the equilibrium without regulation, banks that receive bailouts choose an initial bail-in of $\underline{h} > 0$ to avoid a run, as in panel (b) of Figure 3. The first inequality in equation (20) implies that $\underline{h}$ is smaller than the bail-in the planner would set at these banks and, therefore, the regulator would like to require them to set a larger bail-in. Doing so using either of the policies described above, however, would create significant distortions at banks with small or no losses.

In this situation, the regulator can leverage the disciplining effect of bank runs using a policy that leads banks to effectively self-select. If h_{min} is set slightly above $\underline{h}$, banks that will be bailed out will still prefer h_{min} to setting their bail-in to zero and suffering a run. This fact allows the regulator to add $h(\phi) = 0$ to the choice set without it being used by banks that anticipate being bailed out. Since banks with small or zero losses will choose $h(\phi) = 0$, the minimum bail-in h_{min} no longer distorts the allocation of resources at these banks. This fact, in turn, makes the regulator more willing to set h_{min} above $\underline{h}$ to control the moral hazard problem at banks that will be bailed out.

Panel (a) of Figure 7 depicts the allocation of losses under the optimal optional minimum bail-in policy for the same parameter values as panel (b) of Figure 3. For banks below the equilibrium bailout cutoff ϕ^e , the effect of the policy is as before: it increases the initial bail-in and decreases both the remaining bail-in and bailout. The novel feature comes for banks with the highest values of ϕ , who now choose $h(\phi) = 0$. Because half of the banks in our example have no loss, allowing them to set a zero bail-in brings a substantial welfare gain. It also makes the regulator willing to set h_{min} higher than $\underline{h}$, improving the bail-in choices of banks that are bailed out.

Panel (b) of the figure plots equilibrium welfare as a function of h_{min} for both the mandatory and the optional minimum bail-in policies. In this example, a mandatory minimum bail-in policy (the solid blue curve) always lowers welfare. As h_{min} is increased above zero, the policy is initially binding only for banks with small or no losses, where it introduces a distortion into an allocation that was previously efficient. This distortion is initially second-order, but becomes significant as h_{min} increases. When h_{min} reaches $\underline{h}$, which is approximately 0.25, the policy begins to improve the choices of banks that will be bailed out, which creates

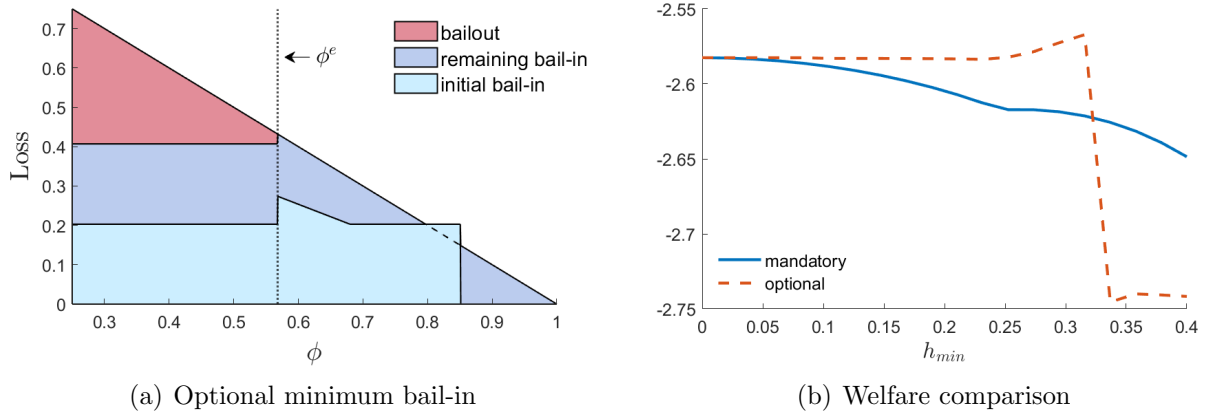


Figure 7: Optional minimum bail-in

an upward kink in the curve in the figure. However, the cost of the distortion it creates at banks with small or no losses is large enough at this point that welfare is still lower than with no policy.

Under the optimal minimum bail-in policy (the dashed red curve), welfare is again initially decreasing as h_{min} is increased above zero. However, since banks with no loss can now choose a bail-in of zero, the welfare loss is much smaller and the curve in the figure appears almost flat. There is again an upward kink as h_{min} crosses $\underline{h}$ and the policy begins to affect the choices of banks that will be bailed out. Because the distortions created by this policy are much smaller than with a mandatory minimum, this benefit now outweighs the costs and adopting the policy raises welfare.

When the optional minimum passes a second threshold, approximately at $h_{min} = 0.32$, banks that will be bailed out would prefer to set $h(\phi) = 0$ even though doing so triggers a run. This shift leads to a sharp decrease in welfare, as bank runs lead to early liquidation of investment and bailouts increase. The optimal policy in this example is to set h_{min} to the highest value for which banks that anticipate a bailout will choose h_{min} over $h(\phi) = 0$.

7 Concluding remarks

Several policy reforms implemented in response to the financial crisis of 2008 aim to give financial intermediaries the ability to more easily impose losses on their investors and/or creditors without declaring bankruptcy or being placed in resolution. Examples include allowing money market mutual funds to restrict withdrawals and impose withdrawal fees,

the introduction of swing pricing in the mutual fund industry more generally,²⁴ and the adoption of rules encouraging the issue of “bail-inable” bank debt. These reforms aim to allow intermediaries to better handle periods of financial stress without the need for bailouts or other forms of public support. While it remains to be seen how effective these reforms will be across a range of situations, the indications to date are not encouraging. At the onset of the COVID-19 crisis in the U.S. in March 2020, the Federal Reserve and U.S. Treasury moved quickly to “assist money market funds in meeting demands for redemptions” by creating a special facility to finance the purchase of assets from these funds.²⁵ The new tools designed for dealing with high redemption demand—restricting withdrawals or imposing withdrawal fees—were not used by any fund. This episode serves as a clear warning that financial-stability policies that rely on intermediaries choosing to quickly bail in their investors in periods of financial stress may be ineffective.

Our model helps illustrate the incentive problems that can undermine the effectiveness of these types of policies and points to a better approach. Banks and other intermediaries anticipate that, when the situation is bad enough, the public sector will respond with bailouts. It does not seem feasible for governments to commit to a strict no-bailout policy, and such a policy may not even be desirable; in our environment, it is optimal for the public sector to absorb some of the tail risk in economic outcomes. Either way, the anticipation of being bailed out undermines the incentive for an intermediary to quickly bail in its investors and creditors, even if it has complete flexibility in choosing the bail in. As a result, equilibrium bail-ins are inefficiently small and bailouts are inefficiently large. Moreover, the unwillingness of intermediaries to bail-in can be a source of fragility: it can lead to runs by investors that deepen the crisis and lead to even larger bailouts.

A regulator can improve financial stability and welfare in our model by mandating system-wide bail-ins at the onset of a crisis. While a “one size fits all” bail-in may create substantial distortions in some intermediaries, the regulator can design the policy to maximize the benefits while keeping these distortions to a minimum. For example, the policy should specify a minimum bail-in, but allow intermediaries to choose a larger bail-in if doing so is in the best interests of their investors. In some cases, the policy may allow intermediaries to opt out of the bail-in if they can do so without causing their investors to lose confidence and run.

These recommendations can be implemented across a range of intermediation arrange-

²⁴See [Chen et al. \(2010\)](#) for evidence of strategic complementarities in the withdrawal decisions of investors in open-end mutual funds where the price is set daily according to the net asset value of the fund. [Jin et al. \(2019\)](#) study the ability of swing pricing to remove these complementarities and prevent runs.

²⁵Detailed information on the Money Market Mutual Fund Liquidity Facility is available at www.federalreserve.gov/monetarypolicy/mmlf.htm.

ments. In general terms, our results provide support for calls to restrict dividend payments and share repurchases by banks in the early stages of a crisis. Banks could also be required to issue debt that is either automatically written down or converted to equity based on a systemic trigger. Similarly, a minimum withdrawal fee could be imposed at all money market mutual funds based on a systemic rather than a fund-specific trigger. One interesting area for future research is adapting our model to the specific institutional features of different intermediation arrangements and deriving the prescriptions for bail-in policy at a finer level of detail.

Appendix: Proofs

Proposition 1. *The efficient plan satisfies $h^*(\phi) = \hat{h}^*(\phi)$ for each $\phi \in \Phi$. That is, all investors within a bank face the same bail-in.*

Proof. To begin, note that the resource constraint (3) will hold with equality for all ϕ at the solution to the planner's problem. The non-negativity restrictions then imply that the planner will set $h(1) = \hat{h}(1) = b(1) = 0$. In other words, investors in banks with zero loss are neither bailed in nor bailed out. For banks with $\phi < 1$, let $\lambda(\phi)f(\phi)$ denote the multiplier on the resource constraint (3), where f is the density function of the distribution F on $[\underline{\phi}, 1)$. We can then write the first-order conditions for the optimal choice of $h(\phi)$ as²⁶

$$u'((1 - h(\phi))c_1^*) \geq \lambda(\phi) \quad \text{and} \quad h(\phi)[u'((1 - h(\phi))c_1^*) - \lambda(\phi)] = 0 \quad (24)$$

and for the optimal choice of $\hat{h}(\phi)$ as

$$u'\left(\left(1 - \hat{h}(\phi)\right)c_2^*\right) \geq \frac{\lambda(\phi)}{R} \quad \text{and} \quad \hat{h}(\phi)\left[u'\left(\left(1 - \hat{h}(\phi)\right)c_2^*\right) - \frac{\lambda(\phi)}{R}\right] = 0. \quad (25)$$

We will show that the solutions to these two sets of equations are necessarily the same, considering the cases of boundary and interior solutions separately.

