

Bail-ins and Bail-outs: Incentives, Connectivity, and Systemic Stability

Benjamin Bernard, Agostino Capponi, and Joseph E. Stiglitz*

October 24, 2021

This paper endogenizes intervention in financial crises as the strategic negotiation between a regulator and creditors of distressed banks. Incentives for banks to contribute to a voluntary bail-in arise from their exposure to financial contagion. In equilibrium, a bail-in is possible only if the regulator's threat to not bail out insolvent banks is credible. Contrary to models without intervention or with government bailouts only, sparse networks enhance welfare for two main reasons: they improve the credibility of the regulator's no-bailout threat for large shocks and they reduce free-riding incentives among bail-in contributors when the threat is credible.

Financial institutions are linked to each other via bilateral contractual obligations and are thus exposed to counterparty risk of their obligors. If one institution defaults on its liabilities, it affects the solvency of its creditors. Since the creditors are also borrowers, they may not be able to repay what they owe and default themselves—problems in one financial institution spread to others in what is known as financial contagion. Large shocks can trigger a cascade of defaults, which impose negative externalities on the economy. The extent of these cascades—the magnitude of the systemic risk—depends on the nature of the linkages, i.e., the structure of the financial system. In the 2008 crisis, it became apparent that the financial system had evolved in a way which enhanced its ability to absorb small shocks but made it more fragile in the face of a large shock. While a few studies called attention to these issues before the crisis, it was only after the crisis that the impact of the network structure on systemic risk became a major object of analysis.¹ Most of the existing

*Bernard: Department of Economics, National Taiwan University, No. 1, Sec. 4, Roosevelt Rd, Da'an District, Taipei City, 10617 Taiwan; benbernard@ntu.edu.tw. Capponi: Department of Industrial Engineering and Operations Research, Columbia University, 500 W 120th St, Mudd Hall 535-G, New York, NY 10027; ac3827@columbia.edu. Stiglitz: Columbia Business School, Columbia University, 3022 Broadway, Uris Hall 212, New York, NY 10027; jes322@gsb.columbia.edu. We are grateful to George Pennacchi, Yiming Ma, Asuman Ozdaglar, Alireza Tahbaz-Salehi, Darrell Duffie, Jakša Cvitanić, Matt Elliott, Douglas Gale, Matthew Jackson, Piero Gottardi, and Felix Corell for interesting discussions and perceptive comments. We would like to thank seminar participants of the Laboratory for Information and Decision Systems at the Massachusetts Institute of Technology, the Cambridge Finance Seminar series, the London School of Economics, Stanford University, New York University, the Center of Operational Research and Econometrics at the University of Louvain, the Fields Institute, the Austrian Central Bank, the National Bank of Belgium, the third annual conference on Network Science and Economics, the Columbia Conference on Financial Networks: Big Risks, Macroeconomic Externalities, and Policy Commitment Devices, the 2018 SFS Cavalcade North America, the 2019 American Finance Association meeting, and the 2018 North American Summer Meeting of the Econometric society for their valuable feedback. The research of Agostino Capponi is supported by a NSF-CMMI: 1752326 CAREER grant. Benjamin Bernard acknowledges financial support from grant P2SKP1.171737 by the Swiss National Science Foundation, from grant 108-2410-H-002-249 by the Ministry of Science and Technology in Taiwan, and from grant 109L900203 by the Center for Research in Econometric Theory and Applications by the Ministry of Education in Taiwan. Joseph Stiglitz acknowledges the support of the Columbia Business School and of the grant on Financial Stability from the Institute for New Economic Thinking.

¹Most notably, Allen and Gale (2000) and Greenwald and Stiglitz (2003). See also Boissay (2006), Elsinger, Lehar and Summer (2006), Castiglionesi (2007), May, Levin and Sugihara (2008), and Nier et al. (2007). One of the reasons for the limited study is the scarce availability of data on interbank linkages.

studies analyze the systemic risk implications of a default cascade, taking into account the network structure, asset liquidation costs, and different forms of inefficiencies that arise at default. Many of these models, however, do not account for the possibility of intervention to stop the cascade. There is either no rescue of insolvent banks or the regulator (or central bank or other government institution) intervenes by following an exogenously specified protocol. The goal of our paper is to endogenize the intervention mechanism as the equilibrium outcome of the strategic interaction between regulator and financial institutions.

The most common default resolution procedure during the 2007–2009 financial crisis was the *bailout*.² In a bailout, the government injects liquidity to help distressed banks service their debt, effectively transferring liabilities from the private sector to the public sector. Some governments, such as Germany’s, have called for private-sector participation through *bail-ins*, in which creditors write down their interbank claims against troubled banks.³ Bail-ins effectively amount to a transfer of liabilities within the private sector, which places the burden of losses on creditors as opposed to taxpayers. A prominent example of a bail-in is the rescue of the hedge fund Long-Term Capital Management (LTCM).⁴ In our paper, we also consider *assisted bail-ins*, in which the regulator provides some liquidity assistance to incentivize the formation of a bail-in.

The negotiation process between the regulator and the banks consists of three stages. In the first stage, the regulator proposes an assisted bail-in allocation, which specifies the contributions by each solvent bank, as well as the liquidity injections provided to each bank. In the second stage, the banks simultaneously decide whether or not to participate in the proposed rescue. If they all participate, the game ends with the proposed rescue consortium. Otherwise, the regulator reacts in the third stage by either (i) proceeding with the residual bail-in at increased taxpayer expense, (ii) proceeding with a bail out, or (iii) avoiding any intervention. After transfers are made, the banks’ liabilities are cleared simultaneously in the spirit of Eisenberg and Noe (2001), possibly leading to a default cascade if the outcome of the negotiation leaves some banks insolvent.

Financial contagion in our model occurs through the two most prominent channels identified by historical events.⁵ First, distressed banks may have to liquidate some of their asset holdings in order to fulfill their obligations. In the liquidation process, the asset is transferred to buyers with lower levels of expertise in managing the asset, causing a drop in its value. Due to financial frictions, addi-

²The Bush administration bailed out large financial institutions (AIG insurance, Bank of America, Bear Stearns and Citigroup) and government sponsored entities (Fannie Mae, Freddie Mac) at the heart of the crisis. The European Commission intervened to bail out financial institutions in Greece and Spain.

³In spite of such calls and the design of instruments to make private-sector participation automatic, there have been few successful bail-ins. Automatic participation is implemented through the use of “bail-inable debt” such as contingent-convertible bonds in order to reduce the banks’ credit risk. The focus of this paper is on the welfare impact of default resolution policies after these risk-mitigating instruments have already been used.

⁴LTCM Portfolio collapsed in the late 1990s. On September 23, 1998, a recapitalization plan of \$3.6 billion was coordinated under the supervision of the Federal Reserve Bank of New York (FRBNY). A total of fourteen banks agreed to participate and two banks (Bear Stearns and Lehman Brothers) rejected the proposal.

⁵Quoting Greenspan (1998): “It was the judgment of officials at the FRBNY, who were monitoring the situation on an ongoing basis, that the act of unwinding LTCM’s portfolio in a forced liquidation would not only have a significant distorting impact on market prices but also in the process could produce large losses, or worse, for a number of creditors and counterparties, and for other market participants who were not directly involved with LTCM. In that environment, it was the FRBNY’s judgment that it was to the advantage of all parties—including the creditors and other market participants—to engender if at all possible an orderly resolution rather than let the firm go into disorderly fire-sale liquidation following a set of cascading cross defaults.”

tional units of the asset are sold to marginally less efficient buyers. Hence, one bank’s liquidation decision affects the value that another bank is able to recover from its asset sale. Second, if a distressed bank defaults on its obligations, its creditors only collect a fraction of their claims. Bankruptcy imposes deadweight losses, which are amplified by feedback loops among defaulting banks.

We investigate the structure of default resolution plans that arise in equilibrium when the regulator cannot ~~credibly~~ commit *ex ante* to an ~~ex-post-suboptimal~~ *fixed* resolution policy. The lack of commitment power has important consequences for the regulator’s negotiation power: if banks are aware that without their participation, the regulator prefers a bailout over no intervention, then they have no incentive to participate in any assisted bail-in. We say that, in this case, the regulator’s *no-intervention threat* fails to be credible.⁶ Only if the threat is credible can an assisted bail-in be organized in equilibrium. Individually, a bank is willing to contribute up to the maximum it would lose in a default cascade. However, because losses are amplified as the shock propagates through the system, *aggregate* losses exceed the required ex-ante contributions. Therefore, it is not necessary that every bank contributes and banks have an incentive to free-ride on the contributions of others. In the equilibrium bail-in, the set of contributing banks minimizes free-riding incentives by consisting of the banks with the largest exposure to contagion. It thus follows that banks are willing to contribute more in sparser networks: because losses are more concentrated, the benefits of a bail-in are more targeted to the contributors than in more diversified networks, thereby reducing free-riding incentives.

A key determinant of the equilibrium outcome is the credibility of the regulator’s no-intervention threat. We show that the threat is credible if and only if the losses generated by the regulator’s inaction—equal to the amplification of the shock as it propagates through the network—do not exceed a given threshold. Whether the shock amplification increases with the size of the initial shock faster than the threshold rises depends on asset illiquidity and on the network structure. We identify a variable, which we call the *total throughput* of defaulting banks, as a sufficient statistic for the dependence of the credibility on the network structure, conditional on the banks’ levels of solvency and their total claims on solvent banks. The total throughput measures the rate of spillover losses transmitted to the solvent banks in the system. Conditional on the banks’ solvency, the total throughput depends only on the network structure and not on the banks’ balance sheet *quantities*, making it a convenient measure to compare the potential to propagate losses among different network structures. We demonstrate that the throughput increases as the connectivity of defaulting banks increases. As a result, in sparsely connected networks, the regulator’s threat may not be credible for small shocks, but the credibility improves as the shock grows larger. Because the total throughput is small, the systemic threat does not increase much with the size of the shock. By contrast, in more diversified network structures, small losses can be well absorbed and the threat not to intervene is credible. However, because the total throughput is large, the threat becomes less credible as the shock size increases. As illustrated in Figure 1, endogeneity of the default resolution plan thus reverses the relative desirability of network structures for intermediate

⁶Several attempted bail-ins failed because the threat of not undertaking a bail-out was not credible; see Stiglitz (2002).

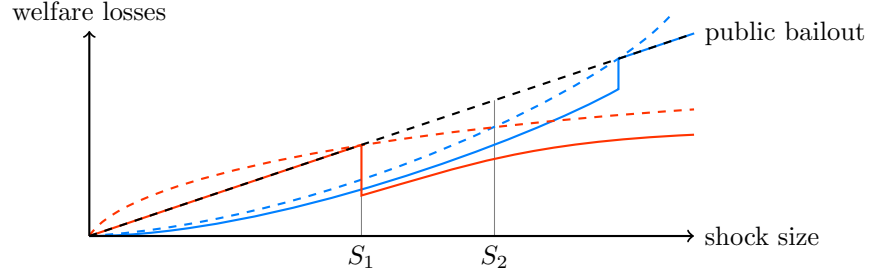


Figure 1: The figure compares equilibrium welfare losses in a diversified (blue) and a concentrated network (red) in the presence (solid lines) and absence (dashed lines) of intervention.⁷ Because the total throughput is small (large) in sparse (dense) networks, welfare losses in absence of intervention are concave (convex). When the no-intervention losses exceed the costs of a public bailout (black dashed line), the government’s threat to not intervene fails to be credible and a public bailout is the only possible equilibrium rescue. If the threat is credible, contributions from the private sector are possible. Contributions are larger in the sparse network because losses without intervention are more concentrated, thereby reducing free-riding incentives. Without intervention or in a model with bailouts only, the diversified network is preferable up until shock size S_2 . Endogeneity of the intervention reverses the relative desirability of the two network structures for intermediate shock sizes between S_1 and S_2 .

shock sizes when compared to models without intervention or models with bailouts only.

We then proceed to investigate the structure of optimal intervention plans, that is, the set of banks that are rescued. A bank is not rescued in a bailout if its creditors have the capacity to absorb a significant portion of the losses due to contagion. In a bail-in, not only the total loss absorption capacity matters but also its distribution among the creditors. When a bank has a single creditor, that creditor cannot free-ride on the contributions of others. Proposing a bail-in for a bank whose losses are absorbed by a single creditor thus leads to an increase in welfare. However, if contagion effects are spread among many creditors, rescuing that bank does not substantially increase contributions from the private sector because of the ~~banks’~~ inherent free-riding incentives ~~among the banks’ many creditors~~. It may then be welfare-enhancing to not rescue such a bank. Banks rescued in the optimal bail-in may thus default in the optimal bailout and vice versa.

Our results uncover the economic forces behind the decision to rescue banks which default due to fundamental shocks versus those failing as a result of contagion. For example, the government opted to rescue AIG as opposed to Goldman Sachs during the global financial crisis.⁸ Because loss absorption capacities were low in the aftermath of Lehman Brother’s default, ~~a bankruptcy filings~~ ~~of~~ ~~by~~ AIG might have had far-reaching consequences, bringing under many of its creditors including Goldman Sachs. If the financial system had been in a more resilient state, in which Goldman Sachs was the only contagiously defaulting bank, it might have been welfare-enhancing to let AIG default and provide liquidity assistance only to Goldman Sachs after the depletion of its capital buffers.

The remainder of the paper is organized as follows. In Section 1, we relate our work to the existing literature. We develop the model in Section 2. We characterize incentives and the equilibrium intervention outcome for any financial network in Section 3. In Section 4, we analyze the impact of the network structure, the shock size, and asset illiquidity on the public bailout and the credibility of

⁸The bailout of AIG was widely speculated to be an indirect bailout of Goldman Sachs, to which it had sold millions of dollars worth of insurance. Because of the large exposure, AIG’s default might have lead to the contagious default of Goldman Sachs, the largest investment bank at the time. ~~Our analysis highlights the factors that rationally should have motivated the decision about whether to bail out AIG or allow AIG to go under, and then possibly bail out Goldman Sachs.~~

the no-intervention threat. In Section 5, we characterize the set of banks to be optimally rescued in bail-ins and bailouts. We discuss the impact of the network structure on the equilibrium intervention plan in Section 6 using a model calibrated to a data set of the European Banking Authority (EBA). Section 7 concludes. The proofs of the main results are contained in the appendix. Supplementary results and auxiliary proofs are in the online appendix; see Bernard, Capponi and Stiglitz (2019).

1 Literature Review

Our paper is related to a vast branch of literature on financial contagion in interbank networks, pioneered by Allen and Gale (2000) and Eisenberg and Noe (2001). Cifuentes, Ferrucci and Shin (2005) have extended those models to include contagion through asset sell-offs, a contagion channel that is also present in our paper. In a model that draws parallels between the financial crisis and a systemic bank run, Uhlig (2010) provides a micro-foundation for fire-sale contagion through uncertainty aversion and adverse selection. The impact of the network structure on the extent of financial contagion has been studied by Gai and Kapadia (2010), Gai, Haldane and Kapadia (2011), Battiston et al. (2012), Elliott, Golub and Jackson (2014), Acemoglu, Ozdaglar and Tahbaz-Salehi (2015), Glasserman and Young (2015), Capponi, Chen and Yao (2016), and Cabrales, Gottardi and Vega-Redondo (2017). The above mentioned works show that in absence of any intervention, dense connections tend to reduce financial contagion for small shocks but can serve as an amplifier of large shocks. While this effect remains present in our model with strategic interventions, dense networks additionally cause free-riding incentives among bail-in contributors, further reducing the desirability of dense networks. We refer to Glasserman and Young (2016), Benoit et al. (2017), and Jackson and Pernoud (2020) for thorough surveys on systemic risk and financial contagion.

The role of the government in stopping financial contagion has been studied since the works of Freixas, Parigi and Rochet (2000) and Gale and Vives (2002). Many papers have focused on the moral hazard problem that bailouts create ex ante, causing banks to take on excessive risks. In order to trigger a bailout more frequently when the market moves against them, banks have an incentive to overborrow (Chari and Kehoe (2016)) or correlate their investments (Acharya and Yorulmazer (2007) and Farhi and Tirole (2012)). Despite those perverse incentives, committing to a no-bailout policy ~~leads~~[can lead](#) to a reduction of welfare ex ante (Keister (2016)). We focus instead on the incentives needed ex post for private-sector involvement in resolving financial distress. These incentives also crucially depend on the government's preferred bailout through the credibility of the no-bailout threat. Our model, like the rest of the literature on financial networks, does not account for the endogenous structure of the interbank network. The exception is Erol (2018), which shows that bailouts cause banks to organize in a core-periphery network. Different from Erol (2018), the regulator in our model is strategic and cannot commit to a bailout or a no-intervention policy.

A few papers have studied private default resolutions or bail-ins in models different from ours. In Rogers and Veraart (2013), banks can prevent a default cascade through mergers. In their paper, however, a merger needs not be incentive compatible for the shareholders of an individual bank, nor does the government take an active role. By contrast, the credibility of the regulator's actions

and the ~~free-riding problem that the shareholders face~~ the rational responses of other banks, including their incentives to free ride, are at the heart of our analysis. Duffie and Wang (2017) consider a bargaining model, in which bail-ins are done contractually rather than through a central planner. They show that private bail-ins reached in a 3-bank network are efficient in the limit as ~~disagreement is~~ bargaining failures are disallowed. Unlike our paper, they abstract away from cross-network externalities and the government’s involvement. In our paper, reducing externalities from asset liquidation and bankruptcy losses are the regulator’s ~~driving force to organize a default resolution plan~~ main motivation to facilitate bail-ins by providing well-designed subsidies. Schilling (2018) shows that if the regulator can impose a mandatory bail-in on depositors in a single-bank model, the depositors will preempt such a policy by running on the bank more frequently ex ante. Because bail-ins and bailouts are financed by solvent financial institutions and the government in our paper, interventions always make depositors better off.

2 Model

We consider an economy consisting of n risk-neutral financial institutions $i = 1, \dots, n$, called “banks”, which lasts for three periods $t = 0, 1, 2$. At the initial period, each bank i is endowed with capital that it can lend to other banks, invest into a liquid asset, called “cash”, or invest into an illiquid asset yielding a random return at time $t = 1$ and a non-pledgeable return at $t = 2$. Because long-term returns are non-pledgeable, liabilities have to be cleared at time $t = 1$. After short-term returns are realized, banks may liquidate their risky asset to help service their debt, but doing so imposes a downward impact on its value because the asset is sold to less efficient users.

Banks negotiate debt contracts with each other at $t = 0$. We denote by L^{ji} bank i ’s liabilities to bank j at time $t = 1$ and use $L^i := \sum_{j=1}^n L^{ji}$ to denote bank i ’s total interbank liabilities. Let p^i denote the market value of bank i ’s liabilities, which may be lower than the notional value L^i if bank i is unable to repay its liabilities in full. If $p^i < L^i$, we say that bank i defaults. All interbank liabilities have equal seniority: if bank i defaults, each creditor j receives $\pi^{ji}p^i$ from bank i , where

$$\pi^{ji} = \begin{cases} L^{ji}/L^i & \text{if } L^i > 0, \\ 0 & \text{otherwise.} \end{cases}$$

The relative liability matrix $\pi = (\pi^{ij})_{i,j=1,\dots,n}$ captures the structure of the financial network. In addition to its liabilities within the financial network, bank i has financial commitments w^i outside the financial sector due at time $t = 1$, which have higher seniority than the interbank liabilities. These commitments include wages, depositors’ claims, and other operating expenses.

For each bank i , we denote by e^i the long-term returns of the illiquid asset held by i . We denote by c^i the sum of i ’s cash holdings and short-term returns realized at time $t = 1$. If a bank i is not able to meet its liabilities out of current income, it will sell a portion $\ell^i \in [0, e^i]$ of its illiquid asset.

Liquidation imposes a downward impact on the value of the asset, yielding a recovery rate of

$$\alpha(\ell) = \exp\left(-\gamma \sum_{i=1}^n \ell^i\right). \quad (1)$$

Equation (1) captures the fact that buyers of the asset have a lower expertise than the seller and that additional units of the asset are sold to marginally less efficient buyers. The parameter $\gamma \geq 0$ captures the rate, at which the pool of efficient asset buyers diminishes.⁹

A defaulting bank liquidates all of its assets to repay the maximal amount **that it can** to its creditors. Because liquidation is costly, a solvent bank liquidates just enough to meet its liabilities. If interbank repayments are p and the asset recovery rate is α , bank i thus liquidates an amount

$$\ell^i(p, \alpha) = \min\left(\frac{1}{\alpha} \left(L^i + w^i - c^i - \sum_{j=1}^n \pi^{ij} p^j\right)^+, e^i\right), \quad (2)$$

where $(\cdot)^+ = \max(\cdot, 0)$ denotes the positive part.¹⁰ If a bank i cannot meet its liabilities even after liquidating all of its assets, it will default. The default creates losses proportional to the banks' asset value: only a fraction $\beta \in (0, 1]$ of the value is paid to the creditors and a fraction $(1 - \beta)$ is lost.¹¹ The market value of bank i 's interbank liabilities is thus equal to

$$p^i = \begin{cases} L^i & \text{if } c^i + \alpha \ell^i + (\pi p)^i \geq L^i + w^i, \\ (\beta(c^i + \alpha \ell^i + (\pi p)^i) - w^i)^+ & \text{otherwise,} \end{cases} \quad (3)$$

where $(\pi p)^i = \sum_{j=1}^n \pi^{ij} p^j$ is the market value of bank i 's interbank claims. To summarize, the financial system in our model is parametrized by $(L, \pi, e, c, w, \gamma, \beta)$, where L, e, c , and w are vectors, whose entries are the corresponding balance sheet quantities of each bank.

How much a bank is able to repay depends on the solvency of the other banks in the system. In a clearing equilibrium, every solvent bank repays its liabilities in full and every insolvent bank pays the entire value (after bankruptcy costs) to its creditors.

Definition 2.1. The triple of repayments, liquidation decisions, and recovery rate (p, ℓ, α) is a *clearing equilibrium* for a financial system $(L, \pi, e, c, w, \gamma, \beta)$ if it satisfies (1)–(3).

For payments p , recovery rate α , and liquidation decisions ℓ , the value of bank i 's equity is

$$V^i(p, \ell, \alpha) := ((\pi p)^i + c^i + e^i - (1 - \alpha)\ell^i - w^i - L^i)^+. \quad (4)$$

If the payment p^i made by bank i is positive, it is divided pro-rata among bank i 's junior creditors and the senior creditors are paid in full. If $p^i = 0$, the junior creditors lose the full amount of their

⁹The elasticity parameter γ typically depends on the asset class: Ellul, Jotikasthira and Lundblad (2011) find that γ is on the order of 10^{-8} for corporate bonds (see Table 8 therein).

¹⁰The model can be adjusted to account for liquidity requirements such as Basel-III (2013). This can be achieved by setting c^i equal to the amount of liquid assets that can be liquidated before hitting the liquidity coverage ratio requirement.

