

Bank Runs, Lender of Last Resort, and Liquidity Regulation*

Toni Ahnert[†]

Kartik Anand[‡]

Guillem Ordoñez-Calafi[§]

July 2026

Abstract

We study the liquidity choice of a bank subject to rollover risk and support from a lender of last resort (LLR) and study the consequences for bank stability, funding costs, and liquidity regulation. The liquidity choice balances forgoing profitable yet illiquid investment with fewer panic runs and cheaper debt. A higher LLR penalty rate increases bank liquidity and has a V-shaped effect on ex-ante bank stability. Turning to normative implications, the availability of the LLR *reduces* welfare for a large social cost of bank failure. The anticipation of ex-post support mitigates panic runs but induces lower bank liquidity ex ante, which increases funding costs and the frequency of bank failure. Liquidity regulation aligns private with social incentives and increases the social value of LLR support.

Keywords: Bank Runs, Rollover Risk, Lender of Last Resort, Liquidity Choice, Liquidity Regulation, Global Games.

JEL classifications: G01, G21, G28.

*We thank seminar participants at Deutsche Bundesbank, University of Nottingham, Bank of Canada, Rotman School of Management, and ECB for helpful comments. The views expressed in this paper are the authors' and do not necessarily reflect those of the ECB, Deutsche Bundesbank, or the Eurosystem.

[†] European Central Bank and CEPR; toni.ahnert@ecb.europa.eu

[‡] Deutsche Bundesbank; kartik.anand@bundesbank.de

[§] University of Bristol; g.ordonez-calafi@bristol.ac.uk

1 Introduction

Liquidity crises remain a defining feature of financial intermediation. During the U.S. banking turmoil of March 2023, for instance, banks experienced large outflows of uninsured deposits (Luck et al., 2023; Cipriani et al., 2024). The outflows subsided only after the Federal Reserve introduced emergency liquidity facilities, which policymakers argued were critical in restoring confidence and preventing disorderly contagion (Powell, 2023). The notion that central banks should step in during crises has a long historical pedigree, but its modern implementation is far from straightforward. While emergency support can stem panic in the short run, critics argue that generous liquidity provision risks weakening banks' incentives to hold sufficient reserves (Congressional Research Office, 2023). This raises fundamental questions: how does the lender of last resort (LLR) reshape a bank's incentives to self-insure with liquidity, and can this institutional design itself become a source of instability?

Despite the centrality of these questions in policy debates, theoretical foundations remain incomplete. Answering them requires a framework in which a bank's funding cost and liquidity are jointly determined. Holding more liquidity, for example, makes the bank more stable, which lowers the price of debt required by investors, which in turn shapes how much liquidity the bank wants to hold in the first place. This joint determination is precisely where the terms of a LLR become pivotal. Whether such support strengthens financial stability depends on how its anticipation interacts with liquidity regulation, such as minimum liquidity requirements (MLRs), and feeds into the bank's ex-ante choices. Thus, public liquidity provision, banks' private liquidity choices, and the disciplining role of MLRs are best understood not in isolation but within a single framework, as argued by a senior policymaker (Bowman, 2024).

Building on Rochet and Vives (2004) and Vives (2014), we develop a parsimonious model in which a risk-neutral bank raises demandable debt (uninsured deposits)

from competitive risk-neutral investors. The bank can make profitable but illiquid investment and hold liquidity. The bank is subject to rollover risk as funding may be withdrawn before the maturity of investment. Upon large interim withdrawals, the bank can use its liquidity, liquidate investment at a cost, or borrow from the LLR at a penalty rate. We take this penalty rate as well as restricting LLR support to solvent but illiquid banks as exogenous institutional parameters throughout the paper.

The model has a unique equilibrium in which rollover risk, debt pricing, and bank liquidity are jointly determined. We use the global-games methods (Carlsson and van Damme, 1993; Morris and Shin, 2003; Vives, 2005) to solve for the equilibrium at the rollover stage. Funding is rolled over whenever a private signal about the health of the bank’s balance sheet is sufficiently strong, and the bank fails due to a run whenever its realized balance sheet health falls below a threshold. A lower LLR penalty rate mitigates panic runs ex post and reduces this failure threshold. The threshold also depends on the liquidity choice and funding costs determined ex ante.

At the initial stage, the bank’s liquidity choice balances forgoing profitable investment with higher stability and cheaper funding. In terms of comparative statics, a higher LLR penalty rate induces the bank to hold more liquidity (higher self-insurance). A key feature of the equilibrium is that a more liquid bank is more stable ex ante, so investors break even for a lower price of debt. Consistent with this feature, Aldasoro et al. (2026) offer causal evidence that higher bank liquidity leads to cheaper wholesale funding (which are part of the uninsured deposits of our model).

A remarkable implication of the model is that LLR support, while stabilising ex post, can *destabilise* the bank ex ante. Bank stability is V-shaped in the LLR penalty rate because the penalty rate has two opposing effects. Holding the bank’s liquidity fixed, a lower penalty rate makes emergency borrowing cheaper and therefore strengthens bank stability. This is the direct stabilising effect shown by Rochet and Vives (2004). At the same time, a lower penalty rate makes reliance on the LLR more

attractive, inducing the bank to hold fewer liquid reserves. This self-insurance effect raises the probability of bank failure and increases the face value of debt required by investors. For low penalty rates, the bank forgoes holding liquidity altogether, so the direct effect dominates and stability falls. Once the penalty rate is high enough to induce positive liquidity, the self-insurance effect dominates and stability increases. The minimum of the V-shape occurs precisely when LLR borrowing is costly enough to weaken solvency, but not yet costly enough to induce sufficient private liquidity.

Our first normative result is that the availability of the LLR can *reduce* welfare. Although LLR support shrinks the range of shocks for which panic runs occur, it also weakens banks' incentives to self-insure ex ante. Anticipating LLR support, the profit-maximizing bank holds excessively low liquidity, which raises both funding costs and the probability of bank failure. When the social cost of bank failure is large (e.g. due to lost output, contagion across institutions, or fiscal burdens from emergency lending), this ex-ante moral-hazard distortion dominates the direct benefit of ex-post stabilisation. Therefore, our result speaks to the long-standing criticism of public liquidity provision: by reducing the bank's incentives to self-insure, the LLR undermines the very stability it was meant to protect.

Our second normative result is that liquidity regulation can remedy the distortions of the LLR. By obliging banks to hold a minimum of liquid reserves, the regulation closes the wedge between the private and social marginal benefits of liquidity. There are two channels at work. First, holding more liquidity lets the bank meet more interim withdrawals without relying on the LLR, which reduces withdrawal incentives. Second, by raising stability, the MLR lowers the face value of debt that investors demand. This, in turn, lowers the debt burden of potential LLR borrowing, which further improves stability and feeds back into a still-lower face value. This debt-pricing relief channel is, to the best of our knowledge, novel, and can be a key driver of the welfare gains. When acting in tandem, the MLR and the LLR strictly

improve upon welfare relative to what either tool could have achieved in isolation, and the welfare gap widens with the social cost of bank failure.

Our model distinguishes between the *strictness* and the *stringency* of liquidity regulation. Strictness refers to the level of the regulatory floor, while stringency refers to the distance between that floor and the bank’s unregulated liquidity choice. This distinction matters for both policy and empirical work. The model implies that (i) strictness increases in the LLR penalty rate, but (ii) stringency is non-monotonic. This suggests that the relevant measure of regulatory pressure is not simply how much liquidity a bank holds, but also how far regulation pushes the bank above the level chosen privately. In the cross-section, this wedge could be proxied using bank-level measures of exposure to liquidity regulation as in [Sundaresan and Xiao \(2024\)](#).

In sum, our paper identifies a novel micro-prudential rationale for liquidity regulation in that it raises the social value of the LLR and makes it unambiguously desirable. There is a familiar parallel in the classic literature on deposit insurance and capital regulation ([Kim and Santomero, 1988](#); [Keeley, 1990](#)). An ex-post back-stop that mitigates runs weakens the bank’s incentives on the asset side, prompting greater risk-taking, and an ex-ante capital requirement restores discipline. Our result shares the broad shape of this argument, but the mechanism is novel and focuses on bank liquidity. By abstracting from asset-side risk-shifting, we emphasize the bank’s liquidity choice and how it feeds into ex-ante bank stability and funding costs.

Literature. The paper is closely related to four strands of literature. First, the lender of last resort has been studied in [Thornton \(1802\)](#), [Bagehot \(1873\)](#), [Freixas and Parigi \(2008\)](#), and [Tucker \(2016\)](#). [Repullo \(2005\)](#) and [Ratnovski \(2009\)](#) study how ex-post LLR support affects the bank’s ex-ante liquidity and risk choices. By contrast, we endogenize the fragility on the bank’s liability side and focus on risk taking via illiquid investment.

Second, our model builds on [Rochet and Vives \(2004\)](#) and [Vives \(2014\)](#). [Rochet and Vives \(2004\)](#) examine the impact of the LLR on a bank subject to runs. Greater LLR leniency mitigates the strategic complementarity in withdrawal decisions and can eliminate panic runs altogether. [Vives \(2014\)](#) studies how prudential tools (solvency / liquidity requirements and disclosure rules) affect the probability of bank insolvency and illiquidity. Our contribution is to endogenize bank liquidity and the face value of short-term debt. This introduces a feedback from anticipated LLR support into banks’ ex-ante choices. While a more generous LLR reduces ex-post rollover risk, it also induces the bank to hold less liquidity ex ante, which decreases bank stability.

Third, a literature studies liquidity regulation as an ex-ante policy tool. [Santos and Suarez \(2019\)](#) view liquidity requirements as a device that “buys time” for the LLR to learn about bank quality. In particular, we share the view that liquidity regulation increases the social value of the LLR. Our paper is complementary, in that we emphasize the benefit of the LLR for debt pricing and ex-ante bank stability. We share with [Sundaresan and Xiao \(2024\)](#) a focus on liquidity buffers, who show that quantity-based rules reduce measured risk but can crowd out lending. In our model liquidity regulation mitigates the price impact that LLR support has on debt, so MLR acts as complements, rather than substitutes, for central-bank backstops.

Fourth, our paper is related to recent global-games papers on bank runs and policy interventions. [Huang et al. \(2024\)](#) study the joint design of deposit insurance and liquidity injections when the policymaker learns about fundamentals from observed withdrawals. They show that liquidity injection overshadows deposit insurance, so the optimal policy is to set deposit insurance to zero and tolerate more runs to improve allocative efficiency. [Schilling \(2025\)](#) examines emergency liquidity assistance and shows that delayed or poorly financed interventions can redistribute repayment burdens from early to late depositors. Both papers focus on policy timing and information, but treat debt pricing and bank liquidity as exogenous. In contrast, we study

how LLR design shapes the bank’s choice of liquidity and debt pricing, generating a new channel through which ex-post support influences ex-ante bank stability.

Structure. The remainder of the paper is organised as follows. Section 2 presents the model. Section 3 describes the equilibrium (in the absence of liquidity regulation) and the positive implications of the model. Section 4 contains the normative analysis, including the welfare implications of a stand-alone LLR and the impact of liquidity regulation. Section 5 concludes. All proofs are relegated to the appendix.

2 Model

We develop a parsimonious model of a bank subject to rollover risk in the presence of a lender of last resort. The model builds on [Rochet and Vives \(2004\)](#), [Vives \(2014\)](#), and [Ahnert et al. \(2019\)](#).

Agents, preferences, and endowments. There are three dates, $t = 0, 1, 2$. A single-good economy comprises a representative banker who owns and operates a bank and a unit mass of investors. All agents are risk-neutral. While the banker cares only about consuming at $t = 2$, investors are indifferent between consuming at $t = 1$ and $t = 2$. Each investor is endowed with one unit of the consumption good at $t = 0$, while the banker is penniless.

Investment opportunities. Investors have access to a risk-free technology that yields a gross return $r > 0$ at $t = 2$ per unit allocated at $t = 0$. The bank has access to illiquid investments that yield a gross return $R > r$ at $t = 2$ per unit invested at $t = 0$. However, liquidating investments at $t = 1$ yields only a fraction $\psi \in (0, 1)$ of the return. The bank can also hold liquid reserves L at $t = 0$, e.g. central bank reserves or government bonds. We normalise the return on liquid assets

to one and further assume that

$$\psi R < 1 < R. \quad (1)$$

The ordering of returns in Equation (1) implies that while illiquid investments have a higher return than liquid reserves, they are costly to liquidate (as in [Allen and Gale \(2007\)](#), for example). To finance its investment, the bank issues demandable debt (uninsured deposits) to investors at $t = 0$. Each unit of debt promises a face value $F \geq r$, payable upon withdrawal at $t = 1$ or $t = 2$, with F being independent of the withdrawal date.¹ The initial balance sheet of the bank is

$$I + L = D, \quad (2)$$

where I is the amount invested, L is the level of liquid reserves, and $D = 1$ when all investor resources are raised.²

Balance sheet shock. The bank is protected by limited liability and its balance sheet is subject to a shock A at $t = 2$. This shock may improve the balance sheet, $A < 0$, but operational, market, legal, or cyber risks may require writedowns, $A > 0$. The additive shock ensures that changes to the bank's liquid reserves do not induce additional risk-taking on the asset-side of the balance sheet, thus allowing us to isolate the role of illiquidity in shaping the bank's choices. The shock is drawn at $t = 1$ from a continuous probability distribution with a decreasing reverse hazard rate,

¹We assume that debt is demandable for exogenous reasons. This could reflect liquidity preferences (e.g. [Diamond and Dybvig \(1983\)](#)) or a commitment device to overcome agency conflicts (e.g. [Calomiris and Kahn \(1991\)](#) and [Diamond and Rajan \(2001\)](#)). See [Rochet and Vives \(2004\)](#) for such an approach in the context of a global-games bank run model. Accordingly, uninsured deposits refer to any short-term or demandable debt instrument, which includes uninsured retail deposits and insured deposits when deposit insurance is not credible ([Bonfim and Santos, 2023](#)). Three quarters of U.S. commercial bank funding are deposits, half of which uninsured ([Egan et al., 2017](#)).

