

Model Secrecy and Stress Tests*

Yaron Leitner[†] Basil Williams[‡]

February 2018

Abstract

Conventional wisdom holds that the models used to stress test banks should be kept secret to prevent gaming. We show instead that secrecy can be suboptimal, because although it deters gaming, it may also deter socially desirable investment. When the regulator can choose the minimum standard for passing the test, we show that secrecy is suboptimal if the regulator is sufficiently uncertain regarding bank characteristics. When failing the bank is socially costly, then under some conditions, secrecy is suboptimal when the bank's private cost of failure is either sufficiently high or sufficiently low. Finally, we relate our results to several current and proposed stress testing policies.

JEL Classification: D82, G01

Keywords: stress tests, information disclosure, delegation, bank incentives, Fed models

*Preliminary draft. The views expressed in this paper are those of the authors and do not necessarily reflect those of the Federal Reserve Bank of Philadelphia or the Federal Reserve System.

[†]Federal Reserve Bank of Philadelphia, Yaron.Leitner@phil.frb.org

[‡]New York University, basil.williams@nyu.edu

1 Introduction

According to Federal Reserve officials, the models that are used to stress test banks are kept secret to prevent banks from gaming them. Indeed, if a bank knows that the Fed’s models underestimate the risks of some class of assets, the bank can invest in those assets without fear of failing the test. However, banks complain about this secrecy, claiming that even their best efforts to prepare for a test could result in unexpected and costly failure.¹

Our main contribution is to present conditions under which, contrary to conventional wisdom and the statements of some policymakers, fully revealing the stress model to banks is optimal.² The results build on the idea that hidden models make banks cautious about risky investment, which could have two effects: banks may game less, but they may also invest less in socially desirable assets. Revealing the model leads to a better social outcome if the second effect dominates. This idea leads to three main results. First, if banks are sufficiently cautious about risky investment or if failing the test is sufficiently costly to them, revealing the regulator’s model is optimal because it prevents underinvestment in socially desirable assets. Second, even if the regulator can adjust the test to make it easier to pass, revealing may still be optimal if uncertainty about the bank characteristics is sufficiently high, or if the regulator is forced to apply the same test to sufficiently different banks. Third, if there is some social cost when banks fail the test, then the optimal disclosure policy may be nonmonotonic in bank characteristics. For example, revealing could be optimal when the bank’s bias toward risky investment or the bank’s private cost of failure is either sufficiently high or sufficiently low.

In our baseline model, the bank can invest in one of two portfolios: a safe

¹A recent proposal from the Federal Reserve suggests enhanced disclosure of the Fed models, such as revealing key variables and some equations, and illustrating how the Fed model will work on some hypothetical loan portfolios; but even under this proposal, the Fed will not reveal the exact models. See <https://www.gpo.gov/fdsys/pkg/FR-2017-12-15/pdf/2017-26856.pdf>

²See former Fed Governor Tarullo’s speech for arguments against fully revealing the model. <https://www.federalreserve.gov/newsevents/speech/tarullo20160926a.htm>

portfolio, which will surely pass the test, or a risky portfolio, which may or may not pass the test. We assume that the bank always prefers to invest in the risky portfolio, whereas the regulator prefers the risky portfolio only if its value during a crisis is sufficiently high. This value is represented by the state of nature. The bank knows more than the regulator about the value of the risky portfolio, and for simplicity, we assume that the bank observes the state with certainty.

We capture the idea of a hidden stress test model by assuming that the regulator observes a noisy signal of the state, and that the regulator passes a bank that invested in the risky portfolio if and only if the signal realization is above some threshold. If the bank fails the test, the regulator forces the bank to alter its portfolio, which we assume is costly for the bank. If the bank passes the test, the regulator leaves the bank's portfolio unchanged. Because the regulator bases his decision on a noisy signal, he could err by passing a bank that invested in a socially undesirable portfolio or by failing a bank that invested in a socially desirable portfolio.

When the regulator's model is hidden, the bank fears failure and is therefore cautious, investing only when its privately observed state exceeds some threshold. We refer to this threshold as the bank's *cautious threshold*. In contrast, when the regulator reveals his signal, and that signal exceeds the passing threshold, the bank invests in the risky portfolio regardless of its privately observed state. So the bank may invest in the risky portfolio even if it knows that doing so is harmful to society. In other words, the bank may game the test.

We compare between two disclosure regimes: a transparent regime under which the regulator reveals his signal to the bank before the bank selects a portfolio, and a secrecy regime in which the regulator's signal is kept secret. We focus on two cases. In the first case, the regulator must follow an exogenously given threshold for passing or failing the bank.³ In the second case,

³For example, the regulator must ensure that the bank's capital during an adverse stress scenario does not fall below some predetermined level.

the regulator can choose the passing threshold optimally. So in the second case, the regulator has two tools to influence the bank's portfolio decision: the disclosure regime and the standard to which the bank is held. In both cases, the regulator announces and commits to the passing threshold publicly before the bank selects its portfolio.

In the first case with an exogenously given threshold, we show that revealing is optimal if the bank's cautious threshold is sufficiently high. This happens, for example, if the bank's cost of failing the test is sufficiently high. Intuitively, in this case, the bank's fear of failing the test leads to a significant reduction in socially beneficial investment, and this reduction more than offsets the benefits from a reduction in a socially harmful investment.

In the second case in which the regulator can choose the passing threshold optimally, he can reassure an overly cautious bank by lowering the passing threshold, thereby making the test easier to pass. However, the bank's cautious threshold depends not only on the difficulty of the test but also on the bank's characteristics (e.g., cost of failing the test). If the regulator is certain about the bank characteristics, he can precisely calibrate the bank's cautious threshold by adjusting the passing threshold, so it is optimal to not reveal. However, precise calibration is impossible when bank characteristics are unknown. We show that, under some conditions, if the regulator is sufficiently uncertain about the bank's characteristics, then revealing is optimal.

Finally, we focus on another force that increases the benefit of revealing the regulator's model. Failing the test and the resulting change in the bank's portfolio might be costly not only for the bank but also for society. We show that if the social cost of failing the test is sufficiently high, it is optimal to reveal the regulator's signal. If instead, the social cost of failing the test is low, the optimal disclosure regime depends on the bank's cautious threshold, and in particular, on the bank's cost of failing the test. Interestingly, the relationship between the optimal disclosure policy and the bank's cost of failure is not necessarily monotone.