¹¹According to Moody's analysis, the average recovery rate for unsecured corporate bonds ranges from 30% to 43%; see Exhibit 8 in <https://www.moody's.com/sites/products/DefaultResearch/2006600000428092.pdf>.

claims and the senior creditors of bank i suffer a loss of

$$\delta^i(p, \alpha) := (w^i - \beta(c^i + \alpha e^i + (\pi p)^i))^+. \quad (5)$$

We denote by $\mathcal{D}(p, \ell, \alpha) := \{i \mid L^i + w^i > c^i + \alpha \ell^i + (\pi p)^i\}$ the set of defaulting banks. Welfare losses in a clearing equilibrium are defined as the weighted sum of default costs, i.e.,

$$W_\lambda(p, \ell, \alpha) := (1 - \alpha) \sum_{i=1}^n \ell^i + (1 - \beta) \sum_{i \in \mathcal{D}(p, \ell, \alpha)} (c^i + \alpha e^i + (\pi p)^i) + \lambda \sum_{i \in \mathcal{D}(p, \ell, \alpha)} \delta^i(p, \alpha). \quad (6)$$

The first term captures losses due to a misallocation of the asset when it is sold to less efficient users. The second term quantifies deadweight losses from bankruptcy. The last term is the welfare cost of losses borne by the senior creditors. The factor $\lambda \geq 0$ captures the importance the regulator assigns to those losses relative to deadweight losses. A regulator with $\lambda = 0$ views losses of senior creditors simply as transfers of wealth and not as losses to the economy. A higher value of λ indicates a higher priority to the economy outside of the banking sector.

Every financial system admits a clearing equilibrium due to Tarski's fixed-point theorem. *In financial systems with bankruptcy costs ($\beta < 1$) or price impacts due to asset liquidation ($\gamma > 0$), there may exist multiple clearing equilibria.* Following standard practices in the literature, liabilities are then cleared with the unique Pareto-efficient clearing equilibrium.

Lemma 2.1. *For any financial system, there exists a greatest clearing equilibrium $(\bar{p}, \bar{\ell}, \bar{\alpha})$ that is Pareto efficient for any $\lambda \geq 0$, i.e., for any other clearing equilibrium (p, ℓ, α) , it holds that $\bar{\alpha} \geq \alpha$ and $W_\lambda(\bar{p}, \bar{\ell}, \bar{\alpha}) \leq W_\lambda(p, \ell, \alpha)$ as well as $\bar{p}^i \geq p^i$, $\bar{\ell}^i \leq \ell^i$, and $V^i(\bar{p}, \bar{\ell}, \bar{\alpha}) \geq V^i(p, \ell, \alpha)$ for any bank i .*

2.1 Contagion and Default Cascade

We position ourselves at time $t = 1$ when short-term returns have been realized but banks have not yet cleared their liabilities. If the banks and/or the regulator have implemented automatic bail-in triggers such as contingent convertible bonds, then $(L, \pi, e, c, w, \alpha, \beta)$ represents the state of the financial system after accounting for these risk-mitigating actions. Depending on the size of the shock, banks may still need to liquidate their assets to remain solvent and defaults may still occur.

There are two channels of financial contagion in our model. The first channel is the downward price pressure imposed on an asset sold by illiquid banks. Due to financial frictions, additional units of the asset are sold to marginally less efficient buyers as formalized in (1). Since illiquid banks target the same pool of potential buyers, one bank's liquidation decision affects the average recovery value of other distressed banks. This leads to a downward spiral, which converges to the highest recovery rate, for which the asset's demand—given by the inverse of (1)—equals the liquidated amount.

Lemma 2.2. *For any vector p of repayments, there exists a solution (ℓ_p, α_p) to (1) and (2) such that $\alpha \leq \alpha_p$ and $\ell_p^i \leq \ell^i$ for any bank i and any other solution (ℓ, α) to (1) and (2).*

The second channel of contagion in our model is credit contagion. A defaulting bank i does not repay its creditors in full, thereby imposing losses $L^i - p^i$ to the rest of the financial system. Creditors with large interbank exposures may thus default as a consequence of these losses and trigger a cascade of defaults. The cascade starts with the set $\mathcal{F} := \{i \mid c^i + \alpha_L \ell_L^i + (\pi L)^i < L^i + w^i\}$ of *fundamentally defaulting banks*, i.e., the set of banks that are unable to repay their liabilities even if every other bank repays its liabilities in full. If there are no fundamentally defaulting banks, then (L, ℓ_L, α_L) is the Pareto-efficient clearing equilibrium and all interbank liabilities are honored.

2.2 Coordination of Rescues

Financial stability can be ensured by providing distressed banks with sufficient liquidity to meet their liabilities. In a bailout, these injections are funded through taxpayer contributions, whereas in a bail-in, creditors of insolvent banks voluntarily take a haircut on their claims. As an intermediate option, we also consider assisted bail-ins, which include transfers by both banks and taxpayers.

Definition 2.2. An *assisted bail-in* (b, s) specifies, for each bank i , the contribution b^i to be made by i and the size s^i of the subsidy i receives. The government's contribution to the bail-in is $\sum_{i=1}^n (s^i - b^i)$, which is required to be non-negative.¹²

Observe that assisted bail-ins contain bailouts and privately backed bail-ins as special cases. As in any negotiation, the outcome depends crucially on the bargaining dynamics—which party gets to make offers at which point in time. We choose to model the negotiation as a three-stage interaction. First, the regulator makes a take-it-or-leave-it bail-in proposal to the banks. This allows us to capture the regulator's role as a coordinator as well as the fact that the government holds much of the bargaining power. Second, the banks simultaneously decide whether or not to participate in the proposed bail-in. Third, the regulator decides what to do if banks reject the proposal. This captures the fact that the government is the lender of last resort and cannot credibly commit to not bailing out banks if it is welfare-maximizing to do so.

In our model, the regulator knows the financial position of each bank and, thus, he can anticipate the banks' responses to any bail-in proposal. Therefore, he need not make a proposal that is not incentive compatible and the negotiation collapses into a single stage.¹³ In reality, the regulator is not fully informed and the coordination of a bail-in might take the form of a strategic bargaining game instead. Some banks might reject the regulator's initial proposal, after which the regulator revises his proposal to either exclude those banks or to accommodate them.¹⁴

¹²This formulation is equivalent to one, in which creditors write down their claims on insolvent banks such that b^i is the net debt forgiven by bank i and s^i is the sum of government injections and net debt forgiven to bank i by i 's creditors.

¹³Bail-ins are typically organized over short periods of time. For example, the critical negotiations leading to the bail-in of LTCM and to the takeover of Bear Stearns by JPMorgan Chase took place over the span of a weekend. It is thus reasonable to assume that there is no discounting between rounds of negotiation. Therefore, the solution to our model coincides with the Ståhl bargaining solution for any finite number of rounds of negotiation in which the regulator makes the last proposal. Indeed, in any such bargaining game, a bail-in can be implemented without the regulator's approval only if it is backed privately. Since a bank is willing to contribute only if it also has to contribute to the regulator's preferred proposal, the only privately-backed bail-in that can arise as a subgame Pareto-efficient equilibrium is, in fact, the regulator's preferred proposal.

¹⁴The original proposal made by the FRBNY for the rescue of LTCM involved a total of 16 of LTCM's creditors. However, Bear Stearns and Lehman Brothers later declined to participate. Upon the rejection of these two banks, the Fed adjusted its

Organization of a rescue. In our model, rescue are organized with the following steps:

1. The regulator proposes an assisted bail-in (b, s) .
2. The banks in $\mathcal{A}(b) := \{j \notin \mathcal{F} \mid b^j > 0\}$ are considered to be part of the negotiation and each bank $i \in \mathcal{A}(b)$ chooses an action $a^i \in \{0, 1\}$, indicating whether it agrees to contribute b^i .¹⁵
3. The regulator chooses his response r from the following three options:
 - (i) “*bail-in*”: Proceed with the proposed subsidies s , using taxpayer money to make up for the missing contributions. Cash holdings and financial commitments to outside parties of each bank i are then equal to $c^i(s) := c^i + s^i$ and $w^i(b, a) := w^i + b^i 1_{\{a^i=1\}}$, respectively, where 1_A is the indicator function that is equal to 1 if A is true and equal to 0 otherwise. The resulting financial system is cleared as in Section 2, where we denote the Pareto-dominant clearing equilibrium by $(\bar{p}(b, s, a), \bar{\ell}(b, s, a), \bar{\alpha}(b, s, a))$. Bank i ’s equity value is equal to $V^i(\bar{p}(b, s, a), \bar{\ell}(b, s, a), \bar{\alpha}(b, s, a))$. Welfare losses are obtained from (6) by additionally accounting for the social cost of government subsidies, that is,

$$W_\lambda(b, s, a) := W_\lambda(\bar{p}(b, s, a), \bar{\ell}(b, s, a), \bar{\alpha}(b, s, a)) + \lambda \sum_{i=1}^n (s^i - b^i 1_{\{a^i=1\}}).^{16} \quad (7)$$

- (ii) “*bailout*”: Resort to a public bailout $(0, \tilde{s})$ with subsidies $\tilde{s}$ decided by the regulator. Then, cash holdings of bank i are equal to $c^i(\tilde{s}) = c^i + \tilde{s}^i$ and we denote by $(\bar{p}(\tilde{s}), \bar{\ell}(\tilde{s}), \bar{\alpha}(\tilde{s}))$ the Pareto-dominant clearing equilibrium. Each bank i ’s equity value is $V^i(\bar{p}(\tilde{s}), \bar{\ell}(\tilde{s}), \bar{\alpha}(\tilde{s}))$ and welfare losses are denoted by $W_\lambda(\tilde{s}) := W_\lambda(0, \tilde{s}, 0)$.
 - (iii) “*no intervention*”: Abandon the rescue, which results in the default cascade described in Section 2. We denote by (p_N, ℓ_N, α_N) the Pareto-dominant clearing equilibrium and by W_N the welfare losses in the default cascade in the absence of intervention.

Remark 2.1. Some banks may be left with zero equity after they are bailed in or bailed out. Such banks cease to exist as a separate entity after the intervention and their bail-ins or bailouts should be understood as an orderly liquidation through takeovers by the bail-in contributors or the government.¹⁷ Because these banks have zero equity value after the intervention, we do not model how the assets are distributed among contributors.

Our solution concept is that of a subgame Pareto-efficient equilibrium, defined as follows.

proposal so that the contributions of Bear Stearns and Lehman Brothers were covered by the remaining 14 banks. The fact that these two banks decided not to cooperate shows that participation in the bail-in was at least partially voluntary.

¹⁵We assume that a bank i with $b^i = 0$ is simply not part of the negotiation and hence has no power to reject the proposal. For ease of notation, we write $a^i = 1$ for such a bank.

¹⁶We assume that the regulator assigns the same weight to taxpayer contributions as to losses of senior creditors. One way to interpret this is that the senior creditors are the depositors of the banks, who are protected by a deposit insurance scheme that is financed through taxpayer contributions. ~~This assumption is made for notational convenience and our~~ Our results would hold if the regulator used different weights.

¹⁷Examples of such takeovers from the global 2007–2009 financial crisis are plentiful. Among the most prominent ones are the takeovers of Bear Stearns and Merrill Lynch by JP Morgan Chase and Bank of America, respectively, or the federal takeovers of Fannie Mae and Freddie Mac.

Definition 2.3. A strategy profile (b, s, a, r) is *subgame Pareto efficient* if it is subgame perfect and after any proposal (b, s) , there is no other continuation equilibrium $(\tilde{a}, \tilde{r})$ of the accepting/rejecting subgame that Pareto dominates (a, r) for the banks in $\mathcal{A}(b)$ and the regulator.

This equilibrium concept is meant to capture that the coordination of a bail-in is a negotiation between the regulator and the banks in $\mathcal{A}(b)$. For example, during the bail-in of LTCM, Peter Fisher of the FRBNY sat down with representatives of LTCM’s creditors to find an appropriate solution; and it is implausible that they would have agreed on a bail-in that is Pareto dominated. Note that banks in the complement of $\mathcal{A}(b)$ are potentially worse off than in an alternative continuation equilibrium of the accepting/rejecting subgame because they are not part of the discussion.

Any proposal of the regulator admits an equilibrium response by the banks due to the following lemma. This result also implies existence of subgame Pareto-efficient equilibria in our model.

Lemma 2.3. *After any proposal (b, s) , the resulting accepting/rejecting subgame has a subgame Pareto-efficient continuation equilibrium (a, r) in pure strategies.*

Because subgame Pareto efficiency is a refinement of subgame perfection, it eliminates the non-credible threat by the regulator to abandon the rescue in the third stage if a public bailout leads to lower welfare losses than a default cascade in absence of intervention. This inability to commit to a no-intervention policy limits the regulator’s ability to incentivize banks to contribute, as we discuss in Section 3.2. Because the state of the financial system is common knowledge among the banks and the regulator in our model, we may assume without loss of generality that the regulator proposes only so-called feasible bail-ins, in which every bank can afford the proposed contribution.

Definition 2.4. A bail-in proposal (b, s) is *feasible* if $b^i = 0$ for any fundamentally defaulting bank $i \in \mathcal{F}$ and $L^i + w^i + b^i \leq c^i + s^i + \bar{\alpha}(b, s, 1)\bar{\ell}^i(b, s, 1) + \sum_{j=1}^n \pi^{ij}\bar{p}^j(b, s, 1)$ for any bank $i \notin \mathcal{F}$, where $1 = (1, \dots, 1)$ is the response vector of unanimous agreement.

While a bank i can refuse to make the proposed contribution b^i , in our model it cannot reject the subsidies s^i it is supposed to receive. Thus, if the regulator chooses option “*bail-in*,” the same subsidies are paid regardless of the banks’ responses. Since less taxpayer money flows into the financial system in $(b, s, 1)$ than in (b, s, a) for any $a \neq (1, \dots, 1)$, each bank is better off in the latter. Feasibility thus guarantees that each bank can afford the proposed contribution *in any* response vector.¹⁸

3 Incentives and Credibility of Intervention Plans

To highlight the primary economic forces at play, we focus on the case of *complete interventions* in this section, where the regulator considers only bailouts and bail-ins that rescue every bank in the system. Section 5 treats the more general case with partial interventions.

¹⁸We show in Lemma E.3 in the online appendix that for any given (b, s) , the payments $\bar{p}(b, s, a)$ and recovery rate $\bar{\alpha}(b, s, a)$ are weakly decreasing in a^i . A bail-in proposal (b, s) thus includes a guarantee that clearing payments and asset recovery are *at least* $\bar{p}(b, s, 1)$ and $\bar{\alpha}(b, s, 1)$ —backed by government contributions if necessary. Price guarantees have played a prominent role in government interventions such as, for example, in TARP or in the acquisition of Merrill Lynch by Bank of America.

3.1 Public Bailout

In a complete bailout, the regulator provides subsidies so that every bank repays its liabilities in full and the resulting clearing payment vector is L . The smallest provided subsidies are equal to the shortfall $s_L := (L + w - c - \alpha_L \ell_L - \pi L)^+$ after banks liquidate the maximal feasible amount ℓ_L , where

(ℓ_L, α_L) is given by Lemma 2.2. If asset liquidation is more costly than taxpayer contributions, the regulator will want to provide additional subsidies to cover the banks' shortfall before liquidation

$$s_0 := (L + w - c - \pi L)^+. \quad (8)$$

Note that subsidies s_L and s_0 support the clearing equilibria (L, ℓ_L, α_L) and $(L, 0, 1)$, respectively. The welfare-maximizing subsidies s in a complete bailout are such that the marginal losses from liquidation are as close to the marginal welfare cost λ of taxpayer contributions as possible, given the constraints $s_L^i \leq s^i \leq s_0^i$ which guarantee solvency of every bank in the system.

Lemma 3.1. *Suppose that the price elasticity γ is positive and let $g(\alpha) := ((1 + \lambda)\alpha - 1) \ln(\alpha)/\gamma$. In any complete bailout with subsidies $s^i \leq s_0^i$ for every bank i , welfare losses are equal to*

$$W_\lambda(s) = \lambda \sum_{i=1}^n s_0^i + g(\alpha). \quad (9)$$

~~Given that every bank is rescued~~, the regulator's choice of subsidies affects welfare only through the induced asset recovery rate. Any bailout that induces recovery rate α requires banks to liquidate an aggregate amount $-\ln(\alpha)/\gamma$ as seen in (1), reducing the required subsidies by the market value of those liquidated assets. Compared to subsidies s_0 , choosing subsidies that induce recovery rate α lowers welfare cost of subsidies by $-\lambda\alpha \ln(\alpha)/\gamma$ but increases liquidation losses by $-(1 - \alpha) \ln(\alpha)/\gamma$. The function g captures this welfare trade-off between liquidation costs and taxpayer contributions. It is strictly convex with the global minimum attained at the *indifference recovery rate* α_{ind} . The indifference recovery rate is decreasing in λ , but does not depend on the rate γ , ~~at which the buyers' efficiency decreases~~.

Lemma 3.2. *The asset recovery rate α_P in the welfare-maximizing complete bailout is 1 if $\gamma = 0$ and $\alpha_P = \max(\alpha_{\text{ind}}, \alpha_L)$ otherwise. A welfare-maximizing complete bailout awards subsidies s_L if $\gamma = 0$. Otherwise, it awards subsidies s with $s_L^i \leq s^i \leq s_0^i$ for every bank i such that*

$$\sum_{i=1}^n s^i = \sum_{i=1}^n s_0^i + \frac{\alpha_P \ln(\alpha_P)}{\gamma}. \quad (10)$$

We denote by W_P the corresponding welfare losses.

If $\alpha_P = \alpha_L$ or $\alpha_P = 1$, subsidies in the bailout are uniquely determined and equal to s_L and s_0 , respectively. However, subsidies are not uniquely determined if $\alpha_L < \alpha_P < 1$: since the asset recovery rate depends only on the total amount of liquidation, welfare depends only on the total subsidies awarded, but not on how those are distributed among banks. Note that the optimal complete bailout does not depend on the network structure because all defaults are prevented.

3.2 Credibility of the Regulator's Threat

If welfare losses in the public bailout of Lemma 3.2 are lower than in absence of intervention, banks know that ~~without their participation~~, the regulator's preferred option is a bailout. The regulator thus has no credible threats to punish recalcitrant banks ~~and, as a consequence, banks have no incentive to participate in a bail-in. The regulator then has no choice but to resort to a public bailout.~~ Consequently, the banks have no incentive to participate in any bail-in and the regulator has no choice but to resort to a public bailout.

Lemma 3.3. *If $W_P < W_N$, then the unique subgame Pareto-efficient equilibrium outcome is the public bailout described in Lemma 3.2, where W_P and W_N denote the welfare losses under the public bailout and no-intervention policies, respectively.*

We say that the regulator's no-intervention threat is credible if and only if $W_N \leq W_P$. We argue in the following three subsections that when the threat is credible, the regulator can incentivize banks to participate in a bail-in. In the remainder of this section, we show that the credibility of the threat is tightly linked to the amplification of the shock in absence of intervention.¹⁹

Since losses to senior creditors are not amplified through the financial system, we exclude those from consideration. The part of the initial shock that is amplified through the network is $S_0 := \sum_{i=1}^n s_0^i - \sum_{i \in \mathcal{D}(p_N, \ell_N, \alpha_N)} \delta^i(p_N, \alpha_N)$, i.e., the difference of aggregate shortfall, defined in (8), and losses to senior creditors. The losses incurred by junior creditors after the amplification are equal to the decrease in the market value of the banks $S_N := \sum_{i=1}^n (V^i(L, 0, 1) - V^i(p_N, \ell_N, \alpha_N))$, where $V^i(L, 0, 1)$ is the book value of bank i 's equity and $V^i(p_N, \ell_N, \alpha_N)$ is the value of i 's equity after clearing liabilities when there is no intervention.

Lemma 3.4. *The regulator's threat is credible if and only if*

$$S_N - S_0 \leq \lambda S_0 - \sum_{i=1}^n \min(e^i, s_0^i) + g(\alpha_P). \quad (11)$$

The credibility threshold depends on the welfare cost λ of taxpayer contributions, the size of the initial shortfall, and the distribution of such shortfall across the system captured by the second and third term on the right-hand side of (11). The second term is a measure of the amount of illiquid assets that can be used to absorb the initial shock: the larger this amount is, the more credible is the threat. The last term in (11) captures the trade-off between liquidation costs and taxpayer contributions in the optimal bailout for a given distribution of shocks s_0 . While this trade-off is minimized at the indifference recovery rate α_{ind} , the regulator may not be able to attain α_{ind} . The difference $g(\alpha_P) - g(\alpha_{\text{ind}})$ is thus a measure of how close to attaining α_{ind} the regulator can tailor a bailout for a given distribution of shocks s_0 ; see also the discussion after Lemma 3.1.

Lemma 3.4 establishes a link between the credibility of the no-intervention threat and existing literature on financial networks without intervention, which often ranks the desirability of network structures according to the welfare loss criterion $S_N - S_0$. The studies of Allen and Gale (2000) and

¹⁹Note that the credibility of the threat is a function of exogenous variables: the welfare losses W_P in the optimal bailout is the result of a minimization problem solved by the regulator alone and W_N are the welfare losses in absence of any action.

Acemoglu, Ozdaglar and Tahbaz-Salehi (2015) show that dense connections between banks may serve as an amplifier for large initial shocks. For most sizes of the initial shock, the right-hand side of (11) does not depend on the network structure.²⁰ Therefore, Lemma 3.4 indicates that dense connections are detrimental to the credibility of the threat when the initial shock is large.

3.3 Regulator's Response in the Last Stage

A consequence of Lemma 3.2 is that, for a given bail-in proposal (b, s) and a response vector a by the banks, the regulator's best response in stage 3 is

$$r(b, s, a) = \begin{cases} \text{"no intervention"} & \text{if } W_N \leq \min(W_\lambda(b, s, a), W_P), \\ \text{"bailout"} & \text{if } W_P < \min(W_N, W_\lambda(b, s, a)), \\ \text{"bail-in"} & \text{otherwise.} \end{cases} \quad (12)$$

The regulator chooses the action that minimizes welfare losses if such an action is unique. Ties are broken according to *"no intervention"* $\succ$ *"bail-in"* $\succ$ *"bailout"* so that (a) taxpayer money is used only if it is strictly welfare increasing and (b) unilateral deviations by banks in stage 2 can be discouraged when $W_N = W_\lambda(b, s, a)$; see Lemma 3.5 and Footnote 22 below for details.

3.4 Banks' Equilibrium Responses

In this section, we analyze how the banks respond to a given proposal when the regulator's threat is credible. A crucial feature of the banks' response vector is whether a sufficient proportion of banks accepts the proposal for the regulator to implement the residual bail-in.