²We abstract from bank equity to focus on the interaction between ex-ante liquidity choice and ex-post LLR support. Introducing exogenous equity $E \geq 0$ on the liability side, so that $I + L = D + E$, would not change the qualitative interaction between the LLR and minimum liquidity requirements.

$\frac{d}{dA} \frac{g(A)}{G(A)} < 0$, where $G(A)$ denotes the cumulative distribution and $g(A)$ the probability density. This property ensures a unique liquidity choice of the bank.³

Debt rollover. Investors delegate rollover decisions at $t = 1$ to professional fund managers, as in [Rochet and Vives \(2004\)](#) and [Vives \(2014\)](#), for example. These managers are indexed $i \in [0, 1]$, who are rewarded for making the right decision: if the bank does not fail, a manager's payoff difference between withdrawing and rolling over is $-c < 0$; if the bank fails, the differential payoff is $b > 0$.⁴ A conservatism ratio, $\gamma \equiv \frac{b}{b+c} \in (0, 1)$, summarises these payoffs, with more conservative managers (higher γ) being less inclined to roll over.⁵ Fund managers base their withdrawal decisions on noisy private signals,

$$x_i = A + \epsilon_i, \quad (3)$$

for all $i \in [0, 1]$, where ϵ_i is a mean-zero noise term that is independent of the shock A and distributed identically and independently distributed across fund managers according to a continuous distribution H with support $[-\epsilon, \epsilon]$ for $\epsilon > 0$. The realization of A is subsequently publicly observed at $t = 2$.

LLR and bank failure. If a fraction $\ell \in [0, 1]$ of fund managers choose to withdraw at $t = 1$, the bank can serve them as long as it has enough liquidity, $\ell DF \leq L$. But when liquid reserves are low, $L < \ell DF$, the bank must either liquidate assets or request liquidity assistance from the LLR. Drawing on [Tucker \(2014\)](#), the LLR operates according to two key principles:

1. *The LLR's policy should be clear and leave no doubt.* In modern payments

³This property is shared by many distributions, including the exponential, normal, log-normal, uniform, and Weibull distributions.

⁴As an example, assume the cost of withdrawal is c ; the benefit from getting the money back or withdrawing when the bank fails is $b + c$; the payoff for rolling over when the bank fails is zero.

⁵Reviewing debt markets during the financial crisis, [Krishnamurthy \(2010\)](#) argues that investor conservatism was an important determinant of short-term lending behaviour. See also [Vives \(2014\)](#).

systems, banks settle claims against each other using central bank reserves. Thus, the central bank – and by extension the LLR – is at the apex of the payments system. It is, therefore, not credible for the LLR to deny liquidity support for solvent banks and allow clearing to fail.

2. *The LLR should only lend to solvent banks.* The main argument in favour of this stems from the view that providing bailouts are the purview of elected governments. For a LLR to do so would violate the boundary between legislators and an independent central bank and should therefore be avoided.⁶

Accounting for these principles, the LLR in our model is deep-pocketed and credibly commits at $t = 0$ to providing funding at a *penalty rate* ρ to solvent banks at $t = 1$, where

$$R < \rho < 1/\psi. \quad (4)$$

The set of inequalities in Equation (4) implies that (i) the LLR offers more favourable terms of funding than liquidating assets ($\rho < 1/\psi$) and (ii) the opportunity cost of borrowing from the LLR instead of holding more liquidity is strictly positive ($\rho > R$). That is, borrowing from the LLR at $t = 1$ to serve withdrawals is costly for the bank from an ex-ante perspective.⁷

We treat ρ as an exogenous institutional parameter throughout, set by long-standing central-bank practice rather than chosen jointly with the regulatory instru-

⁶Additionally, it would significantly undermine market incentives to monitor and ration risky firms. And, if it was believed that the LLR lends to insolvent banks, it engenders adverse selection since no solvent bank would accept LLR funding for fear of being mistaken as being insolvent.

⁷The assumption $\rho > R$ reflects the long-standing Bagehot–Tucker convention (Bagehot, 1873; Tucker, 2014) that central-bank emergency lending should be priced at a penalty above the rate of return on the bank’s assets, so as to discourage routine reliance on the facility. Discount-window and emergency-facility rates at distressed institutions have at times exceeded the yield on assets in precisely this manner. During the U.S. banking turmoil of March 2023, for example, First Republic Bank borrowed up to \$109 billion from the Federal Reserve’s discount window and \$13.5 billion from the newly established Bank Term Funding Program; the interest rates it paid on these and other emergency loans exceeded the rates it was earning on its assets, placing the bank at continued risk of insolvency (Congressional Research Office, 2023).

ments we study. This mirrors current governance arrangements whereby the central bank, operating in the Bagehot–Tucker tradition (Bagehot, 1873; Tucker, 2014), sets the terms of emergency liquidity provision, while the regulatory authority subsequently sets liquidity requirements taking the LLR’s design as given.

We suppose that the LLR perfectly observes the amount of liquidity on the bank’s balance sheet, L , the face value of debt, F , as well as the shock, A , and the withdrawal volume, ℓ .⁸ Consequently, liquidity assistance is provided as long as the bank can repay all investors who did not withdraw and the LLR:

$$\rho(\ell DF - L) + (1 - \ell)DF \leq RI - A. \quad (5)$$

But, if the condition in Equation (5) is violated, borrowing from the LLR would render it insolvent.⁹ In this case, the LLR cannot provide funding. Instead, the bank must liquidate its assets at $t = 1$ at the market rate $1/\psi$ to serve withdrawals. Insofar that $\rho < 1/\psi$, this also renders the bank insolvent. Thus, the denial of LLR support precipitates bank failure. When the bank fails, for simplicity we assume that its equity is completely wiped out and that the recovery rate for investors is zero.¹⁰

Table 1 summarises the timeline of events.

⁸For a global-games analysis of bank-runs when the policymaker can not tell insolvency and illiquidity apart, see Li and Ma (2022).

⁹Equation (5) can also be thought to reflect the amount of collateral that the bank has on its balance sheet. A necessary condition for the bank to receive liquidity assistance is that it has enough collateral to ensure that the LLR does not suffer any loss.

¹⁰This assumption is consistent with evidence that bankruptcy costs are large. James (1991), for example, measures the losses associated with bank failure as the difference between the book value of assets and the recovery value less direct costs associated with failure. These losses amount to about 30% of failed banks’ assets.

$t = 0$	$t = 1$	$t = 2$
1. LLR and liquidity policies set 2. Debt issuance & liquidity choice	1. Balance sheet shock realizes 2. Private signals about shock 3. Withdrawal decisions 4. LLR funding sought 5. Consumption	1. Investment matures 2. Shock materializes 3. Investors & LLR repaid 4. Consumption

Table 1: Timeline.

3 Equilibrium without liquidity regulation

This section characterises the private choices given the institutional LLR penalty rate ρ . We solve for the symmetric perfect Bayesian Nash equilibrium in pure threshold strategies. Fund managers roll over at $t = 1$ if and only if their private signals indicate that the bank is solvent, $x_i \leq x^f$, and the bank fails if and only if the balance sheet shock is large, $A > A^f$. We refer to higher values of A^f as higher bank stability.

Definition 1. *The symmetric pure-strategy perfect Bayesian equilibrium comprises thresholds x^* and A^* , liquid reserves L^* , and face value of debt F^* such that*

- *at $t = 0$, the bank chooses liquid reserves L^* and issues debt with face value F^* to maximize expected profits, taking the LLR penalty rate ρ as given and subject to the participation of investors who require an expected return of r ;*
- *at $t = 1$, the rollover decisions of all fund managers constitute mutual best responses, with rollover threshold $x^* \equiv x^f(L^*, F^*)$ and failure threshold $A^* \equiv A^f(L^*, F^*)$ such that the bank fails for any shock $A > A^*$, given L^* and F^* .*

3.1 Rollover risk

The rollover subgame at $t = 1$ is solved using global games with vanishing private noise about the balance sheet shock, $\epsilon \rightarrow 0$. We take this limit henceforth, which

eases the exposition and allows us to focus on the bank's ex-ante liquidity choice. As such, the rollover threshold converges to the failure threshold, $x^f \rightarrow A^f$, and fund managers face only strategic uncertainty about the behaviour of other managers.

Lemma 1. *Bank stability.* Denote $\bar{A} \equiv RI + L - DF$. There exists a unique bank failure threshold

$$A^f(L, F) \equiv \begin{cases} \bar{A} & \text{if } L \geq \gamma DF \\ \bar{A} - (\gamma DF - L)(\rho - 1) & \text{if } L < \gamma DF \end{cases}. \quad (6)$$

The failure threshold is single-peaked in L , attaining its maximum at $L = \gamma DF$. When liquidity is scarce, $L \in [0, \gamma DF)$, A^f increases in L with slope $\rho - R > 0$. While, when it is abundant, $L \in (\gamma DF, D]$, A^f decreases in L with slope $1 - R < 0$. Bank stability decreases in the penalty rate when liquid reserves are scarce, $\frac{\partial A^f}{\partial \rho} \leq 0$.

Proof. See Appendix A.1. ■

The failure threshold A^f depends on the amount of liquid reserves L , as shown in Figure 1. When liquid reserves are plentiful, $L \geq \gamma DF$, the bank never requires additional funding from the LLR in equilibrium. Thus, fund managers withdraw only when the bank is fundamentally insolvent, $A > \bar{A}$, where $\bar{A}$ is the upper dominance bound. In other words, high liquid reserves rule out panic-based runs whenever $L \geq \gamma DF$. In this range, more liquidity actually reduces bank stability: it does not mitigate panic runs but carries the opportunity cost of profitable long-term investment, resulting in a negative slope, $1 - R < 0$.

When liquid reserves are scarce, $L < \gamma DF$, the bank can, instead of liquidating its investments, approach the LLR for liquidity assistance to serve withdrawals. The greater is the mass of fund managers who withdraw, the more is borrowed from the LLR, the larger the probability of bank failure at $t = 2$ and thus, the higher the incentives of fund managers to withdraw (so withdrawal actions are strategic

complements). The slope $\rho - R > 0$ reflects the marginal gains from more liquidity: a lower cost of borrowing from the LLR minus the opportunity cost of investment.

As a consequence, the failure threshold increases in the penalty rate ρ for scarce liquid reserves. A higher penalty rate implies that the bank must pay more to the LLR for interim funding, which impinges on its solvency at $t = 2$. The higher penalty rate, moreover, exacerbates the strategic complementarity between fund managers thereby incentivizing them further to withdraw even though the bank is fundamentally solvent. These two effects reinforce each other, reducing bank stability.

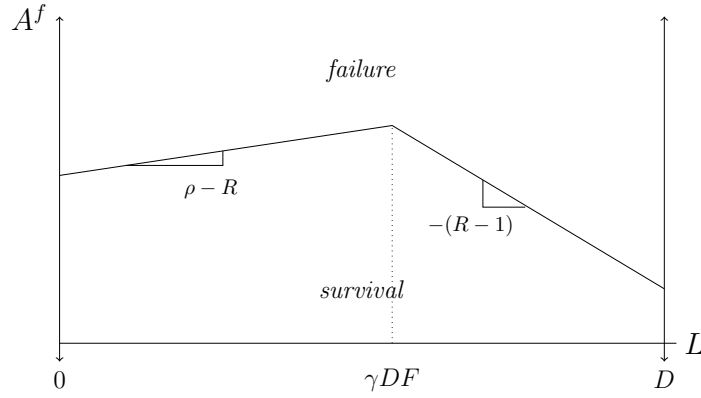


Figure 1: Failure threshold A^f and liquid reserves L . Stability is highest at $L = \gamma DF$.

3.2 Liquidity choice and debt pricing

Given the failure threshold, $A^f(L, F)$, the bank chooses its liquid reserves L at $t = 0$ and invests the remainder, $I = D - L$. It maximizes expected profits subject to offering a face value F that satisfies the participation constraint of investors:

$$\max_{L, F} \pi(L, F) \equiv \int_{-\infty}^{A^f(L, F)} [R(D - L) + L - A - DF] dG(A) \quad \text{s.t.} \quad r \leq F G(A^f(L, F)). \quad (7)$$

While the LLR influences the bank's failure threshold (and thus its expected profits), it is never used in equilibrium. This is because under vanishing private noise, rollover

decisions are perfectly correlated. For $A < A^f(L, F)$, all fund managers rollover their debt and the bank does not require any liquidity assistance. While for $A > A^f(L, F)$, the LLR correctly anticipates insolvency and refuses liquidity assistance. The fund managers, in turn, all withdraw and the bank fails. The LLR therefore operates entirely via *anticipation*: it credibly promises to lend at rate $\rho < 1/\psi$, which weakens the strategic complementarity between fund managers and shifts the failure threshold. This, in turn, influences the liquidity choice and face value of debt, as we see here.