First, suppose the solution has $h(\phi) = 0$ for any given ϕ . Then equation (24) implies

$$u'(c_1^*) \geq \lambda(\phi).$$

The reference allocation (c_1^*, c_2^*) is characterized by the standard optimality condition in the Diamond-Dybvig framework,

$$u'(c_1^*) = Ru'(c_2^*).$$

Combining these two equations yields

$$u'(c_2^*) \geq \frac{\lambda(\phi)}{R}$$

and, therefore, the unique $\hat{h}(\phi)$ satisfying the conditions in equation (25) is $\hat{h}(\phi) = 0$.

Next, suppose the solution has $h(\phi) > 0$ for any given ϕ . Then equation (24) implies

$$u'((1 - h(\phi))c_1^*) = \lambda(\phi).$$

²⁶Note that the Inada conditions on the function u imply that the upper bounds on $h(\phi)$ and $\hat{h}(\phi)$ in equation (3) will never bind at the solution to the problem.

Given that the utility function u is of the constant-relative-risk-aversion form, the ratio of marginal utilities depends only on the ratio of consumption levels, that is, we have

$$\frac{u'(\alpha c_1^*)}{u'(\alpha c_2^*)} = \frac{u'(c_1^*)}{u'(c_2^*)} = R \quad (26)$$

for any $\alpha > 0$. These last two equations imply that we have

$$u'((1 - h(\phi))c_2^*) = \frac{\lambda(\phi)}{R}$$

and, therefore, setting $\hat{h}(\phi) = h(\phi)$ is the unique solution to equation (25). Combining these two cases, we have shown that $\hat{h}(\phi) = h(\phi)$ holds for all ϕ , which establishes the result. $\square$

Proposition 2. *There exists $\phi^* \in \Phi$ such that the efficient plan (h^*, b^*) has the form:*

$$h^*(\phi) = \begin{cases} 1 - \phi \\ 1 - \phi^* \end{cases} \quad \text{and} \quad b^*(\phi) = \begin{cases} 0 \\ \phi^* - \phi \end{cases} \quad \text{as} \quad \phi \begin{cases} \geq \\ < \end{cases} \phi^*.$$

Proof. Letting $z \geq 0$ denote the measure of banks with $\phi = 1$, we can rewrite the objective function in equation (7) as

$$\int_{\underline{\phi}}^1 U(1 - h(\phi)) f(\phi) d\phi + zU(1) + v \left(\tau - \int_{\underline{\phi}}^1 (1 - \phi - h(\phi)) f(\phi) d\phi \right).$$

Let $\underline{\mu}(\phi) f(\phi)$ denote the multiplier on the non-negativity constraint for $h(\phi)$ and let $\bar{\mu}(\phi) f(\phi)$ denote the multiplier on upper bound of $h(\phi)$. Then the first-order conditions

$$U'(1 - h(\phi)) = v'(g) + \underline{\mu}(\phi) - \bar{\mu}(\phi), \quad (27)$$

the complementary slackness conditions

$$\underline{\mu}(\phi) h(\phi) = 0 \quad \text{and} \quad \bar{\mu}(\phi) [h(\phi) - (1 - \phi)] = 0, \quad (28)$$

and non-negativity of the multipliers

$$\underline{\mu}(\phi) \geq 0 \quad \text{and} \quad \bar{\mu}(\phi) \geq 0 \quad \text{for all } \phi \in [\underline{\phi}, 1] \quad (29)$$

are necessary and jointly sufficient for a solution to the problem. We break the analysis into two cases.

Case (i) : Suppose the following inequality holds:

$$v'(\tau) \geq U'(\underline{\phi}). \quad (30)$$

Then concavity of U implies that $v'(\tau) \geq U'(\phi)$ holds for all ϕ . In this case, it is straightforward to check that the following values are the unique solution to equations (27) – (29): $h(\phi) = 1 - \phi$, which implies $b(\phi) = 0$, $\underline{\mu}(\phi) = 0$ and

$$\bar{\mu}(\phi) = v'(\tau) - U'(\phi) \geq 0 \quad \text{for all } \phi.$$

Note that this solution corresponds to the form in the statement of the proposition with $\phi^* = \underline{\phi}$. Intuitively, this case corresponds to a situation in which the government's fiscal capacity τ is small enough that the planner would choose to make no bailout payments to any bank.

Case (ii) : Next, suppose the inequality in equation (30) is reversed. Define the functions

$$g_1(\phi) \equiv U'(\phi)$$

and

$$g_2(\phi) \equiv v' \left(\tau - \int_{\underline{\phi}}^{\phi} (\phi - x) dF(x) \right).$$

It is straightforward to show that g_1 is continuous and strictly decreasing on the interval $[\underline{\phi}, 1]$, and that g_2 is continuous and strictly increasing on the same domain. Moreover, the fact that the inequality in (30) does not hold implies

$$g_1(\underline{\phi}) > g_2(\underline{\phi}),$$

while the upper bound on τ in equation (4) implies

$$g_1(1) < g_2(1).$$

It follows that there exists a unique $\phi^* \in (\underline{\phi}, 1)$ satisfying $g_1(\phi^*) = g_2(\phi^*)$, or

$$U'(\phi^*) = v' \left(\tau - \int_{\underline{\phi}}^{\phi^*} (\phi^* - \phi) dF(\phi) \right). \quad (31)$$

In this case, the unique solution to equations (27) – (29) has the following properties. For

$\phi \geq \phi^*$, it sets $h(\phi) = 1 - \phi$, which implies $b(\phi) = 0$, and sets $\underline{\lambda}(\phi) = 0$ and

$$\bar{\lambda}(\phi) = v' \left(\tau - \int_{\underline{\phi}}^{\phi^*} (\phi^* - \phi) dF(\phi) \right) - U'(\phi) \geq 0,$$

where the non-negativity of this expression follows from the concavity of U . For $\phi < \phi^*$, it sets $h(\phi) = 1 - \phi^*$, which implies $b(\phi) = \phi^* - \phi$, and sets $\underline{\lambda}(\phi) = \bar{\lambda}(\phi) = 0$. Intuitively, equation (31) characterizes the cutoff value ϕ^* by identifying the bank whose investors' marginal value of private consumption with no bailout exactly equals the marginal value of public consumption at the solution. All banks with larger losses than this cutoff receive a bailout that keeps their bail-in equal to the maximum value $1 - \phi^*$. All banks with smaller losses receive no bailout. $\square$

Proposition 3. *The efficient bailout cutoff ϕ^* is strictly increasing in τ whenever $\phi^* \in (\underline{\phi}, 1)$.*

Proof. If $\phi^* \in (\underline{\phi}, 1)$, the first-order condition (31) implicitly defines ϕ^* as a function of the parameter τ . Differentiating this equation with respect to τ yields

$$U''(\phi^*) \frac{d\phi^*(\tau)}{d\tau} = v''(\tau - B(\tau)) \left(1 - \frac{dB(\tau)}{d\tau} \right), \quad (32)$$

where $B(\tau)$ is the aggregate bailout as a function of τ ,

$$B(\tau) \equiv \int_{\underline{\phi}}^{\phi^*(\tau)} (\phi^*(\tau) - \phi) dF(\phi).$$

Differentiating this last equation with respect to τ yields

$$\frac{dB(\tau)}{d\tau} = \frac{d\phi^*(\tau)}{d\tau} F(\phi^*(\tau)).$$

Substituting this expression into equation (32) and solving yields

$$\frac{d\phi^*(\tau)}{d\tau} = \frac{v''(\tau - B(\tau))}{U''(\phi^*) + F(\phi^*(\tau)) v''(\tau - B(\tau))} > 0,$$

which is strictly positive because the functions U and v are strictly concave. $\square$

Proposition 4. *For any (h, y) , there exists a unique $\hat{h}_{max} \in [0, 1]$ such that the fiscal*

authority's bailout policy b will imply

$$\hat{h}(\phi) \leq \hat{h}_{max} \text{ for all } \phi, \text{ with equality if } b(\phi) > 0.$$

Proof. The marginal value of resources in a bank that has ψ units of the good per remaining investor after π withdrawals can be determined by differentiating equation (9) with respect to ψ ,

$$V_1(\psi, y) = \left\{ \begin{array}{c} u'(R\psi) R \\ \pi u'(\psi c_1^*) c_1^* + (1 - \pi) u'(\psi c_2^*) c_2^* \end{array} \right\} \quad \text{as} \quad \left\{ \begin{array}{c} y = 2 \\ y = 1 \end{array} \right\}.$$

Using equation (14), we can rewrite these expressions in terms of the bail-in $\hat{h}$ faced by each of the bank's remaining investors,

$$V_1(\psi, y) = \left\{ \begin{array}{c} u'((1 - \hat{h})c_2^*) R \\ \pi u'((1 - \hat{h})c_1^*) c_1^* + (1 - \pi) u'((1 - \hat{h})c_2^*) c_2^* \end{array} \right\} \quad \text{as} \quad \left\{ \begin{array}{c} y = 2 \\ y = 1 \end{array} \right\}. \quad (33)$$

Since u is of the constant-relative-risk-aversion form, we can use equations (6) and (26) to write

$$u'((1 - \hat{h})c_1^*) = Ru'((1 - \hat{h})c_2^*) = U'(1 - \hat{h}).$$

Using this result, together with the feasibility constraint in (1), the two expressions for $dV/d\psi$ in (33) can be shown to be equal, allowing us to write

$$V_1(\psi, y) = U'(1 - \hat{h}) \quad \text{for } y \in \{1, 2\}. \quad (34)$$

Substituting this expression into equations (12) – (13), we can write the first-order conditions that characterize the fiscal authority's choice of bailout policy as saying that, for each bank ϕ , we have either

$$U'(1 - \hat{h}(\phi)) \leq v'(g) \quad \text{and} \quad b(\phi) [v'(g) - U'(1 - \hat{h}(\phi))] = 0 \quad (35)$$

or

$$U'(1 - \hat{h}(\phi)) > v'(g) \quad \text{and} \quad b(\phi) = \bar{b}(\phi, h(\phi), y(\phi)), \quad (36)$$

where g is determined by the choice of bailouts b as shown in equation (11). We divide the analysis into two cases.