Definition 3.1. Given a bail-in proposal (b, s) , a subgame-perfect equilibrium (a, r) of the continuation game is an *accepting equilibrium* if $r(b, s, a) = \text{"bail-in"}$ and it is a *rejecting equilibrium* otherwise. The banks' response vector a in (a, r) is an *accepting/rejecting equilibrium response*.

For a bank to participate, the bail-in has to be both feasible and incentive-compatible. A complete bail-in is feasible if the net contribution by any bank i does not exceed its budget constraint

$$\eta^i(\alpha, \ell^i) := (c^i + \alpha \ell^i + (\pi L)^i - w^i - L^i)^+, \quad (13)$$

given recovery rate α and liquidation decision ℓ^i .²¹ The incentive-compatibility conditions in a feasible bail-in proposal are stated in the following lemma.

Lemma 3.5. *Suppose the threat is credible, i.e., $W_N \leq W_P$. Let (b, s) be a feasible proposal of a complete bail-in. In any accepting equilibrium response a , bank i with $b^i > 0$ accepts if and only if:*

1. *Welfare losses in the residual bail-in without i satisfy $W_\lambda(b, s, (0, a^{-i})) \geq W_N$, and*

²⁰Only for very large shocks, for which some banks do not repay anything in a clearing equilibrium, the right-hand side depends on the network through losses to senior creditors δ . Most existing literature does not consider shock sizes this large.

²¹This condition is recovered as a special case of Definition 2.4 for complete rescues. The model can easily be adapted to allow for capital requirements of solvent banks by subtracting these requirements from the budget constraint in (13).

2. Bank i 's net contribution to (b, s) satisfies $b^i - s^i \leq b_*^i(\bar{\alpha}(b, s, (1, a^{-i})))$, where

$$b_*^i(\alpha) := \max_{\ell^i \in [0, e^i]} \min \left(\sum_{j=1}^n \pi^{ij} (L^j - p_N^j) + (1 - \alpha_N) \ell_N^i - (1 - \alpha) \ell^i, \eta^i(\alpha, \ell^i) \right). \quad (14)$$

Here, $(\tilde{a}^i, a^{-i})$ indicates the response vector, in which i responds with $\tilde{a}^i$ and $j \neq i$ responds with a^j .

The first condition states that there is no possibility for free-riding: if bank i were to reject the proposal, the regulator would choose not to intervene rather than pay for i 's contribution with taxpayer money.²² In other words, the set of banks which accept the proposal is minimal in any accepting equilibrium response. The second condition states that, in order to prevent a default cascade, bank i is willing to make a net contribution up to its exposure to the default cascade through both channels of contagion. While the second condition is not entirely explicit, Lemma E.3 in the online appendix implies that $b_*^i(\bar{\alpha}(b, s, (1, a^{-i})))$ is non-increasing in $b^i - s^i$, hence the incentive-constraint satisfies a threshold property.

The requirement that equilibria be subgame Pareto-efficient implies that banks will accept an incentive-compatible bail-in proposal. While, in general, accepting equilibria need not be unique, the regulator can preempt any coordination problems by altering the proposed bail-in so that it is incentive compatible for only one consortium of banks to accept the proposal. If the regulator requests zero contributions from any bank outside the selected consortium, unanimous acceptance becomes the unique accepting equilibrium, and hence the unique subgame Pareto-efficient equilibrium, of the revised proposal. We formalize this discussion in Lemmas B.6 and B.7.

3.5 Optimal Proposal by the Regulator

Contributions of banks to a bail-in affect welfare in two ways. First, they reduce the amount of taxpayer contributions needed. Second, if the asset recovery rate in the optimal bailout is higher than the indifference recovery rate, the regulator can use the contributions of banks to enhance welfare by exploiting the trade-off between asset liquidation and taxpayer contributions.

Lemma 3.6. *Let $b_0 := (c + \pi L - w - L)^+$ be the largest feasible contributions without asset liquidation and let g and s_0 be defined as in Lemma 3.1 and (8), respectively. Let (b, s) be a complete feasible bail-in proposal with $b^i s^i = 0$ for every bank i .²³ For any response vector a , welfare losses are equal to*

$$W_\lambda(b, s, a) = W_P + g(\bar{\alpha}(b, s, a)) - g(\alpha_P) + \lambda \sum_{i=1}^n (s^i - s_0^i)^+ - \lambda \sum_{i=1}^n \min(b^i, b_0^i) 1_{\{a^i=1\}}. \quad (15)$$

Equation (15) shows how welfare losses in a bail-in compare to welfare losses in the optimal bailout of Lemma 3.2. A contribution of bank i up to the amount b_0^i does not require asset liquidation, hence it does not impact the asset recovery rate. Thus, each dollar contributed up to b_0^i

²²The regulator prefers “no intervention” over “bail-in” in (12) when they lead to the same welfare losses so that a rejection by bank i can be prevented if the welfare losses $W_\lambda(b, s, (0, a^{-i}))$ in the residual bail-in without i are equal to W_N .

²³The restriction $b^i s^i = 0$ imposes that bank i either receives subsidies or makes contributions to the bail-in, but not both. We show in Lemma C.1 in the appendix that this comes without loss of generality.

improves welfare by λ . Contributions in excess of b_0^i require asset liquidation by the bank, thereby impacting the recovery rate α and the welfare trade-off g . Finally, subsidies beyond bank i 's short-fall s_0^i do not reduce losses from misallocation of the asset because banks can fulfill all obligations without liquidating assets. Each dollar of subsidies awarded in excess of s_0 thus effectively *burns* λ units of welfare. While this generally constitutes a decrease in welfare, we illustrate below how welfare burning can be used by the regulator to deter banks from free-riding.

Next, we analyze how the regulator best implements a rescue plan that satisfies the incentive-compatibility conditions of Lemma 3.5. The no-free-riding constraint in Condition 1 requires that after the rejection by any bank, welfare losses in the residual bail-in are larger than welfare losses without intervention. Using (15), this is equivalent to requiring that for any participating bank i ,

$$W_\lambda(b, s, a) \geq W_N + g(\bar{\alpha}(b, s, a)) - g(\bar{\alpha}(b, s, (0, a^{-i}))) - \lambda \min(b^i, b_0^i). \quad (16)$$

Equation (16) constitutes a lower bound on attainable welfare losses imposed by the no free-riding constraint. It states that welfare losses in an incentive-compatible bail-in cannot be lowered from W_N by more than the welfare impact stemming from the contribution of any participating bank. The no-free-riding constraint thus drives the regulator to include banks, which have a potential for large contributions, or banks whose contributions enhance the welfare trade-off between subsidies and asset liquidation. By choosing an incentive-compatible proposal (b, s) , the regulator implicitly chooses an associated vector of liquidation decisions $\bar{\ell}(b, s, 1)$ and an asset recovery rate $\bar{\alpha}(b, s, 1)$.²⁴ In order to construct a bail-in consortium $\mathcal{C}$ with maximal contributions by its participating banks, the regulator should choose the maximal feasible contribution $\eta^i(\alpha, \ell^i)$ by each bank $i \in \mathcal{C}$, defined in (13), given the desired asset recovery rate α and consistent vector of liquidation decisions ℓ . This contribution is incentive compatible for ℓ^i up to some value $\ell_*^i(\alpha)$, defined as the maximizer in (14). Viewed as a function of α and ℓ , the necessary subsidies to guarantee solvency of every bank are

$$s(\alpha, \ell) := (L + w - c - \alpha\ell - \pi L)^+. \quad (17)$$

For a bail-in with subsidies $s(\alpha, \ell)$ and contributions $\eta(\alpha, \ell)$ for a specific choice of $\mathcal{C}$, α , and a consistent vector ℓ , each bank $i \in \mathcal{C}$ contributes at least $\underline{b}^i := \min((\pi(L - p_N))^i + (1 - \alpha_N)\ell_N^i, b_0^i)$, the maximal incentive-compatible contribution without asset liquidation, and at most $b_*^i(\alpha)$, defined in Lemma 3.5. The following lemma shows that such bail-ins maximize welfare among all bail-ins that are *individually incentive compatible*, i.e., all bail-ins that satisfy Condition 2 of Lemma 3.5.

Lemma 3.7. *Let (b, s) be a complete feasible bail-in such that the response vector $1 = (1, \dots, 1)$ satisfies Condition 2 of Lemma 3.5 for every bank. Denote by $\mathcal{C} := \{i \mid b^i > 0\}$ the set of contributing banks and abbreviate $\alpha = \bar{\alpha}(b, s, 1)$. Welfare losses in this bail-in satisfy*

$$W_\lambda(b, s, 1) \geq W_P - g(\alpha_P) + g(\alpha) - \lambda \sum_{i \in \mathcal{C}} \underline{b}^i. \quad (18)$$

²⁴By Lemma B.6, the regulator can aim to construct bail-ins that can be accepted by all banks without loss of generality.

Equality holds if and only if $b^i - s^i \geq \underline{b}$ for every $i \in \mathcal{C}$ and $s^i \leq s_0^i$ for every $i \notin \mathcal{C}$.

Equation (18) shows that when contributions are of size $\eta(\alpha, \ell)$, their welfare impact depends on the liquidation decision only through the induced asset recovery rate. Thus, similarly to the bailout, the regulator optimizes bail-ins to induce the asset recovery rate, at which he is indifferent between additional taxpayer contributions and asset liquidation.

Suppose now that for a bail-in of the above form, some bank $i \in \mathcal{C}$ has an incentive to free-ride, that is, the regulator proceeds with the residual bail-in even without i 's participation. By Condition 1 in Lemma 3.5, this occurs precisely if $W_N - W_\lambda(b, s, (0, a^{-i})) \geq 0$. If the regulator were able to decrease welfare in the proposed bail-in by this amount across all response vectors, he could eliminate i 's free-riding incentives. It follows from Lemmas 3.6 and 3.7 that welfare in an individually incentive-compatible bail-in (b, s) with contributing banks in $\mathcal{C}$ exceeds the lower bound in (18) by

$$\lambda \sum_{i \notin \mathcal{C}} (s^i - s_0^i)^+ + \lambda \sum_{i \in \mathcal{C}} (\underline{b}^i - b^i)^+. \quad (19)$$

Thus, providing subsidies in excess of s_0 and requesting contributions below $\underline{b}$ are means with which the regulator can decrease or “burn” welfare. Since subsidies in excess of s_0 and contributions below $\underline{b}$ do not affect the asset recovery rate, burning welfare as in (19) does not distort incentives and hence decreases welfare by a constant amount across the banks' responses. We denote by $\chi_{\mathcal{C}}(\alpha)$ the minimal amount of welfare-burning needed to eliminate free-riding incentives from an individually incentive-compatible bail-in with contributing banks $\mathcal{C}$ that induces recovery rate α . The mathematical definition of $\chi_{\mathcal{C}}(\alpha)$ is somewhat convoluted and deferred to Lemma A.1. In Theorem D.2 of the online appendix, we show that welfare burning is used sparingly in equilibrium.

The analysis above brings us to the characterization of the equilibrium intervention plan. The result states that when the threat is credible, the regulator proposes a bail-in which implements the minimum value burning $\chi_{\mathcal{C}}(\alpha)$ for the optimal choice of $\mathcal{C}$ and α . For the sake of reference, we isolate the set of all incentive compatible bail-ins that implement the minimum value burning.

Definition 3.2. Let $z(\alpha) := \alpha \ln(\alpha)$ and let z^{-1} be its inverse on the interval $[\frac{1}{e}, 1]$. The function $g_\alpha(x) := g(z^{-1}(z(\alpha) + \gamma x)) - g(\alpha)$ captures the welfare impact of liquidating x fewer units of the asset at recovery rate α . Let $\Xi(\mathcal{C}, \alpha)$ denote the set of all bail-ins (b, s) satisfying:

- (i) $b^i - s^i \leq b_*^i(\alpha)$ for every $i \in \mathcal{C}$,
- (ii) $b^i = 0$ and $s^i(\alpha, e) \leq s^i$ for every $i \notin \mathcal{C}$,
- (iii) $\sum_{i=1}^n (s_0^i - s^i)^+ + \sum_{i \in \mathcal{C}} (b^i - b_0^i)^+ = -\frac{\alpha \ln(\alpha)}{\gamma}$,
- (iv) $\lambda \sum_{i=1}^n (s^i - s_0^i)^+ + \lambda \sum_{i \in \mathcal{C}} (\underline{b}^i - b^i)^+ = \chi_{\mathcal{C}}(\alpha)$,
- (v) $\lambda \min(b^i, b_0^i) + g_\alpha((b^i - b_0^i)^+) \geq W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C}} \underline{b}^j - \chi_{\mathcal{C}}(\alpha)$ for $i \in \mathcal{C}$.

The first two conditions state that (b, s) is a feasible, complete bail-in with contributing banks in $\mathcal{C}$. Conditions (iii) and (iv) state that the bail-in induces asset recovery rate α and the total amount of welfare burnt is $\chi_{\mathcal{C}}(\alpha)$. To understand Condition (v), recall that a contribution up to $\underline{b}^i$ does not require asset liquidation. A total contribution of size $\underline{b}^i + x^i$ by bank i thus has a total impact on welfare of $\lambda \underline{b}^i + g_{\alpha}(x^i)$. Thus, Condition (v) states that there is no free-riding because the welfare impact of a deviation by bank i (left hand side) is larger than the difference between the welfare losses without intervention and in the bail-in. Conditions (i) and (v) together imply that it is incentive-compatible for any bank $i \in \mathcal{C}$ to accept the proposal.

Theorem 3.8. *For any bail-in (b, s) , let $\ell(b, s)$ denote the vector of liquidated assets if the proposal is accepted by all banks. For any ℓ , let $i_1(\ell), i_2(\ell), \dots$ denote a decreasing order of banks according to $\eta^i(\alpha(\ell), \ell)$. Let $\mathcal{C}(\ell) := \{i_1(\ell), \dots, i_{m(\ell)}(\ell)\}$, where $m(\ell)$ denotes the smallest integer k such that*

$$W_P + (g(\alpha(\ell)) - g(\alpha_P)) - \lambda \sum_{j=1}^k \eta^{i_j(\ell)}(\alpha(\ell), \ell) < W_N.$$

If $W_P < W_N$, then any subgame Pareto-efficient equilibrium outcome is a public bailout with welfare losses W_P as specified by Lemma 3.2. If $W_P \geq W_N$, then there exist generically unique $\mathcal{C}_$ and α_* such that in any subgame Pareto-efficient equilibrium, a bail-in from the set $\Xi(\mathcal{C}_*, \alpha_*)$ is proposed by the regulator and accepted by all banks.²⁵ Welfare losses are equal to*

$$W_E = W_P + (g(\alpha_*) - g(\alpha_P)) - \lambda \sum_{i \in \mathcal{C}_*} \underline{b}^i + \chi_{\mathcal{C}_*}(\alpha_*).$$

Finally, if $\mathcal{C}_$ is unique, then $\mathcal{C}_* = \mathcal{C}(\ell(b, s))$ for all $(b, s) \in \Xi(\mathcal{C}_*, \alpha_*)$.*

As we have highlighted before, a bail-in can be organized in equilibrium if and only if the regulator's no-intervention threat is credible. The set $\mathcal{C}_*$ consists of banks that are most exposed to contagion at the equilibrium recovery rate α_* : after choosing a set of liquidation decisions ℓ that induces α_* , the regulator adds banks into the bail-in consortium in decreasing order of their incentive compatible contributions $\eta(\alpha(\ell), \ell)$ until welfare losses in the bail-in are lower than in the default cascade without intervention. This occurs after adding the $m(\ell)$ most exposed banks. Because of the no-free-riding constraint of Lemma 3.5, no more contributors can be added after that: any additional bank would know that even without its contribution, the regulator will proceed with the residual consortium, hence that bank has no incentive to participate.

The recovery rate and the welfare losses are generically unique in equilibrium. The set of liquidation decisions and the bail-in proposal, however, are not unique in general. Similarly to the public bailout, welfare depends on the liquidation by banks only through the total amount that is being liquidated.²⁶ This gives the regulator some leeway on how to induce recovery rate α_* .

²⁵Generic uniqueness is up to banks i with $\underline{b}_0^i = 0$ because those banks affect welfare only through the welfare trade-off captured in g . Generically unique means that it is unique for an open and dense set of model parameters.

²⁶The only restriction on liquidation by an individual bank is the fifth condition in Definition 3.2, specifying the minimal liquidation amount by bank i for the no-free-riding condition to hold.

Because banks are willing to make larger contributions to a bail-in that guarantees a higher asset recovery rate, the regulator's indifference recovery rate increases in a bail-in, that is, $\alpha_* \geq \alpha_{\text{ind}}$.

In Section D of the online appendix, we highlight the relationship between the equilibrium recovery rate and the amount of welfare burnt in equilibrium: in many situations, the regulator will avoid burning welfare and instead choose to induce a recovery rate, at which the contributions by individual banks are sufficiently large to deter free-riding. Nevertheless, there are scenarios in which welfare is burnt in equilibrium. One such scenario occurs when buyers of liquidated assets are fully efficient, that is, when $\gamma = 0$; see Section 5.

Remark 3.1. We briefly discuss the default resolution outcome if the government had various degrees of commitment power or if a bank had made a pre-commitment to participate in a bail-in.

1. Suppose that the regulator had the power to commit to a no-bailout policy in the third stage. Then the bail-in described in Theorem 3.8 will be organized in equilibrium even if $W_P < W_N$, hence equilibrium welfare losses decrease to W_E in this case.
2. Suppose that the regulator has full commitment power in the third stage, that is, he can commit not only to a no-bailout policy, but also to not proceeding with a residual bail-in if some banks reject the proposal. Then the no-free-riding constraint (Condition 1 in Lemma 3.5) has to be satisfied only if the proposed bail-in does not include any government subsidies. Note that a bail-in still has to be individually incentive compatible (Condition 2 of Lemma 3.5) as, otherwise, a bank would prefer the default cascade over the bail-in. This reflects the fact that in a developed democracy, the government cannot legally appropriate the banks' capital.²⁷ We conclude that the regulator can force his preferred individually incentive-compatible assisted bail-in unless the banks can agree on a private bail-in alternative.²⁸
3. Finally, one could easily adapt the model to allow pre-commitments by banks to participate in a bail-in, which incurs a legal fee C^i on bank i if the bank reneges on its promise. In that case, the legal fee C^i is simply added to the left expression in the minimum in (14). The takeover of Merrill Lynch by Bank of America (BOA) could, perhaps, be interpreted as an assisted bail-in with pre-commitment by BOA.²⁹

4 Shocks, Asset Illiquidity, and Total Throughput

In this section, we analyze the dependence of equilibrium quantities, including asset recovery rate, awarded subsidies, credibility, and welfare losses, on the banks' balance sheet parameters, the

²⁷Without Condition 2 of Lemma 3.5, the government could demand arbitrarily high contributions.

²⁸Formally, such a model would require additional stages of negotiation to give banks the opportunity to coordinate on a private bail-in alternative. It is easy enough to adapt the results in Section 3.4 (by excluding the regulator's incentives) to show that no two private bail-ins are Pareto comparable. Therefore, we expect that the subgame Pareto-efficient equilibrium outcome is a constrained-efficient private bail-in or the regulator's preferred individually incentive-compatible assisted bail-in.

²⁹When BOA wanted to withdraw from the purchase of Merrill Lynch, the Federal Reserve could have claimed that BOA's initial interest stalled a successful default resolution, thereby generating additional losses. To avoid extensive litigation with an uncertain outcome, both parties agreed on a settlement, in which BOA proceeded to purchase the assets, but the price was lowered and backed with a government guarantee. The resulting deal was, presumably, close to incentive compatible for BOA.

network structure, and the recovery rates from asset liquidation and bankruptcy procedures. The results are presented under the assumptions that $\gamma > 0$, $e^i > 0$ for every bank i , and that there is at least one fundamentally defaulting bank. Without this assumption, the results in this section hold when strict monotonicity is replaced with weak monotonicity.

4.1 Optimal Bailout

Since every bank is rescued in a complete bailout, there are no bankruptcy costs and no losses that depend on the network structure. Asset recovery rate, subsidies, and welfare losses in the complete bailout are thus independent of β and π . The dependence of the welfare-maximizing bailout on the rate γ at which the pool of efficient buyers diminishes, the welfare cost λ of a taxpayer dollar, and the size s_0 of the banks' shortfall is given in the following result and illustrated in Figure 2.

Lemma 4.1. *There exist (possibly infinite) thresholds $\gamma_*, \lambda_* > 0$, as well as finite thresholds $s_*^i, e_*^i \geq 0$ for every bank i such that the following conditions hold:³⁰*

- (i) *The recovery rate α_P is decreasing for $\gamma \leq \gamma_*$ and it is constant for $\gamma \geq \gamma_*$. Subsidies and welfare losses are increasing in γ .*
- (ii) *The recovery rate α_P and subsidies are decreasing for $\lambda \leq \lambda_*$ and they are constant for $\lambda \geq \lambda_*$. Welfare losses are increasing in λ .*
- (iii) *For any bank i , the recovery rate α_P , subsidies, and welfare losses in the optimal bailout are decreasing for $e^i \leq e_*^i$ and constant for $e^i > e_*^i$.*
- (iv) *For any bank i , the recovery rate α_P is decreasing for $s_0^i \leq s_*^i$ and constant for $s_0^i > s_*^i$. Subsidies and welfare losses are increasing in s_0^i .*

If the marginal efficiency of buyers decreases slowly, i.e., γ is low, liquidation has a small impact on the asset recovery rate and it is welfare maximizing for the regulator to provide only the minimal amount of subsidies. Under this minimal-intervention policy, the recovery rate of the asset falls as γ increases. The resulting decrease of the banks' equity value requires larger subsidies to restore the system to a going concern. If the marginal efficiency of buyers decreases sufficiently quickly, that is, $\gamma \geq \gamma_*$, the regulator switches from the minimal-intervention policy to a policy that trades off asset liquidation and taxpayer contributions to maintain the indifference recovery rate α_{ind} . As γ increases and liquidation becomes more costly, the regulator has to provide additional subsidies to maintain the indifference recovery rate, further increasing welfare losses.

If the welfare cost λ of a taxpayer dollar is below the threshold λ_* , the regulator balances taxpayer contributions and asset liquidation to maintain the indifference recovery rate. As taxpayer contributions become more costly in terms of welfare, the regulator decreases the size of subsidies until, at level λ_* , he provides only the minimal subsidies necessary to guarantee solvency.

³⁰The thresholds depend on the other model parameters, that is, γ_* depends on λ , s_0 , and e , etc.

³¹Unlike the calibration done in Section 6, buyers of liquidated assets are not assumed to be fully efficient in the calibration that generates Figures 2 and 4. In this calibration, we assume that 90% of the outside assets reported by banks are illiquid and the remaining 10% are perfectly liquid, i.e., cash.