The solution to Problem (7) is given by two equations that characterise the liquidity choice, $L^* \equiv L^*(A^f(L^*, F), F)$, and the face value of debt, $F^* \equiv F^*(A^f(L, F^*))$. We then have $A^* \equiv A^f(L^*(A^f, F^*(A^f)), F^*(A^f))$ and, with a slight abuse of notation, denote the associated system of three equations as $\{L^*, A^*, F^*\}$. In what follows, we introduce the following parametric assumption on the LLR penalty rate.

Assumption 1. *The LLR penalty rate satisfies $\rho < \hat{\rho}$.*

Assumption 1 ensures that the indirect effect induced by the pricing of debt of a marginal increase in liquid reserves on expected profits is limited in scope. This, in turn, ensures that the equilibrium is unique. We characterise $\hat{\rho}$ in Appendix A.2.

Proposition 1. *Liquid reserves.* *The bank's optimal liquid reserves satisfy $L^* \in [0, \gamma DF^*]$. When the solution is interior, $L^* \in (0, \gamma DF^*)$, it is implicitly given by*

$$R - 1 = \left[(\rho - 1)(\gamma DF^* - L^*) + DF^* \right] \frac{g(A^*)}{G(A^*)} \frac{dA^*}{dL}. \quad (8)$$

There exist threshold penalty rates $\underline{\rho}$ and $\bar{\rho}$, with $R < \underline{\rho} < \bar{\rho}$ and $\underline{\rho} < \psi^{-1}$. When $\bar{\rho} < \psi^{-1}$, the corner solution $L^ = 0$ obtain for $\rho \leq \underline{\rho}$, and $L^* = \gamma DF^*$ for $\rho \geq \bar{\rho}$. The interior solution for L^* increases in the penalty rate, $dL^*/d\rho > 0$, and weakly increases in the risk-free rate, $dL^*/dr \geq 0$.*

Proof. See Appendix A.2. ■

The bank's liquidity choice balances the marginal benefit and marginal cost of holding liquid reserves. The marginal cost on the left-hand side of Equation (8) is the foregone return from not investing in illiquid but profitable investments. The right-hand side is the expected gain from a marginal increase in liquidity on the failure threshold A^* . The term $\frac{g(A^*)}{G(A^*)} \frac{dA^*}{dL}$ is the conditional probability that the bank does not fail upon a marginal increase in liquidity, which multiplies into the benefits of staving off bank failure. These benefits are twofold. First, the loss in equity value due to inefficient panic runs, $E_2(A^*) = (\rho - 1)(\gamma DF^* - L^*) > 0$, is avoided. Second, the face value of debt (set for investors to break even in expectation) is lower: higher liquidity raises bank stability and thus the probability of repayment. The bank never chooses to hold more liquidity than γDF^* . Beyond this point, panic runs are already ruled out, so additional reserves only reduce profitability without improving stability.

When the penalty rate is low, $\rho \leq \underline{\rho}$, the marginal benefit of liquidity is small. In this case, both the loss in equity value from inefficient runs, $E_2(A^*)$, and the debt burden incurred from LLR assistance are low, especially in comparison to the foregone revenue from investing in illiquid assets. As a consequence, it is optimally for the bank to hold none, $L^* = 0$. When the penalty rate is high, $\rho \geq \bar{\rho}$, borrowing from the LLR becomes so expensive that the bank prefers to self-insure fully, $L^* = \gamma DF^*$, eliminating panic runs entirely. For intermediate penalty rates $\rho \in (\underline{\rho}, \bar{\rho})$, the bank holds a positive level of reserves. When $1/\psi < \bar{\rho}$, the corner solution is never reached and the optimal liquidity choice is an interior solution. As a consequence, panic runs continue to persist in equilibrium. In what follows, we focus on parameter conditions that ensure $\bar{\rho} < 1/\psi$.¹¹ Figure 2 shows the resulting schedule of bank liquidity, $L^*(\rho)$.

A marginal increase in the risk-free rate r raises the bank's funding cost and thus its debt burden, incentivizing greater self-insurance through liquid reserves. In

¹¹Additional parameters are set at $D = 1$, $R = 2$, $\gamma = 0.2$, $\psi = 0.1$ and $r = 1.2$. The balance sheet shock is normally distributed with mean -1.3 and unit variance. Unless otherwise stated, we use this parameterization in all subsequent numerical exercises.

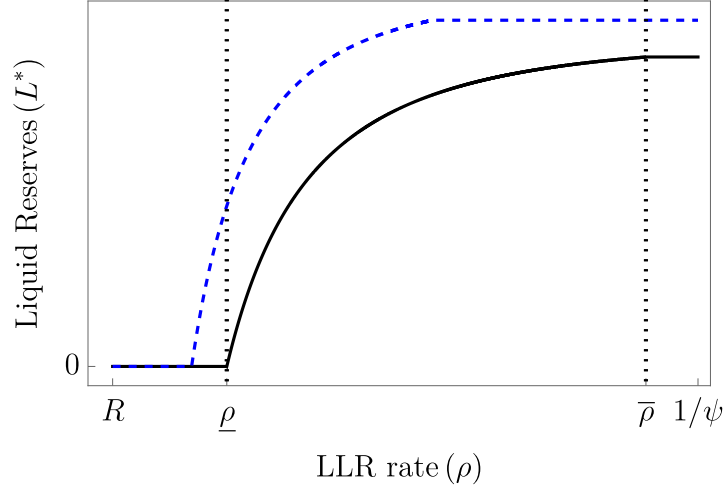


Figure 2: Liquid reserves L^* as a function of the penalty rate ρ . The solid curve is for our baseline parametrization, while the dashed curve features a 10% increase in r .

particular, in the interior region $dL^*/dr \geq 0$, and the interval over which the interior solution applies shifts leftward: $d\rho/dr < 0$ and $d\bar{\rho}/dr < 0$. The dashed curve in Figure 2 illustrates these comparative statics.

3.3 Bank stability and the LLR

We turn to the comparative static for how the LLR shapes bank stability. The effect of the penalty rate decomposes into a direct effect and indirect effects via the bank's liquidity choice and the price of debt. Proposition 2 characterises the total effect.

Proposition 2. *Bank stability and LLR.* *For $\rho < \underline{\rho}$, a marginal increase in the penalty rate reduces bank stability, $\frac{dA^*}{d\rho} < 0$. For $\rho > \underline{\rho}$, by contrast, bank stability weakly increases in the penalty rate, $\frac{dA^*}{d\rho} \geq 0$. In the limit $\rho \rightarrow \bar{\rho}$, panic runs are eliminated, $A^* \rightarrow \bar{A}$. While for $\rho \in (R, \bar{\rho})$, panic runs occur, $A^* < \bar{A}$.*

Proof. See Appendix A.3. ■

Figure 3 shows the V-shaped relationship between bank stability and the penalty rate described in Proposition 2. For a low penalty rate, the marginal cost of holding liquidity outweighs the marginal benefit, so the bank chooses not to hold liquidity on its balance sheet, $L^* = 0$. Thus, to serve withdrawals, the bank relies exclusively on LLR assistance. Thus, following a marginal increase in the penalty rate, the debt burden incurred by the bank from LLR assistance increases, which impinges on its solvency and raises the price of debt. Hence, bank stability decreases, $\frac{dA^*}{d\rho} < 0$.

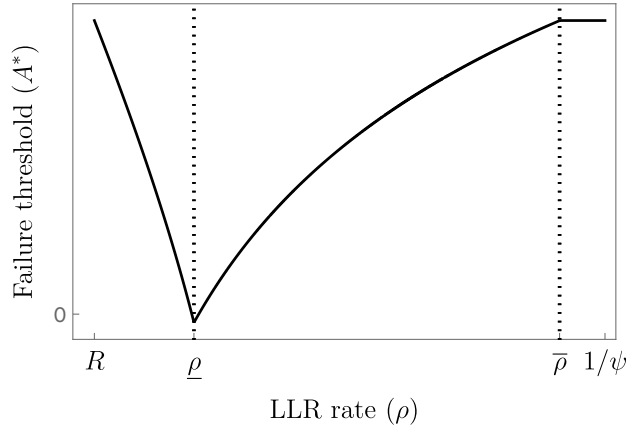


Figure 3: Bank stability depends non-monotonically on the LLR penalty rate.

For a higher penalty rate, $\rho > \underline{\rho}$, the higher debt burden incurred by seeking LLR assistance to serve withdrawals leads the bank to choose a positive level of liquid reserves, $L^* > 0$. In this case, we obtain that

$$\frac{dA^*}{d\rho} > 0 \quad \Leftrightarrow \quad \underbrace{\frac{\partial A^*}{\partial \rho}}_{\text{Direct effect}} + \underbrace{\frac{\partial L^*}{\partial \rho} \frac{\partial A^*}{\partial L}}_{\text{Liquidity effect}} > 0. \quad (9)$$

The total effect of an increase in the penalty rate on the failure threshold can be decomposed into a direct effect and an indirect liquidity effect. Since the face value of debt does not depend directly on ρ , there is no additional effect through the price of debt, $\frac{\partial F^*}{\partial \rho} = 0$. The liquidity effect accounts for how changes in the penalty rate impact the failure threshold by influencing bank liquidity. The direct effect of a

higher penalty rate worsens bank stability, $\frac{\partial A^*}{\partial \rho} < 0$, as in [Rochet and Vives \(2004\)](#). Intuitively, as the cost of liquidity assistance increases, the bank is more likely to fail. However, a higher penalty rate ρ at the same time incentivizes the bank to increase its liquid reserves, which raises stability, $\frac{\partial L^*}{\partial \rho} \frac{\partial A^*}{\partial L} > 0$. When $\rho > \underline{\rho}$, the liquidity effect unambiguously dominates the direct effect. So, bank stability improves with the penalty rate, $\frac{dA^*}{d\rho} \geq 0$. Finally, for $\rho > \bar{\rho}$, the bank chooses $L^* = \gamma DF^*$, thus eliminating panic runs, so $A^* = \bar{A}$. Further increases in ρ have no effect on stability.

4 LLR support, welfare, and liquidity regulation

This section evaluates the welfare impact of LLR support on welfare and derives a role for minimum liquidity regulation. We adopt the same governance split as before: the central bank sets the LLR's terms in the Bagehot–Tucker tradition, while a regulatory authority sets minimum liquidity requirements, taking the LLR policy as given.

4.1 Welfare under isolated LLR support

We consider a social planner who maximizes utilitarian welfare. Unlike the bank, the planner accounts for deadweight losses associated with bank failure — the contraction of credit, the loss of relationship-specific information about firms, and the effects of contagion, among others.¹² Capturing such losses in reduced form by the parameter $\lambda > 0$, utilitarian welfare is the sum of the bank's expected profits, the expected

¹²There is considerable empirical evidence on the social costs of bank failure. For example, [Crawley et al. \(2023\)](#) find that following the failure of SVB in March 2023, local consumer spending in several counties in California was significantly reduced. [Calomiris and Mason \(2003\)](#) show that bank failures reduced loan supply during the Great Depression. Beyond credit-supply considerations, financial-contagion concerns are also important ([Allen and Gale, 2000](#)).

payments to depositors, and the social cost of bank failure:

$$W(L, F) \equiv \pi(L, F) + rD - \lambda \left(1 - G(A^f(L, F)) \right). \quad (10)$$

Three remarks about this social objective function are in order. First, in the limit $\lambda \rightarrow 0$, bank failure exerts no externality and the private equilibrium is constrained efficient. As λ increases, the planner places greater weight on financial stability. Second, since the bank never seeks liquidity assistance from the LLR under vanishing noise, $\epsilon \rightarrow 0$, we do not need to account for transfers between the bank and the LLR. Third, Equation (10) includes the bank profits and expected investor returns but the payoffs to fund managers are absent. This approach is consistent with taking the limits $b \rightarrow 0$ and $c \rightarrow 0$ subject to a constant ratio $\gamma = b/(b+c)$, and follows Ahnert et al. (2019).¹³

To assess the LLR's welfare contribution, we compare two regimes. Let $W_0 \equiv W(L_0^*, F_0^*)$ denote welfare in the absence of the LLR, modelled as $\rho \geq 1/\psi$ so that the bank must liquidate its investments to serve withdrawals. Note that (L_0^*, F_0^*) denotes the bank's equilibrium choices in this benchmark. Let $W_1 \equiv W(L^*, F^*)$ denote welfare when the LLR stands ready to offer liquidity assistance at the penalty rate $\rho < \bar{\rho}$. Again, (L^*, F^*) denotes the equilibrium choices in this case.