For example, under some conditions, revealing is optimal when the bank's

private cost of failing the test is either sufficiently high or sufficiently low. Intuitively, if the cost is high, then fear of failing the test deters the bank from taking a socially desirable risk, and so it is optimal to reveal. If the cost is low, fear of failure does little to deter investment in socially harmful assets. But then it is better to reveal to avoid the social cost of failing the bank. In other words, in this case, providing incentives to the bank via model secrecy is too costly for the regulator.

Before concluding, we discuss additional policy implications from our model. In particular, we relate our results to three specific policies: the current policy of giving banks a short time to revise their capital plans, the proposal to reveal the Fed’s estimated losses on hypothetical portfolios, and the suggestion to accompany greater model transparency with increased capital requirements.

2 Related Literature

The existing literature has focused on disclosure of regulators’ stress test *results* to *investors*.⁴ In contrast, we focus on disclosure of regulators’ stress test *models* to *banks*. To our knowledge, we are the first paper to study this problem.

Our setting is a principal-agent problem in which the principal (the regulator) and agent (the bank) each have private information, the agent takes an action, and the principal can take a follow up action, which is costly both to the agent and to the principal. Our focus is on whether the principal should reveal his private information before the agent takes the action. [Levit \(2016\)](#) also considers a setting in which a principal can reverse the agent’s action. In his basic setting, the principal is more informed than the agent, so intervention can protect the agent from bad outcomes. His paper shows that in some cases the principal can obtain a better outcome by recommending an action to an

⁴See [Goldstein and Leitner \(2017\)](#); [Williams \(2015\)](#); [Goldstein and Sapra \(2014\)](#); [Bouvard, Chaigneau, and Motta \(2015\)](#); [Faria-e Castro, Martinez, and Philippon \(2016\)](#); [Inostroza and Pavan \(2017\)](#); [Orlov, Zryumov, and Skrzypacz \(2017\)](#); [Gick and Pausch \(2012\)](#).

agent and committing not to intervene. In our setting, however, intervention is bad for the agent and is crucial for providing incentives; instead, the principal chooses whether or not to disclose information related to his intervention policy.

As in the delegation literature initiated by [Holmstrom \(1984\)](#)⁵, we rule out transfers between the two parties. The case in which the regulator reveals his information corresponds to a standard delegation problem in which the principal delegates partial authority to the agent. In particular, by revealing his signal, the regulator effectively restricts the bank’s action space to those actions that will surely pass the test. In contrast, the case in which the regulator does not reveal his information is new to this literature and can be thought of as “delegation with hidden evaluation.” The regulator does not restrict the set of actions that the bank can take (i.e., there is full delegation), but the regulator responds to the bank’s action based on an evaluation process (a model) that is hidden from the bank. Our paper provides conditions under which hiding the evaluation process is preferred to revealing it.

The idea that uncertainty regarding the regulator policy can affect incentives appears in other settings. For example, [Lazear \(2006\)](#) shows hidden tests could be a way to induce a socially optimal action, such as studying or not speeding. In his setting, the regulator knows what the socially optimal action is, whereas in our setting the regulator does not know. The possibility of wrongful punishment in our setting can create excessive caution in banks, which is the driving force behind our results. [Freixas \(2000\)](#) offers some justification for “constructive ambiguity” of bank bailout policy by showing that under some conditions, it is optimal for the regulator to use a mixed bailout strategy. In our paper, the regulator follows a deterministic policy rule to pass or fail a bank, but the rule is based on information that could be unknown to the bank.

⁵See also [Dessein \(2002\)](#); [Amador and Bagwell \(2013, 2016\)](#); [Amador, Bagwell, and Frankel \(2017\)](#); [Grenadier, Malenko, and Malenko \(2016\)](#); [Chakraborty and Yilmaz \(2017\)](#); [Harris and Raviv \(2005, 2006\)](#); [Halac and Yared \(2016\)](#).

Finally, there is a large empirical literature that documents how political and regulatory uncertainty can affect the real economy, including reducing investment.⁶ In particular, [Gissler, Oldfather, and Ruffino \(2016\)](#) offer evidence which suggests that uncertainty about the regulation of qualified mortgages caused banks to reduce mortgage lending. The literature is consistent with the idea in our paper that hidden tests could induce the bank to invest less.

3 Model

There are two agents: a bank and a regulator. The bank can invest in either a risky portfolio or a safe portfolio. The payoff from investing in the risky portfolio depends on an unobservable state $\omega \in \Omega \equiv [\underline{\omega}, \bar{\omega}] \subset \mathbb{R}$. The payoff to the bank is $u(\omega)$ and the payoff to the regulator is $v(\omega)$. The payoff to the regulator represents the payoff to society. The payoff from investing in the safe portfolio does not depend on the state, and is normalized to zero for both the bank and regulator. That is, u and v are the relative gains from investing in the risky portfolio, compared to the safe portfolio.

For example, in the context of stress tests, the state ω could represent the value of the risky portfolio in a crisis, and the functions u and v could represent the bank and regulator's expected payoffs, which take into account the probability of a crisis, the resulting losses, the payoffs during normal times, etc.

We assume that u and v are continuous and differentiable. We also assume that:

Assumption 1. *u and v are strictly increasing.*

Assumption 2. *For all $\omega \in \Omega$, $u(\omega) > v(\omega)$.*

⁶For example, [Julio and Yook \(2012\)](#) document that high political uncertainty causes firms to reduce investment during election years. [Fernández-Villaverde, Guerrón-Quintana, Kuester, and Rubio-Ramírez \(2015\)](#) document that temporarily high uncertainty about fiscal policy reduces output, consumption, and investment. See also [Pástor and Veronesi \(2013\)](#) and [Baker, Bloom, and Davis \(2016\)](#).

Assumption 3. $v(\underline{\omega}) < 0 < v(\bar{\omega})$

Assumption 1 implies that both the regulator and the bank prefer higher value. Assumption 2 implies the risky portfolio is more valuable to the bank than to the regulator. This assumption captures the idea that the bank does not internalize the social cost associated with risk. Assumption 3 captures the idea that the regulator may not know which portfolio is best for society. The regulator prefers the risky portfolio only if its value during a crisis is sufficiently high. For use below, we define ω_r to be the unique zero of v .

For simplicity, we assume that the bank always prefers to invest in the risky portfolio.⁷

Assumption 4. For all $\omega \in \Omega$, $u(\omega) > 0$.

The regulator and the bank begin with a common prior belief about the state ω , which prior belief is summarized by a continuous distribution function $G(\omega)$ with support Ω . Before the bank selects its portfolio, the regulator and the bank receive private signals about ω . For simplicity, we assume that the bank receives a perfect signal—i.e., it observes ω . The regulator, however, receives a noisy signal $s \in S \subset \mathbb{R}$ that follows a continuous distribution conditional on ω . We let $F(s|\omega)$ and $f(s|\omega)$ denote the cumulative distribution function and density function of the signal s conditional on ω .