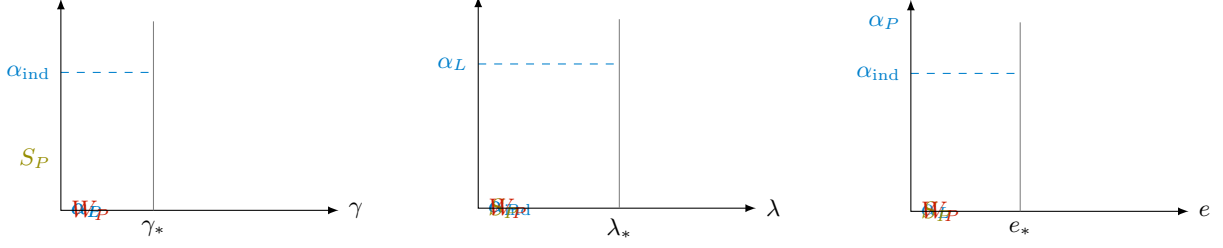


Figure 2: The three charts show the dependence of the optimal bailout on γ , λ , and e for a calibrated model of a dense financial network using the 2018 data of the EBA stress test.³¹ The asset recovery rate $\alpha_P = \max(\alpha_L, \alpha_{\text{ind}})$ is shown in solid lines; the minimal-intervention recovery rate α_L and the indifference recovery rate α_{ind} are shown in dashed lines. The total subsidies S_P awarded and welfare losses W_P are normalized to fit the same scale.

A fundamentally defaulting bank i cannot cover its shortfall by liquidating its assets. If the amount e^i of illiquid asset held by bank i is small, the regulator is thus forced to cover a large portion of the shortfall even if the marginal welfare cost of taxpayer contributions is higher than the marginal welfare impact of asset liquidation. As e^i increases, bank i is able to cover a larger portion of its shortfall by liquidating its assets, which reduces the size of the minimal subsidies required. Below the threshold e_*^i , the benefits of reducing subsidies outweigh additional losses from fire sales, leading to an overall decrease in subsidies and welfare losses.

Finally, both subsidies and welfare losses are increasing in the shortfall s_0 of the banks, which can be understood as a measure of the initial shock size. The asset recovery rate is strictly decreasing in the size of the shock where it exceeds the indifference recovery rate and it is constant otherwise.

4.2 Credibility

For a given set of parameters, the threat is either credible or not. To analyze how the credibility depends on underlying variables, we study the difference $W_P - W_N$ between the welfare losses in the optimal bailout and in absence of intervention. This measures how close to being credible the threat is. We say that the credibility of the threat is increasing or decreasing in a parameter if $W_P - W_N$ is, respectively, increasing and decreasing in that parameter.

A critical measure for the sensitivity analysis is the *total throughput* of a defaulting bank to the solvent members of the economy (i.e., banks and senior creditors). Abbreviate $\mathcal{D}_N = \mathcal{D}(p_N, \ell_N, \alpha_N)$, let $\mathcal{C}_N \subseteq \mathcal{D}_N$ denote the set of defaulting banks which repay their senior creditors in full, and let $\mathcal{I}_N$ denote the set of illiquid but solvent banks, all in absence of intervention. For two sets of banks $\mathcal{S}$ and $\mathcal{C}$, let $\pi^{\mathcal{S}, \mathcal{C}}$ denote the submatrix of π with rows and columns corresponding to banks in $\mathcal{S}$ and $\mathcal{C}$, respectively. The *throughput* of a bank $i \in \mathcal{C}_N$ to a set of banks $\mathcal{S}$ is

$$\theta_{\mathcal{S}}^i(\beta, \pi) := \sum_{j \in \mathcal{S} \setminus \mathcal{D}_N} \pi^{\{j\} \mathcal{C}_N} (I - \beta \pi^{\mathcal{C}_N, \mathcal{C}_N})^{-1} \rho_i^{\mathcal{C}_N} + \beta \sum_{j \in \mathcal{S} \cap \mathcal{D}_N \setminus \mathcal{C}_N} \pi^{\{j\} \mathcal{C}_N} (I - \beta \pi^{\mathcal{C}_N, \mathcal{C}_N})^{-1} \rho_i^{\mathcal{C}_N}. \quad (20)$$

where $\rho_i^{\mathcal{C}_N}$ denotes the unit vector in $\mathbb{R}^{\mathcal{C}_N}$ in the direction of i . The *total throughput* of bank $i \in \mathcal{C}_N$ is then defined as $\theta^i(\beta, \pi) := \theta_{\{1, \dots, N\}}^i(\beta, \pi)$. For a bank $i \in \mathcal{D}_N \setminus \mathcal{C}_N$, we set $\theta_{\mathcal{S}}^i(\beta, \pi) := 1$ for any set of banks $\mathcal{S}$ that contains i and we set $\theta_{\mathcal{S}}^i(\beta, \pi) := 0$ otherwise.

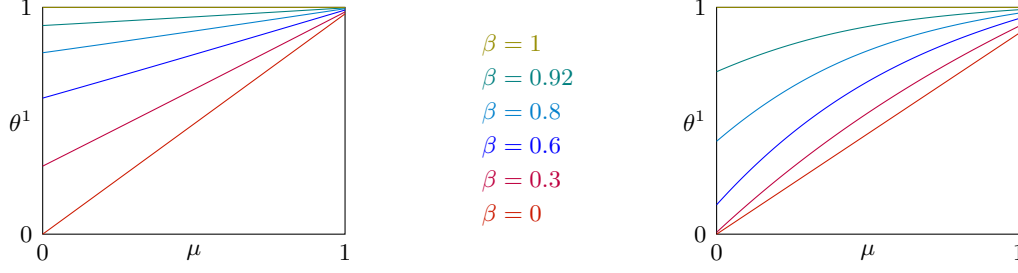


Figure 3: Let π_c and π_r denote the complete and the ring network, respectively, in a network of $n = 36$ banks. The two charts show the total throughput $\theta^1(\beta, \pi_\mu)$ of bank 1 in the network $\pi_\mu = \mu\pi_c + (1-\mu)\pi_r$ for various levels of β , when $\mathcal{C}_N = \mathcal{D}_N = \{1, 2\}$ and $\mathcal{C}_N = \mathcal{D}_N = \{1, 2, 3, 4, 5\}$ in the left and right panel, respectively.

The total throughput of bank i measures the exposure of solvent junior creditors (first term) and senior creditors (second term) to a shock hitting bank i . It quantifies the potential for spill-over losses triggered by defaults. For a bank $i \in \mathcal{C}_N$, the quantity $(I - \beta\pi^{\mathcal{C}_N, \mathcal{C}_N})^{-1} = \sum_{k=0}^{\infty} (\beta\pi^{\mathcal{C}_N, \mathcal{C}_N})^k$ captures the amplification of losses due to feedback effects between defaulting banks: term k in the sum corresponds to the propagation of losses through liability chains in $\mathcal{C}_N$ of length k . After accounting for bankruptcy losses, the exposure of a solvent creditor to a shock on bank i 's assets is π^{ji} for a solvent bank j and $\beta\pi^{ji}$ for the senior creditors of a bank $j \in \mathcal{D}_N \setminus \mathcal{C}_N$.³² The following lemma shows that the total throughput is a normalized measure for the rate of spill-over losses, which condenses all network information needed to determine the credibility of the regulator's threat.

Lemma 4.2. *The total throughput of any bank is non-decreasing in β and it takes values in $[0, 1]$. Conditional on the banks' levels of solvency (the sets $\mathcal{D}_N$, $\mathcal{C}_N$, and $\mathcal{I}_N$) and the total value of their claims on solvent banks, $W_P - W_N$ depends on π only through $\sum_{i \in \mathcal{C}_N} \theta_{\mathcal{I}_N}^i(\beta, \pi)$ and $\sum_{i \in \mathcal{C}_N} \theta^i(\beta, \pi)$.*

Observe from (20) that the throughput depends on the network structure, the location of the shocked bank(s) within the network, the connections of the shocked banks to other defaulting banks, as well as the recovery rate β . Conditional on the banks' levels of solvency, it does not, however, depend on the asset recovery rate or the banks' balance sheet quantities L , c , w , and e . The throughput is increasing in the connectivity between defaulting banks as illustrated in Figure 3.³³

Lemma 4.3.

- (i) *The credibility of the threat is increasing in λ .*
- (ii) *The credibility of the threat is non-decreasing in β and it is strictly increasing on the set $\{\beta \mid \text{there exists } i \in \mathcal{D}_N(\beta) \text{ with } \theta^i(\beta, \pi) > 0\}$.*

³²The total throughput of bank $i \in \mathcal{C}_N$ is related to the bank's Bonacich centrality $B^i = \mathbf{1}_{\mathcal{C}_N}^\top (I - \beta\pi^{\mathcal{C}_N, \mathcal{C}_N})^{-1} \rho_i^{\mathcal{C}_N}$, which captures the total amplification of losses through feedback loops in $\mathcal{C}_N$. The total throughput additionally takes into account how the losses are distributed among the creditors. It is important to note that the Bonacich centrality may diverge to ∞ if $\beta \rightarrow 1$, whereas our notion of total throughput is bounded on the interval $[0, 1]$.

³³The symmetric complete network π_c with $\pi_c^{ij} = 1_{\{i \neq j\}}/(n-1)$ is the most diversified network structure. The ring network π_r with $\pi_r^{ij} = 1_{\{j=i+1 \bmod n\}}$ is the sparsest network structure as measured by the Gini index; see Hurley and Rickard (2009).

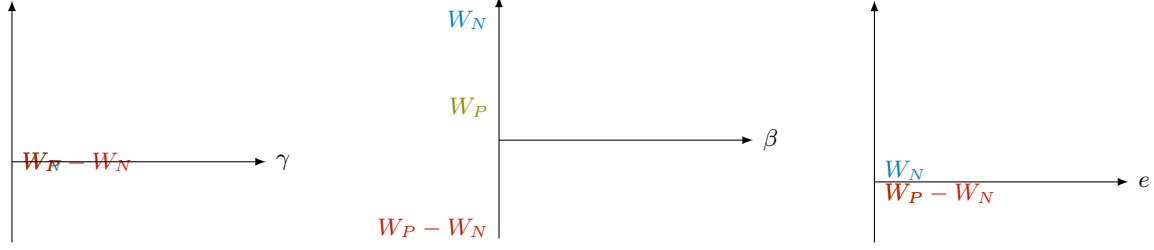


Figure 4: The three charts show the dependence of the optimal bailout on γ , β , and e for a model of a dense financial network calibrated to the 2018 data of the EBA stress test. In the right chart, we scale the size of illiquid assets held by all banks simultaneously, leading to a cluster of thresholds e_*^i , where each bank i becomes solvent and the credibility of the threat increases.

- (iii) *The credibility of the threat is non-monotonic in γ . However, all discontinuities of $W_P - W_N$ with respect to γ are downward discontinuities and the marginal change of $W_P - W_N$ with respect to γ at continuity points is decreasing in β and $\theta^i(\beta, \pi)$ for every defaulting bank i .*
- (iv) *For each i and fixed e^{-i} , there exists e_*^i such that the credibility of the threat is decreasing at all $e^i < e_*^i$, for which $\beta(\theta^i(\beta, \pi) + \lambda \theta_{\mathcal{D}_N}^i(\beta, \pi)) \leq 1$, as well as all $e^i > e_*^i$. At the threshold e_*^i , where bank i becomes solvent, $W_P - W_N$ has an upward discontinuity.*

As the welfare cost λ of taxpayer contribution increases, a bailout becomes more costly and the threat becomes more credible. Similarly, as the recovery rate β increases, bankruptcies become less costly and it becomes more credible that the regulator will not intervene. If the throughput of all defaulting banks is 0, then losses generated by defaulting banks do not spill over to the rest of the system, hence welfare losses are locally constant in β .³⁴ However, if the set of defaulting banks is connected to the rest of the system, then the credibility is strictly increasing in β .

As γ increases, that is, the efficiency of buyers of liquidated assets decreases more quickly, illiquid banks raise a smaller amount of cash from their sales and, if γ exceeds a certain threshold, such banks are unable to meet their liabilities. In absence of intervention, the set of defaulting banks increases at such a threshold, causing a downward discontinuity in the credibility of the threat due to bankruptcy losses. Between these discontinuities, two counteracting forces determine the change in credibility: a larger downward pressure on the asset recovery rate causes larger liquidation losses without intervention, but also mandates larger subsidies in a bailout. Which effect dominates depends on the recovery rate β and the network structure via the total throughput of defaulting banks. As illustrated in Figure 3, the marginal change in credibility is higher in more sparsely connected networks. The left panel of Figure 4 shows that the discontinuous changes dominate the continuous changes in a model calibrated to data from the 2018 EBA stress test.

In absence of intervention, an increase in the amount of illiquid asset held by any bank impacts welfare losses in two ways: liquidation and bankruptcy costs increase, while losses of senior creditors decrease because they are repaid a larger amount with the proceeds from liquidation. If $\theta_{\mathcal{D}_N}^i(\beta, \pi) = 0$, that is, every (direct and indirect) creditor of bank i is able to fully repay its senior creditors, or if

³⁴If all senior creditors are repaid in full and the defaulting banks are only liable to each other but not to the solvent banks in the system, then repayments are just a redistribution of net zero wealth among defaulting banks.

the weight of senior creditor losses in the welfare function is sufficiently small ($\lambda \leq \beta/(1 - \beta)$), then welfare losses in absence of intervention are increasing in asset holdings e^i . Since welfare losses in the public bailout are non-increasing in e^i by Lemma 4.1, the credibility of the threat is decreasing in this case except at the threshold e_*^i , where bank i becomes solvent. An increase in asset holdings e^i can increase the credibility only if a sufficient proportion of the revenue of the liquidated assets accrue to senior creditors and the reduction of senior creditor losses is valued sufficiently highly by the regulator, that is, if $\lambda\theta_{\mathcal{D}_N}^i(\beta, \pi)$ is sufficiently high.

The following result shows how the credibility of the threat changes with the size of the initial shock. For the credibility analysis, we write $c = c_0 - \varepsilon$ as the sum of cash kept in period $t = 0$ and the realization of short-term returns on the illiquid assets in period $t = 1$, where ε is interpreted as the size of the shock to those returns. In contrast, the public bailout depends on the size ε of the initial shock only through the shortfall s_0 , hence Lemma 4.1 simply states the dependency on s_0 .

Lemma 4.4. *For each bank i , there exist $0 < \varepsilon_1^i < \varepsilon_2^i \leq \varepsilon_3^i$ such that the credibility of the threat is constant for $\varepsilon^i \leq \varepsilon_1^i$, decreasing for $\varepsilon^i \in [\varepsilon_1^i, \varepsilon_2^i]$, and increasing for $\varepsilon^i \geq \varepsilon_3^i$. On the interval $[\varepsilon_2^i, \varepsilon_3^i]$, the credibility has only downward discontinuities. The marginal change of $W_P - W_N$ at continuity points in $[\varepsilon_2^i, \varepsilon_3^i]$ is decreasing in β and $\theta^i(\beta, \pi)$.*

For very small shock sizes, bank i is able to honor its liabilities without liquidating its assets. Welfare losses and the credibility of the threat thus remain unaffected. For small shock sizes in the interval $[\varepsilon_1^i, \varepsilon_2^i]$, bank i has to start liquidating its assets in absence of intervention, but not in the public bailout where all interbank claims are honored. The credibility of the threat is thus decreasing in that interval. For intermediate shock sizes in the region $[\varepsilon_2^i, \varepsilon_3^i]$, banks do not have sufficient liquidity to repay their liabilities, both in the bailout and in the no-intervention outcome. Whether this leads to a larger increase of welfare losses without intervention than in a bailout—and hence to a decrease in credibility—depends on the recovery rate β and the total throughput of bank i . Finally, for large shock sizes $\varepsilon^i \geq \varepsilon_3^i$, bank i does not make any payment to its junior creditors in absence of intervention, hence any marginal increase in the shock is not amplified through the network anymore. Consequently, the credibility is increasing in the shock size.

We conclude this section with the following result, which highlights that the credibility of the threat is the most important determinant when comparing welfare losses between two networks.

Lemma 4.5. *For fixed $L, e, c, w, \gamma, \beta$, equilibrium welfare losses after intervention are smaller in network π_1 than in network π_2 if the regulator's threat is credible in π_1 but not in π_2 .*

If the regulator's threat is credible in network π^1 but not in network π^2 , Theorem 3.8 implies that $W_E(\pi^1) < W_N(\pi^1) \leq W_P = W_E(\pi^2)$. If the threat fails to be credible in both networks, the regulator must resort to a public bailout in either network, resulting in identical welfare losses. In Section 6, we compare welfare losses when the threat is credible in both networks. We do so numerically, using a data set from the European Banking Authority's 2018 stress test.

5 Optimal Intervention with Partial Rescues

In this section, we extend the baseline model of Section 3 by enlarging the strategy space of the government from complete rescues to bail-ins and bailouts that may rescue only a subset of the banks. We also refer to these interventions as partial interventions or partial rescues. This analysis reveals the additional forces that contribute to the formation of bail-ins when the regulator does not necessarily rescue every bank in the system. To simplify the analysis and the exposition of our results, we assume that buyers of the asset are fully efficient, i.e., $\gamma = 0$. ~~This is equivalent to assuming that banks do not hold any illiquid asset and all non-interbank claims are cash assets.~~ shuts down one of the channels of contagion and is equivalent to assuming that there are no fire sale effects of liquidation.

5.1 Public Bailout

Without the participation of the banks, the regulator minimizes welfare losses over all possible sets of banks he could bail out. The first lemma describes this minimization procedure.

Lemma 5.1. *For any set of banks $\mathcal{B}$, let $p(\mathcal{B})$ be the greatest fixed-point of*

$$p^i = \begin{cases} L^i & \text{if } i \in \mathcal{B} \text{ or } c^i + \sum_{j=1}^n \pi^{ij} p^j \geq L^i + w^i, \\ \left(\beta(c^i + (\pi p)^i) - w^i \right)^+ & \text{otherwise.} \end{cases}$$

Define the vector of subsidies $s(\mathcal{B})$ by setting $s^i(\mathcal{B}) := (L^i + w^i - c^i - \sum_{j=1}^n \pi^{ij} p^j(\mathcal{B}))^+$ for $i \in \mathcal{B}$ and $s^i(\mathcal{B}) = 0$ otherwise. Let $\mathcal{B}_P := \arg \min_{\mathcal{B}} W_\lambda(s(\mathcal{B}))$. A welfare-maximizing partial bailout awards subsidies from the set $\mathcal{S}_P := \{s(\mathcal{B}_) \mid \mathcal{B}_* \in \mathcal{B}_P\}$ and attains welfare losses $W_P^* := \min_{\mathcal{B}} W_\lambda(s(\mathcal{B}))$.*

The bailout $s(\mathcal{B})$ is the welfare-maximizing bailout among all bailouts that rescue banks in $\mathcal{B}$ by giving subsidies only to banks in $\mathcal{B}$. The regulator thus maximizes welfare by optimally selecting which banks to subsidize. Generically, $\mathcal{B}_P$ is a singleton, that is, welfare is maximized for a unique set of banks to be bailed out. Note that in Lemma 5.1, we do not preclude the possibility that the optimal partial bailout rescues no banks at all. In that case, the optimal “bailout” is no intervention. The following result describes the structure of the partial bailout by characterizing conditions under which it is optimal not to rescue a certain set of banks. Those conditions depend on the shortfall $S(\mathcal{B})$ and the capital buffer $C(\mathcal{B})$ in the bailout $s(\mathcal{B})$, defined as follows:

$$S(\mathcal{B}) := (L + w - c - \pi p(\mathcal{B}))^+, \quad C(\mathcal{B}) := (c + \pi p(\mathcal{B}) - w - L)^+. \quad (21)$$

We denote by $\mathcal{D}(\mathcal{B})$ and $\mathcal{S}(\mathcal{B})$ the set of defaulting and solvent banks, respectively, when liabilities are cleared with $p(\mathcal{B})$. We denote by $\delta(\mathcal{B})$ the corresponding losses to senior creditors.

Lemma 5.2. *For any two sets of banks $\mathcal{B}' \subseteq \mathcal{B}$, we denote by $\zeta := \pi(p(\mathcal{B}) - p(\mathcal{B}'))$ the difference of payments received by banks when liabilities are cleared with $p(\mathcal{B})$ and $p(\mathcal{B}')$, and we denote by $\mathcal{R} := \mathcal{D}(\mathcal{B}') \setminus \mathcal{D}(\mathcal{B})$ the set of banks rescued in bailout $s(\mathcal{B})$ but not in $s(\mathcal{B}')$. Then $W_\lambda(\mathcal{B}') < W_\lambda(\mathcal{B})$*

if and only if

$$\sum_{i \in \mathcal{R}} S^i(\mathcal{B}') + \frac{\lambda}{1 + \lambda} \sum_{i \in \mathcal{S}(\mathcal{B})} \min(\zeta^i, C^i(\mathcal{B})) > \sum_{i \in \mathcal{R}} \delta^i(\mathcal{B}') + \sum_{i \in \mathcal{D}(\mathcal{B})} \min(\beta \zeta^i, \delta^i(\mathcal{B}')) + \sum_{i \in \mathcal{S}(\mathcal{B})} \zeta^i. \quad (22)$$

The left-hand side of (22) captures the benefits for not rescuing banks in $\mathcal{R}$. The first term is proportional to the shortfall of banks in $\mathcal{R}$ if those banks are not rescued. Rescuing them is costly and requires government expenditures in the form of awarded subsidies. The second term captures the potential of solvent banks to absorb losses ζ caused by the default of banks in $\mathcal{R}$. If solvent banks have sufficiently large capital buffers, these losses can be absorbed by the system and there is less benefit in rescuing the defaulting banks using public money. The right-hand side of (22) captures the benefits of rescuing banks in $\mathcal{R}$. These include the direct benefits to senior creditors (first two terms) and solvent junior creditors (third term) if banks in $\mathcal{R}$ are rescued. Benefits to insolvent banks do not explicitly appear in the expression above because those benefits are passed on to creditors of rescued banks through the repayment of liabilities.

Lemma 5.2 implies that a bank is not rescued in the optimal bailout if it is hit by a very large shock or if such a shock can be well-absorbed by the capital buffer of its creditors. By contrast, a bank is rescued in the optimal bailout if, relative to the size of the exogenous shock, the bank's default causes large losses to its creditors that cannot be absorbed by their capital buffers. If a bank is hit by a large shock *and* its bankruptcy causes a default cascade, it may be welfare enhancing to only rescue the contagiously defaulting banks: To rescue the fundamentally defaulting bank, the regulator has to cover the bank's large shortfall entirely using taxpayer money. If he rescues only the contagiously defaulting banks, he can leverage the balance sheet capacity of those banks and cover the residual shortfall only after their capital buffers have been depleted.

Even though we do not explicitly model the network formation stage, our result suggests that a risky bank i has incentives to borrow from other risky banks in the system so that in case of i 's default, its creditors are likely to be distressed as well. Then, their potential for absorbing the losses induced by bank i is small, increasing the chances that bank i is bailed out.

