



FTG Working Paper Series

Rules versus Disclosure: Prudential Regulation and Market Discipline

by

William Fuchs
Daniel Neuhann
Satoshi Fukuda

Working Paper No. 00145-00

Finance Theory Group

www.financetheory.com

*FTG working papers are circulated for the purpose of stimulating discussions and generating comments. They have not been peer reviewed by the Finance Theory Group, its members, or its board. Any comments about these papers should be sent directly to the author(s).

Rules versus Disclosure: Prudential Regulation and Market Discipline*

William Fuchs[†] Satoshi Fukuda[‡] Daniel Neuhann[§]

November 14, 2024

Abstract

We study the joint design of two prominent micro-prudential policy tools: bank regulation that enforces operational standards via rules, and market discipline through information disclosure. Disclosure can be state-contingent but creates a trade-off between incentives and the ex-post protection of weak banks. Hence, regulators use rules to maintain incentives and imperfect disclosure to provide ex-post insurance. In the optimal design, there is *precautionary regulation* to lower the risk of market freezes, and more disclosure in bad times to restore trade. Systemically important banks face more regulation but less disclosure. Banks prefer more disclosure but less regulation.

JEL Classification: G21; G28; D82

Keywords: Bank Regulation; Stress Test; Disclosure; Information Design; Moral Hazard; Banking Crises

*We would like to thank Jean Edouard Colliard, Simon Gervais, Itay Goldstein, Guillermo Ordonez, Paymon Khorrami, George Mailath, Markus Parasca, Rafael Repullo, Andy Skrzypacz, Chester Spatt, Elu von Thadden, João Thereze, Vish Viswanathan, Basil Williams (discussant), Tom Wiseman, Liyan Yang (discussant) and seminar participants at Bocconi University, Conference on Frontier Risks, Financial Innovation and Prudential Regulation of Banks, Duke University, INSEAD, Midwest Economic Theory Conference, Universidad Carlos III de Madrid, UT Austin, Virginia Tech and Wharton Conference on Liquidity and Financial Fragility for very helpful conversations. We thank Kevin Mei for excellent research assistance.

[†]McCombs School of Business, UT Austin, Universidad Carlos III Madrid, CEPR, and FTG. E-mail: william.fuchs@mcombs.utexas.edu.

[‡]Department of Decision Sciences and IGIER, Bocconi University. E-mail: satoshi.fukuda@unibocconi.it.

[§]McCombs School of Business, UT Austin and FTG. E-mail: daniel.neuhann@mcombs.utexas.edu.

1 Introduction

One of the central goals of financial regulation is to reduce the risk of financial crises by ensuring that banks do not hold too many risky assets on their balance sheets. In practice, this task is complicated by the fact that markets for financial assets are typically plagued by adverse selection, making it difficult for banks to sell assets to other investors and creating a natural tendency for banks to retain risky assets. Because the extent of adverse selection depends on the information environment and the ex-ante incentives of banks, this creates a natural role for supervisory regulation and disclosure in managing the allocation of risk.

In line with this view and as a response to the Global Financial Crisis, recent regulatory frameworks such as Basel III assign prominent roles to regulation and supervision on the one hand, and market discipline induced by information disclosure on the other. Existing work has studied the optimal design of individual tools in isolation. Here, our goal is to derive some principles of the optimal *joint* design of regulation and disclosure. Our approach is motivated by the fact that both tools exhibit natural synergies: market discipline through disclosure is valuable when supervision is ineffective, while detailed disclosure about particular banks is less exigent when regulation ensures prudent behavior. Hence, optimally trading off both tools may result in a lower regulatory burden than using each tool individually.

We consider these issues in a model of bank moral hazard in which there is a social benefit to banks selling their assets to other financial investors but asset markets are hampered by adverse selection. Central to our approach is the interaction between fixed *rules* that can directly influence bank behavior by enforcing operational standards, and information disclosure that can *flexibly* respond to exogenous shocks but may engender moral hazard. Specifically, we view regulation and supervision as ex-ante tools that are relatively efficient at deterring bad behavior (e.g., poor operational standards that lead to excessive risk taking), but cannot respond to shocks (i.e., bad results that arise despite sound behavior). Conversely, disclosure can be flexibly adapted to the realized state because it is an ex-post tool, but it also creates a trade-off with respect to incentives: while the regulator wants transparency and market discipline to induce banks to exert effort, she may be tempted to partially obfuscate information to allow even bad banks to sell their assets ex-post. We ask how a welfare-maximizing regulator should use both tools in an overarching regulatory framework.

Our main findings are as follows. In the absence of market discipline through disclosure, the regulator chooses to prescribe excessively high effort. This is because the lack of state-contingency forces the regulator to require “precautionary effort” to ensure that the secondary markets for bank assets do not break down after bad shocks, and banks do not fully internalize the social value of trading assets. Conversely, holding fixed bank effort, the

optimal disclosure policy is to partially obfuscate bank asset quality. More transparency is necessary after bad shocks in order to maintain secondary market liquidity. However, because obfuscation redistributes from healthy to distressed banks, it creates ex-ante moral hazard that may reduce welfare. In the optimal joint design, disclosure serves to maintain market liquidity while rules deter moral hazard. Despite using multiple tools, the optimal joint design reduces the overall regulatory burden on banks because there is less need for precautionary effort. Comparing across banks, we find that more systemically important institutions should face tighter effort regulation but are also optimally subjected to less market discipline via reduced disclosure. These findings lend support to the general approach laid out in the Basel III framework, and go beyond the basic framework by establishing a number of principles of optimal design. They are also in line with the Dodd-Frank Act provisions that call for closer supervision of systemically important financial institutions.

Formally, we consider a model in which a bank first originates assets of uncertain quality (high or low) subject to moral hazard. The probability of obtaining the high quality asset is increasing in unobserved bank effort and a publicly observed aggregate shock that is realized after the effort is sunk. The bank can keep its assets or sell them to a competitive fringe of buyers subject to asymmetric information. There are private gains from trading high quality assets, but not for low quality assets. Hence, the market may break down if average asset quality is too low. However, a social planner would prefer all assets to be traded because there is an additional social value to moving assets off banks' balance sheets. We refer to this consideration simply as the "externality." This assumption captures the idea that the regulator is worried about the risk of bank failure and cannot commit to not bailing out banks in the event of distress. However, if bad assets are held outside of the banking system then there is less risk of financial crisis and therefore a lower likelihood of bailouts. Since non-bank investors are not expected to be bailed out, they value bad assets less than banks. The value of the externality captures the social cost of bailouts not internalized by banks. As such, we will vary the size of the externality to inform how optimal policy should account for the systemic importance of banks.

In this setting, we study the welfare-optimal design of two regulatory tools. The first is a combination of ex-ante rules and supervision which we refer to as *regulation*. This tool allows the regulator to directly influence bank effort subject to some limits which allow us to capture different levels of regulatory capacity.¹ While this tool is effective at reducing

¹We will abstract from the exact tools used to focus on the interactions and trade-offs with the disclosure policy. For instance, the Office of the Comptroller of the Currency defines high-quality bank supervision as assessing "whether each bank has a sound risk management system consisting of policies, processes, personnel, and control systems to measure, monitor, and control risk" (OCC, 2019). Regulatory assessments of bank practices are explicitly codified in the CAMELS rating system.

moral hazard, it is also blunt because it cannot respond to the aggregate shock. The second is *disclosure*, which is the ability to partially or fully reveal the *realized* quality of bank assets to potential buyers. This tool allows the regulator to modulate the degree of market discipline via the information structure. Because the policy pertains to realized asset quality, it can also respond flexibly to aggregate shocks.

We first study each policy instrument in isolation to have a better understanding of their roles. If the regulator must rely only on market discipline then it faces a trade-off between providing ex-post insurance (by not revealing the quality of some bad assets) and ex-ante incentives (effort is less valuable if the market will not learn the asset quality).

The amount of ex-post insurance will be constrained by the state of the economy θ . Fixing the state, the optimal disclosure rule is to honestly reveal when assets are of high quality, but to report ‘good’ with probability $\beta(\theta)$ when assets are in fact of low quality. The value of $\beta(\theta)$ cannot be too high because the market price after a good report must still be sufficiently high to induce banks with high-quality assets to trade. Thus, unlike the standard Bayesian Persuasion setting (e.g., Kamenica and Gentzkow, 2011) we have to contend with both ex-ante effort considerations and the ex-post incentive compatibility of trades for different seller types. We show in Proposition 1 that when the social value of trading bad assets is low, insurance becomes less relevant and the optimal policy calls for full transparency (i.e., $\beta(\cdot) = 0$). This is because the only relevant consideration in this case is ex-post adverse selection, and providing full transparency deals with this in a manner that provides the right effort incentives for the bank. Next, when the social value of trading bad assets is high, we show that the optimal policy will call for partial disclosure.

Next, we consider optimal regulation absent disclosure. The efficacy of regulation depends on the volatility of the aggregate shock. When there is no uncertainty, regulation alone is sufficient to attain trade for all states. In particular, the regulator mandates an effort level such that the average asset quality is sufficiently high, mitigating the adverse selection problem and allowing for efficient trade. As the level of uncertainty increases, however, there is the risk that in particularly bad states the average asset quality would not be sufficient to sustain efficient trade. In order to insure against such realizations, the optimal policy calls for a higher level of precautionary effort to guarantee ex-post efficient trading. Eventually, when uncertainty is too high, insuring against all possible realizations becomes too costly, and the regulator reduces the mandated effort again. As a result, the mandated effort is non-monotonic on the level of uncertainty. Yet, it is always higher than in the effort level that attains full trade under no uncertainty. See Figure 4 and Proposition 2. We also show that the mandated effort level is increasing in the externality. Another way to interpret this finding is that more systemically important institutions should face tighter regulation,

which is consistent with the recent regulatory frameworks such as Basel III and the Dodd-Frank Act. It is also consistent with empirical findings in Schneider et al. (2023), who show that large banks face more scrutiny on risk management practices and governance in the qualitative parts of recent stress tests. In our model, this can be interpreted precisely as demanding higher effort.

Having established these important benchmarks, we study the optimal design problem with both instruments. We show in Lemma 2 that these instruments are substitutes. Namely, the more demanding the effort regulation is, the less disclosure there will be ex-post. This implies that, when there is less regulatory capacity, which limits the ability to enforce higher effort levels, the regulator tends to rely instead more on disclosure. On a related point, when the uncertainty in the environment increases, disclosure will allow for some trade to take place even in bad states. This reduces the level of precautionary effort ex-ante (see Corollary 2), thereby significantly increasing efficiency particularly in very uncertain environments. This rationalizes the use of both instruments together. Thus, when designing optimal regulation, it is important that these pillars of policy be jointly determined. Furthermore, it is worth highlighting that the optimal disclosure rule is more transparent in bad states. While the planner values opacity, in these bad states the planner increases transparency to maintain a conditional expected asset quality that is high enough to induce trade. Theorem 1 provides the characterization of an optimal joint design.

In contrast to the case with only disclosure, Lemma 3 further shows that the optimal policy will *never* entail full disclosure (i.e. $\beta(\cdot) > 0$). Finally, we find that more systemically important institutions should face higher effort regulation and, in turn, should be optimally subjected to less market discipline via reduced disclosure (Corollary 3). An interpretation of the latter result is that the regulator has less incentive to expose systemically important institutions to market discipline.

Lastly, we solve for the optimal *ex-ante* level of regulation from the bank's perspective, under the counterfactual assumption that banks could commit to a certain level of effort ex-ante. Banks do welcome the possibility to subject themselves to some minimal standards, as this helps mitigate the ex-post adverse selection problem. Yet, when the environment is volatile, their willingness to exert precautionary effort is much lower than that which would be demanded by the regulator (see Proposition 3). This means that there is more disagreement between regulators and banks when there is more volatility, and thus also more potential concern regarding regulatory capture.

The rest of the paper is structured as follows. The rest of the Introduction discusses the related literature. Section 2 lays out the model. Section 3 individually studies two policies: information disclosure without regulation (Section 3.1) and regulation without disclosure

(Section 3.2). With these in mind, Section 4 studies the optimal joint design of regulation and disclosure. Section 5 discusses banks’ optimal regulation. Section 6 provides concluding remarks. The proofs are relegated to Appendix A.

Related Literature

This paper is related to three strands of literature that study optimal bank regulation and information disclosure: (i) disclosure of stress test results; (ii) interactions of prudential policy instruments; and (iii) information design and moral hazard.

First, there is a fast-growing literature on disclosure of stress test results as a Bayesian persuasion or an information design problem.² More specifically, the literature has shown that a regulator may optimally obfuscate information for various motives such as ensuring risk-sharing arrangements among banks (Goldstein and Leitner, 2018) and reducing the likelihood of bank runs (e.g., Bouvard et al., 2015; Faria-e Castro et al., 2017; Williams, 2017; Moreno and Tuomas, 2023). In our model, the value of pooling bad and good banks stems from ensuring trade.

Also, in papers such as Bouvard et al. (2015), Williams (2017), and Goldstein and Leitner (2018), a regulator discloses more information during bad times.³ In our paper, absent regulation, increasing opacity has two effects. On the one hand, by pooling bad banks with good ones, trading opportunities increase. On the other hand, the banks are disincentivized to exert effort to improve the quality of their assets. When information disclosure is combined with regulation, the regulator’s main concern is the first effect so that it discloses more during bad times. Also, the higher the externality, the less informative disclosure becomes.

Bouvard et al. (2015) and Parlasca (2023) study a situation in which disclosure by a privately informed regulator may signal the weakness of the financial system. Our paper differs from theirs in two respects. First, we consider the situation in which banks’ private effort determine the soundness of the financial system. This allows us to study the trade-off of ex-post insurance versus ex-ante moral hazard. Second, we study an optimal joint design of regulation and disclosure. With respect to Bouvard et al. (2015), we show that ex-ante regulation can alleviate the tension between the ex-post desire for obfuscation and the ex-ante incentives from disclosure.⁴ With respect to Parlasca (2023), we show that the

²See, for instance, Goldstein and Sapra (2013), Hirtle and Lehnert (2015), and Goldstein and Leitner (2022) for a survey on stress testing.

³See also Parlato (2015), Dang et al. (2017), and Monnet and Quintin (2017) for bank opaqueness.

⁴On a related point, Morrison and White (2013), Shapiro and Skeie (2015), and Shapiro and Zeng (2024) focus on regulator’s reputational concerns for forbearing a failing bank. In Rhee and Dogra (2024), while the regulator commits to information disclosure, banks choose their risk profile in a way leading to “model monoculture.”

informativeness of stress testing depends on the other instrument available to the regulator.

Williams (2017) studies the interaction between stress tests and bank’s ex-ante portfolio choice. He shows that, since stress tests reduce the likelihood of runs, stress tests and liquidity buffers may be seen as policy substitutes. Alvarez and Barlevy (2021) study mandatory disclosure in a banking network and show that mandatory disclosure can improve welfare when banks are vulnerable to contagion. They argue that bank regulation and mandatory disclosure may be a substitute as regulation can mitigate the potential for contagion and make disclosure beneficial. Our contribution is to provide an optimal joint design of regulation and disclosure.

Finally, while this paper focuses on banks’ incentive to exert private effort, papers such as Leitner and Yilmaz (2019), Dai et al. (2024), and Quigley and Walther (2024) suppose that banks have private information in some ways.

Second, our paper is among one of a few papers that study a joint design of prudential policy instruments. Bhattacharya et al. (2002) and Décamps et al. (2004) study the interaction of policy tools in a dynamic context different from ours. More specifically, Bhattacharya et al. (2002) study optimal capital regulation with Poisson-distributed audit. Décamps et al. (2004) study Basel II Accord, which consists of minimal capital requirements, supervision, and market discipline. They show that market discipline can lower the intensity of capital regulation. However, the trade-off between ex-post desire for information obfuscation and ex-ante moral hazard is absent in their papers.

As surveyed by Hirtle and Kovner (2022), little in the current literature studies bank supervision relative to regulation, let alone its joint design. Eisenbach et al. (2016, 2022) and Agarwal and Goel (2024) consider the possibility that bank supervision may be noisy. Eisenbach et al. (2016, 2022) focus on the resource allocation problem between supervising larger and riskier banks and regulation. Agarwal and Goel (2024) study capital requirement when potentially noisy stress testing may lead to mis-classification of banks.

Biswas and Koufopoulos (2022) study a model with a different focus on ex-post moral hazard (diversion) and ex-ante adverse selection (a bank’s type is private information). They show that, while information disclosure alone can be costly because it pools bad assets, it can improve welfare if it is optimally combined with capital regulation. The mechanisms through which optimal information disclosure and regulation interact are different. In our joint design problem of regulation and disclosure, partial information disclosure provides ex-post insurance and thus it lowers the required level of regulation.

Orlov et al. (2023) study the design of sequential stress tests that can impose contingent capital requirements.⁵ The optimal sequential test consists of a precautionary recapital-

⁵In terms of the stress testing literature, Orlov et al. (2023) also allow for correlation in banks’ portfolio.

ization followed by a more informative test that fails only weak banks. In their model, precautionary capital requirements eliminate the need to fail strong banks and thus enable a more informative stress test. In our model, in contrast, higher ex-ante regulation enables the planner to obfuscate information more during bad times.