Proposition 3. *Let $\hat{\lambda}$ be a critical value of the social cost of bank failure. The availability of the LLR improves welfare, $W_1 \geq W_0$, if and only if $\lambda \leq \hat{\lambda}$.*

Proof. See Appendix A.4. ■

Figure 4 depicts the result of Proposition 3. The bank does not internalize the social costs of failure when choosing its asset side. When λ is small, the efficiency gain

¹³Including these payoffs would increase the weight that the social planner places on bank survival, so the parameter λ would effectively be rescaled to $\lambda + c$. Therefore, all of our normative results would go through qualitatively.

from anticipated LLR support outweighs the additional social cost of the bank's lower liquidity and higher ex-ante probability of failure. So $W_1 > W_0$. But, because the probability of failure is higher in the presence of the LLR, W_1 declines more steeply in λ than W_0 . The two schedules cross at $\hat{\lambda}$. And for an even larger social cost of bank failure, the distortions caused by the LLR are sufficiently large, so $W_1 < W_0$. That is, the LLR is detrimental to welfare.

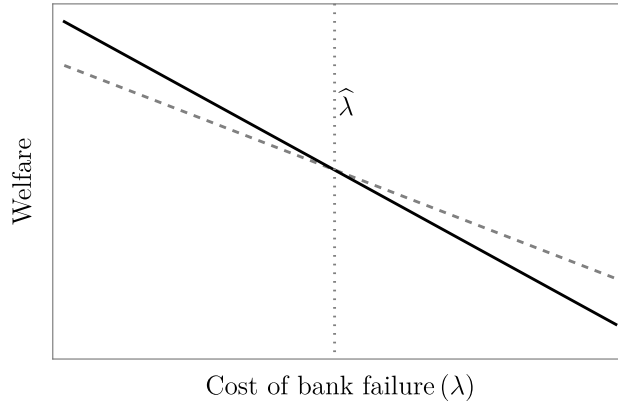


Figure 4: Welfare as a function of the social cost of bank failure λ . The solid line refers to the case with a LLR(W_1), while the dashed line to the case without (W_0).

4.2 The social planner and minimum liquidity requirements

The proximate reason for lower welfare in the presence of the LLR alone is that the bank holds too little liquidity. A natural remedy is to impose a minimum liquidity requirement (MLR). We derive the MLR from the social planner's problem. The planner chooses the socially optimal level of liquidity, L^P , and the corresponding face value of debt, F^P , to maximize welfare $W(L, F)$ subject to the depositors' participation constraint $r \leq G(A^f)F^P$. Using notation analogous to the bank's private solution in Section 3.2, we denote by $\{L^P, A^P, F^P\}$ the solution to the planner's problem. Corollary 1 compares it to the private optimum.

Corollary 1. *There exists a bound $\underline{\rho}^P < \underline{\rho}$ such that, for $\rho \in (\underline{\rho}^P, \bar{\rho})$, the planner*

chooses strictly more liquid reserves than the bank, $L^P > L^*$. Outside this region the bank's choice is already constrained-efficient, $L^P = L^*$. The planner's allocation is implemented by a minimum liquidity requirement $L \geq \underline{L} \equiv L^P$.

Proof. See Appendix A.5. ■

The planner's (interior) choice of liquid reserves equates the bank's opportunity cost of holding liquid reserves with the marginal social benefit and is given by:

$$R - 1 = \left[(\rho - 1)(\gamma DF^P - L^P) + DF^P + \lambda \right] \frac{g(A^P)}{G(A^P)} \frac{dA^P}{dL}. \quad (11)$$

Relative to the bank's private choice in Equation (8), the marginal social benefit from increased liquidity also includes avoiding the social cost of bank failure, λ . This, in turn, leads the planner to choose a higher level of liquid reserves, $L^P \geq L^*$, as depicted in Figure 5. Moreover, the planner's solution can be implemented by introducing minimum liquidity requirements for the bank, $L \geq \underline{L} \equiv L^P$. By choosing $L = \underline{L}$ in the regulated economy, the corresponding face value of debt for the bank is $F = F^P$.

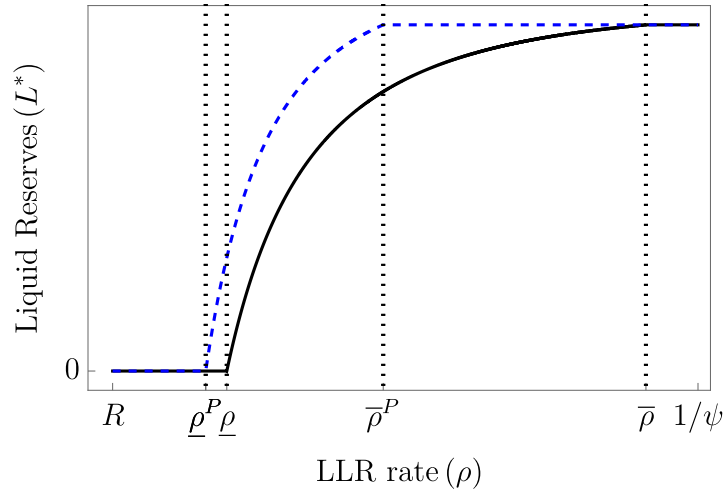


Figure 5: Privately optimal (black solid line) and socially optimal (blue dashed) level of bank liquidity.

Proposition 4. *Liquidity regulation reduces panic runs.* The introduction of

liquidity regulation reduces the range of panic runs

$$\gamma DF^* - L^* \geq \gamma DF^P - L^P, \quad (12)$$

with strict inequality for $\rho \in (\underline{\rho}^P, \bar{\rho})$.

Proof. See Appendix A.5. ■

The mechanism is twofold. First, requiring the bank to hold more liquid reserves, $L^P \geq L^*$, directly improves fund managers' incentives to roll over. Second, by raising ex-ante stability, liquidity requirements lower the equilibrium face value of debt, $F^P \leq F^*$. Cheaper debt, in turn, reduces the bank's ex-post debt burden, which feeds back into stronger rollover incentives. The two effects reinforce each other.

4.3 Liquidity regulation and the social value of the LLR

Returning to the LLR's contribution to welfare, we investigate the role of liquidity regulation. Let $W_1^P \equiv W(L^P, F^P)$ denote welfare under both instruments, that is when the LLR is ready to offer liquidity assistance at the penalty rate $\rho < 1/\psi$ and the minimum liquidity requirement $L \geq \underline{L}$ with $\underline{L} \equiv L^P$ is in place.

Proposition 5. Social value of the LLR. *When the penalty rate satisfies $\rho \in (\underline{\rho}^P, \bar{\rho})$, we have two results:*

- (i) **Levels.** *For any $\lambda > 0$, combining liquidity regulation with the LLR delivers strictly higher welfare than either tool in isolation:*

$$W_1^P > W_0 \geq W_1.$$

- (ii) **Efficacy.** *Liquidity regulation amplifies the welfare contribution of the LLR at*

the margin:

$$\frac{dW_1^P}{d\lambda} > \frac{dW_1}{d\lambda}.$$

Proof. See Appendix A.6. ■

Proposition 5 is the central normative result of the paper. The MLR (liquidity regulation) and the LLR are not substitutes for one another but complements, and the two parts of Proposition 5 capture distinct dimensions of this interaction. First, there is how the interaction impacts the *level* of welfare. By obliging the bank to internalise the social cost of bank failure, liquidity regulation forces the bank to hold a level of liquidity closer to γDF (recall that this is the maximum level that the bank may find privately optimal). This reduces the probability of bank failure and, in turn, reduces the debt burden the bank would carry were it to draw on the LLR. The net effect is to push the critical level $\hat{\lambda}$ to infinity. In words, minimum liquidity regulation ensures that the LLR unambiguously improves welfare.

The second part of the proposition concerns the *efficacy* of the LLR. Without liquidity regulation, the bank does not internalize the social value of its failure, so changes in λ do not affect its private liquidity choice L^* , the associated debt face value F^* , or the failure threshold A^* . Welfare therefore deteriorates as λ rises at rate $\frac{dW_1}{d\lambda} = -(1 - G(A^*))$. With the MLR, by contrast, the regulatory floor is chosen by the planner. A higher λ raises the social marginal benefit of liquidity, leading to a higher level L^P , a higher failure threshold A^P , and a lower failure probability. By the envelope theorem, $\frac{dW_1^P}{d\lambda} = -(1 - G(A^P))$. Since liquidity regulation raises the failure threshold relative to the private allocation, $A^P > A^*$, welfare under the combined MLR–LLR regime declines more slowly in the social cost of bank failure: $dW_1^P/d\lambda > dW_1/d\lambda$. In this sense, liquidity regulation makes the LLR more effective by preserving its anticipatory stabilization benefit while increasing ex-ante bank stability through which the LLR alone becomes socially costly. Figure 6 shows the results.

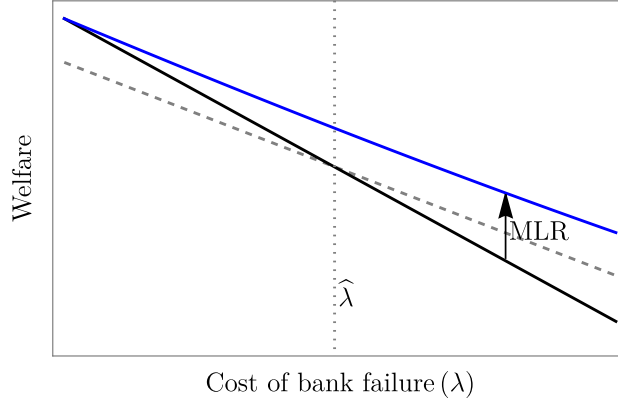


Figure 6: Welfare as a function of the social cost of bank failure λ . The solid black line plots welfare with a LLR (W_1). The blue line is welfare with both the MLR and the LLR (W_1^P). The dashed line is welfare in the absence of both (W_0).

To quantify the relative size of the two channels using our maintained parameter values of the numerical example, Figure 7 reports a Shapley-symmetric decomposition of the welfare gap $W_1^P - W_1$. For each value of λ , we evaluate welfare at four points: the bank's private equilibrium (L^*, F^*) , the planner's allocation (L^P, F^P) , and the two off-equilibrium counterfactuals, (L^P, F^*) and (L^*, F^P) . The decomposition shows that the welfare gain stems entirely from the effect of reducing the banks ex-post debt burden. At the bank's private optimum, L^* is chosen so that the marginal cost of liquidity is balanced against the marginal benefit including the equilibrium re-pricing of F . Stripping out that F-adjustment by holding F at F^* and forcing L up to L^P leaves the bank over-paying for debt. As a consequence, the drop in profits outweighs the inefficient-run cost reduction from the higher failure threshold. The welfare gain therefore comes entirely from the equilibrium re-pricing of debt.

4.4 Comparative statics on the optimal liquidity regulation

Finally, we describe how the optimal MLR moves with the institutional environment. Two concepts are useful. The *strictness* of the MLR, $\underline{L}$, measures the level of regulation: $\underline{L}$ is stricter than $\underline{L}'$ whenever $\underline{L} > \underline{L}'$. The *stringency* of the MLR, $\underline{L} - L^* > 0$,

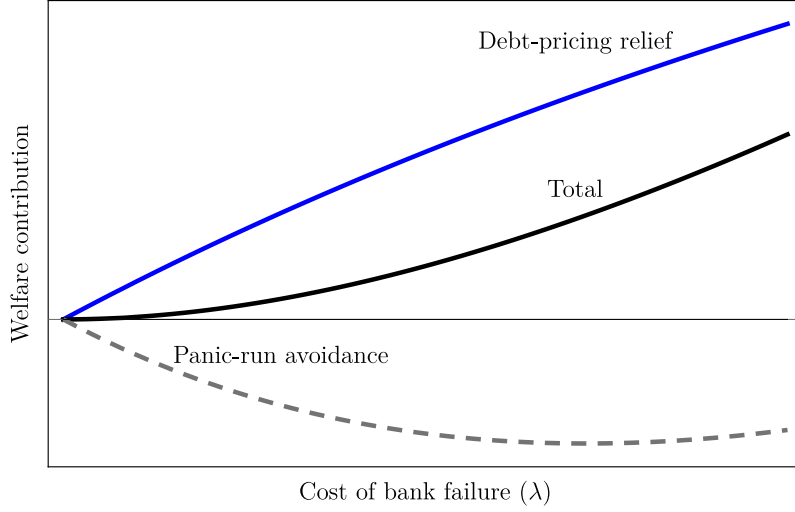


Figure 7: Channel decomposition of the welfare gain $W_1^P - W_1$ at the baseline calibration, as a function of the social cost of bank failure λ . The decomposition is Shapley-symmetric: it averages the L-first and F-first paths through (L, F) space. The solid blue line is the debt-pricing relief channel; the solid black line is the total $W_1^P - W_1$; the dashed gray line is the panic-run avoidance channel.

measures the wedge between the regulatory floor and the bank's unregulated choice—the extent of the correction that liquidity regulation imposes on the bank.

Proposition 6. *Optimal minimum liquidity requirements and LLR penalty rates.* *The strictness of the MLR weakly increases in the LLR penalty rate, $\frac{\partial \underline{L}}{\partial \rho} \geq 0$. Its stringency is non-monotone in ρ , with an inverted-V shape: it rises with ρ over the corner region $\rho \in [\underline{\rho}^P, \underline{\rho}]$ and tends to zero over the interior $\rho \in (\underline{\rho}, \bar{\rho})$.*

Proof. See Appendix A.7. ■

A higher penalty rate raises the ex-post debt burden the bank incurs upon borrowing from the LLR. The planner counteracts this by tightening the MLR because a higher level of $\underline{L}$ improves stability (lowering the face value of debt required by investors) and reduces the bank's reliance on the LLR in the first place. Both effects shrink the ex-post debt burden and thereby the range of shocks where panic runs occur.