5.2 Banks' Equilibrium Responses

When the regulator commits to complete rescues as in Section 3, the threat towards the banks is binary: a bank's assets are either fully protected or, in absence of intervention, they are exposed to the full extent of the default cascade. This is no longer the case when the government allows for partial rescues. A bank's assets may be protected to varying degrees in a welfare-maximizing bailout: while some debtors may be rescued and hence the claims towards those banks are protected, other debtors may still default, thereby inducing losses to the remaining banks in the system.

If the regulator announces that he will implement bailout $s(\mathcal{B})$ when some bank i fails to cooperate, he threatens losses to creditors of banks in the complement $\mathcal{B}^c$ of $\mathcal{B}$. If there is more than one welfare-maximizing (i.e., *credible*) bailout, the regulator can choose which bailout to "threaten" to which banks. Consider a financial system with two identical defaulting banks i and j , where

the welfare-maximizing bailout prescribes the rescue of only one of them. Since rescuing either is a credible action by the regulator, he can threaten the creditors of bank i that, without their participation, he will bail out bank j and vice versa. This is formalized in the following lemma.

Lemma 5.3. *Let (b, s) be a feasible bail-in proposal. In an accepting equilibrium a , bank i with $b^i > 0$ accepts if and only if for some $\mathcal{B}_i \in \mathcal{B}_P$, the following conditions hold:*

1. $W_\lambda(b, s, (0, a^{-i})) \geq W_P^*$,
2. $b^i - s^i \leq \sum_{j=1}^n \pi^{ij}(\bar{p}^j(b, s, (1, a^{-i})) - p^j(\mathcal{B}_i)) - s^i(\mathcal{B}_i)$.

Moreover, if there exists $\mathcal{B}_* \in \mathcal{B}_P$ such that Condition 2 holds simultaneously for every bank i with $\mathcal{B}_i = \mathcal{B}_*$, then rejecting equilibria are subgame Pareto efficient only if $s \in \mathcal{S}_P$. If there exists no such $\mathcal{B}_*$, then rejecting equilibria are not subgame Pareto dominated by accepting equilibria.

Analogous to Lemma 3.5, Condition 1 is a no free-riding constraint, specifying that without i 's participation, the regulator chooses his preferred outside option over the residual bail-in.³⁵ Condition 2 states that bank i 's net contribution to the bail-in (b, s) has to be smaller than or equal to i 's benefits in the bail-in over the threatened bailout $s(\mathcal{B}_i)$. If the regulator threatens different bailouts after a rejection by different banks, the threats cannot be carried out against all banks simultaneously. Because equilibria are robust only to unilateral deviations, this has no effect on accepting equilibria, but it may result in rejecting equilibria also being subgame Pareto efficient.³⁶

5.3 Optimal Proposal of the Regulator

In a partial bail-in, the regulator selects both rescued banks and contributors. Because a bank is more willing to contribute to a bail-in that protects its debtors, the two decisions are interconnected. When the regulator proposes to rescue a set of banks $\mathcal{B}$ and threatens the bailout $s(\mathcal{B}_*)$, the maximal incentive-compatible contribution by bank i , or the *threat level* towards bank i , is equal to

$$\eta_{\mathcal{B}_*}(\mathcal{B}) := \min\left((\pi(p(\mathcal{B}) - p(\mathcal{B}_*)) - s(\mathcal{B}_*))^+, C(\mathcal{B})\right), \quad (23)$$

where $C(\mathcal{B})$ is the vector of the banks' capital buffer defined in (21). The notion of threat levels generalizes the credibility of the threat in Section 3.³⁷ If the optimal bailout happens to be the complete bailout, then threat levels towards all banks are zero and, as in Section 3, the regulator cannot incentivize any banks to participate. If the optimal bailout is not the complete bailout, the regulator minimizes welfare losses over which banks to subsidize, taking into account that it affects

³⁵Note that in the case of partial rescues, the preferred outside option is the threatened partial bailout $s(\mathcal{B}_i)$ leading to welfare losses W_P . By contrast, Lemma 3.5 for complete rescues is formulated under the assumption that the threat is credible, where the regulator's preferred outside option is no intervention and yields welfare losses W_N .

³⁶Whether or not is possible to preclude rejecting equilibria by assigning threats to response vectors $a \neq (1, \dots, 1)$ is a combinatorial argument that depends on the number of participating banks, the number of credible bailout threats, and the rank-order of threats for the individual banks. Since the optimal bailout is generically unique, this analysis is beyond the scope of the paper.

³⁷Because we consider complete rescues in Section 3, it follows that $\mathcal{B} = \{1, \dots, n\}$ and hence $p(\mathcal{B}) = L$. With complete rescues, the threat level towards the banks is thus either 0 or it is given by $b_*^i(\alpha)$ in (14) for any bail-in that induces asset recovery rate α . Because $\ell = 0$ and $\alpha = 1$ in this section, we highlight the dependence on the set $\mathcal{B}$ of subsidized banks.

the contributions he can extract from the private sector. We state the result under the generically satisfied assumption that the optimal bailout is unique. If the welfare-maximizing bailout fails to be unique, the regulator will additionally optimize over which bailouts to threaten.

Theorem 5.4. *Suppose that $\mathcal{B}_P = \{\mathcal{B}_*\}$. For any set of banks $\mathcal{B}$, denote by $i_1(\mathcal{B}), i_2(\mathcal{B}), \dots$ a non-increasing ordering of banks according to $\eta_{\mathcal{B}_*}^i(\mathcal{B})$. For any integer k , define*

$$W^k(\mathcal{B}) := W_\lambda(s(\mathcal{B})) - \lambda \sum_{j=1}^k \eta_{\mathcal{B}_*}^{i_j(\mathcal{B})}(\mathcal{B}).$$

Let $m(\mathcal{B})$ denote the smallest k for which $W_{\mathcal{B}_*}^k(\mathcal{B}) < W_P$ and set

$$W(\mathcal{B}) := \min\left(W^{m(\mathcal{B})}(\mathcal{B}), W_P^* - \lambda \eta_{\mathcal{B}_*}^{i_{m(\mathcal{B})+1}(\mathcal{B})}(\mathcal{B})\right). \quad (24)$$

In the game with partial interventions, welfare losses in any subgame Pareto-efficient equilibrium are equal to $W_E^* := \min_{\mathcal{B}} W(\mathcal{B})$.

As in Theorem 3.8, the no-free-riding constraint drives the regulator to ask for contributions from banks, towards which the threat level is the highest. He includes banks into the bail-in consortium according to the decreasing order $i_1(\mathcal{B}), i_2(\mathcal{B}), \dots$ until welfare losses are lower than in the optimal bailout (left expression in the minimum of (24)). After that, he can only include additional banks into the consortium by burning welfare (right expression in the minimum of (24)).³⁸

In addition to the above forces that also govern the formation of a complete bail-in, the selection of rescued banks affects the size of the contributions the regulator can demand from the private sector. In equilibrium, the regulator will include few banks into the bail-in consortium, each willing to make a large contribution, rather than many banks with small contributions each. It is, therefore, beneficial to rescue banks which have few large creditors rather than banks with many small creditors. This is formalized in the following lemma. It characterizes the structure of equilibrium partial bail-ins by giving conditions, under which it is optimal not to rescue a certain set of banks. We use the same notation as in Lemma 5.2 and Theorem 5.4, and additionally denote by $\mathcal{C}(\mathcal{B}) := \{i_1(\mathcal{B}), \dots, i_{m(\mathcal{B})}(\mathcal{B})\}$ the set of contributing banks towards a rescue of banks in $\mathcal{B}$.

Lemma 5.5. *Suppose that $\mathcal{B}_P = \{\mathcal{B}_*\}$. For any two sets of banks $\mathcal{B}' \subseteq \mathcal{B}$, let $\mathcal{R} := \mathcal{D}(\mathcal{B}') \setminus \mathcal{D}(\mathcal{B})$ as in Lemma 5.2. Then $W_\lambda(s(\mathcal{B}')) - \sum_{i \in \mathcal{C}(\mathcal{B}')} \lambda \eta^i(\mathcal{B}') < W_\lambda(s(\mathcal{B})) - \sum_{i \in \mathcal{C}(\mathcal{B})} \lambda \eta^i(\mathcal{B})$ if and only if*

$$\begin{aligned} & \sum_{i \in \mathcal{R}} S^i(\mathcal{B}') + \frac{\lambda}{1+\lambda} \sum_{i \in \mathcal{S}(\mathcal{B})} \min(\zeta^i, C^i(\mathcal{B})) + \frac{\lambda}{1+\lambda} \left(\sum_{i \in \mathcal{C}(\mathcal{B}')} \eta_{\mathcal{B}_*}^i(\mathcal{B}') - \sum_{i \in \mathcal{C}(\mathcal{B})} \eta_{\mathcal{B}_*}^i(\mathcal{B}') \right) \\ & > \sum_{i \in \mathcal{R}} \delta^i(\mathcal{B}') + \sum_{i \in \mathcal{D}(\mathcal{B})} \min(\beta \zeta^i, \delta^i(\mathcal{B}')) + \sum_{i \in \mathcal{S}(\mathcal{B})} \zeta^i + \frac{\lambda}{1+\lambda} \sum_{i \in \mathcal{C}(\mathcal{B})} \min(\zeta^i, \eta_{\mathcal{B}_*}^i(\mathcal{B})). \end{aligned} \quad (25)$$

³⁸If buyers of liquidated assets are fully efficient, the order $i_1(\mathcal{B}), i_2(\mathcal{B}), \dots$ of banks most exposed to contagion does not depend on the number of banks included in the consortium. The order is fixed for given $\mathcal{B}$ and $\mathcal{B}_*$, and this simplifies the characterization of welfare burning in equilibrium: To include an additional bank i into the consortium, the regulator burns $W_P - W^{m(\mathcal{B})}(\mathcal{B})$ units of welfare so that without i 's participation, welfare in the residual bail-in and the optimal bailout are identical. Then, bank i does not have an incentive to free-ride. Because banks are decreasingly ordered according to their threat levels, the regulator will consider burning welfare only to include bank $i_{m(\mathcal{B})+1}(\mathcal{B})$.

Similarly to the characterization of the optimal bailout in Lemma 5.2, the left-hand side of (25) represents the benefits of not rescuing banks in $\mathcal{R}$, whereas the right-hand side represents the benefits of rescuing banks in $\mathcal{R}$. The majority of terms are identical to (22), but there are two key differences. The first is related to the system's ability to absorb losses transmitted from banks in $\mathcal{R}$ when they are not rescued (second term on the left-hand side of (25)). Losses are absorbed either partially or completely by the capital buffers of $\mathcal{R}$'s creditors. In order to benefit from those capital buffers in a partial bailout, the regulator has an incentive to let banks in $\mathcal{R}$ default. In the bail-in, the regulator benefits from those capital buffers even if he rescues the banks in $\mathcal{R}$ because he can extract larger contributions—up to the amount $\eta(\mathcal{B}) \leq C(\mathcal{B})$ —from contributors in $\mathcal{C}(\mathcal{B})$ (the third term on the right-hand side of (25)). The choice of which banks to rescue is thus *network-dependent*: If banks in $\mathcal{R}$ have many creditors, only some of them will be included in $\mathcal{C}(\mathcal{B})$ due to free-riding incentives. Thus, if capital buffers are large and the defaults of banks in $\mathcal{R}$ can be well absorbed, the second term on the left-hand side of (25) is larger than the third term on the right-hand side, constituting a reason not to rescue banks in $\mathcal{R}$. If, however, banks in $\mathcal{R}$ have only a few large creditors, those are likely included in $\mathcal{C}(\mathcal{B})$, hence the two terms balance out.³⁹ While, in the partial bailout, only the size of the absorbed losses matter, in the partial bail-in the distribution of those losses matters as well because it determines the contributions that can be elicited from the private sector.

The second difference from the case of complete bailouts is the third term on the left-hand side of (25). It captures the structure of rescue consortia in the two alternative bail-ins, stating that it is beneficial to not rescue banks in $\mathcal{R}$ if, by doing so, the regulator does not lose any contributors. Indeed, if the number of banks contributing towards a bail-in rescuing banks in $\mathcal{B}'$ is larger than when rescuing banks in $\mathcal{B}$, i.e., $m(\mathcal{B}') \geq m(\mathcal{B})$, then this term is positive because $\mathcal{C}(\mathcal{B}')$ is the set of banks of size $m(\mathcal{B}')$ that maximizes contributions of size $\eta(\mathcal{B}')$.

Both Lemmas 5.2 and 5.5 imply that it is beneficial for banks to have only a small number of creditors. To be rescued in a bailout, the bank must cause large contagion effects. To be rescued in a bail-in, the bank's creditors need to be among the largest potential contributors to the rescue consortium. Both are more likely to happen when the losses caused by the bank's default are spread only across a few creditors. Because the regulator prefers sparsely connected networks, it follows that ex-ante incentives of banks are better aligned with the regulator's objective when he allows for partial intervention, rather than when restricting himself to complete rescues only.

Our final result relates the structure and welfare losses of partial bail-ins and bailouts.

Lemma 5.6. *Let $\mathcal{B} \subseteq \mathcal{B}'$ with welfare losses $W_\lambda(s(\mathcal{B}')) \leq W_\lambda(s(\mathcal{B}))$ in the corresponding partial bailouts. Then welfare losses in the optimal partial bail-ins, defined in (24), satisfy $W(\mathcal{B}') \leq W(\mathcal{B})$.*

For $\mathcal{B}' \in \mathcal{B}_P$, Lemma 5.6 shows that it cannot be optimal to rescue a smaller set of banks in a bail-in than in the optimal bailout: by Lemma 5.3, no bank would have any incentive to contribute to such a bail-in. Lemma 5.6, however, does not imply that all banks from the optimal bailout are rescued in the optimal bail-in, ~~nor is that statement true in general. Condition 2 in Lemma 5.3 shows that, in~~

³⁹One can show that $\min(\zeta^i, \eta^i(\mathcal{B})) = \min(\zeta^i, C^i(\mathcal{B}))$ for any bank i with a positive threat level $\eta^i(\mathcal{B}')$ when banks in $\mathcal{R}$ are not rescued. The terms in the sums for each such bank are thus identical.

order to incentive banks to participate, the regulator needs to propose bail-ins that rescue banks that are not bailed out. In the online appendix, we provide a numerical example, in which banks rescued in the optimal bailout are not rescued in the optimal bail-in; we provide a counterexample to that claim in Section ?? of the online appendix. In general, the partial bail-in must rescue additional banks to create incentives for banks to contribute, but may give up the rescue of some banks that are protected in the bailout.

6 Equilibrium Welfare Losses and Network Structure

In this section, we analyze the dependency of equilibrium welfare losses on the structure of the interbank network using data from the 2018 stress test of the European Banking Authority (EBA). After eliminating banks with zero reported interbank claims, we are left with 36 banks in the data set. We elaborate in Section I of the online Appendix how we calibrate our model to the data. To highlight more prominently the impact of the network structure on welfare, we assume that buyers of liquidated assets are fully efficient for this calibration exercise. This corresponds to setting $\gamma = 0$ or, equivalently, assuming that all outside assets are held as cash. Detailed information on bilateral exposures is not publicly available. To compare the relative performance of different network structures, we fit a sparse and a dense network structure π_s and π_d , respectively, to the data from the EBA stress test. We then analyze the credibility of the regulator’s threat and the equilibrium welfare losses as a function of the network structure $\pi_\mu := \mu\pi_s + (1 - \mu)\pi_d$ for $\mu \in [0, 1]$.⁴⁰ We generate the dense network π_d with the maximum entropy method developed by Upper and Worms (2004), which distributes interbank liabilities as evenly as possible among the counterparties, yielding a complete network. We generate the sparse network π_s using an iterated greedy algorithm, for which details are provided in Appendix I.⁴¹ We then apply a shock to the assets of HSBC, Barclays, and Deutsche Bank with a shock size equal to their cash holdings, thereby wiping out the value of their non-interbank assets.

The left plot of Figure 5 shows the impact of the network structure on welfare losses under different resolution plans. As the network becomes sparser, liabilities among banks are more concentrated. This makes contagious defaults in absence of intervention more likely and the corresponding welfare losses W_N change discontinuously where this happens. For the chosen size of shocks, the threat is credible in all networks since W_N is smaller than welfare losses W_P in the complete bailout. This allows the coordination of a complete equilibrium bail-in as described in Theorem 3.8 with welfare losses equal to W_E . Contributions to the equilibrium bail-in are illustrated in the right plot of Figure 5, where it is evident that they increase as the network becomes sparser because free-riding incentives are reduced. The three main creditors of the shocked banks

⁴⁰Craig and Von Peter (2014) show that the German interbank network has a core-periphery structure: while the 45 large core banks act as intermediaries and have many counterparties, the periphery banks trade only with core banks, but not with each other. The participating banks in the EBA stress test are 36 of the largest banks in Europe, which are all considered core banks. Therefore, we do not aim to estimate a core-periphery network to this data set but rather analyze the impact of a range of network structures of different sparsity on credibility and equilibrium welfare losses for the subnetwork of core banks.

⁴¹The resulting network π_d has 1260 edges, i.e., each of the 36 banks is connected to every other bank, and a normalized Gini index of 0.4556. The network π_s has 71 edges with a normalized Gini index of 0.9981. See Hurley and Rickard (2009) for a definition of the Gini index. The Gini index is a measure of sparsity, which we normalize to account for the fact that diagonal entries in any relative liability matrix are 0. The normalized Gini index is 0 for the symmetric complete network and 1 for a ring network.

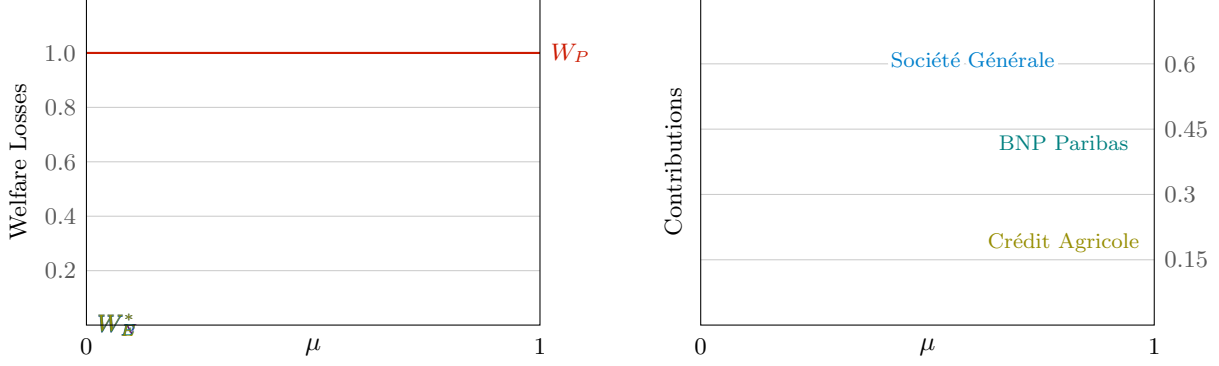


Figure 5: The two plots show how welfare losses under different resolution plans and banks' contributions to the equilibrium bail-in change as a function of the network $\pi_\mu = \mu\pi_s + (1-\mu)\pi_d$ as μ increases from 0 to 1. Thus, at the left edge of each plot, the network is equal to a dense network π_d and it gets progressively sparser as we move to the right until it is equal to a sparse network π_s at the right edge. Welfare losses and welfare impacts of banks' contributions are shown relative to the welfare losses W_P in the complete bailout. Contributions of banks are shown cumulatively so that the contributed amount of a single bank corresponds to the distance between two consecutive lines.

would suffer large losses without intervention and can, therefore, be incentivized to make large contributions. This leads to a continuous decrease in equilibrium welfare losses as the network gets sparser until we observe contagious defaults: due to the no-free-riding constraint, welfare losses in a bail-in can differ from W_N by at most the contribution of any participating bank. Therefore, discontinuous changes in W_N are reflected also in the equilibrium welfare losses W_E with complete rescues. Nevertheless, W_E in the sparsest network is 21.3% lower than in the most dense network despite the fact that without intervention, they would be 17.9% higher.

For the chosen shock sizes, there are no contagiously defaulting banks in the most dense network. This illustrates that even if the regulator's threat is more credible in a dense network, equilibrium welfare losses are typically still decreasing in the sparsity of the network because of the reduced free-riding incentives. For larger shock sizes or a more lowly capitalized financial system, we would observe that the credibility improves as the network becomes sparser because dense connections amplify the shock, leading to a decrease in credibility by Lemma 3.4.

Finally, the left panel of Figure 5 illustrates the findings of Section 5, in which we consider interventions that may target only a subset of banks. For the chosen shock sizes, ~~it turns out that in the optimal bailout,~~ the regulator only rescues the contagiously defaulting banks ~~and lets the fundamentally defaulting banks fail in the optimal bailout.~~ This leads to welfare losses W_P^* , which coincide with W_N when there are no contagiously defaulting banks. In the equilibrium partial bail-in, however, it is optimal to rescue every bank because contributions from the private sector can be solicited for the rescue of fundamentally defaulting banks. This is consistent with the predictions of Lemmas 5.2 and 5.5. Welfare losses W_E^* in the equilibrium bail-in with partial rescues are 41.3% lower than in the complete network structure. In this calibrated model, there are 7,140 ways of applying idiosyncratic shocks to three banks in a network of 36 banks. Taking the average over all combinations of shocks, equilibrium welfare losses in the sparsest network are lower than in the most dense network by 12.5% for complete rescues and by 23.1% for partial rescues. These results ~~highlight the importance of the network structure and~~ suggest that structural policies aiming at

making the financial network more sparse may significantly raise welfare.

7 Concluding Remarks

Various initiatives have been undertaken by central governments and monetary authorities, especially after the global financial crisis, to expand resolution plans and tools. Our paper makes a first step towards a systematic analysis of the incentives that govern alternative resolution plans. At the heart of our analysis is the credibility of the regulator’s no-intervention (or partial rescue) threat, given the desire of each bank to free-ride on the government’s contributions. Selective bail-ins may be preferable to complete bailouts and our analysis provides rational arguments for which banks should be targeted.

The threat **not to intervene** fails to be credible for a given shock size if and only if the shock is heavily amplified through inefficient asset liquidation, bankruptcy costs, and linkages between densely connected banks in distress. The impact of the interbank network structure on the credibility is captured by a measure we call the *total throughput* of defaulting banks, which reflects the rate at which losses spill over to solvent members of the economy, ~~taking into account feedback effects between defaulting banks~~. Conditional on the banks’ levels of solvency, the total throughput depends entirely on the network structure, giving us a metric to rank the desirability of different network structures. In our model of endogenous intervention, sparse networks become relatively more desirable because the threat remains credible for larger shock sizes and free-riding incentives are reduced. For intermediate shock sizes, these effects dominate the diversification benefits of dense networks. Our analysis thus reverses the presumptions concerning the relative desirability of sparse versus dense networks.