Faria-e Castro et al. (2017) consider the interaction between stress testing and a government’s fiscal capacity, i.e., its ability to inject money into banks. They show that a government with a more fiscal capacity conducts a more informative stress test as such a government is less worried about bank runs. Their study sheds light on the difference between the stress tests implemented in the United States and in Europe following the Global Financial Crisis.⁶

Third, more broadly, this paper is related to the literature that studies the trade-off between information and incentive provision in Bayesian persuasion and information design, as this paper studies a bank’s ex-ante effort consideration and the ex-post trades after information disclosure. Firstly, optimal information design has been applied to such numerous settings as schools’ grading standards (e.g., Dubey and Geanakoplos, 2010; Boleslavsky and Cotton, 2015; Zubrickas, 2015) and quality certification (e.g., Albano and Lizzeri, 2001; Zapechelnuk, 2022).⁷ Secondly, papers such as Rodina and Farragut (2018) and Boleslavsky and Kim (2021) develop a methodology for a Bayesian persuasion problem with moral hazard: in addition to the sender and the receiver, there is an agent who exerts a private effort that affects the distribution of an underlying state. Our paper is related to theirs in that the sender (the regulator) also has to take into account the agent’s (bank’s) incentive to exert private effort. With respect to both strands of literature, the key difference from these papers is that our paper focuses on the optimal joint design of an ex-ante regulation and ex-post information disclosure.

See also Gick and Pausch (2012), Huang (2021), Inostroza (2023), Inostroza and Pavan (2023), Leitner and Williams (2023), and Parlato and Philippon (2023) for various forms of correlation.

⁶Spargoli (2012) studies the trade-off for a regulator that disclosing negative stress test results, while it leads banks to replenish their capital and hence to reduce the risk of default, leads banks to downsize or requires capital injections. Huang et al. (2024) study the policy interaction between deposit insurance and liquidity injections.

⁷In the context of grading standards, Boleslavsky and Cotton (2015) study the competition among schools which place graduates by investing in education quality and grading standards. In their extension, they consider a model in which a student makes an unobservable effort. In the context of certification, Zapechelnuk (2022) considers a monopoly producer which decides the quality and price of an indivisible good. While the price is observable to a representative consumer, the quality is not. A regulator sets a certification policy, through which information about the quality is partially revealed to the consumer.

2 Model

We consider the problem of a prudential regulator who is concerned that the banks which hold too many risky assets may experience financial distress at some future date. To reduce this risk, the regulator would like banks to sell their risky assets to other financial institutions who are less likely to contribute to systemic crises. However, the market for risky assets sometimes breaks down, creating a role for a regulator to foster trade by either ex-ante rules and supervision or ex-post disclosure.

Specifically, we consider a two-period model in which a bank originates an asset of uncertain quality and may later sell it to a competitive fringe of buyers (the market).⁸ Asset quality $q \in \{H, L\}$ (i.e., High or Low) is determined by two components: unobservable bank effort $e \in [0, \frac{1}{2})$ that is exerted in the first period, and an exogenous shock θ to the state of the economy that is realized once effort is sunk. We think of state θ as capturing economic conditions that influence the future cash flows generated by bank assets. For example, a low realization of θ might reflect a downturn in which bank loans are less likely to be repaid.⁹

The probability of producing a high quality asset is

$$\text{Prob}(q = H \mid e) = \theta e.$$

Because θ is random, asset quality is uncertain even conditional on effort. Hence, the ability to influence e via regulation may not be sufficient to ensure that all assets are traded.

We assume that θ is uniformly drawn from $[1 - \varepsilon, 1 + \varepsilon]$, where $\varepsilon \in (0, 1)$ captures the uncertainty of the environment. The uniform distribution of θ does not play a crucial role and is for ease of analysis. The assumption that effort is bounded by $\frac{1}{2}$ ensures that the probability θe is well-defined (i.e., between 0 and 1).¹⁰

Effort is costly and the cost function $c : [0, \frac{1}{2}) \rightarrow [0, \infty)$ is increasing, convex, twice continuously differentiable, and satisfies $c(0) = 0$, $c'(0) = 0$, $\lim_{e \rightarrow \frac{1}{2}} c(e) = \infty$, and $\lim_{e \rightarrow \frac{1}{2}} c'(e) = \infty$. For instance, the following two cost functions satisfy our assumptions:

$$c(e) = -k(2e + \log(1 - 2e)) \text{ or } c(e) = \frac{ke^2}{1 - 2e}, \text{ where } k > 0.$$

⁸One could interpret the single bank as representing the whole banking system. Our results readily extend to a setting with a continuum of banks. Yet, to avoid the subscript for individual banks, we consider the single bank.

⁹This, however, does not imply that a bank with low quality assets is currently in financial distress.

¹⁰One could also consider an alternative additive specification where the probability of obtaining a high-quality asset is $e + \theta$ instead of θe . Given some technical modifications so that the conditional probability $e + \theta$ remains well-defined, the only qualitative result affected by this change is the non-monotonicity of optimal disclosure without regulation that is illustrated in the right panel of Figure 3. See footnote 16.

At the beginning of the second period, assets can be traded. Assets of quality $q \in \{H, L\}$ have a value v_q for buyers and ρ_q for the seller. There are private gains from trade for high quality assets, but not for low quality assets. Buyer and seller values are ranked as:

$$v_H > \rho_H > \rho_L > v_L.$$

This makes it difficult to trade bad assets. We introduce a role for policy by assuming that a social planner would want all assets to ultimately be held by the market rather than the bank. In particular, there is an additional social value g to trading each asset that is not captured by the traders, and it is large enough that the planner would like to maximize trade:

$$v_L + g > \rho_L.$$

This payoff structure captures the idea that the regulator is worried about the risk of bank failure and cannot commit to not bailing out banks in the event of distress. However, if the risky assets (or non-performing loans) are held outside of the banking system, there is less risk of financial crises and therefore a lower likelihood of bailouts. Since non-bank investors are not expected to be bailed out, they value bad assets less than banks. The value of g then captures the social cost of bailouts not internalized by banks.¹¹ Section 2.1 discusses in more detail how g might depend on the bank's systemic importance, its capitalization, and the quality of its assets.

There is asymmetric information because the seller is privately informed of the realization of q . As we discuss below, the degree of asymmetric information can be modulated by regulatory disclosure. Buyers form expectations about the quality of the asset and offer a price. Given the highest offer for the asset, which we denote by p , the seller decides to accept or reject. The expected payoff for the seller is:

$$\begin{cases} p - c(e) & \text{if the asset trades} \\ e\rho_H + (1 - e)\rho_L - c(e) & \text{if the asset does not trade} \end{cases}.$$

Policy instruments. In practice, bank regulators have access to a wide array of potential policy instruments. To isolate the key elements of an optimal regulatory framework, we group these instruments into two categories.

The first is a set of *ex-ante* rules and regulations governing bank operational standards including risk management and lending policies. By this we mean that the planner can set

¹¹One could also attain similar trade-offs even without an externality in a model with a continuum of types and imperfect disclosure. There are gains from trade for all assets but the adverse selection problem is particularly severe with low quality assets, and thus they would not trade absent an intervention.

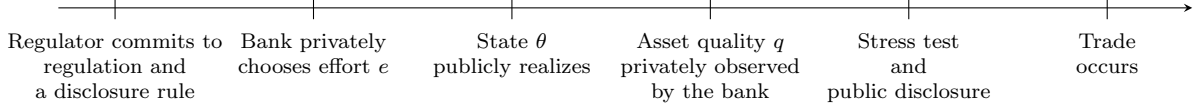


Figure 1: Timeline for Two Policy Instruments

up, before the state of the economy θ is realized, a system that bounds from below the level of effort that the bank must exert (i.e., the minimum effort level). We refer to this minimum effort level as a *regulation* and denote by e_M .¹² We impose a possible limitation on how much effort the regulator can induce: $e_M \leq \bar{e}_M$, in addition to $e_M \in [0, \frac{1}{2})$. The parameter $\bar{e}_M$ allows us to capture different levels of regulatory capacity. While the planner can directly target the minimum effort level, the effort level cannot respond to economic shocks. In order to guarantee trade, it may therefore be necessary to induce excessively high effort.

The second is an *ex-post* instrument, namely, information disclosure. In particular, the planner can commit to a *disclosure* rule which publicly reveals some information about the bank's asset quality conditional on the aggregate state of the economy. In practice, such disclosure might occur in the context of stress testing. Formally, the disclosure rule is a pair (π_H, π_L) where $\pi_H, \pi_L : [1 - \varepsilon, 1 + \varepsilon] \rightarrow \Delta(\{h, \ell\})$. Thus, when the quality of the asset is H (resp. L), given a state $\theta \in [1 - \varepsilon, 1 + \varepsilon]$ of the economy, the planner stochastically announces either h (a good report) or ℓ (a bad report) according to the policy π_H (resp. π_L). To promote trade, the planner has an incentive to obfuscate information, i.e., she has an incentive to announce h even though the bank's asset quality is low. However, expecting information obfuscation, the bank is less likely to exert effort. Thus, partial information disclosure creates the ex-ante moral-hazard problem. Figure 1 illustrates the timeline for these two policy instruments.

While stylized, this timeline captures central elements of current regulatory frameworks. For example, the current CCAR (Comprehensive Capital Analysis and Review) specifies a number of regulations and allows for bank-specific disclosure of stress test results.

2.1 Model Discussion

To isolate the key forces underlying optimal prudential policy, our model assumes a stark dichotomy between ex-ante regulation and ex-post disclosure. In this regard, capital requirements are an example of an ex-ante regulation because higher equity capital would correspond to higher effort. Capital requirements may also have a broader impact, for exam-

¹²Alternatively, one could interpret requiring e_M as requiring the bank to spend some lower bound c_M on the resources such that $c_M = c(e_M)$.

ple, by directly reducing the likelihood of financial distress and lowering the cost of bailouts. In particular, countercyclical capital buffers are intended for this purpose. While we do not explicitly model these effects, they can be captured through comparative statics on the externality parameter g . Since a well-capitalized bank is less likely to face financial distress, and even in that eventuality be easier to rescue, higher capital requirements would thus correspond to a lower g . Also, the size of the externality could capture the systemic importance of a financial institution, with a more systemically important bank having a higher g .

In our analysis, we conduct comparative statics with respect to g to understand how optimal policy relates to the bank size and capitalization. However, our model is not intended for a full analysis of optimal capital requirements since we do not explicitly model the potential costs of capital requirements.

3 Optimal Design of Individual Policies

Before we tackle the joint design problem, we first study each tool on its own to understand their own strengths and weaknesses. Of particular interest in this section are our results on the trade-off between ex-ante incentives and ex-post insurance induced by disclosure in an environment with risky asset quality.

3.1 Information Disclosure without Regulation

We study an optimal information disclosure policy without regulation. We suppose that the planner can commit to an information disclosure policy ex-ante to partially disclose the quality of the assets. The planner is trading off ex-post trades and ex-ante effort costs. We start with the no-information and full-information benchmarks. In the no-information case, the bank would have too little incentive to exert effort. In the full-information case, effort would be high but trade would be suboptimal from the planner's perspective because low types would not trade. We end by studying the optimal information disclosure policy and show that it may call for some obfuscation of information.

3.1.1 No Information Benchmark

Suppose that buyers are uninformed about asset quality while the bank knows the realized quality of the asset, and that there is no regulation or disclosure. Under these assumptions, there is no equilibrium in which trade always occurs for all θ . This is because, in such an equilibrium, prices must be independent of asset quality and the bank would have no incentive to exert any effort. Hence, buyers would be willing to offer at most v_L , which is

less than bank's reservation value ρ_H . Conversely, depending on the degree of uncertainty ε and the relative valuations of buyers and sellers, there may exist an equilibrium where trade never occurs.

Going forward, we focus on the more interesting case in which trade occurs for some realizations of θ but not for others. In Appendix A.1, we show that the necessary and sufficient condition for this to be the case is:

$$(c')^{-1}(\rho_H - \rho_L) > \frac{\rho_H - v_L}{v_H - v_L} \frac{1}{1 + \varepsilon}. \quad (1)$$

This condition requires that the volatility of θ , which is determined by ε , is large relative to the difference in buyer and seller valuations.

The bank's optimal decision is then determined by how effort e and state θ jointly determine the probability of trade. Fixing some effort level e , trade occurs if and only if the state θ is sufficiently favorable. Let $\theta^*(e)$ be the state at which the quality of the asset, $\theta^*(e)v_H + (1 - \theta^*(e))v_L$, is equal to ρ_H . We can pin down this threshold value as

$$\theta^*(e) := \frac{e^*}{e} \text{ with } e^* := \frac{\rho_H - v_L}{v_H - v_L}, \quad (2)$$

where e^* is an effort level for which the *average* asset quality is equal to ρ_H . Note that for all $\theta < \theta^*(e)$ there is no trade. The probability of $\theta \in [1 - \varepsilon, \theta^*(e))$ is important since in these cases the bank will retain the asset and thus has incentives to exert effort.

The equilibrium effort level e^{NI} with no information is an effort level that maximizes the bank's expected payoff given that trade occurs if and only if $\theta \geq \theta^*(e^{\text{NI}})$:

$$\max_{e \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e^{\text{NI}})} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta + \frac{1}{2\varepsilon} \int_{\theta^*(e^{\text{NI}})}^{1+\varepsilon} (\theta e^{\text{NI}} v_H + (1 - \theta e^{\text{NI}}) v_L) d\theta - c(e). \quad (3)$$

The first term captures the bank's payoff under no trade. The second term captures the bank's payoff when trade occurs at price $\theta e^{\text{NI}} v_H + (1 - \theta e^{\text{NI}}) v_L$ when $\theta \geq \theta^*(e^{\text{NI}})$. The no-information equilibrium effort level e^{NI} solves Problem (3), which is characterized as follows:

Lemma 1. *A unique solution $e^{\text{NI}} \in (\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon})$ exists and is characterized by:*

$$\frac{\rho_H - \rho_L}{4\varepsilon} \left(\left(\frac{e^*}{e^{\text{NI}}} \right)^2 - (1 - \varepsilon)^2 \right) = c'(e^{\text{NI}}). \quad (4)$$

The left-hand side of Expression (4) is the marginal benefit of exerting effort. It is given by the expected conditional marginal increase in asset quality times the probability of the

asset not being traded. Note that given buyer's beliefs, neither prices nor the probability of trade are affected by effort. The right-hand side is the marginal cost of effort.

3.1.2 Full Information Benchmark

Next, we consider the bank's problem in the full-information benchmark. When the quality of the asset is high, trade occurs at the price v_H . When the quality of the asset is low, no trade occurs. Thus, the bank would choose the full-information effort level e^{FI} that maximizes its payoff:

$$\mathbb{E}[\theta e v_H + (1 - \theta e) \rho_L] - c(e) = e v_H + (1 - e) \rho_L - c(e).$$

Given our assumptions on the cost function, the first-order condition uniquely pins down the full-information effort level:

$$e^{\text{FI}} = (c')^{-1}(v_H - \rho_L).$$

Note that Lemma 1 implies $e^{\text{NI}} < (c')^{-1}(\rho_H - \rho_L) < e^{\text{FI}}$. Although full information will not necessarily be optimal, it provides incentives to exert higher effort. In the full information case, the planner ex-post may want to obfuscate information. This is because the planner would want to cross-subsidize the low types to generate more trades. Of course, if this were anticipated, then it would diminish the incentives to exert effort. The optimal policy must strike a balance between these two considerations.

3.1.3 Optimal Disclosure without Regulation

We consider the optimal disclosure policy without regulation. To that end, Appendix A.2 shows that, without loss, we can restrict attention to the disclosure policies of the following form: for any θ , (i) if the type (i.e., the asset quality) is H , it reports h with probability 1; and (ii) if the type is L , it reports h with probability $\beta(\theta)$ and ℓ with probability $1 - \beta(\theta)$. Intuitively, since the price under full information satisfies $p^{\text{FI}} = v_H > \rho_H$, the planner can induce some bad assets to trade by reporting them as high quality even though they are of low quality. Ex-post this increases trade, but lowers the bank's ex-ante incentive to exert effort since the price decreases.

We denote by $p(\theta \mid e, \beta)$ the price in state θ after the planner has announced the good signal. By Bayes rule:

$$p(\theta \mid e, \beta) := \frac{\theta e v_H + (1 - \theta e) \beta(\theta) v_L}{\theta e + (1 - \theta e) \beta(\theta)}.$$

Note that $p(\theta \mid e, \beta)$ must be at least as high as ρ_H . Otherwise, there would be no trade.

With this in mind, the planner's problem is:

$$\begin{aligned} & \max_{e, \beta} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e(v_H + g) + (1 - \theta e)(\beta(\theta)(v_L + g) + (1 - \beta(\theta))\rho_L)) d\theta - c(e) \\ & \text{subject to } p(\theta \mid e, \beta) \geq \rho_H \text{ for each } \theta \in [1 - \varepsilon, 1 + \varepsilon] \text{ and} \quad (5) \\ & e \in \operatorname{argmax}_{\hat{e} \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} ((\theta \hat{e} + (1 - \theta \hat{e})\beta(\theta))p(\theta \mid e, \beta) + (1 - \theta \hat{e})(1 - \beta(\theta))\rho_L) d\theta - c(\hat{e}). \quad (6) \end{aligned}$$

Condition (5), which requires the price after the good report to be at least as high as ρ_H , ensures that trade takes place after the good report. Condition (6) is the bank's IC (or Obedience) constraint, which ensures that the bank has an incentive to exert the effort level of e that the planner has specified. Denote by (e^D, β^D) the optimal Disclosure policy.

The optimal policy may take the following two forms. First, full information is optimal: $(e^D, \beta^D) = (e^{\text{FI}}, 0)$. We show below that for any cost function, full information is optimal when the externality g is sufficiently small.

Second, it is optimal to partially disclose information. Since information obfuscation creates a moral hazard problem and the bank does not have an incentive to exert a higher effort level than e^{FI} , in this case the optimal disclosure policy satisfies $e^D < e^{\text{FI}}$. We show below that there exists a cost function such that information obfuscation (i.e., $\beta^D(\theta) > 0$ for a set of θ with positive measure) is optimal when g is sufficiently large.

Proposition 1. *1. For any cost function c , there exists $\underline{g} > \rho_L - v_L$ such that if $g \in (\rho_L - v_L, \underline{g})$, then full information is optimal: $(e^D, \beta^D) = (e^{\text{FI}}, 0)$.*

2. There exist a cost function c and $\bar{g} > \rho_L - v_L$ such that if $g > \bar{g}$, then information obfuscation is optimal: $e^D < e^{\text{FI}}$ and $\beta^D(\theta) > 0$ for some set of θ with positive measure.