Figure 5 gives the graphical intuition for the stringency result. When $\rho < \underline{\rho}^P$, LLR borrowing is so cheap that not even the planner finds it worthwhile for the bank to hold liquidity. Both choose zero reserves, the bank's choice is constrained-efficient, so stringency is zero. In the region $\rho \in [\underline{\rho}^P, \underline{\rho}]$, a divergence emerges: the social marginal benefit of liquidity now justifies positive reserves, but the private benefit does not. A higher penalty rate therefore raises the planner's choice $\underline{L}$ while the bank remains at the corner $L^* = 0$, and stringency rises one-for-one with the penalty rate. Over $\rho \in (\underline{\rho}, \bar{\rho})$, the logic reverses. LLR borrowing is now sufficiently expensive that both the bank and the planner choose to hold liquid reserves. As a consequence, the distance between L^P and L^* shrinks. Higher ρ creates incentives for the bank to further self-insure, increasing stability and thus reducing the gap between private and social optima, so regulation becomes less stringent even though the strictness of the requirement continues to rise. Hence stringency declines even though the strictness of the requirement continues to rise. Finally, for $\rho \geq \bar{\rho}$, the two allocations coincide at the upper bound γDF^* (zero stringency) and panic runs vanish.

5 Conclusion

This paper studies a global-games model of bank runs in which a lender of last resort (LLR) shapes the bank's liquidity choice and the price of uninsured deposits. Anticipating LLR support reduces the bank's incentive to self-insure ex ante by holding liquid reserves. This, in turn, increases the ex-ante probability of bank failure and the face value of debt required by investors. A key feature of the model is that both the bank's liquid reserves and the face value of debt are endogenous and jointly determined.

Our main positive result is that bank stability is V-shaped in the LLR penalty rate. There are two opposing effects. First, holding bank liquidity fixed, the direct

effect of a higher penalty rate is that ex-post emergency liquidity assistance becomes more costly, which weakens bank stability, as in [Rochet and Vives \(2004\)](#). The second effect relates to anticipating LLR support, so the bank holds fewer liquid reserves, increasing the ex-ante probability of failure and reducing debt costs. Taken together, at low penalty rates, the bank holds no liquidity and the direct effect dominates: stability falls as the penalty rate rises. At high penalty rates, liquid reserves are positive and the indirect effect dominates, so stability increases in the penalty rate.

We turn to the normative analysis of anticipating LLR support. Since the bank does not internalize the social cost of its failure, the very availability of the LLR can reduce welfare when this social cost is large. Our result speaks to a long-standing criticism of public liquidity provision: by reducing banks' incentives to self-insure, it can undermine the very stability it was meant to protect. Minimum liquidity requirements resolve the puzzle by closing the wedge between the private and social marginal benefits of liquidity. Thus, the central normative messages are that (i) liquidity regulation and the LLR are complements, rather than substitutes; and (ii) a micro-prudential role of liquidity regulation is to raise the social value of the LLR.

Our results echo the classic argument that an ex-post backstop blunts incentives and, therefore, requires an ex-ante tool to restore discipline. While familiar from the literature on deposit insurance and capital regulation ([Kim and Santomero, 1988](#); [Keeley, 1990](#)), the mechanism in our paper is distinct. We shut down asset-side risk-shifting, so the distortion is not about the riskiness of bank assets but about how little liquidity the bank holds and how it feeds into fragility and the price of its debt.

Several dimensions remain for future work. We abstract from bank equity to isolate the liquidity-side friction; adding capital and capital requirement would extend the analysis to the joint regulation of liquidity and solvency. We also do not model collateral in central-bank lending, although eligibility rules and haircuts determine in practice how much a bank can borrow from the LLR; introducing collateral might

broaden the trade-off between ex-post flexibility and ex-ante incentives. Finally, we treat the LLR's design as exogenous, set by long-standing central-bank practice. Endogenising it within a broader policy game—where the central bank anticipates bank behaviour and responds to political-economy pressures—could be a natural next step.

A Proofs

A.1 Proof of Lemma 1

To derive the failure conditions, it is important to recognise that the LLR has perfect information about the balance sheet shock, A . This means that the LLR knows whether or not the amount of liquidity assistance requested by the bank would render it insolvent. Moreover, insofar that $\rho < 1/\psi$, the refusal of liquidity assistance implies that the bank must liquidate its assets, which would also render it insolvent.

Suppose that the shock is common knowledge among fund managers. In this case, there are two distinct regions with strictly dominant actions for fund managers. First, when all debt is rolled over, $\ell = 0$, the bank still fails if $A > \bar{A} = RI - (DF - L)$. In this case, there is no need for any liquidity assistance from the LLR and so the bank fails whenever the value of its assets is insufficient to cover its debt repayments at $t = 2$. Consequently, fund managers have strictly dominant strategies to withdraw. And second, when no debt is rolled over, $\ell = 1$, the bank does not fail whenever $A < \underline{A} \equiv RI - (DF - L)\rho < \bar{A}$. In this case, the bank requires liquidity assistance from the LLR. Insofar as this additional borrowing at $t = 1$ does not render the bank insolvent, then each fund manager has a strictly dominant strategies to roll over.

Turning to the equilibrium with incomplete information, we solve for the signal and shock thresholds (x^f, A^f) . Suppose that all managers use the threshold strategy x^f . Given A , the proportion of managers who do not roll over is $\ell(A, x^f) = \text{Prob}(x_i > x^f | A) = \text{Prob}(\epsilon_i > x^f - A) = 1 - H(x^f - A)$. When withdrawals can be satisfied by selling liquid reserves, $\ell(A, x^f) \leq L/DF$, failure occurs if $A > \bar{A}$, which does not directly depend on the withdrawal volume. Conversely, if withdrawals lead the bank to require assistance from the LLR, $\ell(A, x^f) > L/DF$, the bank fails whenever it is refused liquidity assistance, i.e., $A > RI - (\ell DF - L)\rho - (1 - \ell)DF$ or,

equivalently, $A > \bar{A} - (\ell DF - L)(\rho - 1)$. Thus, the failure threshold A^f is

$$A^f = \begin{cases} RI - DF + L & \text{if } \ell(A^f, x^f) < \frac{L}{DF} \\ RI \left(1 - \left(\frac{\ell(A^f, x^f) DF - L}{RI} \right) \rho \right) - \left(1 - \ell(A^f, x^f) \right) DF & \text{if } \ell(A^f, x^f) \geq \frac{L}{DF} \end{cases} \quad (13)$$

The threshold A^f exists and is unique. To see this, note that the left-hand side of (13) increases in A over $\mathbb{R}$. Moreover, the right-hand side of (13) is continuous in A , is bounded within $[\underline{A}, \bar{A}]$, and decreases in A over $[\underline{A}, \bar{A}]$.

We use Bayes' rule to derive fund managers' posteriors about the balance sheet shock A and pin down the signal threshold x^f . A manager i who observes $x_i = x^f$ is indifferent between rolling over and withdrawing, hence $1 - \gamma = \Pr(A > A^f | x_i = x^f)$. It follows that $\gamma = \Pr(A < A^f | x_i = x^f) = 1 - H(x^f - A^f)$. The indifference condition therefore implies $x^f - A^f = H^{-1}(1 - \gamma)$. Inserting this into $\ell(A^f, x^f)$, the withdrawal proportion at the threshold A^f becomes $\ell(A^f, x^f) = 1 - H(x^f - A^f) = 1 - H(H^{-1}(1 - \gamma)) = \gamma$. The failure threshold in Equation (6) follows immediately.

A proof of the absence of other equilibria follows directly from previous arguments (Morris and Shin, 2003; Vives, 2005). It is based on iterated elimination of strictly dominated strategies and is skipped for brevity.

A.2 Proof of Proposition 1

We solve Problem (7) assuming $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$ and show that it has a unique solution. We then characterise threshold $\rho = \hat{\rho}$ and demonstrate that assuming $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$ is equivalent to assuming $\rho < \hat{\rho}$.

The participation constraint of investors, $r \leq F G(A^f(L, F))$ in Problem (7), is binding in equilibrium. We, thus, incorporate $F^* = r/G(A^*(L, F^*))$ in the objective function and solve for liquid reserves, L^* .

Interior solution. Total differentiation of (7) yields $\frac{d\pi}{dL} = \frac{\partial\pi}{\partial L} + \frac{\partial\pi}{\partial A^*} \frac{dA^*}{dL} + \frac{\partial\pi}{\partial F^*} \frac{dF^*}{dL}$.

Using expansion $\frac{dF^*}{dL} = \frac{\partial F^*}{\partial A^*} \frac{dA^*}{dL}$ we obtain

$$\frac{d\pi}{dL} = \frac{dA^*}{dL} g(A^*) \left[R(D - L) + L - A^* - DF^* \right] - G(A^*) \left[(R - 1) + D \frac{\partial F^*}{\partial A^*} \frac{dA^*}{dL} \right]. \quad (14)$$

We separately consider the two pieces of the failure threshold A^* given by Equation (6) and take into account that function A^* is not differentiable at $L = \gamma DF^*$.

Consider $L > \gamma DF^*$. Expanding A^* and using $\frac{\partial F^*}{\partial A^*} = -\frac{g(A^*)}{G(A^*)} F^*$ in (14) yields

$$\begin{aligned} \left. \frac{d\pi}{dL} \right|_{L > \gamma DF^*} &= -G(A^*) \left[(R - 1) - \frac{g(A^*)}{G(A^*)} DF^* \frac{dA^*}{dL} \right] \\ &= -G(A^*) \left[(R - 1) + \frac{g(A^*)}{G(A^*)} \left(\frac{(R - 1) DF^*}{1 - \frac{g(A^*)}{G(A^*)} DF^*} \right) \right] = -\frac{(R - 1) G(A^*)}{1 - \frac{g(A^*)}{G(A^*)} DF^*} < 0, \end{aligned} \quad (15)$$

where the second line expands

$$\frac{dA^*}{dL} = \frac{\partial A^*}{\partial L} + \frac{\partial A^*}{\partial F^*} \frac{dF^*}{dL} = \frac{\frac{\partial A^*}{\partial L}}{1 - \frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*}}, \quad (16)$$

using $\frac{\partial A^*}{\partial L} = -(R - 1)$ and $\frac{\partial A^*}{\partial F^*} = -D$. Assumption $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$ guarantees that Equation (16) is positive for $L = L^*$. Hence, A^* and π decrease in L for $L > \gamma DF^*$, resulting in $L^* \leq \gamma DF^*$.

Consider $L \leq \gamma DF^*$. Expanding A^* and using $\frac{\partial F^*}{\partial A^*} = -\frac{g(A^*)}{G(A^*)} F^*$ yields

$$\begin{aligned} \left. \frac{d\pi}{dL} \right|_{L < \gamma DF^*} &= (\gamma DF^* - L)(\rho - 1)g(A^*) \frac{dA^*}{dL} - G(A^*) \left[(R - 1) - \frac{g(A^*)}{G(A^*)} DF^* \frac{dA^*}{dL} \right] \\ &= \left[DF^*(1 + \gamma(\rho - 1)) - L(\rho - 1) \right] g(A^*) \frac{dA^*}{dL} - G(A^*)(R - 1) \\ &= \left[DF^*(1 + \gamma(\rho - 1)) - L(\rho - 1) \right] g(A^*) \left(\frac{\rho - R}{1 - (1 + \gamma(\rho - 1)) DF^* \frac{g(A^*)}{G(A^*)}} \right) \\ &\quad - G(A^*)(R - 1), \end{aligned} \quad (17)$$

where the third line expands $\frac{dA^*}{dL}$ as in Equation (16) with $\frac{\partial A^*}{\partial L} = \rho - R$ and $\frac{\partial A^*}{\partial F^*} = -D(1 + \gamma(\rho - 1))$. Assumption $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$ guarantees that the term in parenthesis in the third line is positive for $L = L^*$. Therefore, $\frac{dA^*}{dL} > 0$, $\frac{d^2 A^*}{dL^2} < 0$, and π increases (decreases) in L when the expression in Equation (17) is positive (negative). Equation (8) follows directly from setting Equation (17) equal to zero. Since the RHS of Equation (8) decreases in L , the equation has one solution at most.