Case (i): First, suppose

$$U'(1) \geq v' \left(\tau - \int_{\underline{\phi}}^1 \bar{b}(\phi, h(\phi), y(\phi)) dF(\phi) \right) \quad (37)$$

holds. Then the conditions in equation (36) will hold at the solution for all ϕ . In other words, the fiscal authority will set the bailout $b(\phi)$ to its upper bound, which corresponds $\hat{h}(\phi) = 0$, for all banks. In this case, the unique value of $\hat{h}_{max}$ is zero.

Case (ii): If the inequality in equation (37) is reversed, then the upper bound $\bar{b}$ will not bind for any ϕ and the bailout $b(\phi)$ will satisfy the conditions in equation (35) for all banks. It follows directly from these conditions that all banks with $b(\phi) > 0$ will have the same bail-in $h(\phi)$. Define $\hat{h}_{max}$ to be this common value, that is, the bail-in satisfying

$$U'(1 - \hat{h}_{max}) = v'(g). \quad (38)$$

Equation (35) and the concavity of U then imply $\hat{h}(\phi) \leq \hat{h}_{max}$ for all ϕ with $b(\phi) = 0$. To show that there is a unique value $\hat{h}_{max}$ with this property, we use equations (8) and (14) to write the bailout $b(\phi)$ received by each bank in terms of its bail-in for remaining investors $\hat{h}(\phi)$,

$$b(\phi) = \begin{cases} 1 - \phi - h(\phi)\pi c_1^* - \hat{h}(\phi)(1 - \pi)\frac{c_2^*}{R} \\ 1 - \phi + \pi(1 - \pi)\left(c_1^* - \frac{c_2^*}{R}\right) - h(\phi)\pi c_1^* - \hat{h}(\phi)(1 - \pi) \end{cases} \quad \text{as } y(\phi) = \begin{cases} 2 \\ 1 \end{cases}.$$

The results above then imply that we can write the bailout for each bank in terms of the maximum bail-in $\hat{h}_{max}$ as

$$b(\phi; \hat{h}_{max}) = \max \left\{ 1 - \phi - h(\phi)\pi c_1^* - \hat{h}_{max}(1 - \pi)\frac{c_2^*}{R}, 0 \right\}$$

for banks with $y(\phi) = 2$ and as

$$b(\phi; \hat{h}_{max}) = \max \left\{ 1 - \phi + \pi(1 - \pi)\left(c_1^* - \frac{c_2^*}{R}\right) - h(\phi)\pi c_1^* - \hat{h}_{max}(1 - \pi), 0 \right\} \quad (39)$$

for banks with $y(\phi) = 1$. Note that $b(\phi, \hat{h}_{max})$ is a continuous, weakly decreasing function of $\hat{h}_{max}$ for all ϕ . The level of the public good is then related to $\hat{h}_{max}$ by

$$g(\hat{h}_{max}) = \tau - \int_{\underline{\phi}}^1 b(\phi, \hat{h}_{max}) dF(\phi).$$

The properties of $b(\phi, \hat{h}_{max})$ imply that g is a continuous, increasing function of $\hat{h}_{max}$. Substituting this function into equation (38) and solving for $\hat{h}_{max}$ on the left-hand side, we can say that any $\hat{h}_{max}$ associated with the fiscal authority's choice of bailouts will satisfy

$$\hat{h}_{max} = 1 - U'^{-1} \left[v' \left(g \left(\hat{h}_{max} \right) \right) \right] \equiv z \left(\hat{h}_{max} \right). \quad (40)$$

The function z defined in this equation is continuous and decreasing for all $\hat{h}_{max} \in [0, 1]$. It follows that there exists a unique solution for $\hat{h}_{max} \in [0, 1]$. $\square$

Proposition 5. *For each profile of initial haircuts h , there is unique minimal-run equilibrium of the overall withdrawal game.*

We begin by establishing two lemmas and then use these preliminary results to construct a sequence of withdrawal profiles y that converges to the unique minimal-run equilibrium. The profile of initial bail-ins h is held fixed throughout. For any profile y of withdrawal behavior across banks, define $\mathbb{P}(y)$ to be the set of banks that experience a run, that is,

$$\mathbb{P}(y) = \{\phi \in \Phi : y(\phi) = 1\}.$$

Proposition 4 establishes that, for any profile y , the fiscal authority's bailout policy generates a maximum bail-in for remaining investors $\hat{h}_{max}(y)$. Our first step is to establish that when a run occurs at a larger set of banks, this maximum bail-in increases.

Lemma 1. The maximum bail-in $\hat{h}_{max}$ established in Proposition 4 satisfies

$$\mathbb{P}(y_2) \supseteq \mathbb{P}(y_1) \quad \Rightarrow \quad \hat{h}_{max}(y_2) \geq \hat{h}_{max}(y_1). \quad (41)$$

Proof. The resource constraint in equation (3) for a bank whose investors do not run will hold with equality when the bank distributes its resources optimally. Solving this equality for the bailout received by the bank yields

$$b(\phi) = 1 - \phi - h(\phi)\pi c_1^* - \hat{h}(\phi)(1 - \pi)\frac{c_2^*}{R}.$$

Using Proposition 4, we can rewrite this expression as a function of $\hat{h}_{max}$,

$$b(\phi) = \max \left\{ 1 - \phi - h(\phi)\pi c_1^* - \hat{h}_{max}(1 - \pi)\frac{c_2^*}{R}, 0 \right\}. \quad (42)$$

Intuitively, the bail-in of remaining investors in all banks that receive a bailout is equal to $\hat{h}_{max}$, while any bank with a smaller bail-in receives zero bailout. Applying the same logic to

equation (23), which is the resource constraint for a bank whose investors do run, we have

$$b(\phi) = \max \left\{ 1 - \phi + \pi(1 - \pi) \left(c_1^* - \frac{c_2^*}{R} \right) - h(\phi)\pi c_1^* - \hat{h}_{max}(1 - \pi), 0 \right\}. \quad (43)$$

Holding the bail-ins $h(\phi)$ and $\hat{h}_{max}$ fixed, it is straightforward to show that the expression in equation (43) is larger than that in equation (42), and strictly larger if it is strictly positive.²⁷ In other words, a bank will receive a larger bailout if it experiences a run, since the extra withdrawals at $t = 1$ would otherwise decrease the consumption of its remaining investors. In addition, it is easy to see that, holding h and the withdrawal behavior $y(\phi)$ fixed, the bailout $b(\phi)$ in both expressions is decreasing in $\hat{h}_{max}$, and strictly decreasing if the bailout is positive.

Let $b(\phi; \hat{h}_{max}, y(\phi))$ denote the bailout amount from equation (42) when $y(\phi) = 2$ and from equation (43) when $y(\phi) = 1$. Let $g(\hat{h}_{max}, y)$ denote the level of public good when the bail-in of all remaining investors is fixed at $\hat{h}_{max}$ and the withdrawal behavior across banks is given by y , that is,

$$g(\hat{h}_{max}, y) = \tau - \int_{\phi}^1 b(\phi; \hat{h}_{max}, y(\phi)) dF(\phi).$$

This discussion above establishes that (i) g is increasing in $\hat{h}_{max}$ and (ii) $\mathbb{P}(y_2) \supseteq \mathbb{P}(y_1)$ implies $g(\hat{h}_{max}, y_2) \leq g(\hat{h}_{max}, y_1)$. In other words, when runs occur at more banks, holding $\hat{h}_{max}$ fixed, more public resources are spent on bailouts and fewer public goods are provided.

To establish the lemma, we combine these results with the characterization $\hat{h}_{max}$ given in equation (38) in the proof of Proposition 4,

$$u' \left((1 - \hat{h}_{max}) c_1^* \right) = v' \left(g(\hat{h}_{max}, y) \right).$$

When $\mathbb{P}(y_2) \supseteq \mathbb{P}(y_1)$, the level of the public good g would tend to decrease if we hold $\hat{h}_{max}$ fixed, raising the marginal utility of public consumption. For this condition to remain satisfied, the maximum bail-in $\hat{h}_{max}$ must increase. Therefore, we have $\hat{h}_{max}(y_2) \geq \hat{h}_{max}(y_1)$, as desired. $\square$

Our second step is to establish that if a bank experiences a run for a given maximum bail-in, it will also experience a run if the maximum bail-in increases. For any value of $\hat{h}_{max}$, define $\Lambda(\hat{h}_{max})$ to be the set of banks that would experience a run if we hold the maximum bail-in

²⁷This result follows from the fact that the efficient allocation satisfies $1 < c_1^* < c_2^* < R$ and the Inada condition on the utility function u implies $\hat{h}_{max} < 1$.

fixed at this level, that is

$$\Lambda(\hat{h}_{max}) = \left\{ \phi \in \Phi : (1 - \hat{h}(\phi)) c_2^* < (1 - h(\phi)) c_1^* \right\}, \quad (44)$$

where

$$\hat{h}(\phi) = \min \left\{ \hat{h}_{NB}(\phi), \hat{h}_{max} \right\}$$

and $\hat{h}_{NB}(\phi)$ is as defined in equation (16).