Figure 2 illustrates Proposition 1: the solid curve depicts the optimal effort level under the information disclosure problem, relative to the full-information effort level e^{FI} .

As long as $g > \rho_L - v_L$ is sufficiently low, as shown in Proposition 1 and depicted in Figure 2, full disclosure is optimal: $(e^D, \beta^D) = (e^{\text{FI}}, 0)$. Note first that, in the limit $g = \rho_L - v_L$, full disclosure is optimal, since there is no welfare gain from trading low quality assets. When g is close to this limit, the planner could in principle reduce disclosure, which would create a second-order gain. Yet, if it were to do that, it would need to contend with the induced moral hazard problem, which would imply a first-order cost in terms of average asset quality. Thus, there is an interval of values of g under which full disclosure is optimal. However, when g is large, the planner may want to cross-subsidize some of low realizations to capture g .

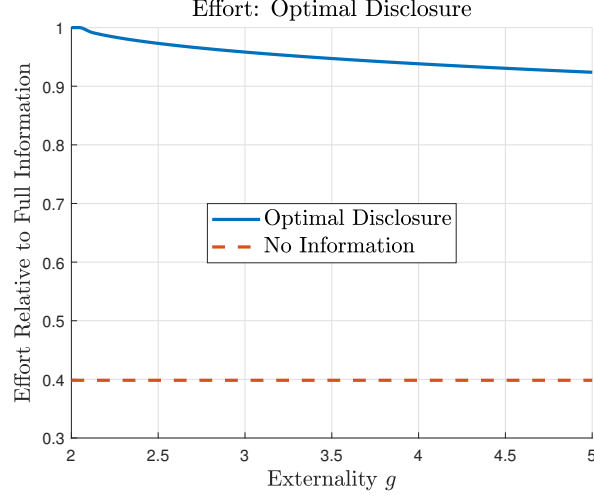


Figure 2: Full versus Partial Information Disclosure (Proposition 1)

The trade-off is that, for incentive compatibility, the proposed effort level must be lowered. Thus, there exists $\bar{g}$ (under certain cost functions) such that if $g > \bar{g}$ then $e^D < e^{FI}$ and $\beta^D(\theta) > 0$ for some set of θ with positive measure.

As g increases, the planner faces a more intense trade-off between an ex-ante effort provision and ex-post insurance. Thus, the planner would be willing to decrease the effort level from the full-information benchmark to increase β more (in an incentive-compatible fashion). If ρ_H is high enough so that Condition (5), the one on the price, starts to bind, then the solution (e^D, β^D) no longer depends on g . Otherwise, e^D is decreasing in g . If we think that g is related to how systemically important an institution is, this suggests that the planner would be more opaque about the assets of large institutions, and as a result, these institutions would have lower standards (worse portfolios). As we will show later, it would then mean that if the planner had other tools available to induce effort, it would have larger incentives to induce higher effort from more systemically important institutions by requiring additional processes. For instance, the “final tailoring rules” to tailor Federal Reserve’s Enhanced Prudential Standards and Basel III, which became effective in December 2019, apply prudential regulations to banking organizations, with increasingly stringent requirements for larger and more complex (i.e., systemically important) ones.¹³

To get a better sense of the optimal disclosure policy for a given value of g , we consider the optimal disclosure policy when externality g is high so that information obfuscation is

¹³The opening statement by Federal Reserve Chair Jerome H. Powell emphasizes that “all of our rules keep the toughest requirements on the largest and most complex firms, because they pose the greatest risks to the financial system and our economy” (<https://www.federalreserve.gov/newsevents/pressreleases/powell-opening-statement-20191010.htm>; Date of Access: April 16, 2024).

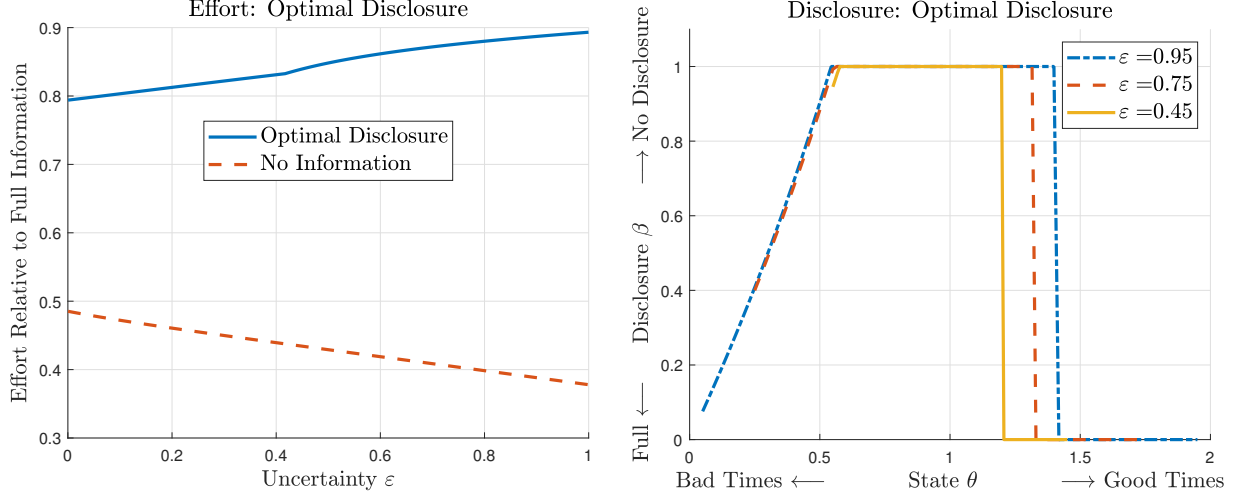


Figure 3: Optimal Information Disclosure Policy (e^D, β^D)

optimal.¹⁴ Figure 3 depicts the optimal disclosure policies for various uncertainty levels ε . The left panel illustrates the effort level given by the optimal information disclosure and the no-information effort levels, normalized by the full-information effort level. The right panel depicts the information obfuscation probability β for various ε .

As illustrated in the figure, the function β is not monotone. For low θ (i.e., in “bad times”), Condition (5) is binding: in order for the price $p(\theta | e, \beta)$ to be at least as high as ρ_H , information has to be disclosed to the extent that the condition is satisfied.¹⁵ In contrast, for high θ (i.e., in “good times”), Condition (5) is not binding. Then, decreasing $\beta(\theta)$ (more information disclosure) relaxes the IC constraint (6) or incentivizes the bank to exert a higher level of effort. In our model in which the probability of a high quality asset is multiplicatively given by θe , lowering $\beta(\theta)$ for high θ relaxes incentives and the loss the planner incurs from foregoing trades for low quality assets is least.¹⁶ Thus, as shown in Proposition 4 in Appendix A.2, there exists $\bar{\theta}$ such that (i) $\beta(\theta) = 0$ when $\theta > \bar{\theta}$ and (ii) $\beta(\theta)$ is the maximum obfuscation probability constrained by Condition (5) when $\theta < \bar{\theta}$. That is, the optimal disclosure policy takes a “bang-bang” disclosure policy.

¹⁴For the purpose of illustration, we take $g = 10$ in Figure 3 so that the optimal effort level e^D is lower than the one depicted in Figure 2.

¹⁵The kink on the curve e^D in the left panel of the figure at around $\varepsilon = 0.42$ corresponds to the fact that Condition (5) starts to bind around that point.

¹⁶If effort and the state entered additively as discussed in footnote 10, then this effect would be absent. We note that this non-monotonicity also disappears under the multiplicative specification once we allow for regulation in Section 4. That is, if effort and the state entered additively, then the right panel of Figure 3 would then look similar to the right panel of Figure 5.

3.2 Regulation without Disclosure

As we saw above, the planner might want to obfuscate information yet worries that doing so creates a moral hazard problem. Thus, a separate policy that can induce effort, namely regulation, may be useful. For simplicity, we directly assume that the planner can choose the minimum effort level $e_M \in [0, \frac{1}{2})$ under the regulatory capacity constraint $e_M \leq \bar{e}_M$. Requiring e_M would correspond to a certain amount of paperwork or conditions to be checked in order to guarantee a minimal loan quality, while the regulatory capacity constraint $\bar{e}_M$ captures a regulatory limitation. As regulation turns out to be binding, the bank would end up taking the effort level of e_M . For ease of notation, therefore, we suppose that the planner directly chooses the effort level $e \in [0, \frac{1}{2})$ under $e \leq \bar{e}_M$. For ease of exposition, we first consider the case in which the regulatory capacity constraint is not binding and then we move on to the case in which the regulatory capacity constraint is binding.

We first consider the optimal regulation when there is no uncertainty: $\varepsilon = 0$. We show that the planner chooses the efficient effort from the productive point of view assuming that all assets are always traded. Since the corresponding social welfare is

$$\mathbb{E} [\theta e(v_H + g) + (1 - \theta e)(v_L + g)] - c(e) = ev_H + (1 - e)v_L + g - c(e),$$

the efficient effort level $e^\diamond$ is given by the first-order condition: $e^\diamond := (c')^{-1}(v_H - v_L)$. This efficient effort level is higher than the full-information effort level $e^{\text{FI}} = (c')^{-1}(v_H - \rho_L)$ because, for the planner, although the value for trading the low quality asset $v_L + g$ is higher than the bank's valuation ρ_L , the externality g does not affect the efficient effort level because assets are always traded. Now, the average quality of the asset satisfies $ev_H + (1 - e)v_L \geq \rho_H$ if and only if $e \geq e^*$, where $e^* = \frac{\rho_H - v_L}{v_H - v_L}$. Thus, when the planner directly targets an effort level, since $(c')^{-1}(v_H - v_L) > e^*$ by Condition (1) and thus trade occurs, the optimal Regulation e^{R} is:

$$e^{\text{R}} = (c')^{-1}(v_H - v_L),$$

provided that the regulatory capacity constraint is not binding (otherwise, $e^{\text{R}} = \bar{e}_M$). In the absence of uncertainty, regulation can attain the efficient level and trade occurs. However, since regulation is determined ex-ante, in a stochastic environment (i.e., $\varepsilon > 0$), it may not be a sufficient tool. Henceforth, we consider the case in which $\varepsilon > 0$.

The planner then must decide how much to insure against bad states (for which trade would not occur) by requiring more cumbersome regulation. To see this, when the planner chooses an effort level e , there exists a unique cutoff $\theta^*(e) = \frac{e^*}{e}$ such that $\theta ev_H + (1 - \theta e)v_L \geq \rho_H$ if and only if $\theta \geq \theta^*(e)$. Using the median $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon)$, trade occurs if and

only if $\theta \in [\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon), 1 + \varepsilon]$.¹⁷ Then, by choosing an effort level e , the planner determines the set of states for which trade occurs.

Formally, the planner's problem is:

$$\begin{aligned} \max_{e \in [0, \frac{1}{2})} & \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta \\ & + \frac{1}{2\varepsilon} \int_{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e). \end{aligned} \quad (7)$$

Below we characterize optimal regulation $e = e^R$ (which we show exists uniquely). Figure 4 depicts $\frac{e^R}{e^{\text{FI}}}$ as a function of uncertainty ε . First note that for even low levels of ε (when trade takes place in all states) the planner still sets regulation above the effort induced by full information. Thus, market discipline by itself is not sufficient to induce the optimal level of effort. Next, when uncertainty increases, the planner increases the level of effort even further to guarantee trade after all realizations of θ . We refer to this (precautionary) regulation above the productively efficient effort as *prudential regulation*. At first, optimal regulation increases in a way so that trades continue to occur for all θ . As ε further increases, full insurance is too costly, and “luck” plays a larger role. Thus, optimal regulation is decreasing in ε . Yet, we show that regulation is still above the efficient effort level (and consequently the full-information effort level) to increase the probability of trade even in bad states.

The next proposition formally characterizes e^R in the relaxed problem without the regulatory capacity constraint. The proof of the proposition implies that the optimal regulation is given by $\bar{e}_M$ when the regulatory capacity constraint is binding at e^R , as the social welfare is increasing on $e \in [0, e^R]$.

Proposition 2. *The optimal regulation e^R is uniquely given by:*

$$e^R = \begin{cases} e^\diamond = (c')^{-1}(v_H - v_L) & \text{if } \varepsilon \leq 1 - \frac{e^*}{e^\diamond} \\ \min(e^\dagger, \frac{e^*}{1-\varepsilon}) & \text{if } \varepsilon > 1 - \frac{e^*}{e^\diamond} \end{cases}, \quad (8)$$

where $e^\dagger \in (e^\diamond, \frac{1}{2})$ is a unique solution satisfying

$$\frac{v_H - v_L}{4\varepsilon}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L}{4\varepsilon}(1 - \varepsilon)^2 + \frac{(v_H - \rho_H + \rho_L - v_L)e^* + 2(v_L + g - \rho_L)}{4\varepsilon} \frac{e^*}{(e^\dagger)^2} = c'(e^\dagger). \quad (9)$$

The proposition implies that $e^R \geq e^\diamond$ for any ε . The strict inequality follows whenever

¹⁷Technically, we would need a slight modification at the end points $\theta \in \{1 - \varepsilon, 1 + \varepsilon\}$ when $\theta^*(e)$ is outside of $[1 - \varepsilon, 1 + \varepsilon]$. Since the planner's value at the end points would not change her objective function, we use this convenient notation.

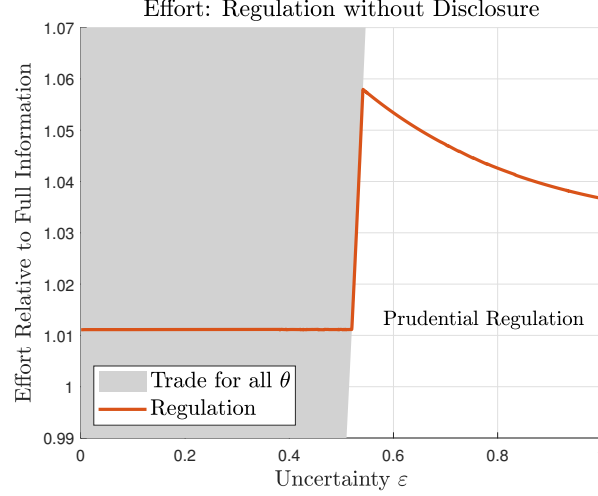


Figure 4: Optimal Regulation (Proposition 2)

$\varepsilon > 1 - \frac{e^*}{e^\dagger}$. Under $e^\dagger$, for some realizations of θ , trades may not take place. The next result establishes this fact.

Corollary 1. *There exists $\bar{\varepsilon} \in (0, 1)$ with the following two properties:*

1. *If $\varepsilon \leq \bar{\varepsilon}$, then $\theta^*(e^R) \leq 1 - \varepsilon$, i.e., trade always occurs.*
2. *If $\varepsilon > \bar{\varepsilon}$, then $\theta^*(e^R) \in (1 - \varepsilon, 1 + \varepsilon)$, i.e., trade takes place for some and only some θ .*

Next, we provide a comparative-static result. It can be shown that the optimal regulation e^R is non-decreasing in g . This implies that the planner would impose higher standards for a systemically important bank (i.e., when the bank has high g).¹⁸ This is indeed consistent with the Basel III framework and the Dodd-Frank Act. As to the Basel III framework, since 2011, the Financial Stability Board, in consultation with Basel Committee on Banking Supervision and national authorities, requires more scrutiny on global systemically important banks (G-SIBs). The provisions of the Dodd-Frank Act state: “Designated [systemically important] FMUs will become subject to the heightened prudential and supervisory provisions of Title VIII, which promote robust risk management and safety and soundness, including conducting their operations in compliance with applicable risk-management standards; providing advance notice and review of changes to their rules, procedures, and operations that could materially affect the nature or level of their risks; and being subject to relevant examination and enforcement provisions.”¹⁹ In the context of the Dodd-Frank Act, we could

¹⁸Corollary 4 in Appendix A.3 formally presents this result.

¹⁹<https://home.treasury.gov/policy-issues/financial-markets-financial-institutions-and-fiscal-service/fsoc/designations> (Date of Access: February 19, 2024).

interpret higher standards for systemically important institutions not only as higher risk-management standards and capital requirements but also the creation of resolution plans known as “living wills” in order to alleviate a moral hazard problem that information obfuscation creates. This is also consistent with empirical findings in Schneider et al. (2023), who show that large banks face more scrutiny on risk management practices and governance in the qualitative parts of recent stress tests. In our model, this can be interpreted precisely as demanding higher effort.

4 Joint Design of Regulation and Disclosure

Regulation is a powerful tool but since it is set ex-ante it cannot be tailored to the realized state. Thus, if the state realization is sufficiently low, then the adverse selection problem is severe and trade completely freezes even though there are some good assets. Thus, the planner can improve by information disclosure. In addition, once the planner is allowed to use information disclosure, she can reduce the amount of regulation required from the bank. We highlight how these two policy tools would be used together and the interactions between them.

We thus consider the case in which the planner can conduct a stress test and disclose its result in addition to regulation (recall Figure 1 for the timeline). While regulation deals with the moral-hazard problem, disclosure allows for adaptation to the states of the economy. As in Section 3.1.3, it is without loss to consider disclosure policies β such that $\beta(\theta)$ is the probability of reporting h when the asset quality is L in state θ .