In our model, the regulator and the banks have complete information about the underlying financial system, which greatly simplifies the negotiation process. In future research, it would be interesting to study the situation, in which the regulator and/or the banks have only partial information about the financial system. Then, banks may have an incentive to reject subsidies in order to signal financial strength. Consequently, banks need to be incentivized not only to participate in a bail-in, but also to provide truthful information about their interbank linkages.

Our results have obvious implications for the design of regulations that affect the network structure and pave the way for future research on endogenous network formation. In such a model, banks anticipate how their ex-ante risk-taking behavior and their choice of counterparties affect the credibility of future rescue plans as well as their expected gains from those rescues. This adds an important dimension to the moral hazard literature: through their interbank linkages, banks can also control the likelihood of a public bailout as well as the structure of the prevailing bail-in. Our preliminary results in this direction suggest that it is beneficial for banks to borrow only from few other banks so that contagion effects in case of the bank’s default are highly concentrated. Then, free-riding incentives among creditors are smaller, making a rescue more appealing to the regulator. Accounting for endogenous network formation will lead to a comprehensive framework for the ~~analysis of welfare-maximizing policies that the regulator can~~

impose;welfare analysis of various policies, combining the bail-out and bail-in strategies analyzed here with policies directed at affecting the network structure ex ante, including exposure limits, imposing taxes on exposures beyond a certain limit, or setting limits on intermediation volume.

References

- Acemoglu, D., A. Ozdaglar, and A. Tahbaz-Salehi.** 2015. “Systemic Risk and Stability in Financial Networks.” *American Economic Review*, 105(2): 564–608.
- Acharya, V.V., and T. Yorulmazer.** 2007. “Too Many to Fail—An Analysis of Time-Inconsistency in Bank Closure Policies.” *Journal of Financial Intermediation*, 16(1): 1–31.
- Acharya, V.V., H.S. Shin, and T. Yorulmazer.** 2011. “Crisis Resolution and Bank Liquidity.” *Review of Financial Studies*, 24(6): 2166–2205.
- Allen, F., and D. Gale.** 2000. “Financial Contagion.” *Journal of Political Economy*, 108(1): 1–33.
- Basel-III.** 2013. “The Liquidity Coverage Ratio and liquidity risk monitoring tools.” *Basel Committee on Banking Supervision*.
- Battiston, S., D. Delli Gatti, M. Callegati, B. Greenwald, and J.E. Stiglitz.** 2012. “Default Cascades: When Does Risk Diversification Increase Stability?” *Journal of Financial Stability*, 8(3): 138–149.
- Benoit, S., J.E. Colliard, C. Hurlin, and C. Perignon.** 2017. “Where the Risks Lie: A Survey on Systemic Risk.” *Review of Finance*, 21: 109–152.
- Bernard, B., A. Capponi, and J.E. Stiglitz.** 2019. “Online Appendix to “Bail-ins and Bail-outs: Incentives, Connectivity, and Systemic Stability”, available at <http://benjamin-bernard.com>.”
- Boissay, F.** 2006. “Credit Chains and the Propagation of Financial Distress, available at http://ssrn.com/abstract_id=872543.” *ECB Working Paper Series, No. 573*.
- Cabrales, A., P. Gottardi, and F. Vega-Redondo.** 2017. “Risk-Sharing and Contagion in Networks.” *Review of Financial Studies*, 30(9): 3086—3127.
- Capponi, A., P.C. Chen, and D. Yao.** 2016. “Liability Concentration and Systemic Losses in Financial Networks.” *Operations Research*, 64(5): 1121–1134.
- Castiglionesi, F.** 2007. “Financial Contagion and the Role of the Central Bank.” *Journal of Banking and Finance*, 31(1): 81–101.
- Chari, V.V., and P.J. Kehoe.** 2016. “Bailouts, Time Inconsistency, and Optimal Regulation: A Macroeconomic View.” *American Economic Review*, 106(9): 2458–2493.
- Cifuentes, R., G. Ferrucci, and H.S. Shin.** 2005. “Liquidity Risk and Contagion.” *Journal of the European Economic Association*, 3(2): 556–566.
- Craig, B., and G. Von Peter.** 2014. “Interbank Tiering and Money Center Banks.” *Journal of Financial Intermediation*, 23(3): 322–347.
- Duffie, D., and C. Wang.** 2017. “Efficient Contracting in Network Financial Markets.” *Working Paper, Graduate School of Business, Stanford University*.
- Eisenberg, L., and T.H. Noe.** 2001. “Systemic Risk in Financial Systems.” *Management Science*, 47(2): 236–249.
- Elliott, M., B. Golub, and M.O. Jackson.** 2014. “Financial Networks and Contagion.” *American Economic Review*, 104(10): 3115–3153.
- Ellul, A., C. Jotikasthira, and C.T. Lundblad.** 2011. “Regulatory Pressure and Fire Sales in the Corporate Bond Market.” *Journal of Financial Economics*, 101(3): 596–620.
- Elsinger, H., A. Lehar, and M. Summer.** 2006. “Risk Assessment for Banking Systems.” *Management Science*, 52(9): 1301–1314.
- Erol, S.** 2018. “Network Hazard and Bailouts.” *Working paper, University of Pennsylvania*.

- Farhi, E., and J. Tirole.** 2012. “Collective Moral Hazard, Maturity Mismatch, and Systemic Bailouts.” *American Economic Review*, 102(1): 60–93.
- Freixas, X., B.M. Parigi, and J.-C. Rochet.** 2000. “Systemic Risk, Interbank Relations and Liquidity Provision by the Central Bank.” *Journal of Money Credit and Banking*, 32(3): 661–638.
- Gai, P., A. Haldane, and S. Kapadia.** 2011. “Complexity, Concentration and Contagion.” *Journal of Monetary Economics*, 58(5): 453–470.
- Gai, P., and S. Kapadia.** 2010. “Contagion in Financial Networks.” *Proceedings of the Royal Society A*, 466: 2401–2423.
- Gale, D., and X. Vives.** 2002. “Dollarization, Bailouts, and the Stability of the Banking System.” *Quarterly Journal of Economics*, 117(2): 467–502.
- Glasserman, P., and H.P. Young.** 2015. “How Likely is Contagion in Financial Networks?” *Journal of Banking and Finance*, 50: 383–399.
- Glasserman, P., and H.P. Young.** 2016. “Contagion in Financial Networks.” *Journal of Economic Literature*, 54(3): 779–831.
- Greenspan, A.** 1998. “Testimony by the Chairman of the Federal Reserve Board before the Committee on Banking and Financial Services of the US House of Representatives.”
- Greenwald, B., and J.E. Stiglitz.** 2003. *Towards a New Paradigm in Monetary Economics*. Cambridge University Press.
- Hurley, N., and S. Rickard.** 2009. “Comparing Measures of Sparsity.” *IEEE Transactions on Information Theory*, (10): 4723–4741.
- Jackson, M., and A. Pernoud.** 2020. “Systemic Risk in Financial Networks: A Survey.” Available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3651864.
- Keister, T.** 2016. “Bailouts and Financial Fragility.” *Review of Economic Studies*, 83(2): 704–736.
- May, R.A., S.A. Levin, and G. Sugihara.** 2008. “Complex Systems: Ecology for Bankers.” *Nature*, 451: 893–895.
- Nier, E., J. Yang, T. Yorulmazer, and A. Alentorn.** 2007. “Network Models and Financial Stability.” *Journal of Economic Dynamics and Control*, 31(6): 2033–2066.
- Rogers, L.C.G., and L.A.M. Veraart.** 2013. “Failure and Rescue in an Interbank Network.” *Management Science*, 59(4): 882–898.
- Schilling, L.** 2018. “Optimal Forbearance of Bank Resolution.” *Working Paper, Becker Friedman Institute for Research in Economics*, 15.
- Stiglitz, J.E.** 2002. *Globalization and its Discontents*. W.W. Norton.
- Uhlig, H.** 2010. “A model of a systemic bank run.” *Journal of Monetary Economics*, 57(1): 78–96.
- Upper, C., and A. Worms.** 2004. “Estimating Bilateral Exposures in the German Interbank Market: Is There a Danger of Contagion?” *European Economic Review*, 48: 827–849.

A Welfare Burning

In this appendix, we define the minimal amount $\chi_{\mathcal{C}}(\alpha)$ of welfare burning needed to eliminate free-riding incentives from an individually incentive-compatible bail-in with contributing banks in $\mathcal{C}$.

Lemma A.1. *As in Definition 3.2, let $z(\alpha) := \alpha \ln(\alpha)$, let z^{-1} be its inverse on the interval $[\frac{1}{e}, 1]$, and set $g_{\alpha}(x) := g(z^{-1}(z(\alpha) + \gamma x)) - g(\alpha)$. Note that the function g_{α} is invertible for $\alpha \geq \alpha_{\text{ind}}$. For a set of banks $\mathcal{C}$, let $\chi_{\mathcal{C}}^i(\alpha) := (W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j - g_{\alpha}(b_{*}^i(\alpha) - \underline{b}^i))^+$ for any bank $i \in \mathcal{C}$. Moreover, let $\hat{\chi}_{\mathcal{C}}(\alpha)$ denote the unique non-negative solution χ to*

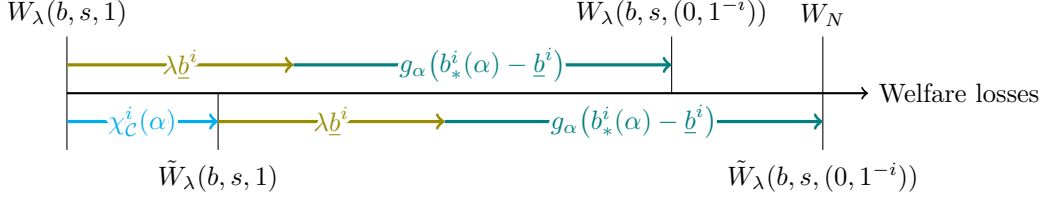


Figure 6: Consider an individually incentive-compatible bail-in (b, s) that induces asset recovery rate α , in which bank i contributes the largest incentive-compatible amount $b_*^i(\alpha)$. A rejection by bank i has an impact of $\lambda \underline{b}^i + g_\alpha(b_*^i(\alpha) - \underline{b}^i)$ on welfare (top). If the regulator burns $\chi_C^i(\alpha) = (W_N - W_\lambda(b, s, (0, 1^{-i})))^+$ units of welfare in a proposal that is otherwise identical, he becomes indifferent between the residual bail-in without bank i and no intervention (bottom). Thus, bank i can no longer free-ride on the contributions of the remaining banks.

$$-\frac{\alpha \ln(\alpha)}{\gamma} = \sum_{i \in \mathcal{C}} g_\alpha^{-1} \left(\left(W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j - \chi \right)^+ \right) \quad (26)$$

if it exists and let $\hat{\chi}_C(\alpha) = 0$ otherwise. Define $\chi_C(\alpha) := \hat{\chi}_C(\alpha) \vee \max_{i \in \mathcal{C}} \chi_C^i(\alpha)$. In any bail-in (b, s) with equilibrium response $1 = (1, \dots, 1)$, set of contributing banks $\mathcal{C} = \{i \mid b^i > 0\}$, and induced asset recovery rate $\alpha = \bar{\alpha}(b, s, 1)$, the amount of welfare burnt is bounded below by $\chi_C(\alpha)$, i.e.,

$$\lambda \sum_{i \notin \mathcal{C}} (s^i - s_0^i)^+ + \lambda \sum_{i \in \mathcal{C}} (\underline{b}^i - b^i)^+ \geq \chi_C(\alpha). \quad (27)$$

For the no-free-riding incentives to hold, the rejection by any bank must have an impact on welfare in excess of $W_N - W_\lambda(b, s, 1)$. It follows from (15) that the impact on welfare has two components: contributions up to $\underline{b}$ affect welfare directly by an amount $\lambda \underline{b}$, whereas contributions that exceed $\underline{b}$ by an amount x require asset liquidation and affect welfare through the trade-off $g_\alpha(x)$. A contribution of size $\underline{b}^i + x^i$ by bank i thus has a total impact on welfare of $\lambda \underline{b}^i + g_\alpha(x^i)$. Among individually incentive-compatible contributions, the welfare impact by bank i is maximized when i contributes $b_*^i(\alpha)$. If the maximal welfare impact is smaller than $W_N - W_\lambda(b, s, 1)$, the remaining value $\chi_C^i(\alpha) = W_N - W_\lambda(b, s, 1) - \lambda \underline{b}^i - g_\alpha(b_*^i(\alpha) - \underline{b}^i)$ has to be burnt; see Figure 6 for an illustration. Burning an amount equal to $\max_{i \in \mathcal{C}} \chi_C^i(\alpha)$ thus deters free-riding by all banks when each bank liquidates the maximal incentive-compatible amount. However, asking each bank to contribute $b_*^i(\alpha)$ may require asset liquidation in excess of $-\ln(\alpha)/\gamma$, which would depress the asset recovery rate below α by (1). In that case, the regulator has to ask for lower contributions from banks that require aggregate liquidation of only $-\ln(\alpha)/\gamma$ and burn additional welfare to balance the reduced welfare impact of the contributing banks. The lowest amount of welfare burnt in that case is $\hat{\chi}_C(\alpha)$; see Figure 7 for an illustration. The proof of Lemma A.1 is in Appendix C.

B Proofs of Results in Sections 2 and 3

B.1 Existence and Monotonicity of Clearing Equilibria

This appendix provides the proofs of Section 2, asserting existence and monotonicity of clearing equilibria, together with Pareto dominance of the greatest clearing equilibrium. We begin with the

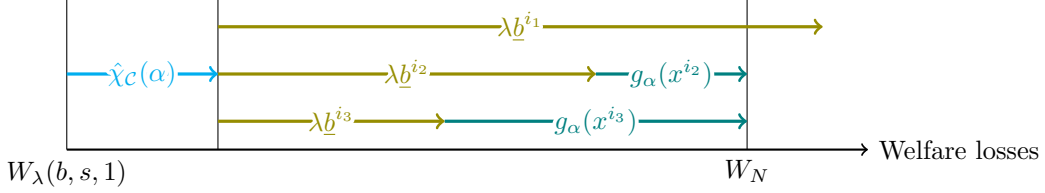


Figure 7: The construction of $\hat{\chi}_C(\alpha)$ is illustrated for $C = \{i_1, i_2, i_3\}$. It is the minimal amount of welfare burnt such that the welfare impact $\lambda \underline{b}^i + g_\alpha(x^i)$ caused by the rejection of any bank $i \in C$ exceeds $W_N - W_\lambda(b, s, 1) - \hat{\chi}_C(\alpha)$, subject to the constraint $\sum_{i \in C} x^i \leq -\alpha \ln(\alpha)/\gamma$, which guarantees that the total amount liquidated does not depress the asset recovery rate below α . Note that $x^{i_3} = 0$ because the welfare impact of i_3 's rejection is large enough to deter free-riding without asset liquidation. If $\hat{\chi}_C(\alpha) > 0$, then the constraint binds, that is, $x^{i_2} + x^{i_3} = -\alpha \ln(\alpha)/\gamma$. Otherwise, the welfare impacts of banks i_1 and i_2 could be increased and the amount of welfare burnt decreased.

following auxiliary result, which provides us with an alternative expression for welfare losses.

Lemma B.1. *For any clearing equilibrium (p, ℓ, α) , the following identity holds*

$$W_\lambda(p, \ell, \alpha) = \sum_{i=1}^n (c^i + e^i - w^i - V^i(p, \ell, \alpha)) + \sum_{i \in \mathcal{D}(p, \ell, \alpha)} (1 + \lambda) \delta^i(p, \alpha). \quad (28)$$

The proof follows from (3), (4), and the fact that $\sum_i (\pi x)^i = \sum_i x^i$ for any vector x by row-stochasticity of π . Since the equity value of any bank is monotonically increasing in the clearing equilibrium, a direct consequence of Lemma B.1 is the fact that welfare losses are monotonically decreasing in the clearing equilibrium. For the sake of reference, we state this property as a lemma.

Lemma B.2. *Any bank i 's value of equity $V^i(p, \ell(p, \alpha), \alpha)$ is non-decreasing and welfare losses $W_\lambda(p, \ell(p, \alpha), \alpha)$ are non-increasing in α and p^j for any j .*

Proof. By definition in (2), $\ell^i(p, \alpha)$ is non-decreasing in p and α . Since bank i 's value in (4) is equal to $V^i(p, \ell, \alpha) = (\pi p + c + e - (1 - \alpha)\ell(p, \alpha) - w - L)^i 1_{\{p^i = L^i\}}$, it is non-decreasing in p . Because the weak derivative of $-(1 - \alpha)\ell(p, \alpha)$ is non-negative by the product rule, $V^i(p, \ell, \alpha)$ is also non-decreasing in α . Monotonicity of $V^i(p, \ell, \alpha)$ implies that the first term in (28) is non-increasing in (p, α) . The second term in (28) is non-increasing in (p, α) by definition of $\delta(p, \alpha)$ in (5). $\square$

Proof of Lemma 2.2. Let $\mathcal{L} := [0, e^1] \times \cdots \times [0, e^n]$ denote the set of possible liquidation decisions by the banks. Fix a vector p of interbank repayments and define the operator $\Phi_p : \mathcal{L} \rightarrow \mathcal{L}$ by setting $\Phi_p^i(x) := \ell^i(\alpha(x), p)$ for $i = 1, \dots, n$, where $\alpha(x)$ and $\ell(\alpha, p)$ are defined in (1) and (2). By construction, a pair (ℓ, α) is a solution to (1) and (2) if and only if it is of the form $(x, \alpha(x))$ for a fixed point x of Φ_p . Since both α and ℓ^i are non-increasing, Φ_p^i is non-decreasing. Therefore, Tarski's fixed-point theorem implies that the set of Φ_p 's fixed points forms a complete lattice. In particular, there exists a fixed point $\underline{x}$ such that $\underline{x}^i \leq x^i$ for any other fixed point x of Φ_p and each i . Let $\ell_p = \underline{x}$ and $\alpha_p = \alpha(\ell_p)$. By construction, (ℓ_p, α_p) satisfies (1) and (2) and for any other solution $(\tilde{\ell}, \tilde{\alpha})$, we have $\tilde{\alpha} = \exp(-\gamma \sum_{i=1}^n \tilde{\ell}^i) \leq \exp(-\gamma \sum_{i=1}^n \ell_p^i) = \alpha_p$. $\square$

Before proving existence of clearing equilibria, we show the following comparison result for fixed points that will be used many times throughout our analysis.

Lemma B.3. *Let f and g be two non-decreasing functions mapping a compact set $\mathcal{X}$ into itself. Let $\bar{x}_f$ ($\underline{x}_f$) and $\bar{x}_g$ ($\underline{x}_g$) denote the greatest (least) fixed points of f and g , respectively, that exist by Tarski's fixed point theorem. If $f(\bar{x}_g) \geq g(\bar{x}_g)$, then $\bar{x}_f \geq \bar{x}_g$. If $f(\underline{x}_g) \leq g(\underline{x}_g)$, then $\underline{x}_f \geq \underline{x}_g$.*

Proof. Let $\bar{x}_n := f^{(n)}(\bar{x}_g)$ define the n -fold application of f to $\bar{x}_g$. Since $f(\bar{x}_g) \geq g(\bar{x}_g) = \bar{x}_g$ and f is non-decreasing, it follows that $(\bar{x}_n)_{n \geq 1}$ is non-decreasing in each component. Because of compactness, $(\bar{x}_n)_{n \geq 1}$ converges to some fixed point $\bar{x}_\infty$ of f . Therefore, $\bar{x}_g \leq \bar{x}_\infty \leq \bar{x}_f$ because $\bar{x}_f$ is the greatest fixed point of f . The analogous argument shows $\underline{x}_g \geq \lim_{n \rightarrow \infty} f^{(n)}(\underline{x}_g) \geq \underline{x}_f$. $\square$

Proof of Lemma 2.1. Let $\mathcal{P} := [0, L^1] \times \cdots \times [0, L^n]$ denote the set of all repayment vectors. Observe first that in any clearing equilibrium of the form (p, ℓ_p, α_p) , in which (ℓ_p, α_p) is given by Lemma 2.2, the vector of repayments p is a fixed point of the operator $\Phi : \mathcal{P} \rightarrow \mathcal{P}$, defined by

$$\Phi^i(p) := \begin{cases} L^i & \text{if } c^i + \alpha_p e^i + (\pi p)^i \geq L^i + w^i, \\ \left(\beta(c + \alpha_p e + \pi p)^i - w^i \right)^+ & \text{otherwise.} \end{cases} \quad (29)$$

We proceed to show that Φ is monotone. It follows directly from (2) that, for any i , $\ell^i(\alpha, p)$ is non-increasing in p^j for any j . Using the definition of Φ_p given in the proof of Lemma 2.2, we deduce that $\Phi_{(\bar{p}^j, p^{-j})}^i(\ell_p) \leq \Phi_p^i(\ell_p)$ for any $\bar{p}^j > p^j$. Therefore, Lemma B.3 shows that $\ell_{(\bar{p}^j, p^{-j})}^i \leq \ell_p^i$ and hence $\alpha_{(\bar{p}^j, p^{-j})} \geq \alpha_p$. This shows that $p \mapsto \alpha_p$ is non-decreasing, hence so is Φ^i . Tarski's fixed point theorem thus implies the existence of a fixed point $\bar{p}$ with $\bar{p} \geq p$ for any fixed point p . Monotonicity of the maps $p \mapsto \alpha_p$ and $p \mapsto \ell_p$ shows that any clearing equilibrium of the form (p, ℓ_p, α_p) for some p is dominated by $(\bar{p}, \bar{\ell}, \bar{\alpha})$ for $\bar{\alpha} = \alpha_{\bar{p}}$ and $\bar{\ell} = \ell_{\bar{p}}$. Maximality of (ℓ_p, α_p) in Lemma 2.2 and monotonicity in (3) show that $(\bar{p}, \bar{\ell}, \bar{\alpha})$ also dominates any other clearing equilibrium. Monotonicity of the banks' equity value and welfare losses now follows from Lemma B.2. $\square$

B.2 Bailouts

This appendix shows that welfare-maximizing bailouts are of the form given in Lemma 3.2. We begin this appendix with the following auxiliary lemma, whose elementary proof is omitted.