Assumption 5 (Monotone Likelihood Ratio Property). *If $\omega' > \omega$, then the ratio $f(s|\omega')/f(s|\omega)$ is strictly increasing in s .*

Assumption 5 implies that $1 - F(s|\omega)$ is strictly increasing in ω .⁸ That is, the regulator is more likely to observe higher signals when the state ω is higher.

After the bank chooses its portfolio, the regulator assesses the value of the bank's portfolio—i.e., the regulator performs a stress test. If the bank

⁷This assumption helps us focus on the main tradeoff in our paper. It is easy to relax this assumption, but relaxing this assumption does not provide any interesting insights.

⁸See Milgrom (1981).

chooses a safe portfolio, we assume the regulator passes the bank. If the bank chooses a risky portfolio, the regulator uses his private information s to decide whether to pass the bank. We assume that the bank passes if the signal s exceeds some threshold; we consider not only exogenous thresholds but also the case in which the threshold is optimally chosen by the regulator. Passing the test means the regulator leaves the bank's portfolio unchanged, whereas failing the test means the regulator requires the bank to replace the risky portfolio with the safe portfolio. This replacement incurs a cost $c_b > 0$ to the bank and $c_s \geq 0$ to the regulator. For example, these costs could represent the opportunity cost of delaying investment in the safe portfolio.⁹ We allow that the regulator may be uncertain about the bank's cost of failing the test, and we let $H(c) = P(c_b \leq c)$ be the cumulative distribution function which describes the regulator's beliefs about the distribution of c_b .

The focus of the paper is whether the regulator should reveal or not reveal his private signal s to the bank. The sequence of events is as follows.

1. The regulator publicly commits to either reveal or not reveal his private signal.
2. Nature chooses the state ω . The bank privately observes ω , and the regulator privately observes the signal s .
3. In accordance with his prior commitment in step (1), the regulator either reveals or does not reveal his signal.
4. The bank selects a portfolio: risky or safe.
5. The regulator conducts the test, passing or failing the bank.
6. Payoffs are realized.
 - If the bank invested in the safe portfolio, both the bank and regulator receive 0.

⁹We assume that investment in the risky portfolio is available only before the stress test.

- If the bank invested in the risky portfolio, then if the bank passes the test, the bank receives $u(\omega)$ and the regulator receives $v(\omega)$; if the bank fails the test, the bank receives $-c_b$ and the regulator receives $-c_s$.

4 Exogenous Pass/Fail Rule

We begin our analysis with the case in which the regulator follows an exogenously given threshold rule for passing or failing the bank. So the only decisions are in step (1) for the regulator and step (4) for the bank. In this section, we focus on the special case in which: (i) the regulator has no uncertainty regarding the bank's private cost c_b of failing the test, so the regulator's beliefs are a point mass at a particular $c_b > 0$; and (ii) the social cost c_s of failing the bank is zero. In Section 6, we discuss the case in which $c_s > 0$, which will give us more results.

Let s_e be the exogenously given, publicly known threshold such that if the regulator receives a signal $s \geq s_e$, the bank passes the test. We first characterize the bank's investment decision. Then, we compare the regulator payoffs under the two regimes: revealing the signal and not revealing. Assume that if the bank is indifferent between investing and not investing, the bank invests.

If the regulator reveals his signal s to the bank, the bank invests if and only if it expects to pass the test—that is, the bank invests when $s \geq s_e$, irrespective of ω . This follows because the bank's payoff from investing and passing the test is positive for all states $\omega \in \Omega$, the payoff from not investing is zero, and the payoff from investing and failing the test is negative. So when the regulator reveals s , the bank uses its knowledge of the test to act in a way that improves its payoff, regardless of the impact on society; i.e., the bank games the test.

If, instead, the regulator does not reveal his signal, the bank's action de-

depends only on the bank's private information, the state ω . Conditional on ω , the bank's expected payoff from investing is $[1 - F(s_e|\omega)]u(\omega) - F(s_e|\omega)c_b$. In particular, with probability $1 - F(s_e|\omega)$, the bank passes the test and obtains $u(\omega)$, and with probability $F(s_e|\omega)$, the bank fails the test and suffers a cost c_b . If the bank does not invest, its payoff is zero. Hence, the bank invests in state ω if and only if

$$[1 - F(s_e|\omega)]u(\omega) - F(s_e|\omega)c_b \geq 0. \quad (1)$$

Next, we show that the bank follows a threshold investment policy: it invests if and only if the state is sufficiently high. Specifically, if the left-hand side of (1) is negative for all $\omega \in \Omega$, the bank never invests. Otherwise, denote the lowest¹⁰ $\omega \in \Omega$ that satisfies (1) by $\omega_b(s_e)$. Because the left-hand side is strictly increasing in ω , the bank invests if and only if $\omega \geq \omega_b(s_e)$. We refer to $\omega_b(s_e)$ as the bank's *cautious threshold*.

The cautious threshold can be found as follows. If the left-hand side is positive at $\underline{\omega}$, then $\omega_b(s_e) = \underline{\omega}$; otherwise, $\omega_b(s_e)$ is the unique zero of the left-hand side, and by the implicit function theorem must be strictly increasing in s_e . Intuitively, when the threshold for passing the test is higher, the bank is less likely to invest because it is more afraid of failing the test.

The next lemma summarizes the results above:

Lemma 1. *1. When the regulator does not reveal his signal to the bank, the bank invests if and only if $\omega \geq \omega_b(s_e)$.*

2. $\omega_b(s_e)$ is increasing in s_e .

Given the bank's equilibrium strategy, we compare the regulator's payoff under both regimes. If the regulator reveals his signal s to the bank, the bank invests if and only if $s \geq s_e$. So in state ω , the bank invests with probability $1 - F(s_e|\omega)$. Taking the expectation across all states gives the regulator's

¹⁰This exists because the left-hand side is continuous.

expected payoff under the revealing regime:

$$V_r(s_e) \equiv \int_{\omega \geq \underline{\omega}} [1 - F(s_e|\omega)]v(\omega)dG(\omega). \quad (2)$$

If the regulator does not reveal his signal, the bank invests if $\omega \geq \omega_b(s_e)$, and if the bank invests in state ω , the bank passes the test with probability $1 - F(s_e|\omega)$. Hence, the regulator's expected payoff is

$$V_n(s_e) \equiv \int_{\omega \geq \omega_b(s_e)} [1 - F(s_e|\omega)]v(\omega)dG(\omega). \quad (3)$$

Equations (2) and (3) show the effect of revealing the regulator's signal s : the bank invests for more states ω . That is, when the regulator reveals his signal, the bank invests for all $\omega \in \Omega$, but when the regulator does not reveal his signal, the bank invests only if $\omega \geq \omega_b(s_e)$.