The analysis in the previous section implies that, when ε is small, since regulation alone can induce trades for all realizations θ , ex-post information disclosure will not be needed, provided that the regulatory capacity constraint $e \leq \bar{e}_M$ is slack. However, once regulation alone cannot maximize trades, the planner would combine ex-post information disclosure with ex-ante regulation. For a given effort level e , as in Expression (2), let $\theta^*(e) = \frac{e^*}{e}$ be the cutoff state under which the average quality of the asset is ρ_H . Since the planner can choose $\beta(\theta)$ for each θ , she chooses each $\beta(\theta)$ as the maximum obfuscation probability with which trade still occurs. With some abuse of notation, the planner would choose $e \in [0, \frac{1}{2})$ with $e \leq \bar{e}_M$ and $\beta : [1 - \varepsilon, 1 + \varepsilon] \rightarrow [0, 1]$ such that

$$\beta(\theta) = \begin{cases} \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta e}{1 - \theta e} & \text{if } \theta \in [1 - \varepsilon, \text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon)) \\ 1 & \text{if } \theta \in [\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon), 1 + \varepsilon] \end{cases} \quad (10)$$

²⁰Technically, we would need a slight modification at the end points when $\theta^*(e)$ is outside of $[1 - \varepsilon, 1 + \varepsilon]$. Since the value of β at the end points would not change the integral, we define β in this way.

When $\theta \leq \theta^*(e)$, $\beta(\theta)$ in Expression (10) makes the price at θ equal to ρ_H . When $\theta \geq \theta^*(e)$, the price at θ is at least as high as ρ_H for any information obfuscation probability. Thus, we let $\beta(\theta) = 1$.

At this point, we provide three implications of Expression (10) as lemmas. First, if we consider a feasible set of (e, β) , then we can show that regulation and information disclosure are policy substitutes.

Lemma 2. *The more demanding the regulation is, the less disclosure occurs. Formally, for any (e, β) and $(\tilde{e}, \tilde{\beta})$ satisfying Expression (10) and $\tilde{e} \geq e$, $\tilde{\beta} \geq \beta$.*

The formal statement in this lemma states that, for a higher level $\tilde{e}$ of regulation, the corresponding disclosure policy $\tilde{\beta}$ given by Expression (10) discloses less information. Thus, a decrease in $\bar{e}_M$ leads to more information disclosure. That is, a regulator with lower regulatory capacity would naturally rely on more information disclosure.

Second, full information is never revealed for any state θ , irrespective of the regulatory capacity constraint $e \leq \bar{e}_M$.

Lemma 3. *Information will never be fully disclosed for any state θ . Formally, $\beta(\theta) > 0$ for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$.*

To see this, suppose that information is fully disclosed at some state. Since information is fully disclosed, the price at that state (average quality after the announcement of h) has to be v_H . This can never be optimal for the planner since it can announce h when the asset quality is L with some small probability in a way such that the price remains at least as high as ρ_H . This increases trade, and thus is an improvement.

Third, Expression (10) implies that information disclosure is countercyclical in the sense that information is disclosed more during bad times. Formally, β is non-decreasing in θ . This follows because the planner discloses information in a way such that the average quality is equal to ρ_H . In other words, although the planner likes to ex-post obfuscate, when the realized state is sufficiently bad, the market beliefs are so low that, absent a strong signal, the market would freeze. In practice, in order to avoid a market freeze, the planner is forced to provide more detailed information even at the cost of some institutions retaining more of their bad assets.²¹

²¹One could apply this result to detailed disclosure of stress tests during the crisis times. For one example, the 2011 Irish and 2011 Europe-wide EBA (European Banking Authority) stress tests are considered to be detailed, including comparisons of bank and third-party estimates of losses in the Irish case and data in electronic and downloadable form in the EBA case. For another example, while currently the Dodd-Frank Act stress test discloses bank-level results, the CCAR in 2011 disclosed only the macro-scenario as opposed to the SCAP (Supervisory Capital Assessment Program) in 2009 during the crisis time. See, for instance, Schuermann (2014) for more details.

Turning to the planner's problem, we start with the relaxed problem that ignores the regulatory capacity constraint $e \leq \bar{e}_M$. Since e determines β through Expression (10), the planner would choose e to solve:

$$\max_{e \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e(v_H + g) + (1 - \theta e)(\beta(\theta)(v_L + g) + (1 - \beta(\theta))\rho_L)) d\theta - c(e).$$

Substituting β given by Expression (10) into the social welfare, we can write the planner's problem as:

$$\begin{aligned} \max_{e \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)} \left(\frac{e\theta}{e^*}(\rho_H + g) + \left(1 - \frac{e\theta}{e^*}\right) \rho_L \right) d\theta \\ + \frac{1}{2\varepsilon} \int_{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e). \end{aligned} \quad (11)$$

The first term is social welfare when $\theta \leq \theta^*(e)$. With probability $e\theta + (1 - e\theta)\beta(\theta) = \frac{e\theta}{e^*}$, the planner announces h , and trade occurs with price ρ_H which also yields the externality g . With the complementary probability, trade does not occur, and in this case the bank's valuation is ρ_L . The second term is social welfare when $\theta \geq \theta^*(e)$, in which trade occurs with no disclosure (i.e., $\beta(\theta) = 1$).

We solve for the optimal policy $(e^{\text{RD}}, \beta^{\text{RD}})$, where “RD” stands for Regulation and Disclosure. The following proposition formally characterizes the optimal policy, provided that the regulatory capacity constraint is not binding:

Theorem 1. 1. *The optimal regulation e^{RD} is uniquely given by:*

$$e^{\text{RD}} = \begin{cases} e^\diamond = (c')^{-1}(v_H - v_L) & \text{if } \varepsilon \leq 1 - \frac{e^*}{e^\diamond}, \\ e^\dagger & \text{if } \varepsilon > 1 - \frac{e^*}{e^\diamond}, \end{cases} \quad (12)$$

where $e^\dagger \in (e^\diamond, \frac{e^*}{1-\varepsilon})$ is a unique solution satisfying

$$\frac{v_H - v_L}{4\varepsilon}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{4\varepsilon e^*}(1 - \varepsilon)^2 + \frac{v_L + g - \rho_L}{4\varepsilon} \frac{e^*}{(e^\dagger)^2} = c'(e^\dagger). \quad (13)$$

2. *The optimal disclosure policy β^{RD} is given through Expression (10).*

- (a) *If $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, then: $\beta^{\text{RD}}(\theta) = 1$ for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$.*
- (b) *If $\varepsilon > 1 - \frac{e^*}{e^\diamond}$, then: (i) $\beta^{\text{RD}}(\theta) \in (0, 1)$ for all $\theta \in [1 - \varepsilon, \theta^*(e)]$; and (ii) $\beta^{\text{RD}}(\theta) = 1$ for all $\theta \in [\theta^*(e), 1 + \varepsilon]$.*

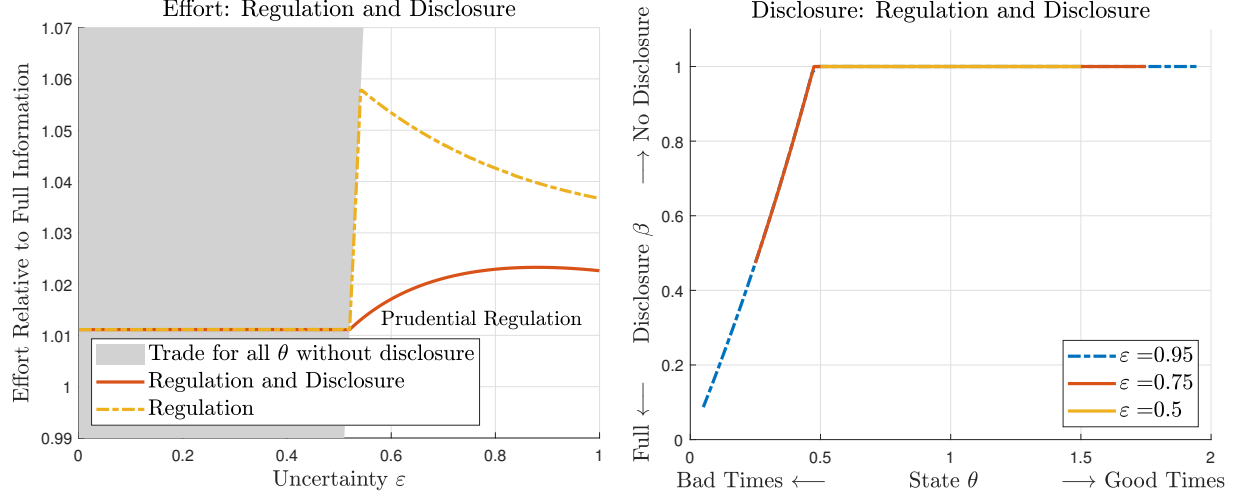


Figure 5: Optimal Regulation and Disclosure (Theorem 1 and Corollary 2)

Figure 5 depicts Theorem 1. The left panel depicts optimal regulation e^{RD} with disclosure. The right panel depicts the optimal information disclosure policy β^{RD} , which is given through Expression (10) from e^{RD} , for various ε .

As shown in the left panel of Figure 5, when uncertainty ε is small, trade always occurs under regulation alone. In this case, information disclosure is not needed (i.e., $e^{\text{RD}} = e^{\text{R}}$), and regulation coincides with the efficient effort level. In the right panel, when $\varepsilon = 0.5$, the information disclosure policy β satisfies $\beta(\cdot) = 1$.

Once ε is sufficiently large (precisely $\varepsilon > 1 - \frac{e^*}{e^{\text{R}}}$), however, trade occurs only for some θ under this effort level. On the one hand, when regulation is the only tool, since losing trade opportunities for some realizations of θ is a first-order loss as compared to a second-order cost of increasing effort, the planner requires prudential regulation so that trade always occurs. As was seen, this corresponds to the increased dashed curve in the left panel of Figure 5. On the other hand, when the planner can utilize both regulation and disclosure, the planner can increase the average asset quality and thus enhance trade either by requiring a higher level of effort or by more detailed information disclosure. Thus, the planner can ensure trade from information disclosure. In other words, information disclosure can reduce the burden of regulation. As we formally show below, the left panel of Figure 5 illustrates that the optimal effort level e^{RD} is lower than e^{R} .

Corollary 2. *For every $\varepsilon \in (0, 1)$, denote by $e^{\text{R}}(\varepsilon)$ and $e^{\text{RD}}(\varepsilon)$ the optimal regulation without disclosure and with disclosure, respectively. Then, $e^{\text{RD}}(\varepsilon) \leq e^{\text{R}}(\varepsilon)$ for all $\varepsilon \in (0, 1)$.*

As shown in the left panel of Figure 5, optimal regulation e^{RD} with disclosure is still at least as high as the efficient effort level and in fact it is strictly higher whenever uncertainty ε

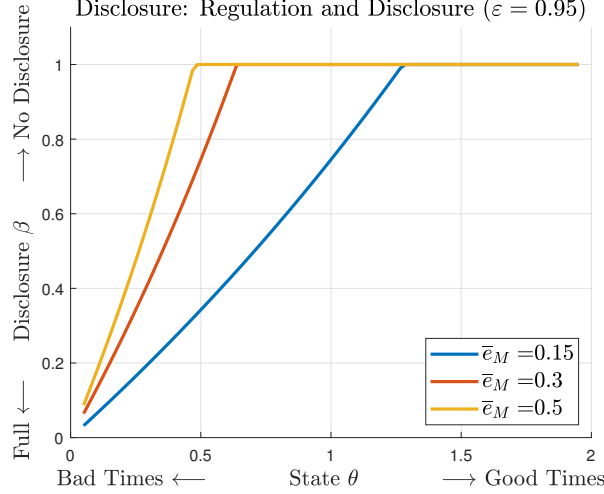


Figure 6: Optimal Disclosure Policy β for Various $\bar{e}_M$

is sufficiently high. Thus, the planner uses both prudential regulation and disclosure together to increase the likelihood of trade.

When the regulatory capacity constraint $e \leq \bar{e}_M$ is binding at e^{RD} characterized in Theorem 1, since the proof implies that the social welfare is increasing on $e \in [0, e^{\text{RD}}]$, the optimal regulation is given by $\bar{e}_M$. The disclosure policy is then determined according to Expression (10). Figure 6 illustrates the optimal disclosure policy β for various $\bar{e}_M$. Since regulation and disclosure are policy substitutes, a lower $\bar{e}_M$ leads to more disclosure.

Also, optimal regulation e^{RD} with disclosure may not be monotone in uncertainty ε . As we discussed before, when uncertainty is too high, luck plays a larger role than effort, and this might lead to a decrease in the level of regulation for sufficiently high uncertainty. The planner can then use higher levels of disclosure (as an ex-post partial substitute) when the realized state is indeed low. Since the values of e^{RD} are similar between $\varepsilon \in \{0.75, 0.95\}$, as shown in the right panel of Figure 5, the optimal disclosure policies β are also similar.

As in Section 3.2, the optimal regulation with disclosure is non-decreasing in externality g , and it implies that the planner would impose higher standards for a systemically important bank (when g is higher).²² Since a disclosure policy is a substitute to regulation, the planner discloses less detailed information. This contrasts with the case in which the regulator is able to design disclosure policies alone. Formally:

Corollary 3. 1. e^{RD} may not be monotone in ε .

2. e^{RD} is non-decreasing in g . Consequently, β^{RD} is non-decreasing in g .

²²Recall the discussion on the “tailoring rules” of the Federal Reserve in Section 3.1.3.

5 Bank-Optimal Regulation

In this section we characterize what regulation and disclosure the bank would choose if it were to “self-regulate.” We refer to it as bank-optimal regulation and disclosure. As in our main analysis, we start with bank-optimal disclosure and regulation individually. Then, we analyze bank-optimal regulation with disclosure. As we will show, although third party regulation is useful from the bank’s perspective, since it provides commitment, it could call for too much effort or too little disclosure.

5.1 Bank-Optimal Disclosure

First, we consider what the optimal disclosure from the bank’s perspective would be. The main difference, vis-à-vis the planner, is that the bank does not care about trading the low quality assets. As a result, the bank would like to commit upfront to full disclosure.

Note that the commitment assumption is not important in this case. This is because, if the bank could not commit ex-ante but could disclose ex-post, then, in equilibrium, the bank would still fully disclose. Of course, after bad realized asset quality, the bank would like buyers to believe the asset is of high quality to command a higher price. The problem is that in equilibrium this cannot be sustained since the bank that has good assets would always like to fully disclose. Thus, any bank that does not fully disclose would be believed to be a bad type. Hence, full disclosure would be the unique equilibrium.

In practice, banks could use credit rating agencies to provide a credible signal. Yet, due to conflicts of interest and rate shopping, the set of signals the bank might credibly deliver might be more restrictive than the ones a planner can commit to. Depending on how credible and precise the set of signals the bank has access to, there could be a limit to the planner’s ability to restrict information. Of course, banks want to be on good terms with regulators so they might still prefer not to disclose even if they could.

5.2 Bank-Optimal Regulation without Disclosure

Next, we consider the bank-optimal regulation, i.e., the effort level that the bank would choose if it can commit, without disclosure. On the one hand, the bank welcomes regulation since it allows to commit to a minimum effort, helping it alleviate the ex-post adverse selection problem. On the other hand, given that the bank does not internalize the externality, it would prefer a lower effort level than the planner.

To characterize the bank-optimal regulation, we make two observations. First, if trade occurs for all θ under the bank-optimal regulation, then the bank’s effort level would be the

minimum effort level under which the average quality of the asset is at least as high as ρ_H for all state θ . Second, suppose that trade never occurs under the bank-optimal regulation. Since the bank's payoff is $e\rho_H + (1 - e)\rho_L - c(e)$, its effort level is $e = (c')^{-1}(\rho_H - \rho_L)$. However, similarly to Section 3.1.1, Condition (1) rules out this no-trade case.

Denoting by e^{BR} the Bank-optimal Regulation, it follows from these two observations that trade occurs if and only if $\theta \geq \theta^*(e^{\text{BR}})$. The bank's problem is:

$$\max_{e \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta + \frac{1}{2\varepsilon} \int_{\text{med}(1-\varepsilon, \theta^*(e), 1+\varepsilon)}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L) d\theta - c(e).$$

The first term corresponds to the payoff from no trade, while the second the payoff from trade. Note that the difference from Expression (7) is that the externality g is absent.

When $\theta \leq \theta^*(e^{\text{BR}})$, no trade takes place as the expected quality of the asset is below ρ_H . Since the second observation implies that some trade occurs, we have to have $\theta^*(e^{\text{BR}}) < 1 + \varepsilon$, i.e., $\frac{e^*}{1+\varepsilon} < e^{\text{BR}}$. In contrast, when $\theta \geq \theta^*(e^{\text{BR}})$, trades take place at the price $\theta e^{\text{BR}} v_H + (1 - \theta e^{\text{BR}}) v_L$. The second observation implies that the bank does not have an incentive to exert e^{BR} with $\theta^*(e^{\text{BR}}) < 1 - \varepsilon$, as trade occurs for all θ once $\theta^*(e^{\text{BR}}) \leq 1 - \varepsilon$. Thus, we have $\theta^*(e^{\text{BR}}) \geq 1 - \varepsilon$, i.e., $\frac{e^*}{1-\varepsilon} \geq e^{\text{BR}}$. Below we formally characterize the bank-optimal effort level e^{BR} .

Proposition 3. *A unique solution e^{BR} exists and satisfies the following:*

$$e^{\text{BR}} = \begin{cases} e^\diamond = (c')^{-1}(v_H - v_L) & \text{if } \varepsilon \leq 1 - \frac{e^*}{e^\diamond} \\ \min(e^\diamond, \frac{e^*}{1-\varepsilon}) & \text{if } \varepsilon > 1 - \frac{e^*}{e^\diamond} \end{cases}, \quad (14)$$

where $e^\diamond \in (\frac{e^*}{1+\varepsilon}, e^{\text{R}})$ is a unique solution satisfying

$$\frac{v_H - v_L}{4\varepsilon} (1+\varepsilon)^2 - \frac{\rho_H - \rho_L}{4\varepsilon} (1-\varepsilon)^2 + \frac{(v_H - \rho_H + \rho_L - v_L)e^* + 2(v_L - \rho_L)}{4\varepsilon} \frac{e^*}{(e^\diamond)^2} = c'(e^\diamond). \quad (15)$$

The left panel of Figure 7 illustrates Proposition 3: the solid curve depicts bank-optimal regulation, compared with the dashed curve which depicts optimal regulation e^{R} (provided $e^{\text{R}} \leq \bar{e}_M$). When the degree of uncertainty is low, the bank could self-regulate since it would pick the same effort level to guarantee high quality as the planner. Instead, when the degree of uncertainty is sufficiently large, supervision by the regulator is necessary since the level chosen by the planner is significantly higher than the one chosen by the bank. This is because, due to the externality, the planner cares more about the probability of trade than the bank. Thus, although both increase the effort once trade is not always guaranteed, the planner would demand more prudential regulation than the bank would like to commit to.