Lemma B.4. *For any $\gamma > 0$, the function $g(\alpha) = ((1 + \lambda)\alpha - 1) \ln(\alpha) / \gamma$ defined in Lemma 3.1 is strictly convex. Moreover, its global minimizer α_{ind} is decreasing in λ and it lies in $(\max(\frac{1}{1+\lambda}, \frac{1}{e}), 1)$.*

Proof of Lemma 3.1. Let s denote a vector of subsidies of a complete bailout with $s^i \leq s_0^i$ for every bank i . Since every bank is rescued when subsidies s are awarded, the definition of s_0 implies that $\ell^i = \frac{1}{\bar{\alpha}(s)}(s_0^i - s^i)$ for any bank i . It follows from (1) that

$$\bar{\alpha}(s) = \exp\left(-\frac{\gamma}{\bar{\alpha}(s)} \sum_{i=1}^n (s_0^i - s^i)\right). \quad (30)$$

Solving (30) for $\sum_{i=1}^n s^i$ and substituting into (7) shows (9). $\square$

Proof of Lemma 3.2. Let s denote a vector of subsidies that maximize welfare in a complete bailout. Because a minimal subsidy of s_L is needed to support the clearing payment vector L , it follows that $s^i \geq s^i(\alpha, e)$ for every bank i . Since any subsidies beyond s_0 have infinitesimal welfare impact $-\lambda$, it follows that $s^i \leq s_0^i$ for every bank i . Therefore, $\ell^i = \frac{1}{\alpha_P}(s_0^i - s^i)$ for any bank i , which implies

$$\alpha_P = \exp\left(-\frac{\gamma}{\alpha_P} \sum_{i=1}^n (s_0^i - s^i)\right). \quad (31)$$

By Lemma 3.1, welfare in the complete bailout depends on the awarded subsidies only through $g(\alpha_P)$. Since g is differentiable, it follows from Lemma B.4 that α_P is either a boundary point or it is equal to α_{ind} . Monotonicity of (31) in the awarded subsidies implies α_P lies between recovery rates α_L and 1 that are attained by subsidies s_L and s_0 , respectively. Since $g(\alpha_{\text{ind}}) < 0 = g(1)$, it follows from monotonicity in (31) that $\alpha_P = \max(\alpha_{\text{ind}}, \alpha_L)$. Inverting (31) for $\sum_{i=1}^n s^i$ yields (10). $\square$

B.3 Incentives and Bail-In Selection

In this appendix we provide the proof of Lemma 3.5 and formalize the discussion in Section 3.4 how the regulator can select among multiple accepting equilibria. We begin with the following auxiliary result, formalizing that a bank is better off rejecting a bail-in proposal if its participation is not needed for the regulator to proceed with the bail-in. It will be convenient to denote by $V^i(b, s, a) = V^i(\bar{p}(b, s, a), \bar{\ell}(b, s, a), \bar{\alpha}(b, s, a))$ bank i 's value of equity in the bail-in (b, s, a) . Similarly, let $V^i(s) = V^i(\bar{p}(s), \bar{\ell}(s), \bar{\alpha}(s))$ denote bank i 's value of equity in the bailout with subsidies s .

Lemma B.5. *Fix a feasible bail-in proposal (b, s) with $b^i > 0$ for some bank i . For any response a^{-i} , we have $V^j(b, s, (0, a^{-i})) \geq V^j(b, s, (1, a^{-i}))$ for any bank j with strict inequality if $j = i$.*

Proof. The two financial systems resulting from $(b, s, (0, a^{-i}))$ and $(b, s, (1, a^{-i}))$ are identical up to the financial commitments by bank i , which are larger by b^i in the latter system. The result thus follows from Statement 3 in Lemma E.3 of the online appendix and monotonicity in Lemma B.2. $\square$

Proof of Lemma 3.5. Fix a feasible proposal (b, s) with accepting equilibrium response a . We first show necessity of the stated conditions. To this end, fix a bank i with $b^i > 0$ and suppose towards a contradiction that $a^i = 1$ but at least one of the two conditions is violated. Suppose first that Condition 1 is violated. Then the regulator proceeds with a bail-in even if bank i rejects the proposal, hence subsidies are the same under a and $(0, a^{-i})$. It follows from Lemma B.5 that bank i is strictly better off under $(0, a^{-i})$, contradicting the assumption that a is an equilibrium. Suppose now that Condition 1 holds but Condition 2 is violated. Then a rejection by i leads to a default cascade without intervention. Feasibility of (b, s) implies $b^i - s^i \leq \eta^i(\alpha_1, \ell^i(L, \alpha_1))$, where we abbreviate $\alpha_1 = \alpha(b, s, (1, a^{-i}))$. Thus, for Condition 2 to be violated, we must have

$$b^i - s^i > IC^i(\alpha_1, \ell^i(L, \alpha_1)), \quad (32)$$

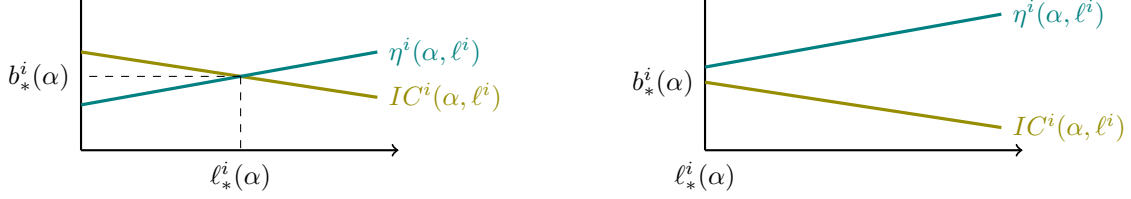


Figure 8: Relation of incentive constraint $IC^i(\alpha, \ell^i)$, budget constraint $\eta^i(\alpha, \ell^i)$, $\ell_*^i(\alpha)$, and $b_*^i(\alpha)$.

where we denote $IC^i(\alpha, \ell^i) := (\pi(L - p_N))^i + (1 - \alpha_N)\ell_N^i - (1 - \alpha)\ell^i$. A straightforward calculation shows that (32) is equivalent to $V_N^i > V^i(b, s, (1, a^{-i}))$, hence bank i has an incentive to deviate.

For sufficiency, note that Condition 2 implies $\ell^i(L, \alpha_1) \leq \ell_*^i(\alpha_1)$ since asset liquidation is monotonic in bank i 's net contribution $b^i - s^i$. This implies $b^i - s^i \leq b_*^i(\alpha) \leq IC^i(\alpha_1, \ell^i(L, \alpha_1))$ as illustrated in Figure 8, which is equivalent to $V^i(b, s, (1, a^{-i})) \geq V_N^i$. By Condition 1, the regulator will choose to not intervene if bank i rejects the bail-in, hence V_N^i is i 's value for rejecting the proposal. It is thus optimal for bank i to accept the proposal. $\square$

The next two results formalize the discussion at the end of Section 3.4, stating that in a subgame Pareto-efficient equilibrium, we can focus on bail-in proposals that can be accepted by every bank.

Lemma B.6. *Let (b, s) be a bail-in with accepting equilibrium responses $\{a_1, \dots, a_m\}$. For any $k = 1, \dots, m$, there exists a proposal $(\tilde{b}, \tilde{s})$ with $W_\lambda(b, s, a_k) = W_\lambda(\tilde{b}, \tilde{s}, 1)$, to which $1 = (1, \dots, 1)$ is the unique accepting equilibrium response.*

Proof. Fix a bail-in (b, s) with an accepting equilibrium response a_k . Let $\mathcal{B} = \{i \mid b^i 1_{\{a_k^i=1\}}\}$ denote the set of banks with a positive contribution in (b, s, a_k) . Define a bail-in $(\tilde{b}, \tilde{s})$ by setting $\tilde{b}^i = b^i 1_{\{i \in \mathcal{B}\}}$ and $\tilde{s}^i = s^i$ for $i = 1, \dots, n$. We will show that $1 = (1, \dots, 1)$ is an accepting equilibrium response to $(\tilde{b}, \tilde{s})$. Note that, by convention, any bank $i \in \mathcal{B}^c$ accepts the proposal; see Footnote 15. Moreover, Lemma 3.5 shows that the stated Conditions 1 and 2 are satisfied in (b, s, a_k) for any $i \in \mathcal{B}$ since a_k is an accepting equilibrium response. Because each bank makes the same contribution in $(\tilde{b}, \tilde{s}, 1)$ as in (b, s, a_k) , it follows that $\bar{\alpha}(\tilde{b}, \tilde{s}, 1) = \bar{\alpha}(b, s, a_k)$, $W_\lambda(b, s, 1) = W_\lambda(b, s, a_k)$, and $W_\lambda(b, s, (0, 1^{-i})) = W_\lambda(b, s, (0, a_k^{-i}))$ for each $i \in \mathcal{B}$. Therefore, Conditions 1 and 2 of Lemma 3.5 are satisfied also in $(\tilde{b}, \tilde{s})$ for any $i \in \mathcal{B}$. It follows from Lemma 3.5 that 1 is an accepting equilibrium response to $(\tilde{b}, \tilde{s})$. Uniqueness follows because the regulator will not proceed with the bail-in if only a proper subset of $\mathcal{B}$ accepts the proposal due to Condition 1 of Lemma 3.5. $\square$

Lemma B.7. *Suppose that a proposal (b, s) admits at least one accepting equilibrium. Among all continuation equilibria, welfare losses are minimized in an accepting equilibrium and at least one of the welfare-minimizing accepting equilibria is subgame Pareto efficient. Moreover, a rejecting equilibrium is subgame Pareto efficient if and only if (b, s) is the bailout given by Lemma 3.2.*

Proof. Fix a complete bail-in proposal (b, s) with accepting equilibrium response a . It follows from Lemma 3.5 that the value of any bank with $b^i > 0$ is at least as high in a as in a rejecting equilibrium response. Since $W_\lambda(b, s, a) \leq W_N \leq W_P$ by definition of an accepting equilibrium,

no rejecting equilibrium can subgame Pareto dominate a and a rejecting equilibrium is subgame Pareto efficient only if it is equivalent to (b, s, a) . Since there are only two possible outcomes in a rejecting equilibrium (public bailout and no rescue), the complete bail-in proposal (b, s) has to coincide with the complete bailout as, by definition, it rescues every bank in the system.

Let A denote the set of accepting equilibria that minimize welfare losses. We show that any $a_* \in \arg \max_{a \in A} \bar{\alpha}(b, s, a)$ is subgame Pareto efficient. Suppose towards a contradiction that some a Pareto dominates a_* , which requires $W_\lambda(b, s, a) \leq W_\lambda(b, s, a_*)$. Since rejecting equilibria cannot Pareto dominate a_* , it follows that $a \in A$. Let $\mathcal{C}$ and $\mathcal{C}_*$ denote the set of banks with positive contributions in a and a_* , respectively. By Condition 1 of Lemma 3.5, the regulator rejects the bail-in if a strict subset of $\mathcal{C}_*$ accepts the proposal, hence $\mathcal{C} \setminus \mathcal{C}_* \neq \emptyset$. Since $\bar{\alpha}(b, s, a) \leq \bar{\alpha}(b, s, a_*)$ by maximality of a_* in A , each bank in $\mathcal{C} \setminus \mathcal{C}_*$ is strictly worse off in a than in a_* , a contradiction. $\square$

B.4 Credibility and Existence of Subgame Pareto-Efficient Equilibria

In this appendix we prove Lemma 2.3, which establishes existence of subgame Pareto-efficient continuation equilibria after the proposal of any bail-in. In order to present the proofs as succinctly as possible, we invoke Lemma 3.3 to deal with the case when the threat fails to be credible. The proof of Lemma 3.3 does not rely on Lemma 2.3: existence of equilibria when the threat fails to be credible follows directly from the existence of strictly dominant strategies.

Proof of Lemma 3.3. Fix a feasible proposal (b, s) and any response vector a . Because the threat fails to be credible, the regulator will never respond with “no intervention”. Since any bank i with $b^i - s^i > 0$ is strictly worse off in (b, s) than in a complete rescue without its participation (e.g., bailout or bail-in by residual consortium) by Lemma B.5, rejection is the strictly dominant action for bank i . Thus, an optimal bailout is the only possible equilibrium outcome by Lemma 3.2. $\square$

Proof of Lemma 2.3. If the threat fails to be credible, Lemma 3.3 establishes that any rejecting equilibrium after any proposal (b, s) is subgame Pareto efficient. Suppose, therefore, that the threat is credible. Fix a feasible proposal (b, s) and let 0 denote the vector of unanimous rejections. The credibility of the threat imposes that $W_N \leq W_P \leq W_\lambda(b, s, 0)$, where we have used that W_P are the lowest-possible welfare losses without contributions by banks. The regulator thus chooses $r(b, s, 0) = \text{“no intervention”}$. If no bank in $\mathcal{A}(b)$ has a profitable deviation, then 0 is rejecting equilibrium response. Suppose, therefore, that some bank $i \in \mathcal{A}(b)$ has a profitable deviation and let a_i denote the corresponding action profile. We will show that a_i is an equilibrium response.

For a_i to be a profitable deviation for bank i , two conditions must hold. First, it is necessary that $r(b, s, a_i) = \text{“bail-in”}$ as otherwise, bank i ’s value of equity would be equal to V_N^i both when i accepts and when i rejects the proposal. Second, we must have that $V^i(b, s, a_i) > V_N^i$. The latter condition is equivalent to Condition 2 of Lemma 3.5 for bank i . Since, by construction, $(0, a_i^{-i})$ is the vector of unanimous rejections, it follows that $W_\lambda(b, s, (0, a_i^{-i})) = W_\lambda(b, s, 0) \geq W_N$, i.e., Condition 1 of Lemma 3.5 is satisfied for bank i as well.

To conclude that a_i is an accepting equilibrium response, it is sufficient to show that Condition 1 of Lemma 3.5 is violated for any bank $j \neq i$. Indeed, since $(0, a_i^{-j}) = a_i$ and $r(b, s, a_i) = \text{"bail-in"}$, it follows that $W_\lambda(b, s, (0, a_i^{-j})) = W_\lambda(b, s, a_i) < W_N$, thereby concluding the proof. $\square$

We conclude this appendix by proving that the regulator's threat is credible if and only if the amplification of the shock through the network is below the threshold given in Lemma 3.4.

Proof of Lemma 3.4. Observe first that $(L^i + w^i - c^i - e^i - (\pi L)^i)^+ = (s_0^i - e^i)^+ = s_0^i - \min(e^i, s_0^i)$. It follows from (28) and the definitions of S_0 and S_N that welfare losses without intervention equal

$$W_N = S_N - S_0 + \sum_{i=1}^n \min(e^i, s_0^i) + \lambda \sum_{i \in \mathcal{D}(p_N, \ell_N, \alpha_N)} \delta^i(p_N, \alpha_N).$$

Lemma 3.2 shows that $W_P = \lambda S_0 + g(\alpha_P) + \lambda \sum_{i \in \mathcal{D}(p_N, \ell_N, \alpha_N)} \delta^i(p_N, \alpha_N)$. Solving the inequality $W_N - W_P \leq 0$ for $S_N - S_0$ using the above expressions for W_N and W_P , we obtain (11). $\square$

C Proof of Theorem 3.8

We start by showing that without loss of generality, we can restrict our attention to bail-ins, in which each bank either makes a contribution or receives a subsidy. Moreover, because the regulator can anticipate the banks' responses, we may also restrict our attention unanimous accepting equilibria.

Lemma C.1. *For any bail-in (b, s) with accepting equilibrium response a , there exists a proposal $(\tilde{b}, \tilde{s})$ with accepting equilibrium response 1 such that $\tilde{b}^i \tilde{s}^i = 0$ and $W_\lambda(\tilde{b}, \tilde{s}, 1) = W_\lambda(b, s, a)$.*

Proof. Fix a bail-in proposal (b, s) with accepting equilibrium response a . The existence of such a implies via Lemma B.7 that either (b, s) is the public bailout of Lemma 3.2 or that the threat is credible. In the former case, the statement holds trivially, hence suppose that the threat is credible. Denote by $\mathcal{C} = \{b^i - s^i > 0, a^i = 1\}$ the set of banks which make a positive net contribution. Define the proposal $(\tilde{b}, \tilde{s})$ by setting $\tilde{b}^i = (b^i 1_{\{a^i=1\}} - s^i)^+$ and $\tilde{s}^i = (s^i - b^i 1_{\{a^i=1\}})^+$, and set $\tilde{a} = (1, \dots, 1)$. It follows straight from the construction of $(\tilde{b}, \tilde{s})$ that $\tilde{b}^i 1_{\{\tilde{a}^i=1\}} - \tilde{s}^i = b^i 1_{\{a^i=1\}} - s^i$ for any bank i , hence each bank's net contribution remains unchanged. This implies that clearing equilibria coincide both in $(\tilde{b}, \tilde{s}, \tilde{a})$ and (b, s, a) and also in $(\tilde{b}, \tilde{s}, (0, \tilde{a}^{-i}))$ and $(b, s, (0, a^{-i}))$ for any $i \in \mathcal{C}$. It follows from (7) that $W_\lambda(\tilde{b}, \tilde{s}, \tilde{a}) = W_\lambda(b, s, a)$ and $W_\lambda(\tilde{b}, \tilde{s}, (0, \tilde{a}^{-i})) = W_\lambda(b, s, (0, a^{-i}))$ for any $i \in \mathcal{C}$. Since a is an accepting equilibrium response to (b, s) , Lemma 3.5 implies that for every bank $i \in \mathcal{C}$,

$$\tilde{b}^i - \tilde{s}^i = b^i - s^i \leq \sum_{j=1}^n \pi^{ij} (\bar{p}^j(b, s, a) - p_N^j) + (1 - \alpha_N) \ell_N^i - (1 - \bar{\alpha}(b, s, a)) \ell^i(L, \bar{\alpha}(b, s, a)).$$

Because the clearing equilibria under $(\tilde{b}, \tilde{s}, \tilde{a})$ and (b, s, a) coincide, this shows that Condition 2 of Lemma 3.5 is satisfied in $(\tilde{b}, \tilde{s}, \tilde{a})$. In particular, $\tilde{a}$ is an accepting equilibrium response to $(\tilde{b}, \tilde{s})$. $\square$

Because every bank is rescued in a complete feasible bail-in, the shortfall of each bank i before liquidation is at most e^i . The liquidated amount by any bank i is thus inversely proportional to the asset recovery rate. This imposes a lower bound on the asset recovery rate of $\frac{1}{e} = \exp(-1)$.

Lemma C.2. *Let (b, s) be a complete feasible bail-in proposal. In any response a , each bank i liquidates $\bar{\ell}^i(b, s, a) = \frac{1}{\alpha}(L^i + w^i + b^i 1_{\{a^i=1\}} - c^i - s^i - (\pi L)^i)^+$ and $\bar{\alpha}(b, s, a) \geq \frac{1}{e}$.*

Proof. Since (b, s) is a complete rescue, Lemma 2.2 implies that $(\bar{\alpha}(b, s, a), \bar{\ell}(b, s, a))$ is the solution to (1) and (2) for $p = L$ with the largest asset recovery rate. By feasibility,

$$L^i + w^i + b^i 1_{\{a^i=1\}} - c^i - s^i - (\pi L)^i \leq \bar{\alpha}(b, s, 1)e^i \leq \bar{\alpha}(b, s, a)e^i.$$

This shows that $\bar{\ell}^i(b, s, a)$ is indeed of the desired form. Therefore, $\bar{\alpha}(b, s, a)$ is a fixed point of the function $f_{x,y}$ in Lemma E.2 for $y = 0$ and $x = \gamma \sum_{i=1}^n (L^i + w^i + b^i 1_{\{a^i=1\}} - c^i - s^i - (\pi L)^i)^+$. Since $\bar{\alpha}(b, s, a)$ is the greatest fixed point of f on $(0, 1]$, Lemma E.2 implies that $\bar{\alpha}(b, s, a) \geq \frac{1}{e}$. $\square$

Due to Lemma C.1, we may restrict attention to bail-in proposals (b, s) , in which $b^i s^i = 0$ for every bank i . Then, liquidation and welfare losses take a simple form as stated in Lemma 3.6.

Proof of Lemma 3.6. Fix a bail-in (b, s) with an accepting equilibrium response a such that $b^i s^i = 0$ for every bank i . Denote $\alpha = \bar{\alpha}(b, s, a)$ for the sake of brevity. Lemma C.2 shows that each bank i liquidates an amount $\bar{\ell}^i(b, s, a) = \frac{1}{\alpha}(s_0^i - b_0^i + b^i 1_{\{a^i=1\}} - s^i)^+$. Using that $b^i s^i = 0$ by assumption, $b_0^i s_0^i = 0$ by definition, and the elementary identity $\min(x, y) = x - (x - y)^+$, it follows that

$$\bar{\ell}^i(b, s, a) = \frac{1}{\alpha}(s_0^i - \min(s^i, s_0^i) + (b^i - \min(b^i, b_0^i))1_{\{a^i=1\}}). \quad (33)$$

Equation (1) implies that liquidation losses in the bail-in are equal to $-(1 - \alpha) \ln(\alpha)/\gamma$. It follows from (1), (33), and $\min(s^i, s_0^i) = s^i - (s^i - s_0^i)^+$ that net subsidies in the bail-in amount to

$$\sum_{i=1}^n (s^i - b^i 1_{\{a^i=1\}}) = \sum_{i=1}^n (s_0^i + (s^i - s_0^i)^+ - \min(b^i, b_0^i)1_{\{a^i=1\}}) + \frac{\alpha \ln(\alpha)}{\gamma}.$$

The result now follows from (7), Lemmas 3.1 and 3.2, and the specific form of g . $\square$

Lemma 3.7 states that among all individually incentive-compatible bail-ins with contributing banks in $\mathcal{C}$ that induce asset recovery rate α , welfare is maximized for bail-ins with contributions $\eta^i(\alpha, \ell)$ by banks $i \in \mathcal{C}$ for any vector ℓ of asset liquidation that induces asset recovery rate α .

Proof of Lemma 3.7. Fix a complete feasible bail-in proposal (b, s) with accepting equilibrium response a . By Lemma C.1 we may assume that $b^i s^i = 0$ and $a^i = 1$ for each bank i . Since (b, s) is a complete rescue, it follows that $s \geq s(\alpha, e)$, where we abbreviate $\alpha = \bar{\alpha}(b, s, 1)$. Condition 2 of Lemma 3.5 and feasibility imply that $b^i \leq b_*^i(\alpha)$ for each $i \in \mathcal{C}$. Define the bail-in $(\tilde{b}, \tilde{s})$ by setting $\tilde{b}^i = \max(b^i, \underline{b}^i)$ for each bank $i \in \mathcal{C}$ and $\tilde{s}^i = \min(s^i, s_0^i)$ for each bank $i \notin \mathcal{C}$. We first show that

$$\min(\tilde{b}^i, b_0^i) = \underline{b}^i. \quad (34)$$

If $\underline{b}^i = (\pi(L - p_N))^i + (1 - \alpha_N)\ell_N^i$, then $b^i \leq b_*^i(\alpha) = \underline{b}^i \leq b_0^i$ and hence (34) follows. If $\underline{b}^i = b_0^i$ instead, then $\min(\tilde{b}^i, b_0^i) = b_0^i = \underline{b}^i$ holds as well. Because subsidies beyond s_0 and contributions below $\underline{b}$ do

not prevent or require liquidation, it follows that $\bar{\alpha}(\tilde{b}, \tilde{s}, 1) = \alpha$. Therefore, Lemma 3.6 implies that subsidies beyond s_0^i and contributions below $\underline{b}^i$ are welfare decreasing, i.e., $W_\lambda(b, s, 1) \geq W_\lambda(\tilde{b}, \tilde{s}, 1)$. Applying Lemma 3.6 to the proposal $(\tilde{b}, \tilde{s})$ and (34) yields

$$W_\lambda(b, s, 1) \geq W_\lambda(\tilde{b}, \tilde{s}, 1) = W_P - g(\alpha_P) + g(\alpha) - \lambda \sum_{i \in \mathcal{C}} \underline{b}^i. \quad (35)$$

This shows the first statement. The final statement follows by observing that the inequality in (35) holds with equality precisely if $\underline{b}^i \leq b^i$ for each bank $i \in \mathcal{C}$ and $s^i \leq s_0^i$ for each bank $i \notin \mathcal{C}$. $\square$

Next, we show that $\chi_{\mathcal{C}}(\alpha)$ in Lemma A.1 is well-defined and that any incentive-compatible bail-in with contributing banks in $\mathcal{C}$ and asset recovery rate α burns at least $\chi_{\mathcal{C}}(\alpha)$ units of welfare.