It is optimal for the regulator to reveal his signal if $V_r(s_e) \geq V_n(s_e)$. Rearranging terms, we obtain that it is optimal to reveal if and only if

$$\int_{\underline{\omega}}^{\omega_b(s_e)} [1 - F(s_e|\omega)]v(\omega)dG(\omega) \geq 0. \quad (4)$$

Hence, whether it is optimal to reveal depends on whether the additional investment in states $[\underline{\omega}, \omega_b(s_e)]$ is socially beneficial. As we explain below, the net effect from this additional investment on the regulator's expected payoff can be either positive or negative.

Specifically, if $\omega_b(s_e) \leq \omega_r$, the left-hand side of (4) is negative, capturing the idea that revealing the signal causes the bank to invest in socially undesirable projects. On the other hand, if $\omega_b(s_e) > \omega_r$, the left-hand side can be written as the sum of two terms:

$$\int_{\underline{\omega}}^{\omega_r} [1 - F(s_e|\omega)]v(\omega)dG(\omega) \quad (5)$$

and

$$\int_{\omega_r}^{\omega_b(s_e)} [1 - F(s_e|\omega)]v(\omega)dG(\omega). \quad (6)$$

Expression (5) is negative and represents the cost of revealing the signal, as mentioned above. We refer to this as the overinvestment effect of revealing the signal. Expression (6) is positive and represents a benefit of revealing the signal, which is that the bank invests in more states in which it is socially desirable to do so. Specifically, if the regulator does not reveal the signal, the bank does not invest in states $\omega \in [\omega_r, \omega_b(s_e))$, which would have given a positive social payoff $v(\omega)$. Revealing the signal avoids this underinvestment effect.

The discussion above suggests that it is optimal to reveal only if the underinvestment effect (6) of not revealing the signal is sufficiently high or the overinvestment effect (5) of revealing the signal is sufficiently low. The next proposition formalizes this intuition.

Proposition 1. *Given a passing threshold s_e such that $V_r(s_e) > 0$, there exists $\bar{\omega}_I \in (\omega_r, \bar{\omega})$ such that the regulator prefers to:*

$$\begin{cases} \text{not reveal} & \text{if } \omega_b(s_e) \in (\underline{\omega}, \bar{\omega}_I) \\ \text{reveal} & \text{if } \omega_b(s_e) > \bar{\omega}_I \\ \text{either} & \text{if } \omega_b(s_e) \in \{\underline{\omega}, \bar{\omega}_I\}. \end{cases}$$

The Proposition shows that whether it is optimal for the regulator to reveal his signal depends on the bank's cautious threshold ω_b . When ω_b is sufficiently high, it is optimal to reveal because not revealing induces the bank to invest too little in socially desirable projects. In contrast, if the cautious threshold ω_b is sufficiently low, but still above $\underline{\omega}$, it is optimal to not reveal because then the bank invests less in socially undesirable projects.

Using Proposition 1, we can derive comparative statics as to how the regulator's optimal disclosure policy changes when model parameters change. For

example, consider ω_b . Observe that the cautious threshold ω_b is increasing in the bank's cost c_b of failing the test. Hence, we have the following.

Corollary 1. *Given a passing threshold s_e such that $V_r(s_e) > 0$, there exists $\bar{c}_b \in (0, \infty)$ such that the regulator prefers to:*

$$\begin{cases} \text{not reveal} & \text{if } c_b \in (0, \bar{c}_b) \\ \text{reveal} & \text{if } c_b > \bar{c}_b \\ \text{either} & \text{if } c_b \in \{0, \bar{c}_b\}. \end{cases}$$

In general, the bank's cautious threshold ω_b , which reflects the bank's reluctance to invest when the regulator's signal is hidden, depends not only on the bank's cost c_b of failing the test, but also on the bank's utility function u . Particular functional forms for u could include parameters that describe various other features, such as risk aversion or the extent of conflict of interest $u(\cdot) - v(\cdot)$ with the regulator. If such parameters have a monotonic relationship to ω_b , they would produce comparative statics similar to Corollary 1.

5 Optimal Pass/Fail Rule

The message of Proposition 1 is that revelation is optimal when the bank is too cautious about investment otherwise. However, if the regulator can choose and commit to an optimal s_e , then he has two ways to reassure an overly cautious bank: he can reveal the signal s , as in the previous section, and he can also make the test easier to pass by decreasing s_e . It is natural to ask whether the freedom to choose the passing threshold s_e obviates signal revelation altogether. In this section, we show that the answer to that question depends on how certain the regulator is about c_b . Intuitively, the more certain the regulator is about c_b , the more precisely he can calibrate the bank's cautious threshold ω_b by adjusting s_e , making revelation unnecessary.

So in contrast to the previous section, suppose the passing threshold s_e can be optimally chosen by the regulator, and recall that the regulator's beliefs about c_b follow cumulative distribution H .

5.1 Optimal passing thresholds

If the regulator reveals his signal, the bank can perfectly avoid test failure, so its investment behavior is independent of c_b . Hence, the regulator's payoff $V_r(s)$ under revealing is identical to (2) in the previous section.

We can show that $V_r(s)$ is strictly quasiconcave and therefore has a unique maximum $s_r \in S$. That is,

$$s_r = \arg \max_{s \in S} V_r(s).$$

To see why, observe that

$$V'_r(s) = - \int_{\Omega} f(s|\omega)v(\omega)dG(\omega) = -f(s) \int_{\Omega} f(\omega|s)v(\omega)d\omega = -f(s)E[v(\omega)|s],$$

so $V'_r(s)$ has the same sign as $-E[v(\omega)|s]$, which by Assumption 5 is strictly decreasing in s .¹¹

To simplify the exposition, we focus on the case in which $E[v(\omega)|\underline{s}] < 0 < E[v(\omega)|\bar{s}]$. Then s_r is interior and is the unique value which satisfies $E[v(\omega)|s_r] = 0$.¹² This condition says that when the regulator observes s_r , he is indifferent between having the bank invest in all states and having the bank not invest at all.