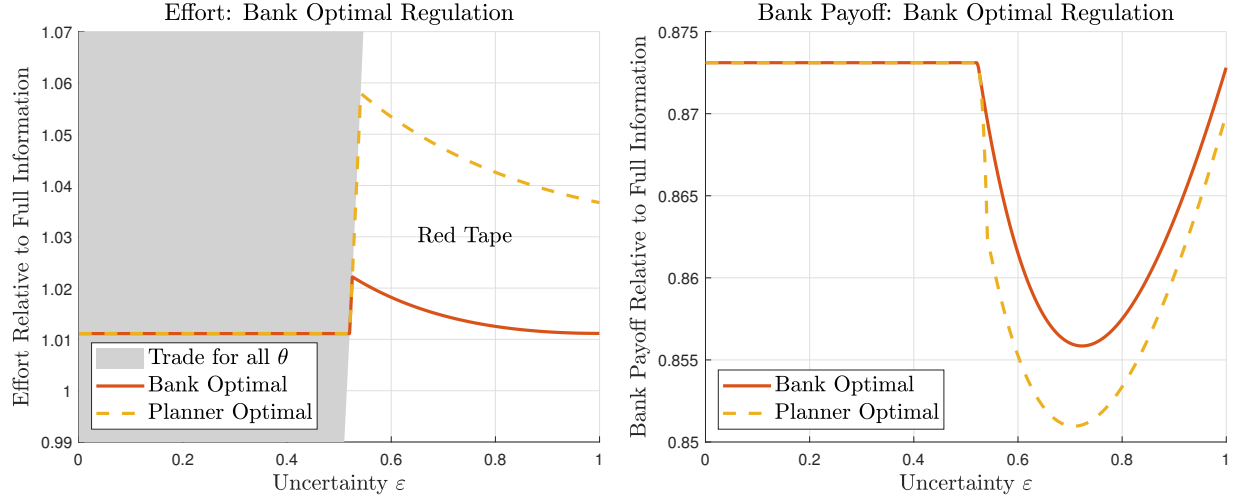


Figure 7: Bank Optimal Regulation (Proposition 3) and Bank's Payoff

In the left panel, we refer to the difference between the optimal regulation and the bank-optimal regulation as “red tape.” As illustrated in the left panel of Figure 7, there exists $\bar{\varepsilon} \in (0, 1)$ such that if $\varepsilon \leq \bar{\varepsilon}$, then trade occurs for all θ ; and if $\varepsilon > \bar{\varepsilon}$, then trade occurs for only some θ .

We move on to the right panel of Figure 7, which depicts the bank's payoff from bank-optimal regulation (solid curve) and from optimal regulation (dashed curve). It is interesting to observe that the bank's payoff is non-monotonic in uncertainty. At first uncertainty is not relevant since trade happens with probability 1. Eventually, as uncertainty increases, there are sufficiently bad states that, absent an increase in effort, there would be no trade after those realizations. At first additional effort is used to reduce said possibility. Of course, this is costly, so payoffs decrease. Payoffs continue to decrease as the probability of trade is reduced. Eventually though, there is an effect in the opposite direction coming from the fact that the price conditional on trade is more sensitive than the value conditional on not trading. This option-like feature leads bank's payoff to increase for high levels of uncertainty.

Finally, we make a technical remark on Proposition 3. While Expression (15) corresponds to Expression (9) with $g = 0$, since the bank's objective function in the bank-optimal regulation problem may not be concave, we show that the unique solution to the bank's optimal regulation problem is indeed characterized by Expression (15).

5.3 Bank-Optimal Regulation and Disclosure

Having analyzed individual policies, we consider bank-optimal regulation with disclosure. As in Section 5.1, the bank would like to commit to full disclosure in equilibrium. This maximizes the gains from trade from the bank's perspective. Note that as long as there is

full information, exerting e^{FI} is optimal for the bank, as shown in Section 3.1.3. Importantly, this implies that it will not be necessary for the bank to be able to commit to hold certain effort standards.

In practice, it is natural to think that the bank would be limited in its ability to commit to the full disclosure policy. Even when relying on third parties such as credit agencies, rate shopping and conflicts of interest can naturally limit the credibility or transparency of the ratings (e.g., Hau et al., 2013; White, 2010). With less than full disclosure, being able to commit to a given effort level would be strictly valuable for the bank.

To illustrate, we focus on the case where the disclosure policy is the planner's optimal disclosure policy β^{RD} characterized in Section 4. First, we can show that there exists a threshold level of uncertainty $\bar{\varepsilon} \in (1 - \frac{e^*}{e^{\text{RD}}}, 1]$ such that for low levels of uncertainty $\varepsilon < \bar{\varepsilon}$ the bank-optimal regulation coincides with the planner's optimal regulation e^{RD} with disclosure. This implies that the bank can self-regulate in environments of low uncertainty. Instead, for high levels of uncertainty $\varepsilon > \bar{\varepsilon}$, the bank would like to commit to lower standards than those imposed by the regulator. In fact, the bank would commit to the bank-optimal regulation e^{BR} without disclosure.²³ In this case, what the planner views as prudential regulation is mostly considered red tape by the bank.

6 Conclusion

We study the optimal joint design of regulation and disclosure. Regulation can ensure prudent behavior but cannot respond to economic shocks. The regulator would choose the optimal regulation higher than the productively efficient effort level, when the degree of uncertainty is high. In contrast, information disclosure can be contingent on economic shocks but leads to the trade-off between providing incentives ex-ante and the ex-post protection of weak banks. When the realized state is sufficiently bad, although the planner likes to obfuscate ex-post, the market beliefs are so low that, without sufficiently detailed disclosure, the market would break down. To prevent a market freeze, the planner would provide more detailed information even at the cost of letting some institutions fail.

Our paper shows the importance of studying the interaction of the financial regulatory tools. While disclosure and regulation are policy substitutes, the joint use of these two policy tools mitigates the overall regulatory burden. Partial information disclosure provides ex-post insurance so that it lowers the required level of regulation. While information obfuscation disincentivizes the banks to behave prudently, regulation maintains ex-ante incentives. Thus,

²³This is because, if the bank's effort level is below the planner's optimal regulation e^{RD} , then trade occurs when the state θ is high enough so that disclosure does not occur under β^{RD} .

our paper provides a justification for the joint use of regulation and supervision on the one hand and market discipline on the other, as in the current Basel III framework. Consistently with the Basel III framework and the Dodd-Frank Act, we also highlight the importance of higher precautionary regulation of systemically important financial institutions. This, in turn, implies that systemically important financial institutions would optimally face less disclosure.

We note that while financial institutions welcome some level of regulation, from their perspective, the regulator imposes too much red tape. This highlights the importance of remaining vigilant since banks have an incentive to capture the regulator and water down regulatory requirements.

Our parsimonious model may also shed light on more broader information disclosure problems such as firm audits in which red tape and inspections interact. Our simple model also admits such extensions as introducing correlated shocks among banks.

References

- AGARWAL, I. AND T. GOEL (2024): “Bank Regulation and Supervision: A Symbiotic Relationship,” *Journal of Banking and Finance*, 163, 107185.
- ALBANO, G. L. AND A. LIZZERI (2001): “Strategic Certification and Provision of Quality,” *International Economic Review*, 42, 267–283.
- ALVAREZ, F. AND G. BARLEVY (2021): “Mandatory Disclosure and Financial Contagion,” *Journal of Economic Theory*, 194, 105237.
- BHATTACHARYA, S., M. PLANK, G. STROBL, AND J. ZECHNER (2002): “Bank Capital Regulation with Random Audits,” *Journal of Economic Dynamics and Control*, 26, 1301–1321.
- BISWAS, S. AND K. KOUFOPOULOS (2022): “Bank Capital Structure and Regulation: Overcoming and Embracing Adverse Selection,” *Journal of Financial Economics*, 143, 973–992.
- BOESLAVSKY, R. AND C. COTTON (2015): “Grading Standards and Education Quality,” *American Economic Journal: Microeconomics*, 7, 248–279.
- BOESLAVSKY, R. AND K. KIM (2021): “Bayesian Persuasion and Moral Hazard,” Working paper.
- BOUVARD, M., P. CHAIGNEAU, AND A. DE MOTTA (2015): “Transparency in the Financial System: Rollover Risk and Crises,” *Journal of Finance*, 70, 1805–1837.

- DAI, L., D. LUO, AND M. YANG (2024): “Disclosure of Bank-Specific Information and the Stability of Financial Systems,” *Review of Financial Studies*, 37, 1315–1367.
- DANG, T. V., G. GORTON, B. HOLMSTRÖM, AND G. ORDOÑEZ (2017): “Banks as Secret Keepers,” *American Economic Review*, 107, 1005–29.
- DÉCAMPS, J.-P., J.-C. ROCHET, AND B. ROGER (2004): “The Three Pillars of Basel II: Optimizing the Mix,” *Journal of Financial Intermediation*, 13, 132–155.
- DUBEY, P. AND J. GEANAKOPOLOS (2010): “Grading Exams: 100, 99, 98, ... or A, B, C?” *Games and Economic Behavior*, 69, 72–94.
- EISENBACH, T. M., D. O. LUCCA, AND R. M. TOWNSEND (2016): “The Economics of Bank Supervision,” Working paper.
- (2022): “Resource Allocation in Bank Supervision: Trade-offs and Outcomes,” *Journal of Finance*.
- FARIA-E CASTRO, M., J. MARTINEZ, AND T. PHILIPPON (2017): “Runs versus Lemons: Information Disclosure and Fiscal Capacity,” *Review of Economic Studies*, 84, 1683–1707.
- GICK, W. AND T. PAUSCH (2012): “Persuasion by Stress Testing—Optimal Disclosure of Supervisory Information in the Banking Sector,” Working paper.
- GOLDSTEIN, I. AND Y. LEITNER (2018): “Stress Tests and Information Disclosure,” *Journal of Economic Theory*, 177, 34–69.
- (2022): “Stress Test Disclosure: Theory, Practice, and New Perspectives,” in *Handbook of Financial Stress Testing*, ed. by J. D. Farmer, A. M. Kleinnijenhuis, T. Schuermann, and T. Wetzler, Cambridge University Press, 208–223.
- GOLDSTEIN, I. AND H. SAPRA (2013): “Should Banks’ Stress Test Results be Disclosed? An Analysis of the Costs and Benefits,” *Foundations and Trends in Finance*, 8, 1–54.
- HAU, H., S. LANGFIELD, AND D. MARQUES-IBANEZ (2013): “Bank Ratings: What Determines their Quality?” *Economic Policy*, 28, 289–333.
- HIRTLE, B. AND A. KOVNER (2022): “Bank Supervision,” *Annual Review of Financial Economics*, 14, 39–56.
- HIRTLE, B. AND A. LEHNERT (2015): “Supervisory Stress Tests,” *Annual Review of Financial Economics*, 7, 39–55.

- HUANG, C., L. SHEN, AND J. ZOU (2024): “Policy Portfolio for Banks: Deposit Insurance and Liquidity Injection,” ”working paper”.
- HUANG, J. (2021): “Optimal Stress Tests in Financial Networks,” Working paper.
- INOSTROZA, N. (2023): “Persuading Multiple Audiences: Strategic Complementarities and (Robust) Regulatory Disclosures,” Working paper.
- INOSTROZA, N. AND A. PAVAN (2023): “Adversarial Coordination and Public Information Design,” Working paper.
- KAMENICA, E. AND M. GENTZKOW (2011): “Bayesian Persuasion,” *American Economic Review*, 101, 2590–2615.
- LEITNER, Y. AND B. WILLIAMS (2023): “Model Secrecy and Stress Tests,” *Journal of Finance*, 78, 689–1223.
- LEITNER, Y. AND B. YILMAZ (2019): “Regulating a Model,” *Journal of Financial Economics*, 131, 251–268.
- MONNET, C. AND E. QUINTIN (2017): “Rational Opacity,” *Review of Financial Studies*, 30, 4317–4348.
- MORENO, D. AND T. TUOMAS (2023): “Precision of Public Information Disclosures, Banks’ Stability and Welfare,” Working paper.
- MORRISON, A. D. AND L. WHITE (2013): “Reputational Contagion and Optimal Regulatory Forbearance,” *Journal of Financial Economics*, 110, 642–658.
- OCC (2019): *Bank Supervision Process: Comptroller’s Handbook*, Office of the Comptroller of the Currency, version 1.1.
- ORLOV, D., P. ZRYUMOV, AND A. SKRZYPACZ (2023): “The Design of Macroprudential Stress Tests,” *Review of Financial Studies*, 36, 4460–4501.
- PARLASCA, M. (2023): “Stress Tests and Strategic Disclosure,” Working paper.
- PARLATORE, C. (2015): “Transparency and Bank Runs,” Working paper.
- PARLATORE, C. AND T. PHILIPPON (2023): “Designing Stress Scenarios,” Working paper.
- QUIGLEY, D. AND A. WALTHER (2024): “Inside and Outside Information,” *Journal of Finance*, 79, 2667–2714.

- RHEE, K. AND K. DOGRA (2024): “Stress Tests and Model Monoculture,” *Journal of Financial Economics*, 152, 103760.
- RODINA, D. AND J. FARRAGUT (2018): “Inducing Effort Through Grades,” Working paper.
- SCHNEIDER, T., P. E. STRAHAN, AND J. YANG (2023): “Bank Stress Testing: Public Interest or Regulatory Capture?” *Review of Finance*, 27, 423–467.
- SCHUERMANN, T. (2014): “Stress Testing Banks,” *International Journal of Forecasting*, 30, 717–728.
- SHAPIRO, J. AND D. SKEIE (2015): “Information Management in Banking Crisis,” *Review of Financial Studies*, 28, 2322–2363.
- SHAPIRO, J. AND J. ZENG (2024): “Stress Testing and Bank Lending,” *Review of Financial Studies*, 37, 1265–1314.
- SPARGOLI, F. (2012): “Bank Recapitalization and the Information Value of a Stress Test in a Crisis,” Working paper.
- WHITE, L. J. (2010): “Markets: The Credit Rating Agencies,” *Journal of Economic Perspectives*, 24, 211–226.
- WILLIAMS, B. (2017): “Stress Tests and Bank Portfolio Choice,” Working paper.
- ZAPECHELNYUK, A. (2022): “Optimal Quality Certification,” *American Economic Review: Insights*, 2, 161–176.
- ZUBRICKAS, R. (2015): “Optimal Grading,” *International Economic Review*, 56, 751–776.

A Proofs

A.1 Section 3.1.1

A.1.1 Derivation of Condition (1)

We first derive Condition (1) under which there is no equilibrium in which trade never occurs. If there is an equilibrium in which trade never occurs, then the bank would maximize the expected private value of the asset net of costs:

$$\mathbb{E}[\theta e \rho_H + (1 - \theta e) \rho_L] - c(e) = e \rho_H + (1 - e) \rho_L - c(e).$$

Given our assumptions on the cost function, the first-order condition uniquely pins down the bank's effort level as

$$e = (c')^{-1}(\rho_H - \rho_L).$$

This condition relates the marginal cost of increasing the probability of obtaining a good asset to the difference in bank valuations between high and low quality assets.

For no trade not to constitute an equilibrium, it must be that at this effort level, the equilibrium price (the average quality of the asset) is as high as ρ_H for some set of θ with positive measure. This condition is exactly Condition (1).

A.1.2 Proof of Lemma 1

Proof of Lemma 1. In the no-information equilibrium, the bank's effort level e^{NI} has to solve Problem (3), which reduces to:

$$\max_{e \in [0, \frac{1}{2})} \frac{\rho_H - \rho_L}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e^{\text{NI}})} \theta d\theta \cdot e - c(e).$$

Given our assumptions on the cost function, the problem has a unique solution, which has to coincide with e^{NI} . The first-order condition at $e = e^{\text{NI}}$ is:

$$c'(e^{\text{NI}}) = \frac{\rho_H - \rho_L}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e^{\text{NI}})} \theta d\theta = \frac{\rho_H - \rho_L}{4\varepsilon} \left(\left(\frac{e^*}{e^{\text{NI}}} \right)^2 - (1 - \varepsilon)^2 \right),$$

which coincides with Expression (4). By the arguments in the main text, $1 - \varepsilon < \theta^*(e^{\text{NI}}) < 1 + \varepsilon$, which implies that $e^{\text{NI}} \in \left(\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon} \right)$. $\square$

A.2 Section 3.1.3

Appendix A.2.1 formulates the planner's optimal disclosure problem without regulation as discussed in the main text. Appendix A.2.2 provides the proof of Proposition 1. Appendix A.2.3 provides the formal characterization of the optimal disclosure policy discussed in the main text.

A.2.1 Planner's Problem

We first show that one can restrict attention to the disclosure policies of the following form.