Lemma 2. The set Λ is expanding in $\hat{h}_{max}$, that is,

$$\phi \in \Lambda(\hat{h}_{max}) \Rightarrow \phi \in \Lambda(\hat{h}'_{max}) \quad \text{for all } \hat{h}'_{max} \geq \hat{h}_{max}. \quad (45)$$

Proof. This result follows directly from the definitions of the set Λ and the bail-in $\hat{h}(\phi)$ given above. When $\hat{h}_{max}$ increases, the bail-in of bank ϕ 's remaining investors $\hat{h}(\phi)$ either increases or remains unchanged. The left-hand side of the inequality in equation (44) is thus weakly decreasing in $\hat{h}_{max}$, while the right-hand side is unchanged. $\square$

Proof of Proposition 5. We now use these two lemmas to construct a sequence of withdrawal profiles that converges to the unique minimal-run equilibrium. This process begins with the profile in which there are no runs on any bank. Define y_0 as

$$y_0(\phi) = 2 \quad \text{for all } \phi \in \Phi. \quad (46)$$

Let $\hat{h}_{max}(y_0)$ denote the maximum bail-in of remaining investors associated with this profile. Lemma 1 implies that $\hat{h}_{max}(y_0)$ is the smallest possible value that the maximum bail-in could take in any equilibrium.

Next, we identify the set of banks whose patient investors would choose to run when the maximum bail-in is $\hat{h}_{max}(y_0)$, which is given by $\Lambda(\hat{h}_{max}(y_0))$. If this set is empty, then y_0 is an equilibrium of the overall withdrawal game and, since there are no runs on any bank, it is clearly the minimal-run equilibrium.

If $\Lambda(\hat{h}_{max}(y_0))$ is not empty, then y_0 is not an equilibrium of the overall withdrawal game. Moreover, since the value of $\hat{h}_{max}$ in any equilibrium is at least $\hat{h}_{max}(y_0)$, Lemma 2 implies that all banks in $\Lambda(\hat{h}_{max}(y_0))$ will experience a run in *any* equilibrium of the overall withdrawal game. Define a new profile of withdrawal behavior in which these (and only these) banks experience a run,

$$y_1(\phi) = \begin{cases} 1 \\ 2 \end{cases} \quad \text{if} \quad \begin{cases} \phi \in \Lambda(\hat{h}_{max}(y_0)) \\ \text{otherwise} \end{cases}.$$

Note that $\mathbb{P}(y_1) \supset \mathbb{P}(y_0) = \emptyset$ and, therefore, by equation (41) we have $\hat{h}_{max}(y_1) \geq \hat{h}_{max}(y_0)$. In addition, since banks in the set $\mathbb{P}(y_1)$ will necessarily experience a run in any equilibrium, we know that $\hat{h}_{max}(y_1)$ is a lower bound on the maximum bail-in in any equilibrium.

We then repeat the process. For $j = 1, 2, \dots$, define

$$y_{j+1}(\phi) = \begin{cases} 1 \\ 2 \end{cases} \quad \text{if} \quad \begin{cases} \phi \in \Lambda(\hat{h}_{max}(y_j)) \\ \text{otherwise} \end{cases}. \quad (47)$$

Note that a fixed point of this equation is necessarily an equilibrium of the overall withdrawal game: a profile y^e such that when the maximum bail-in is equal to $\hat{h}_{max}(y^e)$, the withdrawal behavior in each bank determined by equations (14) and (15) is equal to y^e .

For each ϕ , the sequence $\{y_j(\phi)\}_{j=0}^{\infty}$ either remains equal to 2 for all j or switches to 1 at some j and remains equal to 1 for all $j' > j$. Therefore, the sequence of profiles $\{y_j\}$ converges pointwise to some profile y^e . This profile y^e is a fixed point of equation (47) and, therefore, is an equilibrium of the overall withdrawal game. By construction, the value $\hat{h}_{max}(y_j)$ at each point in the sequence is a lower bound on the value of the maximum bail-in in any equilibrium. It follows that, if the overall withdrawal game has multiple equilibria, the process defined here will converge to the equilibrium with the smallest value of $\hat{h}_{max}$. Lemma 2 above implies that the set of banks experiencing a run in any other equilibrium will contain the set $\mathbb{P}(y^e)$, which establishes that y^e is the unique minimal-run equilibrium. $\square$

Proposition 6. *There exists a pure-strategy equilibrium of the bail-in game.*

Proof. We divide the proof into three steps.

Step (i) : Suppose banks anticipate an arbitrary value for the maximum bail-in of the remaining investors in banks that are bailed out, denoted $\hat{h}_{max}^a$. Find the initial bail-in $h(\phi)$ that each bank will choose and the withdrawal behavior of its investors $y(\phi)$ that will result.

Recall that a bank receiving a bailout can prevent a run by setting $h(\phi) = \underline{h}(\hat{h}_{max}^a)$, as defined in equation (18), but may choose to set $h(\phi) = 0$ and suffer a run if $\underline{h}$ is large enough. Define

$$U_B(\hat{h}_{max}^a) = \max \left\{ \begin{array}{l} \pi u \left(\left(1 - \underline{h}(\hat{h}_{max}^a) \right) c_1^* \right) + (1 - \pi) u \left(\left(1 - \hat{h}_{max}^a \right) c_2^* \right), \\ \pi u(c_1^*) + (1 - \pi) \left[\pi u \left(\left(1 - \hat{h}_{max}^a \right) c_1^* \right) + (1 - \pi) u \left(\left(1 - \hat{h}_{max}^a \right) c_2^* \right) \right] \end{array} \right\}, \quad (48)$$

which measures the expected utility of investors in a bank that is bailed out when the bank chooses $h(\phi)$ optimally. Using equation (18) and the fact that u has constant relative risk aversion, it is straightforward to show that U_B is continuous and strictly decreasing in $\hat{h}_{max}^a$. In addition, comparing the two terms on the right-hand side of equation (48) shows that there exists a scalar $\alpha \in (0, 1)$ such that the maximum is equal to the first line for all $\hat{h}_{max}^a < \alpha$ and to the second line for all $\hat{h}_{max}^a > \alpha$. It follows that banks anticipating a bailout will set $h(\phi)$ in a way that leads to a run if and only if $\hat{h}_{max}^a > \alpha$.

A bank that anticipates not receiving a bailout will set $h(\phi) = 1 - \phi$ and its investors will not run. Let $\hat{\phi}$ denote the unique solution to

$$U_B(\hat{h}_{max}^a) = U(\hat{\phi}),$$

where the function U measures the expected utility of investors in a bank where there is no bailout and all investors share the loss evenly, as defined in equation (6). We can write this solution as

$$\hat{\phi}(\hat{h}_{max}^a) = U^{-1}(U_B(\hat{h}_{max}^a)). \quad (49)$$

Then banks with $\phi < \hat{\phi}(\hat{h}_{max}^a)$ will choose $h(\phi)$ to maximize U_B and receive a bailout, while banks with $\phi \geq \hat{\phi}(\hat{h}_{max}^a)$ will set $h(\phi) = 1 - \phi$ and not be bailed out. It follows directly from the properties of U and U_B that $\hat{\phi}(0) = 1$ and the limit of $\hat{\phi}(\hat{h}_{max}^a)$ as $\hat{h}_{max}^a$ approaches 1 is zero. Moreover, since U is a continuous, strictly increasing function, we have that $\hat{\phi}(\hat{h}_{max}^a)$ is also continuous and strictly decreasing.

Summarizing the results so far, we have that each bank's optimal choice of initial bail-in is given by

$$h(\phi; \hat{h}_{max}^a) = \begin{cases} \underline{h}(\hat{h}_{max}^a) \\ 0 \\ 1 - \phi \end{cases} \quad \text{if} \quad \begin{cases} \phi < \hat{\phi}(\hat{h}_{max}^a) \text{ and } \hat{h}_{max}^a \leq \alpha \\ \phi < \hat{\phi}(\hat{h}_{max}^a) \text{ and } \hat{h}_{max}^a > \alpha \\ \phi \geq \hat{\phi}(\hat{h}_{max}^a) \end{cases}. \quad (50)$$

The resulting withdrawal behavior of each bank's investors is given by

$$y(\phi; \hat{h}_{max}^a) = \begin{cases} 1 \\ 2 \end{cases} \quad \text{if} \quad \begin{cases} \phi < \hat{\phi}(\hat{h}_{max}^a) \text{ and } \hat{h}_{max}^a > \alpha \\ \text{otherwise} \end{cases}. \quad (51)$$

Step (ii) : Combine Step (i) with Proposition 4 to create a mapping from the anticipated maximum bail-in $\hat{h}_{max}^a$ to the actual maximum bail-in $\hat{h}_{max}$ generated by the bailout policy.