Now suppose the regulator does not reveal the signal and selects passing threshold s . Then by (1), the bank invests if and only if $c_b \leq [F(s|\omega)^{-1} - 1]u(\omega)$. So, the probability that the bank invests in state ω is

$$k(s, \omega) \equiv H([F(s|\omega)^{-1} - 1]u(\omega)).$$

¹¹See Milgrom (1981).

¹²If $s_r \in \{\underline{s}, \bar{s}\}$, not revealing weakly dominates revealing, for all parameter values.

Hence, the regulator's payoff under not revealing is

$$V_n(s) \equiv \int_{\Omega} [1 - F(s|\omega)] k(s, \omega) v(\omega) dG(\omega). \quad (7)$$

Observe that in the special case in which the regulator knows the bank's type (i.e., H has all of the mass on a particular c_b), $k(s, \omega)$ is a step function, and (7) reduces to (3).

Let s_n denote an optimal passing threshold if the regulator chooses to not reveal. That is,

$$s_n \in \arg \max_{s \in S} V_n(s).$$

Then it is optimal to reveal if and only if $V_r(s_r) \geq V_n(s_n)$.

5.2 Regulator knows c_b

As a benchmark, we start with the case in which the regulator knows the bank's cost of failing the test. The next lemma shows that in this case, it is optimal to reveal the test.

Lemma 2. *If the regulator can optimally choose with commitment the passing thresholds s_r and s_n , the regulator has no uncertainty about c_b , and $c_s = 0$, then for all $c_b \geq 0$ it is weakly optimal to not reveal.*

Intuitively, suppose the regulator used the optimal revealing threshold under both disclosure regimes. If not revealing were worse than revealing, it must be due to underinvestment. But then, the regulator could simply reduce the not-revealing threshold to eliminate underinvestment without inducing overinvestment, and also pass the bank more frequently than under revealing.

5.3 Regulator does not know c_b

Now suppose the regulator does not know the bank's cost of failing the test. In this case, the regulator cannot eliminate underinvestment without increasing

overinvestment, as he did in Lemma 2, because a test that eliminates underinvestment for one type of bank increases overinvestment for other types. When types are very similar to one another, the regulator can still increase his payoff above $V_r(s_r)$ by making the test easier and not revealing it. So, not revealing is optimal. However, when types are very different from one another, the regulator cannot reach the payoff $V_r(s_r)$ by making the test easier. In this case, revealing becomes optimal.

Hence, one might expect that if the regulator is sufficiently uncertain regarding the bank's cost of failure c_b , it is optimal to reveal, and otherwise, it is optimal not to reveal. In the following example, this intuitive prediction is confirmed.

Example 1. Suppose $\Omega = [0, 1]$, $G(\omega) = \omega$ (i.e., uniform), $f(s|\omega) = 2[(1 - s)(1 - \omega) + s\omega]$, $u(\omega) = \omega + 0.5$, and $v(\omega) = \omega - 0.5$. So investment is socially desirable when $\omega > 0.5$ and socially undesirable when $\omega < 0.5$. Assume that the distribution H over c_b is lognormal with parameters $\mu = \ln 2$ and various values of σ . This amounts to fixing the median of H at 2 and changing the variance.

Figure 1 illustrates the density function of c_b for several values of σ . The figure shows that when σ is low, most of the mass is concentrated at the median of the distribution. When σ is high, the distribution puts a high mass on very low types and very high types.

Figure 2 illustrates the regulator's payoffs $V_r(s_r)$ and $V_n(s_n)$, as a function of various degrees of uncertainty σ about c_b . The payoff from revealing $V_r(s_r)$ does not depend on the level of uncertainty, whereas the payoff from not revealing $V_n(s_n)$ is strictly decreasing in the level of uncertainty. For a very low level of uncertainty, not revealing is strictly optimal. For a very high level of uncertainty, revealing is strictly optimal.

To obtain more intuition for the result in Example 1, observe that from Equation (7), the regulator's payoff $V_n(s)$ depends on the product of two functions: the probability $k(s, \omega)$ that the bank invests in state ω and the prob-

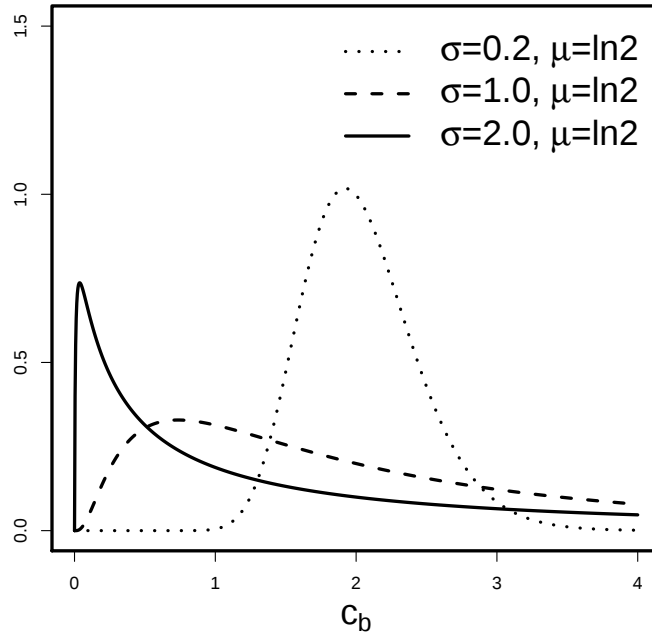


Figure 1: The distribution of c_b is lognormal, with fixed $\mu = \ln 2$, which implies a fixed median of 2. For fixed μ , the variance $[\exp(\sigma^2) - 1] \exp(2\mu + \sigma^2)$ is increasing in σ .