Lemma 4. *It is sufficient to consider a policy (π_H, π_L) with the following properties: (i) π_H reports h with probability $\alpha(\theta)$ and ℓ with probability $1 - \alpha(\theta)$ for each θ ; and (ii) π_L reports h with probability $\beta(\theta)$ and ℓ with probability $1 - \beta(\theta)$ for each θ .*

Proof of Lemma 4. Let $(S, (\pi_H, \pi_L))$ be a signal, i.e., S is a set of signal realizations and $\pi_H, \pi_L : \Theta \rightarrow \Delta(S)$ is a (measurable) policy. We show that, without loss, we can let $S = \{h, \ell\}$. To see this, let T_θ be a set of signal realizations under which trade takes place for a given θ . Then, we can define $(\tilde{S}, (\tilde{\pi}_H, \tilde{\pi}_L))$ by $\tilde{S} = \{h, \ell\}$, $\tilde{\pi}_H(\theta)(h) = \pi_H(\theta)(T_\theta)$, and $\tilde{\pi}_L(\theta)(h) = \pi_L(\theta)(T_\theta)$. Thus, we can consider $(S, (\pi_H, \pi_L))$ such that $S = \{h, \ell\}$, $\pi_H(\theta)(h) = \alpha(\theta)$, and $\pi_L(\theta)(h) = \beta(\theta)$. $\square$

By Lemma 4, the planner's problem is:

$$\begin{aligned} & \max_{e, \alpha, \beta} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e (\alpha(\theta)(v_H + g) + (1 - \alpha(\theta))\rho_H) + (1 - \theta e) (\beta(\theta)(v_L + g) + (1 - \beta(\theta))\rho_L)) d\theta - c(e) \\ & \text{subject to } p(\theta \mid e, \alpha, \beta) := \frac{\theta e \alpha(\theta) v_H + (1 - \theta e) \beta(\theta) v_L}{\theta e \alpha(\theta) + (1 - \theta e) \beta(\theta)} \geq \rho_H \text{ for each } \theta \in [1 - \varepsilon, 1 + \varepsilon] \text{ and} \\ & e \in \operatorname{argmax}_{\hat{e} \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} ((\theta \hat{e} \alpha(\theta) + (1 - \theta \hat{e}) \beta(\theta)) p(\theta \mid e, \alpha, \beta) + \theta \hat{e} (1 - \alpha(\theta)) \rho_H + (1 - \theta \hat{e}) (1 - \beta(\theta)) \rho_L) d\theta - c(\hat{e}). \end{aligned}$$

The policy (e, α, β) is *incentive-feasible* if it satisfies the above two constraints. We show that one can restrict attention to the disclosure policies with $\alpha = 1$.

Lemma 5. *It is without loss to consider the policies with $\alpha = 1$.*

Proof of Lemma 5. If $\alpha = 1$ almost surely, then it is without loss to redefine $\alpha = 1$. Thus, consider an incentive-feasible policy (e, α, β) with $\alpha(\theta) < 1$ on some set with positive measure. We show in two steps that there exists an incentive-feasible policy $(\tilde{e}, 1, \beta)$ which improves the welfare.

First, since (e, α, β) is incentive-feasible, we derive the first-order condition of the bank's IC constraint. The bank's payoff from taking $\hat{e}$ is:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} ((\theta \hat{e} \alpha(\theta) + (1 - \theta \hat{e}) \beta(\theta)) p(\theta \mid e, \alpha, \beta) + \theta \hat{e} (1 - \alpha(\theta)) \rho_H + (1 - \theta \hat{e}) (1 - \beta(\theta)) \rho_L) d\theta - c(\hat{e}),$$

and the derivative with respect to $\hat{e}$ is:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta \{(\alpha(\theta) p(\theta \mid e, \alpha, \beta) + (1 - \alpha(\theta)) \rho_H) - (\beta(\theta) p(\theta \mid e, \alpha, \beta) + (1 - \beta(\theta)) \rho_L)\} d\theta - c'(\hat{e}).$$

Hence, e satisfies:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta \{(\alpha(\theta) p(\theta \mid e, \alpha, \beta) + (1 - \alpha(\theta)) \rho_H) - (\beta(\theta) p(\theta \mid e, \alpha, \beta) + (1 - \beta(\theta)) \rho_L)\} d\theta = c'(e).$$

Second, since $p(\theta \mid e, \alpha, \beta)$ is increasing in α , if $\alpha = 1$ then we have:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta(\theta))(p(\theta \mid e, \beta) - \rho_L) d\theta > c'(e).$$

Then, there exists $\tilde{e} > e$ such that

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta(\theta))(p(\theta \mid \tilde{e}, \beta) - \rho_L) d\theta = c'(\tilde{e}).$$

This is because the right-hand side is increasing in $\tilde{e}$ and diverges to infinity as $\tilde{e} \uparrow \frac{1}{2}$, while the left-hand side is bounded. Moreover, it can be seen that $(\tilde{e}, 1, \beta)$ satisfies the bank's IC constraint (more formally, Lemma 6 shows that the first-order approach is valid). As $(\tilde{e}, 1, \beta)$ improves the welfare, the proof is complete. $\square$

Thus, the planner's problem reduces to the one in Section 3.1.3. We call a policy (e, β) to be *incentive-feasible* if it satisfies Conditions (5) and (6). The rest of Appendix A.2.1 provides some preliminary observations regarding the planner's problem.

We show that the first-order approach is valid.

Lemma 6. *The bank's IC constraint (6) can be replaced with its first-order condition:*

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta(\theta))(p(\theta \mid e, \beta) - \rho_L) d\theta = c'(e). \quad (16)$$

Proof of Lemma 6. First, we show that the first-order condition of the bank's IC constraint yields Expression (16). Given (e, β) , the bank's payoff is $F(\hat{e} \mid e, \beta) - c(\hat{e})$, where

$$F(\hat{e} \mid e, \beta) := \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} ((\theta\hat{e} + (1 - \theta\hat{e})\beta(\theta))p(\theta \mid e, \beta) + (1 - \theta\hat{e})(1 - \beta(\theta))\rho_L) d\theta.$$

Denoting $F_1(\hat{e} \mid e, \beta) = \frac{\partial F}{\partial \hat{e}}(\hat{e} \mid e, \beta)$, the derivative of the bank's payoff with respect to $\hat{e}$ is:

$$F_1(\hat{e} \mid e, \beta) - c'(\hat{e}) = \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta(\theta)) (p(\theta \mid e, \beta) - \rho_L) d\theta - c'(\hat{e}).$$

The first-order condition with respect to $\hat{e}$ at $\hat{e} = e$ is given by Expression (16).

Conversely, assume Expression (16). Hence,

$$F_1(x \mid x, \beta) = c'(x).$$

First, let $e > \hat{e}$. Then,

$$\begin{aligned} c(e) - c(\hat{e}) &= \int_{\hat{e}}^e c'(x) dx = \int_{\hat{e}}^e F_1(x | x, \beta) dx \\ &\leq \int_{\hat{e}}^e F_1(x | e, \beta) dx = F(e | e, \beta) - F(\hat{e} | e, \beta), \end{aligned}$$

where the inequality follows because $F_1(x | e, \beta)$ is non-decreasing in e . This shows that

$$F(e | e, \beta) - c(x) \geq F(\hat{e} | e, \beta) - c(\hat{e}).$$

Similarly, let $e < \hat{e}$. Then,

$$\begin{aligned} c(\hat{e}) - c(e) &= \int_e^{\hat{e}} c'(x) dx = \int_e^{\hat{e}} F_1(x | x, \beta) dx \\ &\geq \int_e^{\hat{e}} F_1(x | e, \beta) dx = F(\hat{e} | e, \beta) - F(e | e, \beta), \end{aligned}$$

where the inequality follows because $F_1(x | e, \beta)$ is non-decreasing in e . This shows that

$$F(e | e, \beta) - c(x) \geq F(\hat{e} | e, \beta) - c(\hat{e}).$$

□

Next, we prove some preliminary observations.

Lemma 7. 1. *The full-information policy $(e^{\text{FI}}, 0)$ is incentive-feasible.*

2. *If a policy (e, β) is incentive-feasible, then $e \leq e^{\text{FI}}$.*

3. *There is no incentive-feasible policy (e, β) with $\beta(\cdot) = 1$. No disclosure for all state realizations, i.e., $\beta(\cdot) = 1$, is never optimal.*

Proof of Lemma 7. 1. We consider the full-information benchmark. To that end, consider a policy (e, β) with $e > 0$ and $\beta(\cdot) = 0$. Since $p(\theta | e, \beta) = v_H$, the bank's IC constraint is:

$$e \in \operatorname{argmax}_{\hat{e} \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta \hat{e} v_H + (1 - \theta \hat{e}) \rho_L) d\theta - c(\hat{e}).$$

This is the same problem as the full-information benchmark. Hence, $(e^{\text{FI}}, 0)$ is incentive-feasible.

2. We show that any effort level $e > e^{\text{FI}}$ cannot be implemented. To see this, we consider the bank's first-order condition. The bank's payoff is given as:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} ((\theta\hat{e} + (1-\theta\hat{e})\beta(\theta))p(\theta | e, \beta) + (1-\theta\hat{e})(1-\beta(\theta))\rho_L) d\theta - c(\hat{e}).$$

The derivative of the bank's payoff with respect to $\hat{e}$ is given as:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1-\beta(\theta))(p(\theta | e, \beta) - \rho_L) d\theta - c'(\hat{e}) \leq v_H - \rho_L - c'(\hat{e}).$$

If $e > e^{\text{FI}}$ is implemented, then at $\hat{e} = e > e^{\text{FI}}$, the bank's marginal payoff is negative.

3. Consider $\beta(\cdot) = 1$. Then the bank's problem is:

$$\max_{\hat{e} \in [0, \frac{1}{2}]} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} p(\theta | e, \beta) d\theta - c(\hat{e}).$$

Thus, the bank would take $\hat{e} = 0$. Hence, if (e, β) is incentive-feasible, then $e = 0$. However, under such (e, β) , the price is $p(\theta | e, \beta) = v_L < \rho_H$ for all θ . Thus, there is no incentive-feasible policy (e, β) such that $\beta(\cdot) = 1$. □

The last part of the lemma states that the stress test that always gives the “passing grade” is not incentive feasible. As discussed in the main text, we will show below that if externality g is low enough then the full information will be optimal.

A.2.2 Proof of Proposition 1

We prove Proposition 1. To prove the second part, we prove the following two lemmas. The first lemma provides a sufficient condition under which full information is not optimal. The second lemma establishes conditions under which the above sufficient condition holds.

Lemma 8. *If*

$$g \frac{v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L)}{c''(e^{\text{FI}})e^{\text{FI}}} < (1 - e^{\text{FI}})(v_L + g - \rho_L), \quad (17)$$

then $(e^{\text{FI}}, 0)$ is not optimal.

Proof of Lemma 8. To derive the sufficient condition (17), we consider a policy (e, β) such that $\beta(\cdot)$ is constant. Since the constraint on the price is slack when $\beta(\cdot) = 0$, we consider a relaxed problem in which the relevant constraint is the IC constraint. By Lemma 6, the

problem is:

$$\begin{aligned} & \max_{e, \beta} e(v_H + g) + (1 - e)(\beta(v_L + g) + (1 - \beta)\rho_L) - c(e) \\ & \text{subject to } \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta)(p(\theta | e, \beta) - \rho_L) d\theta = c'(e). \end{aligned}$$

Then, the Lagrangian is:

$$\begin{aligned} \mathcal{L} = & ev_H + (1 - e)(\beta v_L + (1 - \beta)\rho_L) - c(e) + (e + (1 - e)\beta)g \\ & + \lambda \left(\frac{1 - \beta}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta \left(\frac{\theta ev_H + (1 - \theta e)\beta v_L}{\theta e + (1 - \theta e)\beta} - \rho_L \right) d\theta - c'(e) \right) + \mu\beta, \end{aligned}$$

where μ is the Lagrange multiplier associated with $\beta \geq 0$.

The first-order condition with respect to e is:

$$v_H - (\beta v_L + (1 - \beta)\rho_L) - c'(e) + (1 - \beta)g + \lambda \left(\frac{1 - \beta}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \frac{\theta^2 \beta (v_H - v_L)}{(\theta e + (1 - \theta e)\beta)^2} d\theta - c''(e) \right) = 0.$$

The first-order condition with respect to β is:

$$\begin{aligned} 0 = & (1 - e)(v_L + g - \rho_L) + \mu \\ & - \lambda \left(\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta \left(\frac{\theta ev_H + (1 - \theta e)\beta v_L}{\theta e + (1 - \theta e)\beta} - \rho_L \right) d\theta + \frac{1 - \beta}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta \frac{(1 - \theta e)\theta e (v_H - v_L)}{(\theta e + (1 - \theta e)\beta)^2} d\theta \right). \end{aligned}$$

Suppose that $(e, \beta) = (e^{\text{FI}}, 0)$ is optimal. We derive the first-order necessary conditions. At $(e, \beta) = (e^{\text{FI}}, 0)$, the first-order condition with respect to e reduces to:

$$v_H - \rho_L - c'(e^{\text{FI}}) + g = \lambda c''(e^{\text{FI}}), \text{ that is, } g = \lambda c''(e^{\text{FI}}).$$

At $(e, \beta) = (e^{\text{FI}}, 0)$, the first-order condition with respect to β reduces to:

$$0 = (1 - e^{\text{FI}})(v_L + g - \rho_L) - \lambda \frac{v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L)}{e^{\text{FI}}} + \mu.$$

In order for $(e, \beta) = (e^{\text{FI}}, 0)$ to be an optimum, it is necessary that $\mu \geq 0$. Intuitively, the benefit from infinitesimally increasing e from e^{FI} exceeds the benefit from infinitesimally increasing β from 0. Hence, it is necessary that:

$$0 \leq \mu = g \frac{v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L)}{c''(e^{\text{FI}})e^{\text{FI}}} - (1 - e^{\text{FI}})(v_L + g - \rho_L).$$

Thus, under Expression (17), $(e, \beta) = (e^{\text{FI}}, 0)$ cannot be optimal. $\square$

Moving on to the second lemma:

Lemma 9. *If*

$$v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L) < (1 - e^{\text{FI}})c''(e^{\text{FI}})e^{\text{FI}}, \quad (18)$$

then there exists $\bar{g} > \rho_L - v_L$ such that Condition (17) holds if and only if $g > \bar{g}$. For instance, Expression (18) holds when $c(e) = \frac{ke^2}{1-2e}$ with $k > 0$.

Proof of Lemma 9. First, if Expression (18) holds, then there exists

$$\bar{g} = \frac{v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L)}{(1 - e^{\text{FI}})c''(e^{\text{FI}})e^{\text{FI}} - (v_H - ((1 - e^{\text{FI}})v_L + e^{\text{FI}}\rho_L))}(\rho_L - v_L) > \rho_L - v_L$$

such that Condition (17) holds if and only if $g > \bar{g}$.

Second, assume $c(e) = \frac{ke^2}{1-2e}$. Then, $c'(e) = \frac{k}{2} \left(\frac{1}{(1-2e)^2} - 1 \right)$, $(c')^{-1}(y) = \frac{1}{2} \left(1 - \sqrt{\frac{k}{k+2y}} \right)$, and $c''(e) = \frac{2k}{(1-2e)^3}$. Since

$$e^{\text{FI}} = \frac{1}{2} \left(1 - \sqrt{\frac{k}{k+2(v_H - \rho_L)}} \right) \text{ and } 1 - e^{\text{FI}} = \frac{1}{2} \left(1 + \sqrt{\frac{k}{k+2(v_H - \rho_L)}} \right),$$

Expression (18) reduces to:

$$\begin{aligned} & \frac{1}{2} \left(1 + \sqrt{\frac{k}{k+2(v_H - \rho_L)}} \right) (v_H - v_L) + \frac{1}{2} \left(1 - \sqrt{\frac{k}{k+2(v_H - \rho_L)}} \right) (v_H - \rho_L) \\ & < (v_H - v_L) \left(\frac{k}{k+2(v_H - \rho_L)} \right)^{-\frac{1}{2}}. \end{aligned}$$

Since $\rho_L > v_L$, this inequality holds because

$$\frac{k}{k+2(v_H - \rho_L)} < 1.$$

$\square$

Now, we provide the proof of Proposition 1.

Proof of Proposition 1. 1. We denote by W the social welfare under a policy $(e, \beta(\cdot | e))$, where, since we will vary e , we denote by $\beta(\cdot | e)$. Then, we have:

$$W = e(v_H + g) + (1 - e)\rho_L + \mathbb{E}[(1 - \theta e)\beta(\theta | e)](v_L + g - \rho_L) - c(e).$$

Also, we denote by W^{FI} the social welfare under the full-information benchmark. Since trade yields externality g per unit of trading,

$$W^{\text{FI}} = e^{\text{FI}}(v_H + g) + (1 - e^{\text{FI}})\rho_L - c(e^{\text{FI}}).$$

Then, we have:

$$\frac{W^{\text{FI}} - W}{e^{\text{FI}} - e} = (v_H - \rho_L + g) - \frac{c(e^{\text{FI}}) - c(e)}{e^{\text{FI}} - e} + \frac{\mathbb{E}[(1 - \theta e)(0 - \beta(\theta | e))]}{e^{\text{FI}} - e}(v_L + g - \rho_L).$$

Since c is convex, we have

$$(v_H - \rho_L + g) - \frac{c(e^{\text{FI}}) - c(e)}{e^{\text{FI}} - e} \geq (v_H - \rho_L) - c'(e^{\text{FI}}) + g = g.$$

Thus,

$$\frac{W^{\text{FI}} - W}{e^{\text{FI}} - e} \geq g + \frac{\mathbb{E}[(1 - \theta e)(0 - \beta(\theta | e))]}{e^{\text{FI}} - e}(v_L + g - \rho_L).$$

Hence, letting

$$M = \sup_{e \in [0, e^{\text{FI}}]} -\frac{\mathbb{E}[(1 - \theta e)(0 - \beta(\theta | e))]}{e^{\text{FI}} - e},$$

if $M \leq 1$, then $W^{\text{FI}} \geq W$; if $M > 1$, then $W^{\text{FI}} \geq W$ if

$$g \leq \frac{M}{M - 1}(\rho_L - v_L).$$

2. Lemma 8 establishes a sufficient condition under which full information is not optimal.

Lemma 9 establishes conditions under which Lemma 8 holds, for instance, when the cost function is given by $c(e) = \frac{ke^2}{1-2e}$.