Show that this mapping has a fixed point.

Given any $\hat{h}_{max}^a$ and the resulting profiles h and y derived in step (i), Proposition 4 establishes that there is a unique maximum bail-in $\hat{h}_{max}$ that will be applied to the remaining investors in all banks that are bailed out. Using equations (50) – (51), we can write the bailout received by each bank as

$$b\left(\phi; \hat{h}_{max}^a, \hat{h}_{max}\right) = \begin{cases} 1 - \phi - \underline{h}\left(\hat{h}_{max}^a\right) \pi c_1^* - \hat{h}_{max}(1 - \pi) \frac{c_2^*}{R} \\ 1 - \phi + \pi(1 - \pi) \left(c_1^* - \frac{c_2^*}{R}\right) - \hat{h}_{max}(1 - \pi) \\ 0 \end{cases} \quad \text{if } \begin{cases} \phi < \hat{\phi}\left(\hat{h}_{max}^a\right) \text{ and } \hat{h}_{max}^a \leq \alpha \\ \phi < \hat{\phi}\left(\hat{h}_{max}^a\right) \text{ and } \hat{h}_{max}^a > \alpha \\ \phi \geq \hat{\phi}\left(\hat{h}_{max}^a\right) \end{cases}. \quad (52)$$

The level of public good will equal the fiscal capacity of the government τ minus the total bailout payments to all banks. Using equation (52), together with our assumption that the distribution F has a density function for $\phi < 1$, we can write g as a function of the anticipated and actual maximum bail-ins,

$$g\left(\hat{h}_{max}^a, \hat{h}_{max}\right) = \tau - \int_{\underline{\phi}}^{\hat{\phi}(\hat{h}_{max}^a)} b\left(\phi; \hat{h}_{max}^a, \hat{h}_{max}\right) f(\phi) d\phi. \quad (53)$$

For any given $\hat{h}_{max}$, the fact that $\hat{\phi}$ and $\underline{h}$ are continuous functions of $\hat{h}_{max}^a$ shows that g is continuous at all points except $\hat{h}_{max}^a = \alpha$. In addition, the fact that $\hat{\phi}$ is strictly decreasing and b is weakly decreasing in $\hat{h}_{max}^a$ imply that g is strictly increasing at all points except $\hat{h}_{max}^a = \alpha$, where it jumps down.

For any given $\hat{h}_{max}^a$, the corresponding $\hat{h}_{max}$ from Proposition 4 is given by the solution to

$$u' \left[\left(1 - \hat{h}_{max}\right) c_1^* \right] \leq v' \left[g\left(\hat{h}_{max}^a, \hat{h}_{max}\right) \right], \text{ with equality if } \hat{h}_{max}^a > 0.$$

Equivalently, we can write

$$\hat{h}_{max} \geq 1 - \frac{1}{c_1^*} u'^{-1} \left\{ v' \left[g\left(\hat{h}_{max}^a, \hat{h}_{max}\right) \right] \right\} \equiv z\left(\phi; \hat{h}_{max}^a, \hat{h}_{max}\right) \quad (54)$$

with equality if $\hat{h}_{max}^a > 0$. As discussed in the proof of Proposition 4, given any $\hat{h}_{max}^a$, the function z is continuous and strictly decreasing in $\hat{h}_{max}$. Combined with the boundary

conditions, these properties imply there is a unique solution for $\hat{h}_{max}$, which we denote

$$\hat{h}_{max} = \zeta \left(\hat{h}_{max}^a \right). \quad (55)$$

To establish the properties of the function ζ , recall that equation (53) shows that, for all values of $\hat{h}_{max}$, the function g is continuous and strictly increasing in $\hat{h}_{max}^a$ at all points except $\hat{h}_{max}^a = \alpha$, where it jumps down. Since u' and v' are both continuous, strictly decreasing functions, equation (54) then implies that the function z is continuous and strictly decreasing in $\hat{h}_{max}^a$ at all points except $\hat{h}_{max}^a = \alpha$, where it jumps up. This monotonicity of z in $\hat{h}_{max}^a$ for all $\hat{h}_{max}$ implies monotonicity of the solution in $\hat{h}_{max}^a$. Specifically, the function ζ is continuous and strictly decreasing in $\hat{h}_{max}^a$ at all points except $\hat{h}_{max}^a = \alpha$, where it jumps up. These properties of ζ meet the conditions of Corollary 1 in Milgrom and Roberts (1994), which establishes that there exists a fixed point of the mapping ζ , that is, a value $\hat{h}_{max}^e$ satisfying

$$\hat{h}_{max}^e = \zeta \left(\hat{h}_{max}^e \right). \quad (56)$$

Step (iii) : Show that a fixed point $\hat{h}_{max}^e$ corresponds to an equilibrium of the bail-in game and establish the properties of equilibrium withdrawal behavior.

Given a value $\hat{h}_{max}^e$ satisfying equation (56), calculate the profile of initial bail-ins banks would choose if $\hat{h}_{max}^a = \hat{h}_{max}^e$ using equation (50) and the resulting withdrawal behavior of investors using equation (51). Consider the decision problem of an individual bank ϕ , which takes the initial bail-ins chosen by other banks $h_{-\phi}^e$ as given. Because each bank with $\phi < 1$ has zero measure, its own choice of $h(\phi)$ has no effect on the level of public good g in equation (53) and, therefore, on the maximum bail-in $\hat{h}_{max}^e$ that will be generated by the bailout policy. Since $h^e(\phi)$ is bank ϕ 's best choice when the maximum bail-in is $\hat{h}_{max}^e$, by construction, it is a best response to the profile of bail-ins chosen by other banks, $h_{-\phi}^e$. The profile of initial bail-ins h^e is, therefore, an equilibrium of the bail-in game as defined in Definition 1. $\square$

Proposition 7. *In any equilibrium of the bail-in game, there exists $\phi^e \in \Phi$ such that*

$$h^e(\phi) = \left\{ \begin{array}{c} 1 - \phi \\ \tilde{h} \in \{0, \underline{h}\} \end{array} \right\} \quad \text{and} \quad \left\{ \begin{array}{c} b^e(\phi) = 0 \\ b^e(\phi) > 0 \end{array} \right\} \quad \text{as} \quad \phi \left\{ \begin{array}{c} > \\ < \end{array} \right\} \phi^e.$$

Proof. This result follows directly from the proof of Proposition 6 above by setting $\phi^e = \hat{\phi} \left(\hat{h}_{max}^e \right)$, where the function $\hat{\phi}$ is defined in equation (49) and $\hat{h}_{max}^e$ is the fixed point of equation (55) corresponding to the given equilibrium. The equilibrium values of the initial

bail-ins $h^e(\phi)$ then follow from equation (50) and the equilibrium bailouts $b^e(\phi)$ follow from equation (52) $\square$

Proposition 8. *In any equilibrium of the bail-in game, $h^e(\phi) < h^*(\phi)$ holds for all ϕ with $b^e(\phi) > 0$. In addition, if $b^e(\phi) > 0$ holds for some ϕ , we have*

$$\int_{\underline{\phi}}^1 b^e(\phi) dF(\phi) > \int_{\underline{\phi}}^1 b^*(\phi) dF(\phi).$$

Proof. From Proposition 2, we have

$$h^*(\phi) = \min\{1 - \phi, 1 - \phi^*\} \quad \text{for all } \phi \quad (57)$$

where $\phi^* < 1$. From Proposition 7 and equation (22), we have

$$h^e(\phi) \in \{0, \underline{h}^e\} < 1 - \phi \quad \text{for all } \phi < \phi^e. \quad (58)$$

The proof of the first part of the proposition is by contradiction. Suppose $h^e(\phi) \geq h^*(\phi)$ held for some ϕ with $b(\phi) > 0$, that is, for some $\phi < \phi^e$. We will show that this inequality would imply that the bail-ins of all investors in all banks are larger in equilibrium than in the planner's allocation, which contradicts the optimality conditions for the fiscal authority's choice of bailout payments.

Since equation (58) shows $h^e(\phi) < 1 - \phi$ for all $\phi < \phi^e$, having $h^e(\phi) \geq h^*(\phi)$ for some such ϕ would imply that $h^*(\phi) < 1 - \phi$ also holds. Equation (57) would then imply that this value of ϕ must be strictly less than ϕ^* , meaning that bank ϕ would also receive a bailout in the planner's allocation and that its associated bail-in would be $h^*(\phi) = 1 - \phi^*$. Since $\phi^* < 1$, this bail-in would be strictly positive and, hence, $h^e(\phi)$ must be strictly positive as well. It then follows from equation (58) that $h^e(\phi) = \underline{h}^e > 0$ must hold for all $\phi < \phi^e$ and, therefore, we would have

$$\underline{h}^e \geq 1 - \phi^*. \quad (59)$$

Equation (22) shows $\underline{h}^e < 1 - \phi^e$. Combining these two inequalities yields $\phi^e < \phi^*$, that is, a strictly larger set of banks would be bailed out in the planner's allocation than in equilibrium. It would then follow from equations (57) and (58) that the inequality in equation (59) must be strict, meaning that the initial bail-in is strictly smaller in equilibrium for all $\phi < \phi^e$.