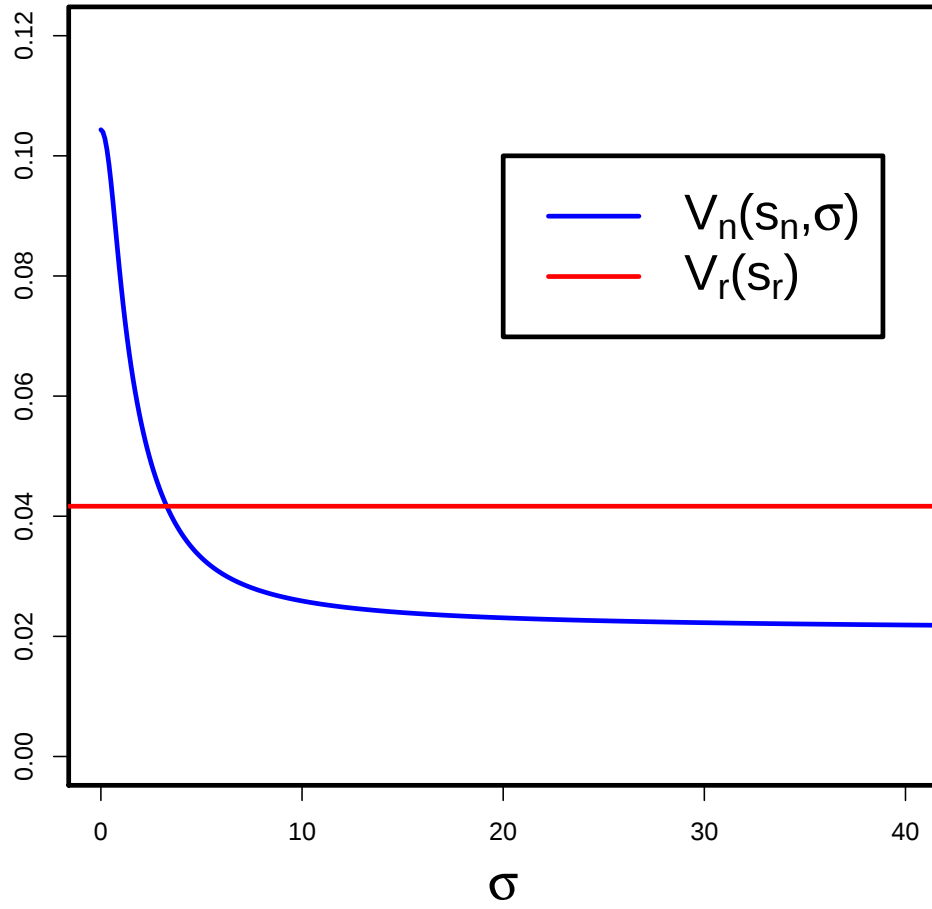


Figure 2: The regulator's payoff $V_n(s_n)$ from not revealing is decreasing in his uncertainty of the bank's type c_b .

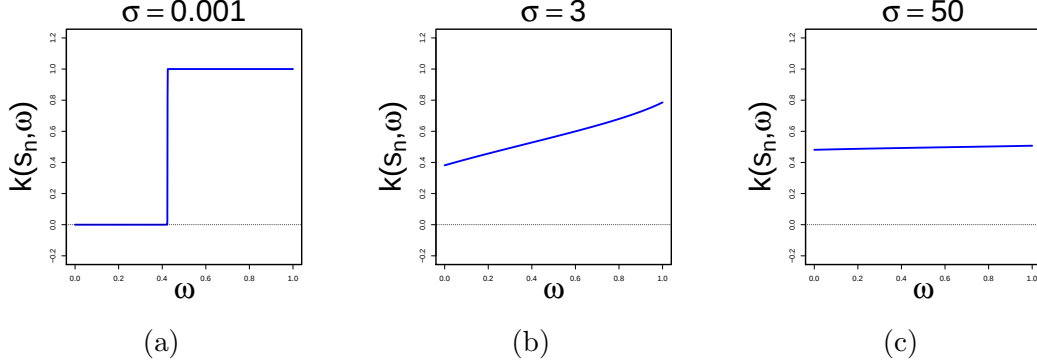


Figure 3: Panels [a](#), [b](#), and [c](#) plot $k(s_n, \omega)$, which is the probability that the bank invests in state ω , for increasing values of σ (i.e., increasing regulator uncertainty about c_b).

ability $1 - F(s|\omega)$ that the regulator passes the bank in state ω . Making the test easier by reducing s shifts both functions upward, which leads to more investment. However, the overall effect is ambiguous because investment increases both in socially desirable states $[\omega_r, \bar{\omega}]$ and socially undesirable states $[\underline{\omega}, \omega_r]$.

Figure [3](#) illustrates the function $k(s, \omega)$ under different levels of uncertainty. To see why not revealing is optimal when uncertainty about c_b is very low, consider panel [3a](#). When uncertainty is close to zero, the regulator can predict with near certainty whether a bank that observes a particular ω will invest or not. This closely approximates the situation in Lemma [2](#). So $k(s_n, \omega)$ is approximately a step function with step at $\omega_b(s_n)$, calculated when $c_b = 2$ (the median of the distribution). In this case, the bank's investment is very sensitive to the bank's private information ω , and so, the regulator can adjust s_e so that the bank almost certainly does not invest in bad states $[\underline{\omega}, \omega_r]$ and almost certainly invests in good states $[\omega_r, \bar{\omega}]$.

To see why revealing is optimal when uncertainty about c_b is very high, consider panel [3c](#). In this case, the function $k(s_n, \omega)$ is nearly horizontal. That is, the probability that the bank invests is very insensitive to the state ω . Intuitively, when a large fraction of types have a very low cost of failure

and a large fraction of types have a very high cost, those with the low cost invest in all states ω , those with high cost invest in no states, and only a very small fraction of types with intermediate cost favor higher states. Therefore, by adjusting the passing threshold, the regulator doesn't change the relative likelihood of investment in good states ω versus bad states ω ; instead he simply scales the probability of investment, which is roughly constant across ω . For example, by lowering the passing threshold s , the regulator can induce some additional types with low cost to invest, but he cannot induce them to invest only when it is socially optimal to invest. So, instead of essentially avoiding investment in all states, these types invest in almost every state, which means that when the regulator reduces underinvestment, he increases overinvestment. In this case, if the regulator wanted to scale up the investment level, he would need to reduce the passing threshold significantly, ending up passing the bank in almost every state. But then revealing the regulator's signal is optimal because the regulator sets a high threshold for passing the test and uses its information to eliminate some socially undesirable investment.

6 Social Cost of Failing the Bank

In this section, we examine whether it is optimal to reveal or not reveal the regulator's signal s when the social cost c_s of failing the bank is strictly positive. We illustrate our result for the case in which the passing threshold s_e is exogenous and the regulator knows the bank's cost c_b of failure.

If the regulator reveals his signal, the bank invests only if it observes a passing signal $s \geq s_e$. Consequently, under revealing, the regulator never fails the bank, so regulator's payoff $V_r(s_e)$ under revealing is independent of c_s and identical to (2).

If the regulator does not reveal his signal, the bank invests if $\omega \geq \omega_b(s_e)$, and if the bank invests in state ω , the bank passes the test with probability $1 - F(s_e|\omega)$. Hence, the regulator's expected payoff is

$$V_n(s_e) \equiv \int_{\omega \geq \omega_b(s_e)} [1 - F(s_e|\omega)]v(\omega)dG(\omega) - c_s \int_{\omega \geq \omega_b(s_e)} F(s_e|\omega)dG(\omega). \quad (8)$$

The first term represents the case in which the bank invests and the regulator passes the bank. The second term represents the case in which the bank invests and fails the test.