□

A.2.3 Characterization of the Optimal Disclosure Policy

As discussed in the main text, we characterize the form of the optimal disclosure policy β^{D} when partial disclosure is optimal (i.e., when full information disclosure is not optimal).

Proposition 4. *Suppose that an optimal policy $(e^{\text{D}}, \beta^{\text{D}})$ entails partial disclosure. Then,*

there exists $\bar{\theta}$ such that

$$\beta^D(\theta) = \begin{cases} \min\left(1, \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta e^D}{1 - \theta e^D}\right) & \text{if } \theta < \bar{\theta} \\ 0 & \text{if } \theta \geq \bar{\theta} \end{cases}.$$

Proof of Proposition 4. We consider the best disclosure policy β given e . Given e , we consider disclosure policies β that satisfy the constraint on the price:

$$\beta(\theta) \leq \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta e}{1 - \theta e} \text{ for all } \theta \leq \theta^*(e).$$

The planner's objective then amounts to maximizing

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (1 - \theta e) (\beta(\theta)(v_L + g) + (1 - \beta(\theta))\rho_L) d\theta$$

among e and β of the above form subject to the IC constraint

$$c'(e) = \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta(1 - \beta(\theta))(p(\theta | e, \beta) - \rho_L) d\theta.$$

For the planner, decreasing $\beta(\theta)$ leads to an instantaneous loss of $(1 - \theta e)(v_L + g - \rho_L)$, which is decreasing in θ and e . Decreasing $\beta(\theta)$ relaxes the constraints (for the IC constraint, the right-hand side is decreasing in $\beta(\theta)$, as $p(\theta | e, \beta)$ is higher when θ is higher and $\beta(\theta)$ is lower). Thus, the optimal disclosure policy takes a bang-bang solution of the form

$$\beta(\theta) = \begin{cases} \min\left(1, \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta e}{1 - \theta e}\right) & \text{if } \theta < \bar{\theta}(e) \\ 0 & \text{if } \theta \geq \bar{\theta}(e) \end{cases}.$$

To see this, for any given incentive-feasible disclosure policy $(\tilde{e}, \tilde{\beta})$, there exists an incentive-feasible disclosure policy (e, β) such that $e \geq \tilde{e}$ and β is of the above form. Moreover, if $(\tilde{e}, \tilde{\beta})$ is not of this form for a set of θ with strictly positive measure, then the new disclosure policy (e, β) leads to a strict improvement. Note that the value of β at $\theta = \bar{\theta}(e)$ does not affect the planner's objective and the constraints. The proof is complete by letting $\bar{\theta} = \bar{\theta}(e^D)$. $\square$

We remark that, in principle, Proposition 4 implies that finding an optimal disclosure policy reduces to a finite-dimensional problem of finding $(e, \bar{\theta})$ that maximizes the planner's objective function subject to the bank's IC constraint. Yet, since the first-order conditions of the planner's problem are convoluted and add little economic insights, we characterize the form of the optimal disclosure policy in terms of the threshold $\bar{\theta}$.

A.3 Section 3.2

Proof of Proposition 2. The proof consists of five steps. In the first step, since the planner's objective function depends on $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon)$, we categorize the following two cases. In the first case, the objective function is a piece-wise continuous function over three intervals. In the second case, the objective function is a piece-wise continuous function over two intervals.

1. The first case is when $\varepsilon \in (0, 1)$ satisfies:

$$\frac{e^*}{1 - \varepsilon} < \frac{1}{2}, \text{ that is, } \varepsilon < 1 - 2e^*. \quad (19)$$

In this case, we need to consider the following three sub-cases on $e \in [0, \frac{1}{2})$:

- (a) $e \in [0, \frac{e^*}{1+\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 + \varepsilon$;
- (b) $e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = \theta^*(e)$; and
- (c) $e \in [\frac{e^*}{1-\varepsilon}, \frac{1}{2})$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 - \varepsilon$.

2. The second case is when $\varepsilon \in (0, 1)$ satisfies:

$$\frac{1}{2} \leq \frac{e^*}{1 - \varepsilon}, \text{ that is, } \varepsilon \geq 1 - 2e^*. \quad (20)$$

In this case, we need to consider the following four sub-cases on $e \in [0, 1]$:

- (a) $e \in [0, \frac{e^*}{1+\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 + \varepsilon$; and
- (b) $e \in [\frac{e^*}{1+\varepsilon}, \frac{1}{2})$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = \theta^*(e)$.

Since the objective function is a piece-wise continuous function for each case, optimal regulation e^R exists.

The second step shows that $e^R \geq \frac{e^*}{1+\varepsilon}$, that is, Cases (1a) and (2a) are never optimal. In either case, Problem (7) reduces to:

$$\max_{e \in [\frac{e^*}{1+\varepsilon}, 1]} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta - c(e),$$

that is,

$$\max_{e \in [\frac{e^*}{1+\varepsilon}, 1]} e \rho_H + (1 - e) \rho_L - c(e).$$

Since the objective function is concave and since Condition (1) implies $\rho_H - \rho_L > c' \left(\frac{e^*}{1+\varepsilon} \right)$, it follows that the objective function is increasing on $\left[0, \frac{e^*}{1+\varepsilon}\right]$, and the proof of the second step is complete.

The third step shows that if $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$ then a unique optimal regulation is $e^R = e^\diamond$, where $e^\diamond = (c')^{-1}(v_H - v_L)$. Take $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$. Then, $\frac{e^*}{1-\varepsilon} \leq e^\diamond$ implies that we are in Case 1. Especially, we consider Case (1c), in which Problem (7) reduces to:

$$\max_{e \in \left[\frac{e^*}{1-\varepsilon}, \frac{1}{2}\right)} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e),$$

that is,

$$\max_{e \in \left[\frac{e^*}{1-\varepsilon}, \frac{1}{2}\right)} e v_H + (1 - e) v_L + g - c(e).$$

As the objective function is strictly concave, the unique solution is

$$\max \left(\frac{e^*}{1-\varepsilon}, e^\diamond \right) = e^\diamond.$$

This is a solution of the entire problem, as the efficient value $e^\diamond v_H + (1 - e^\diamond) v_L + g - c(e^\diamond)$, which is an upper bound of the entire problem, is attained.

The fourth step analyzes Cases (1b) and (2b). In each case, Problem (7) reduces to:

$$\max_{e \in \left[\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}\right] \cap \left[0, \frac{1}{2}\right)} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e)} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta + \frac{1}{2\varepsilon} \int_{\theta^*(e)}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e),$$

that is,

$$\begin{aligned} \max_{e \in \left[\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}\right] \cap \left[0, \frac{1}{2}\right)} & \frac{(v_L + g)(1 + \varepsilon) - \rho_L(1 - \varepsilon)}{2\varepsilon} + \frac{1}{2\varepsilon} \left(\frac{v_H - v_L}{2} (1 + \varepsilon)^2 - \frac{\rho_H - \rho_L}{2} (1 - \varepsilon)^2 \right) e \\ & - \frac{1}{2\varepsilon} \left(\frac{v_H - \rho_H + \rho_L - v_L}{2} e^* + (v_L + g - \rho_L) \right) \frac{e^*}{e} - c(e). \end{aligned}$$

The objective function is a strictly concave function because the first two terms define an affine function in e and the third and fourth terms are a strictly concave function.

The first-order condition with respect to e is given by Expression (9):

$$\left(\frac{v_H - v_L}{2} (1 + \varepsilon)^2 - \frac{\rho_H - \rho_L}{2} (1 - \varepsilon)^2 \right) + \left(\frac{v_H - \rho_H + \rho_L - v_L}{2} e^* + (v_L + g - \rho_L) \right) \frac{e^*}{e^2} = 2\varepsilon c'(e).$$

For the right-hand side, $c'(0) = 0$, c' is increasing, and $\lim_{e \rightarrow \frac{1}{2}} c'(e) = \infty$. For the left-hand side, it diverges to infinity as $e \downarrow 0$ and it is decreasing. Hence, there is a unique $e^\dagger \in (0, \frac{1}{2})$ which satisfies the first-order condition. In fact, we show that $e^\dagger \in (e^\diamond, \frac{1}{2})$. To see this, it suffices to show:

$$\left(\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L}{2}(1 - \varepsilon)^2 \right) + \left(\frac{v_H - \rho_H + \rho_L - v_L}{2}e^* + (v_L + g - \rho_L) \right) \frac{e^*}{(e^\diamond)^2} > 2\varepsilon c'(e^\diamond).$$

Since $c'(e^\diamond) = v_H - v_L$, the above inequality reduces to:

$$\frac{v_H - \rho_H + \rho_L - v_L}{2}(1 - \varepsilon)^2 + \left(\frac{v_H - \rho_H + \rho_L - v_L}{2}e^* + (v_L + g - \rho_L) \right) \frac{e^*}{(e^\diamond)^2} > 0,$$

which follows because $v_H > \rho_H$ and $v_L + g > \rho_L > v_L$.

The fifth step shows that

$$e^R = \min \left(e^\dagger, \frac{e^*}{1 - \varepsilon} \right) \text{ when } 1 - \frac{e^*}{e^\diamond} \leq \varepsilon, \text{ i.e., } e^\diamond \leq \frac{e^*}{1 - \varepsilon}.$$

We start with Case 1. We have seen from the analysis of Case (1a) that the objective function is increasing on $[0, \frac{e^*}{1 + \varepsilon}]$. For Case (1c), since $e^\diamond \leq \frac{e^*}{1 - \varepsilon}$, the objective function is decreasing on $[\frac{e^*}{1 - \varepsilon}, \frac{1}{2})$. Thus, the objective function is maximized on $[\frac{e^*}{1 + \varepsilon}, \frac{e^*}{1 - \varepsilon}]$. Since it is a strictly concave function on this interval, it has a unique maximizer $\min(e^\dagger, \frac{e^*}{1 - \varepsilon})$.

Next, we consider Case 2. We have seen from the analysis of Case (2a) that the objective function is increasing on $[0, \frac{e^*}{1 + \varepsilon}]$. For Cases (2b), as $e^\dagger \in [\frac{e^*}{1 + \varepsilon}, \frac{1}{2})$, the objective function is uniquely maximized at $e^\dagger = \min(e^\dagger, \frac{e^*}{1 - \varepsilon})$, as $e^\dagger < \frac{1}{2} \leq \frac{e^*}{1 - \varepsilon}$.

Thus, $e^R = \min(e^\dagger, \frac{e^*}{1 - \varepsilon})$ holds for each possibility. $\square$

Proof of Corollary 1. Proposition 2 implies the following. When $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, trades occur for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$. When $1 - \frac{e^*}{e^\diamond} \leq \varepsilon$, as long as $\frac{e^*}{1 - \varepsilon} \leq e^\dagger$, trades occur for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$. On the one hand, $\frac{e^*}{1 - \varepsilon}$ is increasing in ε and diverges to infinity as $\varepsilon \uparrow 1$. On the other hand, at $\varepsilon = 1 - \frac{e^*}{e^\diamond}$ (i.e., $e^\diamond = \frac{e^*}{1 - \varepsilon}$), we have $e^\dagger > \frac{e^*}{1 - \varepsilon}$. We show that $e^\dagger$ is decreasing in ε . Since $e^\dagger$ is characterized by Expression (9), differentiating both sides of Expression (9) with respect to ε yields:

$$\begin{aligned} & (v_H - v_L)(1 + \varepsilon) + (\rho_H - \rho_L)(1 - \varepsilon) - 2c'(e^\dagger) \\ &= \left(2 \left(\frac{v_H - \rho_H + \rho_L - v_L}{2}e^* + (v_L + g - \rho_L) \right) \frac{e^*}{(e^\dagger)^3} + 2\varepsilon c''(e^\dagger) \right) \frac{\partial e^\dagger}{\partial \varepsilon}. \end{aligned}$$

Since the left-hand side satisfies

$$\begin{aligned}
& (v_H - v_L)(1 + \varepsilon) + (\rho_H - \rho_L)(1 - \varepsilon) - 2c'(e^\dagger) \\
& < (v_H - v_L)(1 + \varepsilon) + (\rho_H - \rho_L)(1 - \varepsilon) - (v_H - v_L)\frac{(1 + \varepsilon)^2}{2} + (\rho_H - \rho_L)\frac{(1 - \varepsilon)^2}{2} \\
& = -\frac{(1 + \varepsilon)(1 - \varepsilon)}{2\varepsilon}(v_H - \rho_H + \rho_L - v_L) < 0,
\end{aligned}$$

it follows that $e^\dagger$ is decreasing in ε .

Thus, there exists a unique $\bar{\varepsilon} \in (0, 1)$ such that $\frac{e^*}{1-\bar{\varepsilon}} = e^\dagger$. The statement of the proposition holds with this $\bar{\varepsilon}$. $\square$

Finally, as discussed in the main text, we formulate and prove the following comparative-statics result on the optimal regulation e^R .

Corollary 4. 1. e^R is not monotone in ε .

2. e^R is non-decreasing in g .

Proof of Corollary 4. 1. First, when $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, $e^R = e^\diamond$ does not depend on ε . Second, when $e^\dagger \geq \frac{e^*}{1-\varepsilon}$, $e^R = \frac{e^*}{1-\varepsilon}$ is increasing in ε . Third, when $e^\dagger \leq \frac{e^*}{1-\varepsilon}$, $e^R = e^\dagger$ is decreasing in ε .

2. Since $e^\diamond$ and $\frac{e^*}{1-\varepsilon}$ do not depend on g , it suffices to show that $e^\dagger$ is non-decreasing in g . Take $\tilde{g}$ with $\tilde{g} > g$. Suppose to the contrary that $\tilde{e}^* \leq e^\dagger$. Then, while the left-hand side of Expression (9) is strictly increased, the right-hand side of Expression (9) is weakly decreased. This is a contradiction. $\square$

A.4 Section 4

Proof of Lemma 2. It is without loss to consider the case in which $\bar{e}_M \geq \tilde{e} \geq e$. It follows from Expression (2) that $\theta^*(\tilde{e}) \leq \theta^*(e)$. For any $\theta \in [1 - \varepsilon, 1 + \varepsilon]$ with $\theta \geq \theta^*(\tilde{e})$, we have $\tilde{\beta}(\theta) = 1 \geq \beta(\theta)$. For any $\theta \in [1 - \varepsilon, 1 + \varepsilon]$ with $\theta < \theta^*(\tilde{e})$, it follows from

$$\frac{\partial \beta(\theta)}{\partial e} = \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta}{(1 - \theta e)^2} > 0$$

that

$$\tilde{\beta}(\theta) = \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta \tilde{e}}{1 - \theta \tilde{e}} \geq \frac{v_H - \rho_H}{\rho_H - v_L} \frac{\theta e}{1 - \theta e} = \beta(\theta).$$

The proof is complete. $\square$

Proof of Lemma 3. Suppose that $\beta(\theta) = 0$ for some $\theta \in [1 - \varepsilon, 1 + \varepsilon]$. Then, the price at θ satisfies

$$v_H = \rho_H,$$

which is a contradiction. Hence, $\beta(\theta) > 0$ for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$. $\square$

Proof of Theorem 1. The structure of the proof resembles that of Proposition 2. The proof consists of five steps. In the first step, since the planner's objective function depends on $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon)$, we categorize the following two cases. In the first case, the objective function is a piece-wise continuous function over three intervals. In the second case, the objective function is a piece-wise continuous function over two intervals.

1. The first case is when $\varepsilon \in (0, 1)$ satisfies:

$$\frac{e^*}{1 - \varepsilon} < \frac{1}{2}, \text{ that is, } \varepsilon < 1 - 2e^*. \quad (21)$$

In this case, we need to consider the following three sub-cases on $e \in [0, \frac{1}{2}]$:

- (a) $e \in [0, \frac{e^*}{1+\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 + \varepsilon$;
- (b) $e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = \theta^*(e)$; and
- (c) $e \in [\frac{e^*}{1-\varepsilon}, \frac{1}{2})$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 - \varepsilon$.

2. The second case is when $\varepsilon \in (0, 1)$ satisfies:

$$\frac{1}{2} \leq \frac{e^*}{1 - \varepsilon}, \text{ that is, } 1 - 2e^* \leq \varepsilon. \quad (22)$$

In this case, we need to consider the following two sub-cases on $e \in [0, \frac{1}{2}]$:

- (a) $e \in [0, \frac{e^*}{1+\varepsilon}]$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = 1 + \varepsilon$; and
- (b) $e \in [\frac{e^*}{1+\varepsilon}, \frac{1}{2})$, in which $\text{med}(1 - \varepsilon, \theta^*(e), 1 + \varepsilon) = \theta^*(e)$.

Since the objective function is a piece-wise continuous function for each case, optimal regulation e^{RD} exists.

The second step shows that $e^{\text{RD}} \geq \frac{e^*}{1+\varepsilon}$, that is, Cases (1a) and (2a) are never optimal. In either case, Problem (11) reduces to:

$$\max_{e \in [0, \frac{e^*}{1+\varepsilon}]} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \left(\rho_L + \frac{\rho_H - \rho_L + g}{e^*} \theta e \right) d\theta - c(e),$$

that is,

$$\max_{e \in [0, \frac{e^*}{1+\varepsilon}]} \rho_L + \frac{\rho_H - \rho_L + g}{e^*} e - c(e).$$

Since the objective function is strictly concave and Condition (1) implies

$$\frac{\rho_H - \rho_L + g}{e^*} > v_H - v_L > c' \left(\frac{e^*}{1 + \varepsilon} \right),$$

it follows that the objective function is increasing on $[0, \frac{e^*}{1+\varepsilon}]$, and the proof of the second step is complete.