Proof of Lemma A.1. Fix $\alpha \geq \alpha_{\text{ind}}$. We start by showing that (26) has a unique non-negative solution if one exists. Lemma B.4 shows that g is increasing and hence invertible on the interval $[\alpha_{\text{ind}}, \infty)$. Let g^{-1} denote the inverse on $[\alpha_{\text{ind}}, \infty)$ and define the function $\hat{\alpha}(x) := g^{-1}(x + g(\alpha))$ for $x \geq 0$. It is easy to check that $\hat{\alpha}(x) \geq \alpha$ and $g_\alpha^{-1}(x) = \frac{1}{\gamma}(z(\hat{\alpha}(x)) - z(\alpha))$ for any $x \geq 0$. Moreover, for $x > 0$, we get $\hat{\alpha}(x) > \alpha$, hence the formula for the inverse of the derivative implies that

$$\hat{\alpha}'(x) = \frac{1}{g'(g^{-1}(x + g(\alpha)))} = \frac{1}{g'(\hat{\alpha}(x))} > 0,$$

where we have used that $\hat{\alpha}(x) > \alpha \geq \alpha_{\text{ind}}$. Since $\hat{\alpha}(x) > \alpha_{\text{ind}} \geq \frac{1}{e}$ and z is increasing on $[\frac{1}{e}, \infty)$, it follows from the chain rule that $(g_\alpha^{-1})'(x) = \frac{1}{\gamma} z'(\hat{\alpha}(x)) \hat{\alpha}'(x) > 0$. Let $f(\chi)$ denote the right-hand side of (26). Since g_α^{-1} is strictly increasing, f is strictly decreasing where it is positive. Thus, there exists a unique non-negative solution if $f(0) \geq -\alpha \ln(\alpha)/\gamma$ and there exists no non-negative solution otherwise. The fact that $\chi_{\mathcal{C}}(\alpha)$ is a lower bound for welfare burning in an accepting bail-in (b, s) with contributing banks $\mathcal{C}$ and $\alpha = \bar{\alpha}(b, s, 1)$ now follows from Lemma C.3 below. $\square$

Lemmas 3.6 and 3.7 together imply that any bail-in satisfying Conditions (i)–(iv) in Definition 3.2, if accepted, induces welfare losses

$$W_{\mathcal{C}}(\alpha) := W_P - g(\alpha_P) + g(\alpha) - \lambda \sum_{i \in \mathcal{C}} \underline{b}^i + \chi_{\mathcal{C}}(\alpha).$$

The following lemma establishes that this is a lower bound for welfare losses that can be attained by a bail-in with contributing banks $\mathcal{C}$ that induces asset recovery rate α .

Lemma C.3. *Let (b, s) be a complete feasible bail-in proposal with accepting equilibrium response a . Let $\mathcal{C} := \{i \mid b^i 1_{\{a^i=1\}} > 0\}$ and denote $\alpha = \bar{\alpha}(b, s, a)$ for the sake of brevity. Then $W_\lambda(b, s, a) \geq W_{\mathcal{C}}(\alpha)$ and the inequality binds if and only if $(b, s) \in \Xi(\mathcal{C}, \alpha)$.*

Proof. Let (b, s) be a complete feasible bail-in proposal with accepting equilibrium response a . By Lemma C.1, we may assume without loss of generality that $b^i s^i = 0$ and $a^i = 1$ for every bank i . It follows in the same way as in the proof of Lemma 3.7 that $b^i \leq b_*^i(\alpha)$ for each $i \in \mathcal{C}$ and $s^i \geq s_L^i$ for each $i \notin \mathcal{C}$. Therefore, Conditions 1 and 2 of Definition 3.2 are satisfied.

Define the bail-in proposal $(\tilde{b}, \tilde{s})$ by setting $\tilde{b}^i = \max(b^i, \underline{b}^i)$ for each $i \in \mathcal{C}$ and $\tilde{s}^i = \min(s^i, s_0^i)$ for each $i \notin \mathcal{C}$. Let us denote $\chi := \lambda \sum_{i=1}^n (s^i - s_0^i)^+ + \lambda \sum_{i \in \mathcal{C}} (\underline{b}^i - b^i)^+$. It follows as in the proof of Lemma 3.7 that $\bar{\alpha}(\tilde{b}, \tilde{s}, 1) = \alpha$. Together with (34) and Lemma 3.7, this yields

$$W_\lambda(b, s, 1) = W_\lambda(\tilde{b}, \tilde{s}, 1) + \chi = W_P - g(\alpha_P) + g(\alpha) - \lambda \sum_{i \in \mathcal{C}} \underline{b}^i + \chi.$$

To show the first statement, it thus remains to show that $\chi \geq \chi_{\mathcal{C}}(\alpha)$.

For any bank $i \in \mathcal{C}$, let a_{-i} denote the response vector in which every bank but bank i accepts the proposal and set $\alpha_{-i} = \bar{\alpha}(b, s, a_{-i})$. By Lemma 3.6, welfare losses in this response are equal to

$$\begin{aligned} W_\lambda(b, s, a_{-i}) &= W_\lambda(b, s, 1) - g(\alpha) + g(\alpha_{-i}) + \lambda \min(b^i, b_0^i) \\ &= W_P - g(\alpha_P) + g(\alpha_{-i}) - \lambda \sum_{j \in \mathcal{C}} \underline{b}^j + \lambda \min(b^i, b_0^i) + \chi. \end{aligned}$$

Solving this equation for χ , applying Condition 1 of Lemma 3.5, and using the fact that (34) implies $\min(b^i, b_0^i) \leq \underline{b}^i$, we obtain a lower bound for χ given by

$$\chi \geq W_N - W_P + g(\alpha_P) - g(\alpha_{-i}) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j. \quad (36)$$

Equation (33) implies that for any accepting equilibrium response a ,

$$\bar{\ell}^j(b, s, a) = \frac{1}{\bar{\alpha}(b, s, a)} ((s_0^j - s^j)^+ + (b^j - b_0^j)^+ 1_{\{a^j=1\}}). \quad (37)$$

Since $s_0^j = 0$ for $j \in \mathcal{C}$ and $b^j = 0$ for $j \notin \mathcal{C}$, it follows from (1) that

$$-\frac{\alpha_{-i} \ln(\alpha_{-i})}{\gamma} = \sum_{j \in \mathcal{C} \setminus \{i\}} (b^j - b_0^j)^+ + \sum_{j \notin \mathcal{C}} (s_0^j - s^j)^+ = -\frac{\alpha \ln(\alpha)}{\gamma} - (b^i - b_0^i)^+. \quad (38)$$

Recall that z^{-1} denotes the inverse of $z(\alpha) = \alpha \ln(\alpha)$ on $[\frac{1}{e}, \infty)$. Multiplying (38) by $-\gamma$ and applying $g \circ z^{-1}$, we obtain $g(\alpha_{-i}) - g(\alpha) = g_\alpha((b^i - b_0^i)^+)$. In conjunction with (36), this yields

$$\chi \geq W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j - g_\alpha(b_*^i(\alpha) - \underline{b}^i), \quad (39)$$

where we have used that g_α is increasing. Note that $g(\alpha_{-i}) - g(\alpha) = g_\alpha((b^i - b_0^i)^+) \geq 0$. Therefore, solving (36) for $g(\alpha_{-i}) - g(\alpha)$, taking the maximum with 0, and applying g_α^{-1} yields

$$(b^i - b_0^i)^+ \geq g_\alpha^{-1} \left(\left(W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j - \chi \right)^+ \right). \quad (40)$$

Summing (40) over all $i \in \mathcal{C}$ yields

$$\sum_{i \in \mathcal{C}} g_{\alpha}^{-1} \left(\left(W_N - W_P + g(\alpha_P) - g(\alpha) + \lambda \sum_{j \in \mathcal{C} \setminus \{i\}} \underline{b}^j - \chi \right)^+ \right) \leq \sum_{i \in \mathcal{C}} (b^i - b_0^i)^+ \leq -\frac{\alpha \ln(\alpha)}{\gamma}, \quad (41)$$

where we have used (38) in the second inequality. Since $\chi_{\mathcal{C}}(\alpha)$ is the smallest value $\chi' \geq 0$ that satisfies (39) for all $i \in \mathcal{C}$ and (41), it follows that $\chi_{\mathcal{C}}(\alpha) \leq \chi$. This shows $W_{\lambda}(b, s, a) \geq W_{\mathcal{C}}(\alpha)$.

This lower bound holds with equality if and only if $\chi = \chi_{\mathcal{C}}(\alpha)$, i.e., Condition (iv) in Definition 3.2 is satisfied. Summing (37) over all banks shows that Conditions (iii) holds as well. Finally, Condition 1 of Lemma 3.5 for $i \in \mathcal{C}$ implies that Condition (v) holds. This concludes the proof. $\square$

Lemma C.3 shows that welfare losses $W_{\mathcal{C}}(\alpha)$ are attained only by bail-ins in $\Xi(\mathcal{C}, \alpha)$. The following lemma shows that the converse is true as well if $\alpha \geq \frac{1}{e}$, which is satisfied by all Pareto-efficient clearing equilibria; see Lemma E.2 in the online appendix.

Lemma C.4. *For any $\mathcal{C}$ and $\alpha \geq \frac{1}{e}$, any $(b, s) \in \Xi(\mathcal{C}, \alpha)$ is a complete feasible bail-in proposal with $W_{\lambda}(b, s, 1) = W_{\mathcal{C}}(\alpha)$ such that $1 = (1, \dots, 1)$ is an accepting equilibrium response if $W_{\mathcal{C}}(\alpha) < W_N$.*

Proof. Fix $(b, s) \in \Xi(\mathcal{C}, \alpha)$. It follows along the same lines as in the proof of Lemma F.1 that (b, s) is a complete feasible bail-in with $\bar{\alpha}(b, s, 1) = \alpha$. Condition (i) in Definition 3.2 implies that Condition 2 in Lemma 3.5 is satisfied for every bank $i \in \mathcal{C}$ in the response vector 1. It follows from Lemma 3.6 and Condition (iv) in Definition 3.2 that $W_{\lambda}(b, s, 1) = W_{\mathcal{C}}(\alpha)$. Condition (v) in Definition 3.2 thus implies that Condition 1 of Lemma 3.5 is satisfied for every bank $i \in \mathcal{C}$. Therefore, an application of Lemma 3.5 shows that 1 is an accepting equilibrium response if $W_{\mathcal{C}}(\alpha) < W_N$. $\square$

Lemma C.5 shows that the equilibrium bail-in contributors are the banks with the largest exposure to contagion effects.

Lemma C.5. *For any bail-in proposal (b, s) , let $\ell(b, s)$ denote the induced vector of liquidation decisions when every bank accepts the proposal. For any vector ℓ , let $\mathcal{C}(\ell)$ be defined as in Theorem 3.8. Suppose that there exists $\mathcal{C}$ such that in any subgame Pareto-efficient equilibrium, a bail-in from $\Xi(\mathcal{C}, \alpha)$ is implemented. Then $\mathcal{C} = \mathcal{C}(\ell(b, s))$ for any $(b, s) \in \Xi(\mathcal{C}, \alpha)$.*

Proof. Let $\mathcal{C}$ be such that in any subgame Pareto-efficient equilibrium, a bail-in from $\Xi(\mathcal{C}, \alpha)$ is implemented. Fix $(b, s) \in \Xi(\mathcal{C}, \alpha)$ and abbreviate $\ell = \ell(b, s)$. Suppose towards a contradiction that there exists a pair of banks $(i_0, j_0) \in \mathcal{C}^c \times \mathcal{C}$ such that $\eta^{i_0}(\alpha, \ell) > \eta^{j_0}(\alpha, \ell)$. Let $\tilde{\mathcal{C}} := \mathcal{C} \cup \{i_0\} \setminus \{j_0\}$ and define a bail-in $(\tilde{b}, \tilde{s})$ by setting $\tilde{b}^{i_0} = \eta^{i_0}(\alpha, \ell)$, $\tilde{s}^{i_0} = 0$, $\tilde{b}^{j_0} = 0$, $\tilde{s}^{j_0} = 0$, as well as $\tilde{b}^i = b^i$ and $\tilde{s}^i = s^i$ for any other bank $i \notin \{j_0\} \cup \{i_0\}$.

We will first show that $(\tilde{b}, \tilde{s})$ is a complete, feasible bail-in proposal. By definition of ℓ , the shortfall of any bank i in the bail-in (b, s) is $\alpha \ell^i$. Let $\tilde{\sigma}^i := (L^i + w^i + \tilde{b}^i - c^i - \tilde{s}^i - (\pi L)^i)^+$ denote the shortfall of bank i in bail-in $(\tilde{b}, \tilde{s})$. Since $\tilde{b}^i = b^i$ and $\tilde{s}^i = s^i$ for every $i \notin \{i_0, j_0\}$, it follows that $\tilde{\sigma}^i = \alpha \ell^i$ for every such bank i . Since j_0 makes a positive net contribution $b^{j_0} - s^{j_0}$ to the bail-in (b, s) by Condition 1 of Lemma 3.5, it follows that $\tilde{\sigma}^{j_0} \leq (L^{j_0} + w^{j_0} + b^{j_0} - c^{j_0} -$

$s^{j_0} - (\pi L)^{j_0})^+ = \alpha \ell^{j_0}$ with strict inequality if $\ell^{j_0} > 0$. The definition of $\eta(\alpha, \ell)$ in (13) implies that $\tilde{b}^{i_0} \leq (c^{i_0} + \alpha \ell^{i_0} + (\pi L)^{i_0} - w^{i_0} - L^{i_0})^+$, and hence $\tilde{\sigma}^{i_0} \leq \alpha \ell^{i_0}$ also for bank i_0 . We conclude that $\tilde{\sigma}^i \leq \alpha \ell^i \leq e^i$ for every bank i , hence $(\tilde{b}, \tilde{s})$ is a complete feasible bail-in proposal.

Since the total shortfall is smaller in $(\tilde{b}, \tilde{s})$ than in (b, s) , it follows that $\tilde{\alpha} := \alpha(\tilde{b}, \tilde{s}, 1) \geq \alpha$. Observe further that $\eta^{i_0}(\alpha, \ell) > \eta^{j_0}(\alpha, \ell) \geq b^{j_0} > 0$ implies $s^{i_0}(\alpha, \ell) = 0$ and hence $s^{i_0} = 0 = \tilde{s}^{j_0}$. It also implies that $\tilde{b}^{i_0} > b^{j_0}$, which yields

$$W_\lambda(\tilde{b}, \tilde{s}, 1) = W_\lambda(b, s, 1) + \frac{1 - \tilde{\alpha}}{\tilde{\alpha}} \sum_{i=1}^n \tilde{\sigma}^i - \frac{1 - \alpha}{\alpha} \sum_{i=1}^n \alpha \ell^i - \lambda(\tilde{b}^{i_0} - b^{j_0}) < W_\lambda(b, s, 1),$$

where we have used that $\tilde{\sigma}^i \leq \alpha \ell^i$ for every bank i , that $x \mapsto (1-x)/x$ is positive and decreasing, and that $\tilde{\alpha} \geq \alpha$. Because 1 is an accepting equilibrium response for (b, s) , this shows $W_\lambda(\tilde{b}, \tilde{s}, 1) < W_N$.

For any $i \in \tilde{\mathcal{C}}$, let a_{-i} denote the response vector by the banks where every bank but i accepts the proposal. Observe that $W_\lambda(\tilde{b}, \tilde{s}, a_{-i_0}) = W_\lambda(b, s, a_{-j_0}) \geq W_N$ by Condition 1 of Lemma 3.5. For any $i \in \tilde{\mathcal{C}} \setminus \{i_0\}$, let $\tilde{\alpha}_{-i} = \tilde{\alpha}(\tilde{b}, \tilde{s}, a_{-i})$ and $\alpha_{-i} = \alpha(b, s, a_{-i})$ and observe that

$$-\frac{\tilde{\alpha}_{-i} \ln(\tilde{\alpha}_{-i})}{\gamma} + \frac{\alpha_{-i} \ln(\alpha_{-i})}{\gamma} = \sum_{j \neq i} \tilde{\sigma}^j - \sum_{j \neq i} \alpha \ell^j = \tilde{\sigma}^{i_0} + \tilde{\sigma}^{j_0} - \alpha(\ell^{i_0} + \ell^{j_0}) \leq 0. \quad (42)$$

Equation (42) implies that $\tilde{\alpha}_{-i} \geq \alpha_{-i}$. Moreover, since the difference in shortfall is the same as the difference of shortfalls between $(\tilde{b}, \tilde{s}, 1)$ and $(b, s, 1)$, it follows from concavity of $x \mapsto -x \ln(x)$ that $\tilde{\alpha}_{-i} - \alpha_{-i} \leq \tilde{\alpha} - \alpha$ and hence $\tilde{\alpha}_{-i} - \tilde{\alpha} \leq \alpha_{-i} - \alpha$. Convexity of $x \mapsto (1-x)/x$ thus yields

$$\frac{1 - \tilde{\alpha}}{\tilde{\alpha}} - \frac{1 - \tilde{\alpha}_{-i}}{\tilde{\alpha}_{-i}} \leq \frac{1 - \alpha}{\alpha} - \frac{1 - \alpha_{-i}}{\alpha_{-i}}.$$

Together with the fact that $\alpha \ell^j \geq \tilde{\sigma}^j$ for every bank j , this implies

$$W_\lambda(b, s, a_{-i}) - W_\lambda(\tilde{b}, \tilde{s}, a_{-i}) = \frac{1 - \alpha_{-i}}{\alpha_{-i}} \sum_{j \neq i} \alpha \ell^j - \frac{1 - \tilde{\alpha}_{-i}}{\tilde{\alpha}_{-i}} \sum_{j \neq i} \tilde{\sigma}^j \leq W_\lambda(b, s, 1) - W_\lambda(\tilde{b}, \tilde{s}, 1). \quad (43)$$

Define now a vector of subsidies $\hat{s} \geq \tilde{s}$ that burns additional welfare precisely equal to

$$\chi := \max_{i \in \tilde{\mathcal{C}}} (W_\lambda(b, s, a_{-i}) - W_\lambda(\tilde{b}, \tilde{s}, a_{-i})).$$

Condition 1 of Lemma 3.5 for bail-in (b, s) yields $W_\lambda(\tilde{b}, \hat{s}, a_{-i}) \geq W_\lambda(b, s, a_{-i}) \geq W_N$, showing that Condition 1 of Lemma 3.5 is also satisfied in $(\tilde{b}, \hat{s})$. Since $\alpha(\tilde{b}, \hat{s}, 1) \geq \tilde{\alpha} \geq \alpha$, it follows that also Condition 2 of Lemma 3.5 is satisfied in $(\tilde{b}, \hat{s})$ for every bank i . Finally, it follows from (43) that

$$W_\lambda(\tilde{b}, \hat{s}, 1) = W_\lambda(\tilde{b}, \tilde{s}, 1) + \chi \leq W_\lambda(b, s, 1) < W_N. \quad (44)$$

This shows that $a = (1, \dots, 1)$ is an accepting equilibrium response for $(\tilde{b}, \hat{s})$. Condition 2 of Lemma 3.5 and Lemma B.7 imply that it is the unique subgame Pareto efficient continuation

equilibrium. Since the regulator has no profitable deviations from (b, s) by assumption, (44) implies that $(\tilde{b}, \hat{s})$ is part of a subgame Pareto-efficient equilibrium as well. This contradicts the assumption that contributions have to come from banks in $\mathcal{C}$. Therefore, we conclude $\mathcal{C} = \mathcal{C}(\ell)$. $\square$

Proof of Theorem 3.8. If the threat fails to be credible, then it follows from Lemma 3.3 that the regulator will implement an optimal public bailout as given in Lemma 3.2.

Suppose, therefore, that the threat is credible. For any set of banks $\mathcal{C}$, Lemma D.1 in the online appendix characterizes the minimum asset recovery rate $\alpha_{\mathcal{C}}$ that can be sustained in a complete bail-in with contributing banks in $\mathcal{C}$. Let $A(\mathcal{C})$ denote the set of recovery rates $\alpha \in [\alpha_{\mathcal{C}}, 1]$ that minimize $W_{\mathcal{C}}(\alpha)$. Since $W_{\mathcal{C}}$ is continuous, $A(\mathcal{C})$ is non-empty. Lemma F.5 implies that there exists at least one set $\mathcal{C}$, for which $\min_{\alpha \in A(\mathcal{C})} W_{\mathcal{C}}(\alpha) < W_N$. For any such set $\mathcal{C}$ and any $\alpha \in A(\mathcal{C})$, Lemma C.4 implies that any proposal $(b, s) \in \Xi(\mathcal{C}, \alpha)$ admits an accepting equilibrium response $(1, \dots, 1)$ that attains welfare losses $W_{\mathcal{C}}(\alpha)$. Moreover, no bail-in with contributing banks in $\mathcal{C}$ can attain lower welfare losses by Lemma C.3. Condition 1 in Lemma 3.5 implies that $(1, \dots, 1)$ is the unique accepting equilibrium response, hence also the unique subgame Pareto-efficient equilibrium by Lemma B.7. Thus, the regulator must propose a bail-in $(b, s) \in \Xi(\mathcal{C}_*, \alpha)$ for $\alpha \in A(\mathcal{C}_*)$ and contributing banks $\mathcal{C}_*$ that minimizes welfare losses. We show in Theorem D.2 in the online appendix that only bail-ins from $\Xi(\mathcal{C}_*, \alpha(\mathcal{C}_*))$ for $\alpha(\mathcal{C}_*) = \max A(\mathcal{C}_*)$ are subgame Pareto efficient and that $\mathcal{C}_*$ is generically unique. Finally, Lemma C.5 shows that $\mathcal{C}_* = \mathcal{C}(\ell(b, s))$ for any $(b, s) \in \Xi(\mathcal{C}_*, \alpha_*)$ if $\mathcal{C}_*$ is unique. $\square$