It is optimal for the regulator to reveal his signal if $V_r(s_e) \geq V_n(s_e)$. Rearranging terms, we obtain that it is optimal to reveal if and only if

$$\int_{\underline{\omega}}^{\omega_b(s_e)} [1 - F(s_e|\omega)]v(\omega)dG(\omega) + c_s \int_{\omega \geq \omega_b(s_e)} F(s_e|\omega)dG(\omega) \geq 0. \quad (9)$$

Comparing (9) to (4), revealing the signal now has two effects on the regulator's payoff: not only does it cause the bank to invest in more states $[\underline{\omega}, \omega_b(s_e)]$, but it also avoids the social cost c_s of failing the bank. This suggests that if the social cost c_s of failing the bank is sufficiently high, it is optimal to reveal the signal. Otherwise, it is optimal to reveal only if the underinvestment effect (6) of not revealing the signal is sufficiently high or the overinvestment effect (5) of revealing the signal is sufficiently low. The next proposition formalizes this intuition.

Proposition 2. *Given a passing threshold s_e and regulator signal distribution $F(s|\omega)$, such that $V_r(s_e) > 0$, there exists a social cost of failure $\bar{c}_s > 0$ such that:*

1. *If $c_s > \bar{c}_s$, revealing is strictly preferred to not revealing.*
2. *If $c_s \leq \bar{c}_s$, then there exist $\underline{\omega}_I, \bar{\omega}_I \in \Omega$, with $\underline{\omega}_I \leq \bar{\omega}_I$ (with strict inequality if $c_s < \bar{c}_s$), such that:*
 - (a) *If $\omega_b(s_e) \in (\underline{\omega}_I, \bar{\omega}_I)$, not revealing is strictly preferred to revealing.*
 - (b) *If $\omega_b(s_e) < \underline{\omega}_I$ or $\omega_b(s_e) > \bar{\omega}_I$, revealing is strictly preferred to not revealing.*

(c) If $\omega_b(s_e) = \underline{\omega}_I$ or $\omega_b(s_e) = \bar{\omega}_I$, the regulator is indifferent between revealing and not revealing.

3. As c_s decreases from $\bar{c}_s$ to 0, $\underline{\omega}_I$ strictly decreases to $\underline{\omega}$ and $\bar{\omega}_I$ strictly increases to a value strictly less than $\bar{\omega}$.

Part 1 captures the idea that when the social cost of failure is high, the regulator would like to prevent failure by revealing the signal. In part 2, the social cost of failure is low, and whether revealing is optimal depends not only on the cost of failure but also on how not revealing affects the bank's investment decision. There are two circumstances in which it is optimal to reveal. The first case is when not revealing the signal does very little to prevent investment in bad projects—i.e., $\omega_b(s_e)$ is very low. In that case, revealing the signal induces only slightly worse investment behavior but avoids the cost of failure. The second case is when not revealing the signal deters not only bad investment but also much good investment—i.e., $\omega_b(s_e)$ is very high. In that case, revealing the signal permits this good investment and also avoids the cost of test failure.

Because $\omega_b(s_e)$ is increasing in c_b , we obtain the following immediate corollary, which is illustrated in Figure 4.

Corollary 2. *Given passing threshold s_e , if $c_s > \bar{c}_s$, revealing is optimal for all $c_b \geq 0$. If $c_s \leq \bar{c}_s$, there exist $\underline{c}_b, \bar{c}_b \in \mathbb{R}_+$, with $\underline{c}_b \leq \bar{c}_b$ (strict inequality if $c_s < \bar{c}_s$) such that the regulator's optimal policy is to*

$$\begin{cases} \text{Reveal} & \text{if } c_b \in [0, \underline{c}_b) \cup (\bar{c}_b, \infty] \\ \text{Not reveal} & \text{if } c_b \in (\underline{c}_b, \bar{c}_b) \\ \text{Either} & \text{if } c_b \in \{\underline{c}_b, \bar{c}_b\}. \end{cases}$$

Furthermore, $\underline{c}_b$ and $\bar{c}_b$ are respectively strictly increasing and strictly decreasing functions of c_s .¹³

¹³The appendix works out the case in which the passing threshold may be optimally

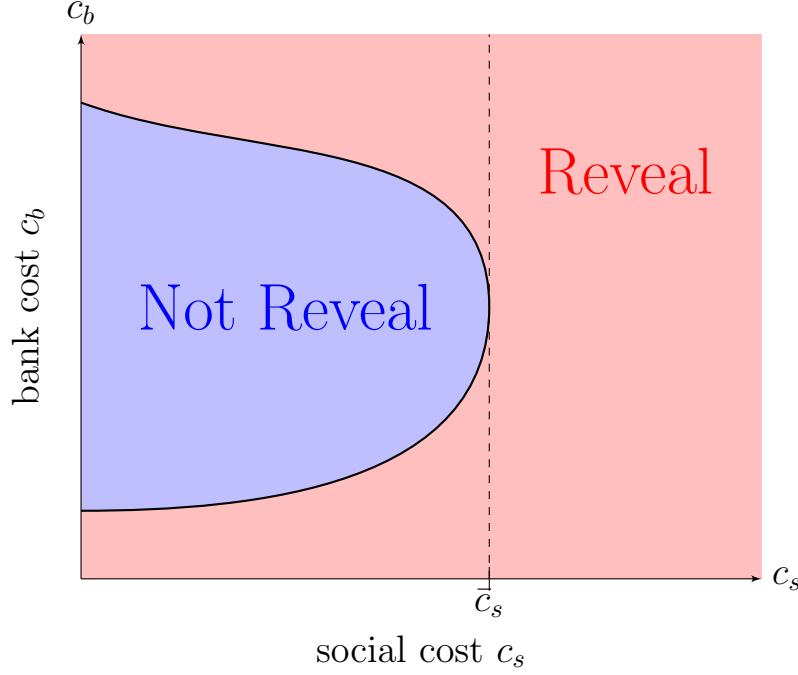


Figure 4: The regulator's optimal policy for the bank cost c_b and social cost c_s of stress test failure. For social cost $c_s < \bar{c}_s$, the optimal policy is nonmonotonic in the bank's private cost c_b of test failure.

7 Policy Implications

In this section, we show how the insights of our model can shed light on current and suggested stress testing policies.