Equation (22) also shows that the equilibrium bail-in of the remaining investors in a bank with $\phi < \phi^e$ satisfies $1 - \hat{h}_{max}^e < \phi^e$. Combining this inequality with the results above would

yield

$$1 - \hat{h}_{max}^e < \phi^e < \phi^* = 1 - h^*(\phi) \quad \text{for all } \phi < \phi^*,$$

which would imply $\hat{h}_{max}^e > h^*(\phi)$ for all $\phi < \phi^*$. In other words, for all banks that are bailed out in both allocations, the bail-in of the remaining investors would be larger in equilibrium than in the planner's allocation. For banks that are not bailed out, the bail-in of all investors is $h(\phi) = 1 - \phi$ in both allocations. The discussion so far has established, therefore, that if $h^e(\phi) \geq h^*(\phi)$ held for some $\phi < \phi^e$, we would have

$$h^e(\phi) \geq h^*(\phi) \text{ and } \hat{h}^e(\phi) \geq h^*(\phi) \text{ for all } \phi, \quad (60)$$

with strict inequalities for $\phi < \phi^e$. Using these inequalities in the feasibility constraint in equation (3), which holds with equality in both allocations, we would then have

$$b^e(\phi) \leq b^*(\phi) \text{ for all } \phi,$$

with strict inequality for $\phi < \phi^e$. In other words, if the bail-ins were larger for all banks in the equilibrium allocation, the bailouts must be smaller in equilibrium. Equation (11) would then imply that the level of the public good must be higher in equilibrium, that is, $g^e > g^*$.

For the final step, we look at the first-order conditions that determine the bailout payments in each allocation. For the planner's allocation, equation (31) can be written as

$$u'((1 - h^*(\phi)) c_1^*) = v'(g^*) \quad \text{for all } \phi < \phi^*. \quad (61)$$

For the equilibrium allocation, equation (38) can be written as

$$u'\left(\left(1 - \hat{h}^e(\phi)\right) c_1^*\right) = v'(g^e) \quad \text{for all } \phi < \phi^e. \quad (62)$$

Using these two equations, $g^e > g^*$ would imply $\hat{h}^e(\phi) < h^*(\phi)$ for all banks that are bailed out in both allocations, which contradicts equation (60) above. Intuitively, if the planner's allocation had a smaller level of the public good than the equilibrium allocation, the planner would impose *larger* bail-ins on the remaining investors, not smaller bail-ins as derived above. Therefore, the conjecture that $h^e(\phi) \geq h^*(\phi)$ holds for some $\phi < \phi^e$ cannot be true.

For the second part of the proposition, we break the proof into two parts. First, suppose that $\phi^* \geq \phi^e$ holds, that is, that the set of banks bailed out in the planner's allocation is no smaller than the set bailed out in equilibrium. Using equations (22) and (57), we then have

$$h^*(\phi) = 1 - \phi^* \leq 1 - \phi^e < \hat{h}^e(\phi) \quad \text{for all } \phi < \phi^*.$$

In other words, the bail-in imposed on the remaining investors in banks that receive a bailout is larger in equilibrium. Using the first and last terms on this line in equations (61) and (62), respectively, it follows that the level of the public good is smaller in equilibrium, that is, $g^* > g^e$. Using equation (11) to relate the level of the public good to the total amount of bailout payments then delivers the desired result.

Now suppose instead that $\phi^* < \phi^e$. The proof for this case is by contradiction. Suppose the result were not true, that is, suppose instead that

$$\int_{\underline{\phi}}^1 b^e(\phi) dF(\phi) \leq \int_{\underline{\phi}}^1 b^*(\phi) dF(\phi) \quad (63)$$

held. Then by equation (11), the level of public good would be higher in equilibrium than in the planner's allocation, $g^e > g^*$. Using equations (61) – (62) above, we would then have

$$\hat{h}^e(\phi) \leq h^*(\phi)$$

for all banks that are bailed out in both allocations, that is, for all $\phi < \phi^*$. Intuitively, if more of public good were provided in equilibrium, then the bail-ins of the remaining investors would be set smaller in equilibrium. Recall that the first part of the proposition established that the initial bail-in is also smaller in equilibrium, $h^e(\phi) < h^*(\phi)$, for all $\phi < \phi^e$. Using these two inequalities in the feasibility constraint (3), which holds with equality, we would then have

$$b^e(\phi) > b^*(\phi) \quad \text{for all } \phi < \phi^*. \quad (64)$$

In other words, if the bail-ins for a given bank are smaller in equilibrium, the bailout must be larger. Integrating this inequality across banks, it would then follow that

$$\begin{aligned} \int_{\underline{\phi}}^1 b^e(\phi) dF(\phi) &\geq \int_{\underline{\phi}}^{\phi^*} b^e(\phi) dF(\phi) \\ &> \int_{\underline{\phi}}^{\phi^*} b^*(\phi) dF(\phi) = \int_{\underline{\phi}}^1 b^*(\phi) dF(\phi), \end{aligned}$$

which contradicts equation (63) above. Therefore, equation (63) cannot hold and the second part of the proposition has been established. $\square$

Proposition 9. *Either equilibrium in the bail-in game is unique or there are exactly two pure-strategy equilibria, one in which no bank runs occur and one in which a run occurs on all banks that are bailed out.*

Proof. This result follows from the proof of Proposition 6. That proof established that a pure-strategy equilibrium of the bail-in game corresponds to a fixed point of the mapping ζ defined in equation (55). It also established that ζ is a decreasing function at all points except $\hat{h}_{max} = \alpha$, where it jumps up. These results imply that ζ has at most two fixed points and, if two fixed points exist, one must have $\hat{h}_{max}^e \leq \alpha$ and the other must have $\hat{h}_{max}^e > \alpha$. Equation (51) then implies that, if multiple equilibria exist, there is one pure-strategy equilibrium in which no investors run on their bank and another in which investors run on all banks with $\phi < \phi \left(\hat{h}_{max}^e \right)$, that is, on all banks that receive a bailout in equilibrium. $\square$

Proposition 10. *The equilibrium bailout cutoff ϕ^e is strictly increasing in τ whenever $\phi^e \in (\underline{\phi}, 1)$.*

Proof. We first examine how an increase in τ affects the bail-in applied to the remaining investors in a bank that is bailed out, $\hat{h}_{max}$. Holding τ fixed, $\hat{h}_{max}$ is defined as the unique fixed point of equation (40). When ϕ^e is interior, this equation can be written as

$$\begin{aligned} \hat{h}_{max} &= 1 - \frac{1}{c_1^*} u'^{-1} \left[v' \left(\tau - \int_{\underline{\phi}}^1 b(\phi; \hat{h}_{max}) dF(\phi) \right) \right] \\ &\equiv z(\hat{h}_{max}; \tau). \end{aligned}$$

The function z is strictly decreasing in τ for all $\hat{h}_{max}$. It follows that the unique solution to the equation, $\hat{h}_{max}$ is strictly decreasing in τ . Next, a change in $\hat{h}_{max}$ affects the equilibrium bailout cutoff ϕ^e according to equation (49), which we can write as

$$\phi^e = U^{-1} \left(U_B(\hat{h}_{max}) \right). \quad (65)$$

Since U_B is strictly decreasing and U is strictly increasing, we have that ϕ^e is strictly decreasing in $\hat{h}_{max}$. Combining these two results, we have that an increase in τ causes $\hat{h}_{max}$ to strictly decrease, which causes the cutoff ϕ^e to strictly increase, as desired. $\square$

Proposition 11. *If the equilibrium of the economy without regulation has $h^e(\phi) = 0$ for those ϕ with $b^e(\phi) > 0$, then $h_{min}^* > 0$.*

Proof. The method of proof is to derive an expression for equilibrium welfare as a function of h_{min} and show that this function is strictly increasing at $h_{min} = 0$. We first need to characterize equilibrium play in the bail-in game with regulation. The proposition applies if, in the absence of regulation, the equilibrium bailout cutoff ϕ^e was strictly larger than $\underline{\phi}$ and

banks with $\phi < \phi^e$ chose $h^e(\phi) = 0$. A continuity argument can then be used to establish that, when h_{min} is sufficiently small, the equilibrium with regulation will be characterized by a cutoff $\phi^e(h_{min}) > \underline{\phi}$ such that all banks with $\phi < \phi^e(h_{min})$ receive bailouts and set their bail-in at the regulatory minimum, h_{min} . Banks with $\phi > \phi^e(h_{min})$ will not receive a bailout and will set their initial bail-in to $1 - \phi$ (as in Proposition 7) if allowed by the regulation, otherwise they will choose the minimum, h_{min} .

Next, we derive the expected utility from private consumption of investors in a bank in each of these situations. For any bank with $\phi < \phi^e(h_{min})$, we have

$$U_B \equiv \pi u[(1 - h_{min})c_1^*] + (1 - \pi)u\left[\left(1 - \hat{h}_{max}(h_{min})\right)c_2^*\right], \quad (66)$$

where $\hat{h}_{max}(h_{min})$ is determined as in Proposition 4. Banks with $\phi^e(h_{min}) < \phi < 1 - h_{min}$ will set $h(\phi) = 1 - \phi$ and the expected utility from private consumption of their investors will follow the function U defined in equation (6),

$$U(\phi) \equiv \pi u(\phi c_1^*) + (1 - \pi)u(\phi c_2^*).$$

For banks with $\phi > 1 - h_{min}$, the mandatory minimum will bind. Using the resource constraint in equation (3) with $b(\phi)$ set to zero, we can write the expected utility from private consumption of investors in such a bank as

$$U_N(\phi) \equiv \pi u[(1 - h_{min})c_1^*] + (1 - \pi)u\left[\frac{(1 - \pi)c_2^* - R(1 - \phi - h_{min}\pi c_1^*)}{1 - \pi}\right]. \quad (67)$$

Using these three expressions, we can write equilibrium welfare as a function of the minimum bail-in when h_{min} is sufficiently small as

$$\begin{aligned} \mathcal{W}(h_{min}) \equiv & F(\phi^e) U_B + \int_{\phi^e}^{1-h_{min}} U(\phi) f(\phi) d\phi + \int_{1-h_{min}}^1 U_N(\phi) f(\phi) d\phi \\ & + zU_N(1) + v(g), \end{aligned} \quad (68)$$

where z is the measure of banks with zero loss. While not explicit in the notation above, keep in mind that U_B , $U_N(\phi)$, ϕ^e and g are all functions of h_{min} .