Corner solutions. The FOC in Equation (8) can be rearranged to

$$L^* = \left(\frac{1 + \gamma(\rho - 1)}{\rho - R} \right) DF^* - \left(\frac{1}{\rho - 1} \right) \left(\frac{R - 1}{\rho - R} \right) \frac{G(A^*)}{g(A^*)}. \quad (18)$$

Solution $L^* = 0$ arises whenever the RHS of Equation (18) is non-positive. We define function

$$A^* \equiv A^*(L^* = 0) = DR - (1 + \gamma(\rho - 1))DF^*, \quad \text{with} \quad F^* = \frac{r}{G(A^*)}. \quad (19)$$

Noting that the RHS of Equation (18) is weakly negative when $\rho \leq \underline{\rho}$, we obtain

$$\underline{\rho} - 1 + \gamma(\underline{\rho} - 1)^2 = \left(\frac{R - 1}{rD} \right) \frac{G(A^*)^2}{g(A^*)}. \quad (20)$$

From the Implicit Function Theorem (IFT) and assumption $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$ it follows that $\frac{dA^*}{d\rho} < 0$. Therefore, the RHS of Equation (20) decreases in $\underline{\rho}$, and so $\underline{\rho}$ is unique. Threshold $\underline{\rho}$ satisfies $R < \underline{\rho}$. To see this, use Equation (20) to note that inequality $R - 1 + \gamma(R - 1)^2 < \left(\frac{R-1}{rD} \right) \frac{G(A^*)^2}{g(A^*)}$ must hold under assumption $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$.

Solution $L^* = \gamma DF^*$ arises whenever the RHS of Equation (18) is weakly larger than γDF^* . We define function

$$A^{\bar{*}} \equiv A^*(L^* = \gamma DF^*) = DR - [1 + \gamma(R - 1)]DF^{\bar{*}}, \quad \text{with} \quad F^{\bar{*}} = \frac{r}{G(A^{\bar{*}})}. \quad (21)$$

Noting that the RHS of Equation (18) is weakly greater than γDF^* when $\rho \geq \bar{\rho}$, we obtain

$$\bar{\rho} - 1 = \frac{1}{1 + \gamma(R - 1)} \left(\frac{R - 1}{rD} \right) \frac{G(A^*)^2}{g(A^*)}. \quad (22)$$

Threshold A^* is independent of ρ , and so it is the RHS of Equation (22). Hence, $\bar{\rho}$ is unique and corner solution $L^* = \gamma DF^*$ is independent of ρ .

In section A.2.2 we show that $\frac{dL^*}{d\rho} > 0$ for interior solutions, $L^*[0, \gamma DF^*)$. It follows immediately that $\underline{\rho} < \bar{\rho}$.

Uniqueness assumption and existence of corner solutions. We showed that a solution to problem (7) exists and is unique when $\frac{\partial A^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} < 1$. Expansion of the condition yields $\frac{g(A^*)}{G(A^*)^2} rD(1 + \gamma(\rho - 1)) < 1$. It follows that the condition is satisfied for $\rho < \hat{\rho}$, with

$$\hat{\rho} - 1 = \frac{1}{\gamma} \left[\frac{1}{rD} \frac{G(A^*)^2}{g(A^*)} - 1 \right], \quad (23)$$

where $A^* \equiv A^*(\rho = \hat{\rho})$. Consider the case where $\rho = \hat{\rho}$ is such that $L^* = \gamma DF^*$, which we verify below. Then, the failure threshold is independent of ρ , so it satisfies $A^* = A^*$, with A^* characterised in (21). Two conclusions follow. First, threshold $\rho = \hat{\rho}$ in (23) is unique. Second, the characterizations of thresholds $\bar{\rho}$ and $\hat{\rho}$, in (22) and (23) respectively, are directly comparable. It is straightforward to verify that $\bar{\rho} < \hat{\rho}$ when $\frac{g(A^*)}{G(A^*)^2} rD(1 + \gamma(\rho - 1)) < 1$.

A.2.1 A benchmark for comparative statics

The equilibrium is given by system of equations $\{L^*, A^*, F^*\}$. Consider the system of total differentials $\left\{ \frac{dL^*}{di}, \frac{dF^*}{di}, \frac{dA^*}{di} \right\}$ where i is some exogenous parameters and $\frac{dL^*}{di} = \frac{\partial L^*}{\partial i} + \frac{\partial L^*}{\partial A^*} \frac{dA^*}{di} + \frac{\partial L^*}{\partial F^*} \frac{dF^*}{di}$, $\frac{dF^*}{di} = \frac{\partial F^*}{\partial i} + \frac{\partial F^*}{\partial A^*} \frac{dA^*}{di}$, and $\frac{dA^*}{di} = \frac{\partial A^*}{\partial i} + \frac{\partial A^*}{\partial L^*} \frac{dL^*}{di} + \frac{\partial A^*}{\partial F^*} \frac{dF^*}{di}$. We

omit superscripts from now. Solving the system yields:

$$\begin{aligned}
\frac{dL}{di} &= \left[\frac{\partial L}{\partial i} \left(1 - \frac{\partial A}{\partial F} \frac{\partial F}{\partial A} \right) + \frac{\partial A}{\partial i} \left(\frac{\partial L}{\partial A} + \frac{\partial L}{\partial F} \frac{\partial F}{\partial A} \right) + \frac{\partial F}{\partial i} \left(\frac{\partial L}{\partial F} + \frac{\partial L}{\partial A} \frac{\partial A}{\partial F} \right) \right] \mathcal{Q}, \\
\frac{dA}{di} &= \left[\frac{\partial A}{\partial i} + \frac{\partial L}{\partial i} \frac{\partial A}{\partial L} + \frac{\partial F}{\partial i} \left(\frac{\partial A}{\partial F} + \frac{\partial A}{\partial L} \frac{\partial L}{\partial F} \right) \right] \mathcal{Q}, \\
\frac{dF}{di} &= \left[\frac{\partial F}{\partial i} \left(1 - \frac{\partial A}{\partial L} \frac{\partial L}{\partial A} \right) + \frac{\partial A}{\partial i} \frac{\partial F}{\partial A} + \frac{\partial L}{\partial i} \frac{\partial F}{\partial A} \frac{\partial A}{\partial L} \right] \mathcal{Q},
\end{aligned} \tag{24}$$

with $\mathcal{Q} \equiv \left[1 - \frac{\partial F}{\partial A} \frac{\partial A}{\partial F} - \frac{\partial L}{\partial A} \frac{\partial A}{\partial L} - \frac{\partial F}{\partial A} \frac{\partial A}{\partial L} \frac{\partial L}{\partial F} \right]^{-1}$. The cross partial derivatives between interior equilibrium outcomes satisfy

$$\begin{aligned}
\frac{\partial A}{\partial L} &= \rho - R > 0, & \frac{\partial L}{\partial A} &= - \left(\frac{1}{\rho - 1} \right) \left(\frac{R - 1}{\rho - R} \right) \frac{\partial \left\{ \frac{G(A)}{g(A)} \right\}}{\partial A} < 0, \\
\frac{\partial A}{\partial F} &= -D(1 + \gamma(\rho - 1)) < 0, & \frac{\partial F}{\partial A} &= -r \frac{g(A)}{G(A)^2} < 0, \\
\frac{\partial L}{\partial F} &= \left(\frac{1 + \gamma(\rho - 1)}{\rho - R} \right) D > 0.
\end{aligned} \tag{25}$$

From Equation (25) and Assumption 1 it follows that $\mathcal{Q} > 0$.

A.2.2 Liquid reserves as a function of the penalty rate.

Interior solution. We use our benchmark for comparative statics in Section A.2.1.

First, consider the partial derivatives of $\{L^*, A^*, F^*\}$ with respect to ρ . Using expression for L^* in Equation (18) we obtain

$$\begin{aligned}
\frac{\partial L^*}{\partial \rho} &= - \left(\frac{1}{\rho - R} \right) \left[\left(\frac{1 + \gamma(R - 1)}{\rho - R} \right) DF^* - \frac{R - 1}{(\rho - 1)^2} \left(\frac{2\rho - R - 1}{\rho - R} \right) \right] \\
&= - \left(\frac{1}{\rho - R} \right) \left[\left(\frac{1 + \gamma(\rho - 1)}{\rho - R} - \gamma \right) DF^* - \left(\frac{R - 1}{(\rho - 1)(\rho - R)} + \frac{R - 1}{(\rho - 1)^2} \right) \frac{G(A^*)}{g(A^*)} \right] \\
&= - \left(\frac{1}{\rho - R} \right) \left[-(\gamma DF^* - L^*) - \frac{R - 1}{(\rho - 1)^2} \frac{G(A^*)}{g(A^*)} \right] > 0.
\end{aligned} \tag{26}$$

Moreover, we have $\frac{\partial A^*}{\partial \rho} = -(\gamma DF^* - L^*) < 0$ and $\frac{\partial F^*}{\partial \rho} = 0$. Plugging these expressions and those in (25), into system (24), we obtain $\frac{dL^*}{d\rho} > 0$.

Corner solutions. Corner solutions $L^* = 0$ and $L^* = \gamma DF^*$ are independent of ρ . The former is trivial; for the latter note that both A^* and F^* defined in Section A.2 are independent of ρ .

A.2.3 Liquid reserves as a function of the risk-free rate.

Interior solution. We use our benchmark for comparative statics in Section A.2.1. First, consider the partial derivatives of $\{L^*, A^*, F^*\}$ with respect to r . We have $\frac{\partial L^*}{\partial r} = \frac{\partial A^*}{\partial r} = 0$ and $\frac{\partial F^*}{\partial r} = 1/G(A^*) > 0$. From Equation (24) it follows that $\frac{dL^*}{dr} = \frac{\partial F^*}{\partial r} \left(\frac{\partial L^*}{\partial F^*} + \frac{\partial L^*}{\partial A^*} \frac{\partial A^*}{\partial F^*} \right) \mathcal{Q}$. Using Equation (25) we obtain $\frac{dL^*}{dr} > 0$.

Corner solutions. Consider threshold $\underline{\rho}$, which is implicitly characterised by Equation (20). The LHS of Equation (20) is independent of r . Meanwhile, from the IFT and Assumption 1 it follows that $\frac{dA^*}{dr} < 0$, so the RHS of Equation (20) decreases in r and we have $\frac{d\underline{\rho}}{dr} < 0$. The associated corner solution $L^* = 0$ is independent of r .

Consider threshold $\bar{\rho}$, which is implicitly characterised by Equation (22). The LHS of Equation (22) is independent of r . Meanwhile, from the IFT and Assumption 1 it follows that $\frac{dA^*}{dr} < 0$, so the RHS of Equation (22) decreases in r and we have $\frac{d\bar{\rho}}{dr} < 0$. The associated corner solution $L^* = \gamma DF^*$ increases in r . To see this, note that $\frac{dL^*}{dr} = \frac{\partial L^*}{\partial F^*} \frac{dF^*}{dr}$, where $\frac{dF^*}{dr} = \frac{\partial F^*}{\partial r} + \frac{\partial F^*}{\partial A^*} \frac{dA^*}{dr}$. Expanding terms yields $\frac{dL^*}{dr} = \gamma D \left(\frac{1}{G(A^*)} - r \frac{g(A^*)}{G(A^*)^2} \frac{dA^*}{dr} \right) > 0$.

A.3 Proof of Proposition 2

A.3.1 The marginal effects of the penalty rate on bank stability.

Interior solution. For $\rho \in (\underline{\rho}, \bar{\rho})$, our benchmark for comparative statics (Section A.2.1) reveals that $\frac{dA^*}{d\rho} \geq 0 \Leftrightarrow \frac{\partial A^*}{\partial \rho} + \frac{\partial L^*}{\partial \rho} \frac{\partial A^*}{\partial L^*} \geq 0$, i.e Equation (9), because $\frac{\partial F^*}{\partial \rho} = 0$. From equation (6) we obtain $\frac{\partial A^*}{\partial \rho} = -(\gamma DF^* - L^*)$, whereas $\frac{\partial A^*}{\partial L^*}$ and $\frac{\partial L^*}{\partial \rho}$ are respectively given by (25) and (26). It follows that

$$\begin{aligned} \frac{dA^*}{d\rho} > 0 &\Leftrightarrow 0 \leq -(\gamma DF^* - L^*) + \left(\frac{1}{\rho - R} \right) \left[\gamma DF^* - L^* + \left(\frac{R-1}{(\rho-1)^2} \right) \frac{G(A^*)}{g(A^*)} \right] (\rho - R) \\ &\Leftrightarrow 0 < \frac{R-1}{(\rho-1)^2} \frac{G(A^*)}{g(A^*)}, \end{aligned} \quad (27)$$

which is always satisfied. Therefore, $\frac{dA^*}{d\rho} > 0$ for $\rho \in (\underline{\rho}, \bar{\rho})$.

Corner solutions. With $\rho \leq \underline{\rho}$, stability is given by failure threshold A^* characterised in equation (19). From the IFT and Assumption 1 it follows that $\frac{dA^*}{d\rho} < 0$. With $\rho \geq \bar{\rho}$, stability is given by failure threshold A^* characterised in equation (21), which is independent of ρ .

A.3.2 Bank stability with and without LLR.

Section A.3.1 reveals that the failure threshold A^* depicts a V-shape as a function of the penalty rate $\rho \in (R, \bar{\rho})$ with a minimum level attained for $\rho = \underline{\rho}$.

In absence of LLR, the bank incurs a liquidation cost $1/\psi > \bar{\rho}$. We know that the bank's failure threshold for corner solution $L = \gamma DF^*$ is independent of ρ , i.e. $\frac{dA^*}{d\rho} = 0$, so bank stability is given for any $\rho \geq \bar{\rho}$ or in absence of LLR is given by A^* in equation (21). This is the stability level achieved at the 'right' end of the V-shape.