The third step shows that if $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$ then a unique optimal regulation is $e^{\text{RD}} = e^\diamond$. Take $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$. Then, $\frac{e^*}{1-\varepsilon} \leq e^\diamond$ and $e^\diamond < \frac{1}{2}$ imply that we are in Case 1. Especially, we consider Case (1c), in which Problem (11) reduces to:

$$\max_{e \in [\frac{e^*}{1-\varepsilon}, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e),$$

that is,

$$\max_{e \in [\frac{e^*}{1-\varepsilon}, \frac{1}{2})} v_L + g + (v_H - v_L)e - c(e).$$

Thus, the analysis reduces to that of Case (1c) in the proof of Proposition 2. As the objective function is strictly concave, the unique solution is $e^\diamond$. This is a solution of the entire problem, as the first-best value $e^\diamond v_H + (1 - e^\diamond) v_L + g - c(e^\diamond)$, which is an upper bound of the entire problem, is attained.

The fourth step analyzes Cases (1b) and (2b). In each case, Problem (11) reduces to:

$$\max_{e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}] \cap [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e)} \left(\rho_L + \frac{\rho_H - \rho_L + g}{e^*} e \theta \right) d\theta + \frac{1}{2\varepsilon} \int_{\theta^*(e)}^{1+\varepsilon} (\theta e v_H + (1 - \theta e) v_L + g) d\theta - c(e).$$

Thus, the problem is:

$$\begin{aligned} \max_{e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}] \cap [0, \frac{1}{2})} & \frac{(v_L + g)(1 + \varepsilon) - \rho_L(1 - \varepsilon)}{2\varepsilon} + \frac{1}{2\varepsilon} \left(\frac{v_H - v_L}{2} (1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{2e^*} (1 - \varepsilon)^2 \right) e \\ & - \frac{1}{2\varepsilon} \frac{v_L + g - \rho_L}{2} \frac{e^*}{e} - c(e). \end{aligned}$$

The objective function is a strictly concave function because the first two terms define an

affine function in e and the third and fourth terms are a strictly concave function.

The first-order condition is given by Expression (13), or:

$$\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{2e^*}(1 - \varepsilon)^2 + \frac{v_L + g - \rho_L}{2} \frac{e^*}{e^2} = 2\varepsilon c'(e). \quad (23)$$

The right-hand side is increasing in e , approaches 0 as $e \downarrow 0$, and diverges to infinity as $e \uparrow \frac{1}{2}$. The left-hand side goes to infinity as $e \downarrow 0$, and is decreasing in e . Hence, there is a unique solution $e^\dagger$ of Expression (13) in $(0, \frac{1}{2})$. In fact, we show that $e^\dagger \geq e^\diamond$ when $\varepsilon \geq 1 - \frac{e^*}{e^\diamond}$. To see this, it suffices to show:

$$\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{2e^*}(1 - \varepsilon)^2 + \frac{v_L + g - \rho_L}{2} \frac{e^*}{(e^\diamond)^2} \geq 2\varepsilon c'(e^\diamond).$$

Since $c'(e^\diamond) = v_H - v_L$, the above inequality reduces to:

$$\frac{(1 - \varepsilon)^2}{2} \left(v_H - v_L - \frac{\rho_H - \rho_L + g}{e^*} \right) + \frac{v_L + g - \rho_L}{2} \frac{e^*}{(e^\diamond)^2} \geq 0,$$

that is,

$$\frac{v_L + g - \rho_L}{2e^*} \left(\left(\frac{e^*}{e^\diamond} \right)^2 - (1 - \varepsilon)^2 \right) \geq 0,$$

which follows when $\varepsilon \geq 1 - \frac{e^*}{e^\diamond}$. In fact, when $\varepsilon = 1 - \frac{e^*}{e^\diamond}$, $e = e^\diamond$ satisfies the first-order condition.

We also show that $e^\dagger \geq \frac{e^*}{1 - \varepsilon}$ when $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$. To that end, it suffices to show:

$$\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{2e^*}(1 - \varepsilon)^2 + \frac{v_L + g - \rho_L}{2} \frac{(1 - \varepsilon)^2}{e^*} \geq 2\varepsilon c' \left(\frac{e^*}{1 - \varepsilon} \right),$$

that is,

$$\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - v_L}{2e^*}(1 - \varepsilon)^2 \geq 2\varepsilon c' \left(\frac{e^*}{1 - \varepsilon} \right).$$

When $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, we have $e^\diamond \geq \frac{e^*}{1 - \varepsilon}$. Thus, it suffices to show:

$$\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - v_L}{2e^*}(1 - \varepsilon)^2 \geq 2\varepsilon c'(e^\diamond) = 2\varepsilon(v_H - v_L),$$

that is,

$$\frac{v_H - v_L}{2\varepsilon}(1 - \varepsilon)^2 \geq \frac{\rho_H - v_L}{2\varepsilon e^*}(1 - \varepsilon)^2,$$

which holds with equality. Similarly, $e^\dagger \leq \frac{e^*}{1-\varepsilon}$ when $\varepsilon \geq 1 - \frac{e^*}{e^\diamond}$.

The fifth step shows:

$$e^{\text{RD}} = e^\dagger \text{ when } 1 - \frac{e^*}{e^\diamond} \leq \varepsilon, \text{ i.e., } e^\diamond \leq \frac{e^*}{1 - \varepsilon}.$$

We start with Case 1. We have seen from the analysis of Case (1a) that the objective function is increasing on $[0, \frac{e^*}{1+\varepsilon}]$. For Case (1c), since $e^\diamond \leq \frac{e^*}{1-\varepsilon}$, the objective function is decreasing on $[\frac{e^*}{1-\varepsilon}, \frac{1}{2})$. Thus, the objective function is maximized on $[\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$. Since it is a strictly concave function on this interval, it has a unique maximizer $\min(e^\dagger, \frac{e^*}{1-\varepsilon}) = e^\dagger$.

Next, we consider Case 2. We have seen from the analysis of Case (2a) that the objective function is increasing on $[0, \frac{e^*}{1+\varepsilon}]$. Thus, the objective function is maximized on $[\frac{e^*}{1+\varepsilon}, \frac{1}{2})$. Since it is a strictly concave function on this interval, it has a unique maximizer $e^\dagger$. $\square$

Proof of Corollary 2. First, when $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, we have $e^{\text{R}}(\varepsilon) = e^\diamond = e^{\text{RD}}(\varepsilon)$. Second, when $1 - \frac{e^*}{e^\diamond} \leq \varepsilon$, to show $e^{\text{R}}(\varepsilon) \geq e^{\text{RD}}(\varepsilon)$, it suffices to show that $e^\dagger > e^\dagger$. To that end, we show:

$$\begin{aligned} & \left(\frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L}{2}(1 - \varepsilon)^2 \right) + \left(\frac{v_H - \rho_H + \rho_L - v_L}{2}e^* + (v_L + g - \rho_L) \right) \frac{e^*}{(e^\dagger)^2} \\ & > 2\varepsilon c'(e^\dagger) \\ & = \frac{v_H - v_L}{2}(1 + \varepsilon)^2 - \frac{\rho_H - \rho_L + g}{2e^*}(1 - \varepsilon)^2 + \frac{v_L + g - \rho_L}{2} \frac{e^*}{(e^\dagger)^2}. \end{aligned}$$

The above inequality follows because it reduces to:

$$\frac{(\rho_H - \rho_L) + (v_H - v_L)g}{\rho_H - v_L} \frac{(1 - \varepsilon)^2}{2} + \left(\frac{v_H - \rho_H + \rho_L - v_L}{2}e^* + \frac{v_L + g - \rho_L}{2} \right) \frac{e^*}{(e^\dagger)^2} > 0.$$

$\square$

Proof of Corollary 3. 1. First, when $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, $e^{\text{RD}} = e^\diamond$ does not depend on ε . Second, when $\varepsilon \geq 1 - \frac{e^*}{e^\diamond}$, we have $e^{\text{RD}} = e^\dagger$. Differentiating Expression (23) with respect to ε yields:

$$(v_H - v_L)(1 + \varepsilon) + \frac{\rho_H - \rho_L + g}{e^*}(1 - \varepsilon) - 2c'(e) = \left((v_L + g - \rho_L) \frac{e^*}{e^3} + 2\varepsilon c''(e) \right) \frac{\partial e}{\partial \varepsilon}.$$

Then, the left-hand side of the above equation satisfies:

$$\begin{aligned}
& (v_H - v_L)(1 + \varepsilon) + \frac{\rho_H - \rho_L + g}{e^*}(1 - \varepsilon) - 2c'(e) \\
&= (v_H - v_L)(1 + \varepsilon) + \frac{\rho_H - \rho_L + g}{e^*}(1 - \varepsilon) - \frac{v_H - v_L}{2\varepsilon}(1 + \varepsilon)^2 + \frac{\rho_H - \rho_L + g}{2\varepsilon e^*}(1 - \varepsilon)^2 \\
&\quad - \frac{v_L + g - \rho_L}{2\varepsilon} \frac{e^*}{e^2} \\
&= \frac{v_L + g - \rho_L}{2\varepsilon e^*} \left((1 + \varepsilon)(1 - \varepsilon) - \left(\frac{e^*}{e} \right)^2 \right) \\
&\geq \frac{v_L + g - \rho_L}{2\varepsilon e^*} \left((1 + \varepsilon)(1 - \varepsilon) - \left(\frac{e^*}{e^\diamond} \right)^2 \right).
\end{aligned}$$

Thus, e^{RD} is initially increasing when ε is close to $1 - \frac{e^*}{e^\diamond}$. In contrast, when ε is close to 1, $e^\dagger$ is locally decreasing.

2. Let g and $\tilde{g}$ be such that $\tilde{g} > g$. Denote by $\Delta := \tilde{g} - g$. Let (e_g, β_g) be a solution to Problem (11) under g , and let $(e_{\tilde{g}}, \beta_{\tilde{g}})$ be a solution to Problem (11) under $\tilde{g}$. Observe that (e_g, β_g) and $(e_{\tilde{g}}, \beta_{\tilde{g}})$ are feasible, i.e., satisfy Expression (10). On the one hand, we have:

$$\begin{aligned}
& \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e_g(v_H + g) + (1 - \theta e_g)(\beta_g(\theta)(v_L + g) + (1 - \beta_g(\theta))\rho_L)) d\theta - c(e_g) \\
& \geq \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e_{\tilde{g}}(v_H + g) + (1 - \theta e_{\tilde{g}})(\beta_{\tilde{g}}(\theta)(v_L + g) + (1 - \beta_{\tilde{g}}(\theta))\rho_L)) d\theta - c(e_{\tilde{g}}).
\end{aligned}$$

On the other hand, we have:

$$\begin{aligned}
& \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e_{\tilde{g}}(v_H + \tilde{g}) + (1 - \theta e_{\tilde{g}})(\beta_{\tilde{g}}(\theta)(v_L + \tilde{g}) + (1 - \beta_{\tilde{g}}(\theta))\rho_L)) d\theta - c(e_{\tilde{g}}) \\
& \geq \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e_g(v_H + \tilde{g}) + (1 - \theta e_g)(\beta_g(\theta)(v_L + \tilde{g}) + (1 - \beta_g(\theta))\rho_L)) d\theta - c(e_g).
\end{aligned}$$

Adding both inequalities yields:

$$\frac{\Delta}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \{(\theta e_{\tilde{g}} - \theta e_g)(1 - \beta_g(\theta)) + (1 - \theta e_{\tilde{g}})(\beta_{\tilde{g}}(\theta) - \beta_g(\theta))\} d\theta \geq 0. \quad (24)$$

Suppose to the contrary that $e_g > e_{\tilde{g}}$. Then, it follows from Lemma 2 that $\beta_g \geq \beta_{\tilde{g}}$.

Then, it follows from Expression (24) that, for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$,

$$(\theta e_{\tilde{g}} - \theta e_g)(1 - \beta_g(\theta)) + (1 - \theta e_{\tilde{g}})(\beta_{\tilde{g}}(\theta) - \beta_g(\theta)) = 0.$$

Then, we obtain $\beta_g(\theta) = \beta_{\tilde{g}}(\theta) = 1$ for all $\theta \in [1 - \varepsilon, 1 + \varepsilon]$. Under $\beta(\cdot) = 1$, consider the problem under generic g :

$$\max_{e \in [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e(v_H + g) + (1 - \theta e)(v_L + g)) d\theta - c(e).$$

The solution does not depend on g , as the problem can be written as

$$\max_{e \in [0, \frac{1}{2})} ev_H + (1 - e)v_L + g - c(e).$$

Thus, we cannot have $e_g > e_{\tilde{g}}$. Hence, it must be the case that $e_{\tilde{g}} \geq e_g$. □

A.5 Section 5.2

Proof of Proposition 3. The proof consists of three steps. The first step shows $e^{\text{BR}} \geq \frac{e^*}{1+\varepsilon}$ by showing that the bank's payoff is increasing on $[0, \frac{e^*}{1+\varepsilon}]$. When $e \leq \frac{e^*}{1+\varepsilon}$, the bank's payoff is:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} (\theta e \rho_H + (1 - \theta e) \rho_L) d\theta - c(e) = e \rho_H + (1 - e) \rho_L - c(e).$$

Since it follows from Condition (1) that $(c')^{-1}(\rho_H - \rho_L) > \frac{e^*}{1+\varepsilon}$, the bank's payoff is increasing on $[0, \frac{e^*}{1+\varepsilon}]$.

The second step shows that $e^{\text{BR}} \leq \max(e^\diamond, \frac{e^*}{1-\varepsilon})$. The second step also shows that if $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, then $e^{\text{BR}} = e^\diamond$. The proof consists of two sub-steps. In the first sub-step, suppose that $e^\diamond \geq \frac{e^*}{1-\varepsilon}$. Then, if $e = e^\diamond$, the bank's payoff is:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta e v_H + (1 - \theta e) v_L d\theta - c(e) = e v_H + (1 - e) v_L - c(e),$$

which is an upper bound of the bank's payoff. Thus, if $\varepsilon \leq 1 - \frac{e^*}{e^\diamond}$, then $e^{\text{BR}} = e^\diamond$. In the second sub-step, suppose that $e^\diamond < \frac{e^*}{1-\varepsilon}$. If $e \geq \frac{e^*}{1-\varepsilon}$, then the bank's payoff is:

$$\frac{1}{2\varepsilon} \int_{1-\varepsilon}^{1+\varepsilon} \theta e v_H + (1 - \theta e) v_L d\theta - c(e) = e v_H + (1 - e) v_L - c(e).$$

Since this is maximized at $e^\diamond < \frac{e^*}{1-\varepsilon}$, it follows that the bank's payoff is decreasing when $e \geq \frac{e^*}{1-\varepsilon}$.

In sum, the second step implies that if $\varepsilon > 1 - \frac{e^*}{e^\diamond}$, then the bank's problem reduces to:

$$\max_{e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}] \cap [0, \frac{1}{2})} \frac{1}{2\varepsilon} \int_{1-\varepsilon}^{\theta^*(e)} (\theta e \rho_H + (1-\theta e) \rho_L) d\theta + \frac{1}{2\varepsilon} \int_{\theta^*(e)}^{1+\varepsilon} (\theta e v_H + (1-\theta e) v_L) d\theta - c(e). \quad (25)$$

The third step shows that if $\varepsilon > 1 - \frac{e^*}{e^\diamond}$, then Problem (25) has a unique maximizer in $[\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$, which is specified by the lemma. To that end, differentiating the bank's payoff function in Problem (25) with respect to e yields

$$\begin{aligned} & \frac{\rho_H - \rho_L}{4\varepsilon} \left(\frac{(e^*)^2}{e^2} - (1-\varepsilon)^2 \right) + \frac{v_H - v_L}{4\varepsilon} \left((1+\varepsilon)^2 - \frac{(e^*)^2}{e^2} \right) + \frac{\rho_H - \rho_L}{2\varepsilon} \frac{e^*(1-e^*)}{e^2} - c'(e) \\ &= \frac{v_H - v_L}{4\varepsilon} (1+\varepsilon)^2 - \frac{\rho_H - \rho_L}{4\varepsilon} (1-\varepsilon)^2 + \frac{(v_H - \rho_H + \rho_L - v_L)e^* + 2(v_L - \rho_L)e^*}{4\varepsilon} \frac{e^*}{e^2} - c'(e). \end{aligned}$$

Thus, we obtain Expression (15) as the first-order condition. If there is no $e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$ such that Expression (15) holds, then the unique solution of the bank's problem is $e^{\text{BR}} = \frac{e^*}{1-\varepsilon}$. Observe that if this is the case then $\frac{e^*}{1-\varepsilon} < \frac{1}{2}$, as there exists some $e \in (\frac{e^*}{1+\varepsilon}, \frac{1}{2})$ which satisfies the first-order condition. This is because, while $c'(e) \uparrow \infty$ as $e \uparrow \frac{1}{2}$, the right-hand side of the first-order condition is bounded on $[\frac{e^*}{1+\varepsilon}, \frac{1}{2}]$. In fact, such e is unique. Also, comparing Expressions (9) and (15) implies that $e < e^{\text{R}}$.

Thus, suppose that there exists $e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$ such that Expression (15) holds. If

$$(v_H - \rho_H + \rho_L - v_L)e^* \geq 2(\rho_L - v_L),$$

then the left-hand side of the first-order condition is weakly decreasing while the right-hand side is increasing. Thus, e is the unique solution of the bank's problem. If

$$(v_H - \rho_H + \rho_L - v_L)e^* < 2(\rho_L - v_L),$$

then the left-hand side of the first-order condition is increasing and concave while the right-hand side is increasing and convex. Thus, there exists a unique $e \in [\frac{e^*}{1+\varepsilon}, \frac{e^*}{1-\varepsilon}]$ that satisfies the first-order condition, and this e is the unique solution of the bank's problem. The proof is complete. $\square$