Differentiating this function with respect to h_{min} yields

$$\frac{d\mathcal{W}(h_{min})}{dh_{min}} = \left\{ \begin{array}{l} U_B f(\phi^e) \frac{d\phi^e}{dh_{min}} + F(\phi^e) \frac{dU_B}{dh_{min}} \\ -U(\phi^e) f(\phi^e) \frac{d\phi^e}{dh_{min}} - U(1-h_{min})f(1-h_{min}) \\ +U_N(1-h_{min})f(1-h_{min}) + \int_{1-h_{min}}^1 \frac{dU_N(\phi)}{dh_{min}} f(\phi) d\phi \\ +z \frac{dU_N(1)}{dh_{min}} + v'(g) \frac{dg}{dh_{min}} \end{array} \right\}.$$

This expression can be simplified using the following observations. First, a bank whose realization of ϕ falls exactly on the equilibrium bailout cutoff ϕ^e must be indifferent between (i) setting the minimum bail-in and receiving a bailout and (ii) setting a bail-in of $1-\phi$ and not receiving a bailout. In other words, $U_B = U(\phi^e)$ must hold, which implies that the first terms on each of the first two lines of this expression sum to zero. Second, using equation (1) in equation (67) above, it is straightforward to show that $U_N(1-h_{min}) = U(1-h_{min})$, so that the second term on the second line and first term on the third line also sum to zero. We then evaluate the remaining terms at $h_{min} = 0$. When we do so, the second term in the third line becomes zero, as the mandatory minimum is no longer binding for any bank with positive losses. Finally, the first term on the last line becomes zero as well, because the allocation of resources in banks with no losses is efficient when $h_{min} = 0$ and, therefore, the utility loss associated with increasing the initial bail-in is second order. We can then write the derivative of equilibrium welfare evaluated at $h_{min} = 0$ as

$$\left. \frac{d\mathcal{W}(h_{min})}{dh_{min}} \right|_{h_{min}=0} = F(\phi^e) \left. \frac{dU_B}{dh_{min}} \right|_{h_{min}=0} + v'(g) \left. \frac{dg}{dh_{min}} \right|_{h_{min}=0}. \quad (69)$$

Intuitively, the impact of increasing h_{min} on welfare initially depends only on how it affects the allocation of resources in banks that are bailed out and how it affects the level of the public good. To establish the result, it suffices to show that this expression is strictly positive.

The remainder of the proof is divided into two steps. First we show that the second term on the right-hand side of equation (69) is always strictly positive. If the first terms is non-negative, the result then follows immediately. In the second step, we show that if the first terms is negative, the overall expression is still strictly positive.

Step (i): Show $\left. \frac{dg}{dh_{min}} \right|_{h_{min}=0} > 0$.

Suppose this were not true, that is, suppose g were initially (weakly) decreasing in h_{min} .

The first-order conditions for the fiscal authority's bailout policy imply

$$u' \left((1 - \hat{h}_{max}) c_1^* \right) = v'(g). \quad (70)$$

Because u and v are both strictly concave, it would then follow that $\hat{h}_{max}$ must be (weakly) increasing in h_{min} . This fact, in turn, would imply that U_B as defined in equation (66) would be strictly decreasing in h_{min} . Equation (65) would then imply that ϕ^e is strictly decreasing in h_{min} . Intuitively, if increasing h_{min} were to cause $\hat{h}_{max}$ to increase, it would imply that the consumption of all investors in a bank that is being bailed out would decrease. If U_B were to decrease, some banks that had previously chosen $h(\phi) = 0$ and received a bailout would instead choose $h(\phi) = 1 - \phi$ and no longer be bailed out, causing the cutoff ϕ^e to decrease.

Using the feasibility constraint in equation (3), we can write the budget constraint of the fiscal authority in equation (11) as

$$g = \tau - \int_{\underline{\phi}}^{\phi^e} \left(1 - \phi - h_{min} \pi c_1^* - \hat{h}_{max} (1 - \pi) \frac{c_2^*}{R} \right) dF(\phi)$$

where ϕ^e and $\hat{h}_{max}$ both depend on h_{min} . Differentiating with respect to h_{min} and evaluating the result at $h_{min} = 0$ yields

$$\left. \frac{dg}{dh_{min}} \right|_{h_{min}=0} = \left\{ \begin{array}{l} - \left(1 - \phi^e - \hat{h}_{max} (1 - \pi) \frac{c_2^*}{R} \right) f(\phi^e) \frac{d\phi^e}{dh_{min}} \Big|_{h_{min}=0} \\ + F(\phi^e) \left(\pi c_1^* + (1 - \pi) \frac{c_2^*}{R} \frac{d\hat{h}_{max}}{dh_{min}} \Big|_{h_{min}=0} \right) \end{array} \right\} \quad (71)$$

The arguments above established that if g were initially weakly decreasing in h_{min} , ϕ^e would be strictly decreasing and $\hat{h}_{max}$ would be (weakly) increasing. Both lines on the right-hand side of equation (71) would then be strictly positive, implying that g is strictly increasing in h_{min} , a contradiction. It follows that $\frac{dg}{dh_{min}}$ must be strictly positive when evaluated at $h_{min} = 0$.

Looking back at equation (69), if $\frac{dU_B}{dh_{min}}$ is non-negative, then the right-hand side is strictly positive and the proposition has been established. If not, we proceed to the second step.

Step (ii): Show $\left. \frac{dU_B}{dh_{min}} \right|_{h_{min}=0} < 0$ implies $\left. \frac{d\mathcal{W}(h_{min})}{dh_{min}} \right|_{h_{min}=0} > 0$.

Differentiating equation (66) with respect to h_{min} and evaluating at $h_{min} = 0$ yields

$$\left. \frac{dU_B}{dh_{min}} \right|_{h_{min}=0} = -\pi u'(c_1^*) c_1^* - (1 - \pi) u' \left((1 - \hat{h}_{max}) c_2^* \right) c_2^* \frac{d\hat{h}_{max}}{dh_{min}} \Big|_{h_{min}=0}.$$

Using equation (26) with α set to $(1 - \hat{h}_{max})$ and equation (70), we can rewrite this expression as

$$\left. \frac{dU_B}{dh_{min}} \right|_{h_{min}=0} = -u'(c_1^*) \pi c_1^* - v'(g)(1 - \pi) \frac{c_2^*}{R} \left. \frac{d\hat{h}_{max}}{dh_{min}} \right|_{h_{min}=0}. \quad (72)$$

Substituting equations (71) and (72) into equation (69) yields

$$\left. \frac{d\mathcal{W}(h_{min})}{dh_{min}} \right|_{h_{min}=0} = \left\{ \begin{array}{l} F(\phi^e) \left(-u'(c_1^*) \pi c_1^* - v'(g)(1 - \pi) \frac{c_2^*}{R} \left. \frac{d\hat{h}_{max}}{dh_{min}} \right|_{h_{min}=0} \right) \\ -v'(g) \left(1 - \phi^e - \hat{h}_{max}(1 - \pi) \frac{c_2^*}{R} \right) f(\phi^e) \left. \frac{d\phi^e}{dh_{min}} \right|_{h_{min}=0} \\ +v'(g)F(\phi^e) \left(\pi c_1^* + (1 - \pi) \frac{c_2^*}{R} \left. \frac{d\hat{h}_{max}}{dh_{min}} \right|_{h_{min}=0} \right) \end{array} \right\}.$$

Note that the two terms involving $\frac{d\hat{h}_{max}}{dh_{min}}$ cancel out. We can combine the remaining terms to write

$$\left. \frac{d\mathcal{W}(h_{min})}{dh_{min}} \right|_{h_{min}=0} = \left\{ \begin{array}{l} F(\phi^e) \pi c_1^* \left(v'(g) - u'(c_1^*) \right) \\ -v'(g) \left(1 - \phi^e - \hat{h}_{max}(1 - \pi) \frac{c_2^*}{R} \right) f(\phi^e) \left. \frac{d\phi^e}{dh_{min}} \right|_{h_{min}=0} \end{array} \right\}.$$

The first line on the right-hand side of this equation is strictly positive because the assumption in equation (4) together with the fiscal authority's choice of bailouts implies that $v'(g)$ is always larger than $u'(c_1^*)$ in equilibrium. This step assumes U_B is strictly increasing in h_{min} which, using equation (65), implies that ϕ^e is strictly decreasing in h_{min} . Therefore, the second line is also strictly positive, which implies that welfare is strictly increasing in h_{min} when evaluated at $h_{min} = 0$, as desired. $\square$