Consider bank stability under the most benevolent LLR, i.e., the bank's failure threshold at the 'left' end of the V-shape, i.e., for ρ arbitrarily close to R . We know $R < \underline{\rho}$, so the bank's choice is given by corner solution $L^* = 0$ and its stability is characterised by A^* in equation (19). Evaluation of the threshold at $\rho = R$ yields $A^*(\rho = R) = RD - [1 + \gamma(R - 1)]DF^*$, which equals failure threshold A^* in equation (21) corresponding to the corner solution where bank liquidity is plentiful, $L^* = \gamma DF^*$. That is, $A^*(\rho = R) = A^*(\rho \geq \bar{\rho})$. Moreover, it follows that an impactful LLR, $\rho < \bar{\rho}$, decreases stability.

A.4 Proof of Proposition 3

The LLR is characterised by a given penalty rate ρ . Meanwhile, the absence of LLR implies that the bank can only liquidate assets at a premium ψ^{-1} . A LLR has impact on equilibrium outcomes if, and only if, $\rho < \min\{\bar{\rho}, \psi^{-1}\}$; we consider the LLR to be *impactful*.

First, we show that an impactful LLR increases bank profits. Next, we show that when $\bar{\rho} \leq \psi^{-1}$, an impactful LLR always reduces stability. Instead, when $\bar{\rho} > \psi^{-1}$, an impactful LLR reduces stability only for a sufficiently large penalty rate $\rho > \tilde{\rho}$, where $R < \tilde{\rho} < \underline{\rho}$. Last, we prove that for any impactful LLR that reduces stability, there exists a unique critical threshold $\lambda = \hat{\lambda}$ on the social cost of bank failure.

A.4.1 Bank profits as a function of the penalty rate.

We show that bank profits decrease with the LLR's penalty rate $\frac{d\pi^*}{d\rho} < 0$ whenever the LLR is impactful. Given that an impactful LLR satisfies $\rho < \psi^{-1}$, it follows that it increases bank profits.

Consider profits as characterised in (7), in equilibrium. Total differentiation

yields $\frac{d\pi^*}{d\rho} = \frac{\partial\pi^*}{\partial\rho} + \frac{\partial\pi^*}{\partial A^*} \frac{dA^*}{d\rho} + \frac{\partial\pi^*}{\partial L^*} \frac{dL^*}{d\rho} + \frac{\partial\pi^*}{\partial F^*} \frac{dF^*}{d\rho}$:

$$\frac{d\pi^*}{d\rho} = g(A^*) \left[R(D - L^*) + L^* - A^* - DF^* \right] \frac{dA^*}{d\rho} - G(A^*)(R - 1) \frac{dL^*}{d\rho} + g(A^*) DF^* \frac{dA^*}{d\rho} \quad (28)$$

where the last term of (28) uses $\frac{dF^*}{d\rho} = \frac{\partial F^*}{\partial A^*} \frac{dA^*}{d\rho} = -\frac{g(A^*)}{G(A^*)} F^* \frac{dA^*}{d\rho}$.

Interior solution. For $\rho \in (\underline{\rho}, \bar{\rho})$, using the relevant expression of the failure threshold A^* in equation (6) and rearranging terms, we obtain

$$\frac{d\pi^*}{d\rho} = \left[(\rho - 1)(\gamma DF^* - L^*) + DF^* \right] g(A^*) \frac{dA^*}{d\rho} - G(A^*)(R - 1) \frac{dL^*}{d\rho}. \quad (29)$$

We have $\frac{dL^*}{d\rho} > 0$ (Appendix A.2.2) and $\frac{dA^*}{d\rho} > 0$ (Proposition 2). The FOC in (8) reveals that, in equilibrium, $[(\rho - 1)(\gamma DF^* - L^*) + DF^*]g(A^*) = G(A^*)(R - 1)/\frac{dA^*}{dL}$. Using this condition in (29) yields

$$\begin{aligned} \frac{d\pi^*}{d\rho} < 0 &\Leftrightarrow \frac{dA^*}{dL} > \frac{dA^*/d\rho}{dL^*/d\rho} \\ &\Leftrightarrow \frac{\frac{\partial A^*}{\partial L}}{\left(1 - \frac{\partial A^*}{\partial F} \frac{\partial F^*}{\partial A^*}\right)} > \frac{\frac{\partial A^*}{\partial \rho} + \frac{\partial L^*}{\partial \rho} \frac{\partial A^*}{\partial L}}{\frac{\partial L^*}{\partial \rho} \left(1 - \frac{\partial A^*}{\partial F} \frac{\partial F^*}{\partial A^*}\right) + \frac{\partial A^*}{\partial \rho} \left(\frac{\partial L^*}{\partial A^*} + \frac{\partial L^*}{\partial F^*} \frac{\partial F^*}{\partial A^*}\right)} \\ &\Leftrightarrow \frac{\partial A^*}{\partial \rho} \frac{\partial A^*}{\partial L} \left(\frac{\partial L^*}{\partial A^*} + \frac{\partial L^*}{\partial F^*} \frac{\partial F^*}{\partial A^*} \right) > \frac{\partial A^*}{\partial \rho} \left(1 - \frac{\partial A^*}{\partial F} \frac{\partial F^*}{\partial A^*} \right), \end{aligned} \quad (30)$$

where the second line uses equation (16) for the LHS and the system in (24) for the RHS. In the last line of (30), the LHS is positive, whereas Assumption 1 implies that the RHS is negative. Therefore, the inequality holds and $\frac{d\pi^*}{d\rho} < 0$.

Corner solutions. For $\rho \leq \underline{\rho}$, the relevant expression of the failure threshold A^* in equation (6) is the same as at the interior solution, so total differentiation of profits in equilibrium is characterised by Equation (29). Further, we have $\frac{dL^*}{d\rho} = 0$ and $A^* = \bar{A}^*$, with $\frac{dA^*}{d\rho} < 0$ (Proposition 2). Therefore, the expression in (29) simplifies

to $\frac{d\pi^*}{d\rho} = DF^*(1 + \gamma(\rho - 1))g(A^*)\frac{dA^*}{d\rho} < 0$.

For $\rho \geq \bar{\rho}$ (a non-impactful LLR), the relevant expression of the failure threshold in equation (6) is $A^* = \bar{A}$. Further, we have $\frac{dL^*}{d\rho} = 0$ and $A^* = A^*$, with $\frac{dA^*}{d\rho} = 0$ (Proposition 2). From Equation (28) it follows that $\frac{d\pi^*}{d\rho} = 0$.

A.4.2 Welfare with and without LLR.

Consider an impactful LLR and use the main text subscripts to denote equilibrium outcomes in the absence (0) and the presence (1) of the LLR. Appendix A.4.1 shows that it increases bank profits, $\pi_0^* < \pi_1^*$. Moreover, Appendix A.3.2 shows that the LLR decreases bank stability $A_1^* < A_0^*$, when $\bar{\rho} < 1/\psi$, or when $\bar{\rho} \geq 1/\psi$ and $\rho > \tilde{\rho}$; it increases bank stability $A_1^* > A_0^*$, otherwise.

Welfare characterised in (10) is smaller in absence of LLR when $\pi_0^* + rD - \lambda(1 - G(A_0^*)) \leq \pi_1^* + rD - \lambda(1 - G(A_1^*))$ or, equivalently, for

$$\lambda \leq \frac{\pi_1^* - \pi_0^*}{G(A_0^*) - G(A_1^*)} \equiv \hat{\lambda}. \quad (31)$$

When LLR decreases bank stability $A_1^* < A_0^*$, threshold $\hat{\lambda}$ is well-defined and unique. When LLR increases bank stability $A_1^* \geq A_0^*$, threshold does not exist for positive values of λ .

A.5 Proof of Corollary 1

A.5.1 Solution to the planner's problem

We characterise the welfare-maximizing liquidity choice following a derivation process analogous to that in Proof of Proposition 1, Section A.2.

Interior solution. Total differentiation of the planner's objective (10) yields

$$\frac{dW}{dL} = \frac{d\pi}{dL} + \frac{dA^P}{dL}g(A^P)\lambda, \quad (32)$$

with $\frac{d\pi}{dL}$ characterised by (14), and expansion of $\frac{dA^P}{dL}$ given by (16), both with the relevant superscript P .

Consider $L > \gamma DF^P$. From equation (15) it follows that $\frac{d\pi}{dL}|_{L>\gamma DF^P} < 0$. Moreover, we have $\frac{\partial A^P}{\partial L} = -(R-1) < 0$, implying $\frac{dA^P}{dL} < 0$. Hence, it holds that $\frac{dW}{dL}|_{L>\gamma DF^P} < 0$, and therefore $L^P \leq \gamma DF^P$.

Consider $L^P \leq \gamma DF^P$. For now, assume that $\rho < \hat{\rho}$ (Assumption 1) guarantees $\frac{\partial A^P}{\partial F^P} \frac{\partial F^P}{\partial A^P} < 1$; we show subsequently that this is true. Then, we have that $\frac{\partial A^P}{\partial L} = \rho - R > 0$ implies $\frac{dA^P}{dL} > 0$ and the FOC associated to (32) is given by (11). This characterises implicitly the optimal liquidity choice when the solution to the planner's problem is interior.

Corner solutions. The FOC in equation (11) can be rearranged to

$$L^P = \left(\frac{1 + \gamma(\rho - 1)}{\rho - R} \right) DF^P - \left(\frac{1}{\rho - 1} \right) \left(\frac{R - 1}{\rho - R} \right) \frac{G(A^P)}{g(A^P)} + \frac{\lambda}{\rho - 1}. \quad (33)$$

Solution $L^P = 0$ arises whenever the RHS of (33) is non-positive. Analogous to the derivation of Proof of Proposition 1, Appendix A.2, we define $A^P \equiv A^P(L^P = 0) = DR - (1 + \gamma(\rho - 1))DF^P$, where $F^P = r/G(A^P)$. Noting that the RHS of equation (33) equals zero when $\rho = \underline{\rho}^P$, we obtain

$$\underline{\rho}^P - 1 + \gamma(\underline{\rho}^P - 1)^2 = \left[\frac{G(A^P)}{g(A^P)}(R - 1) - \lambda(\underline{\rho}^P - R) \right] \frac{G(A^P)}{rD}. \quad (34)$$

From the IFT and Assumption 1 it follows that $\frac{dA^P}{d\rho} < 0$. Moreover, when the RHS

of (34) is positive, it increases in A^P , so it decreases on $\underline{\rho}^P$, and therefore there is a unique $\underline{\rho}^P > 0$. Further, $\underline{\rho}^P$ exists only when the RHS of (34) is positive, which is true for

$$\lambda < \underline{\lambda} \equiv \left(\frac{R-1}{\underline{\rho}^P - R} \right) \frac{G(A^P)}{g(A^P)}. \quad (35)$$

Solution $L^P = \gamma DF^P$ arises whenever the RHS of (33) is weakly greater than γDF^P . Analogous to the derivation of Proof of Proposition 1, Appendix A.2, we define $A^{\bar{P}} \equiv A^P(L^P = \gamma DF^P) = DR - [1 + \gamma(R-1)]DF^{\bar{P}}$, where $F^{\bar{P}} = r/G(A^{\bar{P}})$. Noting that the RHS of Equation (33) equals $\gamma DF^{\bar{P}}$ when $\rho = \bar{\rho}^P$, we obtain

$$\bar{\rho}^P - 1 = \frac{1}{1 + \gamma(R-1)} \left[\frac{G(A^{\bar{P}})}{g(A^{\bar{P}})}(R-1) - \lambda(\bar{\rho}^P - R) \right] \frac{G(A^{\bar{P}})}{rD}. \quad (36)$$

From the IFT and Assumption 1 it follows that $\frac{dA^{\bar{P}}}{d\rho} < 0$. Moreover, when the RHS of (36) is positive, it increases on $A^{\bar{P}}$, so it decreases on $\bar{\rho}^P > 0$, and therefore there is a unique $\bar{\rho}^P > 0$. Further, $\bar{\rho}^P$ exists only when the RHS of (36) is positive, which is true for

$$\lambda < \bar{\lambda} \equiv \left(\frac{R-1}{\bar{\rho}^P - R} \right) \frac{G(A^{\bar{P}})}{g(A^{\bar{P}})}. \quad (37)$$

Uniqueness. Analogous to the benchmark setting, uniqueness is guaranteed for $\frac{\partial A^P}{\partial F^P} \frac{\partial F^P}{\partial A^P} < 1$. Following the steps to obtain the uniqueness condition in Section A.2, here we obtain uniqueness condition $\rho < \hat{\rho}^P$, with $\hat{\rho}^P - 1 = \frac{1}{\gamma} \left[\frac{1}{rD} \frac{G(A^{\hat{P}})^2}{g(A^{\hat{P}})} - 1 \right]$ and such that $\bar{\rho}^P < \hat{\rho}^P$.

Notice that the bank's liquidity choice and associated failure threshold equal those of the planner for $\rho \geq \bar{\rho}$. That is, we have $L^P = L^* = \gamma DF^*$ and, because failure threshold A^* is independent of ρ , we also have $A^* = A^{\hat{P}}$. It follows that $\hat{\rho}^P = \hat{\rho}$, so Assumption 1 guarantees uniqueness of the planner's solution.