References

- Acharya, Viral, Irvind Gujral, Nirupama Kulkarni, and Hyun Song Shin (2011) “Dividends and bank capital in the financial crisis of 2007-2009,” NBER Working Paper 16896, March.
- Acharya, Viral, Hanh Le, and Hyun Song Shin (2016) “Bank capital and dividend externalities,” *Review of Financial Studies*, 30 (3), 988–1018.
- Acharya, Viral and Tanju Yorulmazer (2007) “Too many to fail—An analysis of time-inconsistency in bank closure policies,” *Journal of Financial Intermediation*, 16 (1), 1–31.
- Allen, Franklin, Elena Carletti, Itay Goldstein, and Agnese Leonello (2018) “Government guarantees and financial stability,” *Journal of Economic Theory*, 177, 518–557.
- Allen, Franklin and Douglas Gale (1998) “Optimal financial crises,” *Journal of Finance*, 53 (4), 1245–1284.
- Allen, Franklin and Douglas Gale (2004) “Financial fragility, liquidity, and asset prices,” *Journal of the European Economic Association*, 2 (6), 1015–1048.
- Bernard, Benjamin, Agostino Capponi, and Joseph E. Stiglitz (2017) “Bail-ins and bail-outs: Incentives, connectivity, and systemic stability,” NBER Working Paper 23747, October.
- Bhattacharya, Sudipto and Douglas Gale (1987) “Preference shocks, liquidity and central bank policy,” in *New Approaches to Monetary Economics*, W. Barnett and K. Singleton, eds., Cambridge University Press, 69–88.
- Bolton, Patrick and Martin Oehmke (2019) “Bank resolution and the structure of global banks,” *Review of Financial Studies*, 32 (6), 2384–2421.
- Brown, Craig O. and I. Serdar Dinc (2005) “The politics of bank failures: Evidence from emerging markets,” *Quarterly Journal of Economics*, 120 (4), 1413–1444.
- Calomiris, Charles W. and Charles M. Kahn (1991) “The role of demandable debt in structuring optimal banking arrangements,” *American Economic Review*, 81 (3), 497–513.
- Castiglionesi, Fabio and Noemi Navarro (2020) “(In) Efficient interbank networks,” *Journal of Money, Credit and Banking*, 52 (2-3), 365–407.
- Chari, V.V. and Ravi Jagannathan (1988) “Banking panics, information, and rational expectations equilibrium,” *Journal of Finance*, 43 (3), 749–761.
- Chen, Qi, Itay Goldstein, and Wei Jiang (2010) “Payoff complementarities and financial fragility: Evidence from mutual fund outflows,” *Journal of Financial Economics*, 97 (2), 239–262.
- Cipriani, Marco, Antoine Martin, Patrick E McCabe, and Bruno Maria Parigi (2014) “Gates, fees, and preemptive runs,” *FRB of New York Staff Report* (670).

- Colliard, Jean-Edouard and Denis Gromb (2018) “Financial restructuring and resolution of banks,” HEC Paris Research Paper No. FIN-2018-1272, May.
- Cooper, Russell and Thomas W. Ross (1998) “Bank runs: Liquidity costs and investment distortions,” *Journal of Monetary Economics*, 41 (1), 27–38.
- Dewatripont, Mathias (2014) “European banking: Bailout, bail-in and state aid control,” *International Journal of Industrial Organization*, 34, 37–43.
- Diamond, Douglas W. and Philip H. Dybvig (1983) “Bank runs, deposit insurance, and liquidity,” *Journal of Political Economy*, 91 (3), 401–419.
- Diamond, Douglas W. and Raghuram G. Rajan (2001) “Liquidity risk, liquidity creation, and financial fragility: A theory of banking,” *Journal of Political Economy*, 109 (2), 287–327.
- Engineer, Merwan (1989) “Bank runs and the suspension of deposit convertibility,” *Journal of monetary Economics*, 24 (3), 443–454.
- Ennis, Huberto M. (2012) “Some theoretical considerations regarding net asset values for money market funds,” *Federal Reserve Bank of Richmond Economic Quarterly*, 98 (4), 231–254.
- Ennis, Huberto M. and Todd Keister (2009a) “Bank runs and institutions: The perils of intervention,” *American Economic Review*, 99 (4), 1588–1607.
- Ennis, Huberto M. and Todd Keister (2009b) “Run equilibria in the Green–Lin model of financial intermediation,” *Journal of Economic Theory*, 144 (5), 1996–2020.
- Ennis, Huberto M. and Todd Keister (2010) “Banking panics and policy responses,” *Journal of Monetary Economics*, 57 (4), 404–419.
- Farhi, Emmanuel, Mikhail Golosov, and Aleh Tsyvinski (2009) “A theory of liquidity and regulation of financial intermediation,” *Review of Economic Studies*, 76 (3), 973–992.
- Flannery, Mark J. (1996) “Financial crises, payment system problems, and discount window lending,” *Journal of Money, Credit and Banking*, 28 (4), 804–824.
- Flannery, Mark J. (2014) “Contingent capital instruments for large financial institutions: A review of the literature,” *Annual Review of Financial Economics*, 6 (1), 225–240.
- Freixas, Xavier and José Jorge (2008) “The role of interbank markets in monetary policy: A model with rationing,” *Journal of Money, Credit and Banking*, 40 (6), 1151–1176.
- Goldstein, Itay, Hao Jiang, and David T Ng (2017) “Investor flows and fragility in corporate bond funds,” *Journal of Financial Economics*, 126 (3), 592–613.
- Goldstein, Itay, Alexandr Kopytov, Lin Shen, and Haotian Xiang (2020) “Bank heterogeneity and financial stability,” NBER Working Paper 27376, June.

- Goldstein, Itay and Ady Pauzner (2005) “Demand–deposit contracts and the probability of bank runs,” *Journal of Finance*, 60 (3), 1293–1327.
- Goodhart, Charles and Emiliios Avgouleas (2015) “A critical evaluation of bail-in as a bank recapitalization mechanism,” in Evanoff, Douglas, Andrew Haldane, and George Kaufman eds. *The New International Financial System: Analyzing The Cumulative Impact Of Regulatory Reform*, Chap. 11, 267–305: World Scientific.
- Heider, Florian, Marie Hoerova, and Cornelia Holthausen (2015) “Liquidity hoarding and interbank market rates: The role of counterparty risk,” *Journal of Financial Economics*, 118 (2), 336–354.
- Henderson, Christopher, William Lang, and William Jackson (2015) “Insider bank runs: Community bank fragility and the financial crisis of 2007,” Federal Reserve Bank of Philadelphia Working Paper 15-09, February.
- Iyer, Rajkamal, Manju Puri, and Nicholas Ryan (2016) “A tale of two runs: Depositor responses to bank solvency risk,” *Journal of Finance*, 71 (6), 2687–2726.
- Jacklin, Charles J. (1987) “Demand deposits, trading restrictions, and risk sharing,” in Prescott, Edward C. and Neil Wallace eds. *Contractual Arrangements for Intertemporal Trade*, 26–47: Univ. of Minnesota Press.
- Jin, Dunhong, Marcin Kacperczyk, Bige Kahraman, and Felix Suntheim (2019) “Swing pricing and fragility in open-end mutual funds,” IMF Working Paper 19/227, November.
- Kareken, John H. and Neil Wallace (1978) “Deposit insurance and bank regulation: A partial-equilibrium exposition,” *Journal of Business*, 413–438.
- Keister, Todd (2016) “Bailouts and financial fragility,” *The Review of Economic Studies*, 83 (2), 704–736.
- Kroszner, Randall S. and Philip E. Strahan (1996) “Regulatory incentives and the thrift crisis: Dividends, mutual-to-stock conversions, and financial distress,” *Journal of Finance*, 51 (4), 1285–1319.
- Li, Yang (2020) “Liquidity regulation and financial stability,” *Macroeconomic Dynamics*, 24 (5), 1240–1263.
- Milgrom, Paul and John Roberts (1994) “Comparing equilibria,” *American Economic Review*, 84 (3), 441–459.
- Mitkov, Yuliyana (2019) “Inequality and financial fragility,” *Journal of Monetary Economics*, forthcoming.
- Peck, James and Karl Shell (2003) “Equilibrium bank runs,” *Journal of Political Economy*, 111 (1), 103–123.
- Rogoff, Kenneth and Anne Sibert (1988) “Elections and macroeconomic policy cycles,” *Review of Economic Studies*, 55 (1), 1–16.

- Schmidt, Lawrence, Allan Timmermann, and Russ Wermers (2016) “Runs on money market mutual funds,” *American Economic Review*, 106 (9), 2625–57.
- Sommer, Joseph (2014) “Why bail-in? And how!,” *Federal Reserve Bank of New York Economic Policy Review*, 20 (2), 207–228.
- Sultanum, Bruno (2014) “Optimal Diamond–Dybvig mechanism in large economies with aggregate uncertainty,” *Journal of Economic Dynamics and Control*, 40, 95–102.
- Voellmy, Lukas (2019) “Preventing Bank Panics with Fees and Gates,” mimeo. <https://dx.doi.org/10.2139/ssrn.3425663>.
- Wallace, Neil (1988) “Another attempt to explain an illiquid banking system: The Diamond and Dybvig model with sequential service taken seriously,” *Federal Reserve Bank of Minneapolis Quarterly Review*, 12 (4), 3–16.
- Walther, Ansgar and Lucy White (2019) “Rules versus discretion in bank resolution,” *Review of Financial Studies*, forthcoming.
- Yorulmazer, Tanju (2014) “Case studies on disruptions during the crisis,” *Federal Reserve Bank of New York Economic Policy Review*, 20 (1), 17–28.