A.5.2 Comparison of solutions to the bank's and the planner's problems

Interior solutions. The systems of equations characterizing solutions to the bank's and the planner's problems feature the same dependencies. We use our benchmark for comparative statics in Section A.2.1 to study the effect of λ on interior equilibrium outcomes. Using (33) we obtain $\frac{\partial L^P}{\partial \lambda} = \frac{1}{\rho-1} > 0$. Moreover, we have $\frac{\partial A^P}{\partial \lambda} = \frac{\partial F^P}{\partial \lambda} = 0$. It follows that $\frac{dL^P}{d\lambda} > 0$ and $\frac{dL^P}{d\lambda} > 0$, so $L^P > L^*$ and $A^P > A^*$ for interior values.

Corner solutions. From (34) and (36) we use the IFT and $\frac{dA^*}{d\rho} = \frac{dA^*}{d\rho} = 0$ to note that both $\underline{\rho}^P$ and $\bar{\rho}^P$ diminish in λ . Therefore, we have $\underline{\rho}^P < \underline{\rho}$ and $\bar{\rho}^P < \bar{\rho}$. It follows that $\{L^*, A^*, F^*\} = \{L^P, A^P, F^P\}$ for any $\rho \leq \underline{\rho}^P$ and $\rho \geq \bar{\rho}^P$.

A.5.3 Minimum liquidity regulation

Equation (8) reveals that bank marginal profits only cross zero once as a function of L , and are positive for $L < L^*$ and negative for $L > L^*$. Therefore, bank profits are quasi-concave in L and reach a unique maximum at $L = L^*$. It follows that (a) regulation $\underline{L}$ binds if and only if $L^* < \underline{L}$; (b) when regulation $\underline{L}$ binds, the bank chooses $L = \underline{L}$.

A.6 Proof of Proposition 5

Welfare is strictly higher under both the LLR and MLR We show separately $W_1^P > W_1$, and $W_1^P > W_0$.

Allocation L^* maximizes bank profits in (7); MLR implement the planner's allocation L^P , which maximizes welfare in (10). Therefore $W_1^P > W_1$ and $W_0^P > W_0$ whenever $L^P \neq L^*$, i.e., for $\rho \in (\underline{\rho}^P, \bar{\rho}^P)$.

To show $W_1^P > W_0$, we prove $W_1^P > W_0^P$ (the LLR improves welfare when MLR are in place) and note that we already know $W_0^P > W_0$. Total differentiation yields

$$\begin{aligned} \frac{W_1^P}{d\rho} &= \frac{d\pi^P}{d\rho} + \lambda g(A^P) \frac{dA^P}{d\rho} \\ &= \left[(\rho - 1)(\gamma DF^P - L^P) + DF^P + \lambda \right] g(A^P) \frac{dA^P}{d\rho} - G(A^P)(R - 1) \frac{dL^P}{d\rho}, \end{aligned} \quad (38)$$

where the second line is obtained by differentiating profits π^P as in (28) and (29). For $\rho \geq \underline{\rho}^P$, FOC (11) reveals that, in equilibrium, $[(\rho - 1)(\gamma DF^P - L^P) + DF^P + \lambda]g(A^P) = G(A^P)(R - 1)/\frac{dA^P}{dL}$. Plugging this in (38) yields $\frac{dW_1^P}{d\rho} < 0 \Leftrightarrow \frac{dA^P}{dL} > \frac{dA^P/d\rho}{dL^P/d\rho}$, and we can verify that $\frac{dW_1^P}{d\rho} < 0$ by developing the associated condition as in (30). For $\rho < \underline{\rho}^P$, we have $\frac{dA^P}{d\rho} < 0$ and $\frac{dL^P}{d\rho} = L^P = 0$. It follows immediately that $\frac{W_1^P}{d\rho} < 0$. It follows that $W_1^P > W_0^P$.

The LLR is more effective in the presence of MLR The bank does not internalize the social cost of failure, so $\frac{dW_1}{d\lambda} = \frac{\partial W_1}{\partial \lambda} = -(1 - G(A^*))$. Meanwhile, for the planner's solution, we have $\frac{dW_1^P}{d\lambda} = \frac{\partial W_1^P}{\partial \lambda} + \frac{\partial W_1^P}{\partial \pi} \frac{d\pi^P}{d\lambda} + \frac{\partial W_1^P}{\partial A} \frac{dA^P}{d\lambda}$. Differentiation of profits yields

$$\frac{d\pi^P}{d\lambda} = \left[(\rho - 1)(\gamma DF^P - L^P) + DF^P \right] g(A^P) \frac{dA^P}{d\lambda} - G(A^P)(R - 1) \frac{dL^P}{d\lambda}, \quad (39)$$

where (39) uses the relevant expression for A^P in a derivation analogous to that in (28) and (29). Then, by expanding the differentiation of welfare evaluated at equilibrium one obtains:

$$\begin{aligned} \frac{dW_1^P}{d\lambda} &= \frac{\partial W_1^P}{\partial \lambda} + \frac{d\pi^P}{d\lambda} + \lambda g(A^P) \frac{dA^P}{d\lambda} \\ &= \frac{\partial W_1^P}{\partial \lambda} + \left[(\rho - 1)(\gamma DF^P - L^P) + DF^P + \lambda \right] g(A^P) \frac{dA^P}{d\lambda} - G(A^P)(R - 1) \frac{dL^P}{d\lambda} \\ &= \frac{\partial W_1^P}{\partial \lambda} + G(A^P)(R - 1) \left[\frac{dA^P/d\lambda}{dA^P/dL} - \frac{dL^P}{d\lambda} \right]. \end{aligned} \quad (40)$$

The first line of (40) uses $\partial W_1^P / \partial A^P = \lambda g(A^P)$; the second plugs in the total differentiation of profits in (39); the third line plugs in the FOC (11) rearranged as $[(\rho - 1)(\gamma DF^P - L^P) + DF^P + \lambda]g(A^P) = G(A^P)(R - 1)/\frac{dA^P}{dL}$. The term in brackets in the third line of (40) equals zero. To see this, expand $\frac{dA^P}{dL}$ as in (16), and expand both $\frac{dA^P}{d\lambda}$ and $\frac{dL^P}{d\lambda}$ using system (24) and noting that $\frac{\partial A^P}{\partial \lambda} = \frac{\partial F^P}{\partial \lambda} = 0$. It follows that, in equilibrium, $\frac{dW_1^P}{d\lambda} = \frac{\partial W_1^P}{\partial \lambda} = -(1 - G(A^P))$. Last, note that $A^P > A^*$ whenever $L^P \neq L^*$, i.e., for $\rho \in (\underline{\rho}^P, \bar{\rho})$. It follows that $\frac{dW_1}{d\lambda} < \frac{dW_1^P}{d\lambda} < 0$.

A.7 Proof of Proposition 6

A.7.1 MLR strictness weakly increases with LLR's penalty rate

MLR equate the planner's solution, $\underline{L} = L^P$ (Corollary 1). An analysis of the planner's solution analogous to that of the bank's private solution in Section A.2.2 reveals $\frac{dL^P}{d\rho} > 0$ for $\rho \in [\underline{\rho}, \bar{\rho})$, and $\frac{dL^P}{d\rho} = 0$ for $\rho < \underline{\rho}$ and $\rho \geq \bar{\rho}$. See Figure 5 for a graphical representation.

A.7.2 MLR stringency is nonmonotonic in LLR's penalty rate

Section A.5.2 shows $\underline{\rho}^P < \underline{\rho}$ and $\bar{\rho}^P < \bar{\rho}$; Corollary 1 shows that MLR are binding for $\rho \in (\underline{\rho}^P, \bar{\rho})$.

For $\rho \in (\underline{\rho}^P, \bar{\rho}^P)$, we have $L^P > 0$ with $\frac{dL^P}{d\rho} > 0$, as well as $L^* = 0$ and $\frac{dL^*}{d\rho} = 0$. Therefore, stringency increases with ρ . For $\rho \in (\bar{\rho}^P, \bar{\rho})$, we have $L^P = \gamma DF^P$ with $\frac{dL^P}{d\rho} = 0$, and $L^* < \gamma DF^P$ with $\frac{dL^*}{d\rho} > 0$. Therefore, stringency decreases with ρ .

References

- Ahnert, T., K. Anand, P. Gai, and J. Chapman (2019). Asset encumbrance, bank funding and fragility. Review of Financial Studies 32(6), 2422–55.
- Aldasoro, I., S. Doerr, and H. Zhou (2026). Liquidity regulation and bank funding costs. Hong Kong University Business School WP 2026/024.
- Allen, F. and D. Gale (2000). Financial contagion. Journal of Political Economy 108(1), 1–33.
- Allen, F. and D. Gale (2007). Understanding Financial Crises. Oxford University Press.
- Bagehot, W. (1873). Lombard Street. H.S. King, London.
- Bonfim, D. and J. A. Santos (2023). The importance of deposit insurance credibility. Journal of Banking & Finance 154, 106916.
- Bowman, M. (2024). Bank Liquidity, Regulation, and the Fed’s Role as Lender of Last Resort.
- Calomiris, C. and C. Kahn (1991). The role of demandable debt in structuring optimal banking arrangements. American Economic Review 81(3), 497–513.
- Calomiris, C. W. and J. R. Mason (2003, November). Fundamentals, Panics, and Bank Distress During the Depression. American Economic Review 93(5), 1615–1647.
- Carlsson, H. and E. van Damme (1993). Global games and equilibrium selection. Econometrica 61(5), 989–1018.
- Cipriani, M., T. M. Eisenbach, and A. Kovner (2024, May). Tracing Bank Runs in Real Time. Staff Reports (Federal Reserve Bank of New York), Federal Reserve Bank of New York.
- Congressional Research Office (2023). Bank Term Funding Program (BTFP) and Other Federal Reserve Support to Banking System in Turmoil.

- Crawley, E., T. Doh, and M. Shin (2023). Failure of Silicon Valley Bank Reduced Local Consumer Spending but Had Limited Effect on Aggregate Spending. Federal Reserve Bank of Kansas City Economic Bulletin.
- Diamond, D. and P. Dybvig (1983). Bank runs, deposit insurance and liquidity. Journal of Political Economy 91, 401–19.
- Diamond, D. and R. Rajan (2001). Liquidity risk, liquidity creation, and financial fragility: A theory of banking. Journal of Political Economy 109(2), 287–327.
- Egan, M., A. Hortacsu, and G. Matvos (2017). Deposit competition and financial fragility: Evidence from the us banking sector. American Economic Review 107(1), 169–216.
- Freixas, X. and B. M. Parigi (2008). Lender of last resort and bank closure policy. Cesifo working paper series.
- Huang, C., L. Shen, and J. Zou (2024). Policy Portfolio for Banks: Deposit Insurance and Liquidity Injection.
- James, C. (1991). The losses realized in bank failures. Journal of Finance 46(4), 1223–1242.
- Keeley, M. C. (1990). Deposit insurance, risk, and market power in banking. The American Economic Review 80(5), 1183–1200.
- Kim, D. and A. M. Santomero (1988). Risk in banking and capital regulation. The Journal of Finance 43(5), 1219–33.
- Krishnamurthy, A. (2010). How Debt Markets Have Malfunctioned in the Crisis. Journal of Economic Perspectives 24, 3–28.
- Li, Z. and K. Ma (2022, December). Contagious Bank Runs and Committed Liquidity Support. Management Science 68(12), 9152–9174.
- Luck, S., M. Plosser, and J. Younger (2023). Bank Funding during the Current Monetary Policy Tightening Cycle.

- Morris, S. and H. S. Shin (2003). Global Games: Theory and Applications. In M. Dewatripont, L. P. Hansen, and S. J. Turnovsky (Eds.), Advances in Economics and Econometrics, pp. 56–114. Cambridge: Cambridge University Press.
- Powell, J. (2023). Financial stability and economic developments.
- Ratnovski, L. (2009). Bank liquidity regulation and the lender of last resort. Journal of Financial Intermediation 18(4), 541–58.
- Repullo, R. (2005). Liquidity, Risk Taking, and the Lender of Last Resort. International Journal of Central Banking 1, 47–80.
- Rochet, J.-C. and X. Vives (2004). Coordination failures and the Lender of Last Resort: was Bagehot right after all? Journal of the European Economic Association 2(6), 1116–47.
- Santos, J. A. and J. Suarez (2019). Liquidity standards and the value of an informed lender of last resort. Journal of Financial Economics 132(2), 351–68.
- Schilling, L. (2025, May). Pitfalls and Optimal Design of Emergency Liquidity Assistance.
- Sundaresan, S. and K. Xiao (2024, January). Liquidity regulation and banks: Theory and evidence. Journal of Financial Economics 151, 103747.
- Thornton, H. (1802). An inquiry into the nature and effects of the paper credit of Great Britain. George Allen and Unwin.
- Tucker, P. (2014). The lender of last resort and modern central banking: principles and reconstruction.
- Tucker, P. (2016). How can central banks deliver credible commitment and be “emergency institutions”? In Central bank governance and oversight reform. Hoover Institute.
- Vives, X. (2005). Complementarities and games: New developments. Journal of Economic Literature 43, 437–79.
- Vives, X. (2014). Strategic complementarity, fragility, and regulation. Review of Financial Studies 27(12), 3547–92.