1. In practice, banks whose capital plans would lead their capital to fall below the required level are given a short time to adjust their plans.¹⁴ In our model, this practice could imply a lower private cost for banks from failing the test.¹⁵ Our model suggests that if the social cost of

chosen and $c_s > 0$, but we omit it in the text because the intuition is similar to Lemma 2. In that case, for any positive c_s , not revealing is optimal for high enough c_b , because the optimal cautious threshold ω_b may be induced with a low passing threshold and therefore a low probability of costly test failure.

¹⁴See <https://www.gpo.gov/fdsys/pkg/FR-2017-12-15/pdf/2017-26856.pdf>

¹⁵Earlier, we thought of c_b as the opportunity cost of delaying investment in the safe portfolio, but c_b could also represent other costs, such as the embarrassment involved with

failing a bank is zero, a lower c_b implies the regulator should disclose less. However, if the social cost of a failing a bank is positive, the impact of a lower c_b on optimal disclosure is more nuanced. In particular, if c_b descends from a high value to a middle value, the regulator should disclose less. But if c_b descends from a middle value to a low value, the regulator should disclose more.

2. A widely expressed concern is that disclosing the Fed’s models could increase correlations in asset holdings among banks subject to the stress tests (i.e., the largest banks), making the financial system more vulnerable to adverse financial shocks. A recent proposal suggests that the Fed reveals the outcome of applying its models to hypothetical loan portfolios.¹⁶ An extension of our model would suggest that such enhanced disclosure could also increase correlations in asset holdings. The proposed hypothetical portfolios could serve as benchmark portfolios in which too many banks invest, leading to correlated investment. So just as in our basic model, in which the bank could underinvest in a socially valuable risky portfolio by choosing the safe portfolio for which the test results are predictable, in practice, banks could underinvest in their idiosyncratic risky portfolios, for which the test results are unpredictable, and overinvest in the benchmark risky portfolio, for which the test results are predictable.
3. In his departing speech, Fed Governor Daniel Tarullo suggested that if the Fed model were to be published, the minimum capital requirements (represented by the passing threshold in our model) would need to materially increase.¹⁷ This suggestion reflects the view that revealing the model would lead to more gaming, which would need to be counteracted by a more stringent test. Our model suggests that this conclusion is

the public objections to a bank’s capital plan.

¹⁶See <https://www.gpo.gov/fdsys/pkg/FR-2017-12-15/pdf/2017-26856.pdf>

¹⁷See <https://www.federalreserve.gov/newsevents/speech/tarullo20170404a.htm>

only partially correct. In particular, as we show in Proposition 4 in the appendix, for some parameter values, revealing is strictly preferred to not revealing, yet, the optimal passing threshold under revealing is lower than the optimal passing threshold under not revealing. Intuitively, this happens when the social cost of failure is high enough, so that revealing is optimal, but the banks private cost of failing the test is low, so that deterring overinvestment under not revealing requires a very high threshold: $s_n > s_r$.

8 Conclusion

We present conditions under which it is socially optimal for a regulator to reveal his stress testing model to the bank. The framework we present allows that banks may game a publicly known model, the chief concern underlying the Federal Reserve’s policy of model secrecy. We show that despite the possibility of gaming, revealing the model may still be optimal, because uncertainty about the regulator’s model may prevent banks from investing in socially valuable assets. In addition, even when the regulator can reassure cautious banks by relaxing the minimum standard to which they are held, revealing may be optimal if the regulator is sufficiently uncertain about bank characteristics. Finally, we show that if causing the bank to fail a stress test is socially costly, the optimal disclosure policy may be nonmonotonic in bank characteristics; that is, revealing may be optimal for banks that have very high or very low private cost of failure.

We are currently working on understanding optimal partial disclosure rules. For example, the regulator may commit to identifying an interval in which his private signal falls. Alternately, the regulator may commit to a stochastic map from his privately observed signal to an arbitrary message space, in the manner of Bayesian persuasion à la [Kamenica and Gentzkow \(2011\)](#). Relatedly, if particular asset classes are more relevant to the stress test model than others,

the regulator may make distinct disclosure decisions regarding different classes of assets.

References

- Amador, Manuel and Kyle Bagwell (2013), “The theory of optimal delegation with an application to tariff caps.” *Econometrica*, 81, 1541–1599.
- Amador, Manuel and Kyle Bagwell (2016), “Money burning in the theory of delegation.” Working paper.
- Amador, Manuel, Kyle Bagwell, and Alex Frankel (2017), “A note on interval delegation.” Working paper.
- Baker, Scott R, Nicholas Bloom, and Steven J Davis (2016), “Measuring economic policy uncertainty.” *The Quarterly Journal of Economics*, 131, 1593–1636.
- Bouvard, Matthieu, Pierre Chaigneau, and Adolfo de Motta (2015), “Transparency in the financial system: Rollover risk and crises.” *The Journal of Finance*, 70, 1805–1837.
- Chakraborty, Archishman and Bilge Yilmaz (2017), “Authority, consensus, and governance.” *The Review of Financial Studies*, 30, 4267–4316.
- Dessein, Wouter (2002), “Authority and communication in organizations.” *The Review of Economic Studies*, 69, 811–838.
- Faria-e Castro, Miguel, Joseba Martinez, and Thomas Philippon (2016), “Runs versus lemons: information disclosure and fiscal capacity.” *The Review of Economic Studies*, rdw060.
- Fernández-Villaverde, Jesús, Pablo Guerrón-Quintana, Keith Kuester, and Juan Rubio-Ramírez (2015), “Fiscal volatility shocks and economic activity.” *The American Economic Review*, 105, 3352–3384.
- Freixas, X. (2000), “Optimal Bail Out Policy, Conditionality and Constructive Ambiguity.” DNB Staff Reports (discontinued) 49, Netherlands Central Bank, URL <https://ideas.repec.org/p/dnb/staffs/49.html>.

- Gick, Wolfgang and Thilo Pausch (2012), “Persuasion by stress testing.” Working paper.
- Gissler, Stefan, Jeremy Oldfather, and Doriana Ruffino (2016), “Lending on hold: Regulatory uncertainty and bank lending standards.” *Journal of Monetary Economics*, 81, 89–101.
- Goldstein, Itay and Yaron Leitner (2017), “Stress tests and information disclosure.” FRB of Philadelphia Working Paper No. 17-28.
- Goldstein, Itay and Haresh Sapra (2014), “Should banks’ stress test results be disclosed? an analysis of the costs and benefits.” *Foundations and Trends® in Finance*, 8, 1–54.
- Grenadier, Steven R, Andrey Malenko, and Nadya Malenko (2016), “Timing decisions in organizations: Communication and authority in a dynamic environment.” *The American Economic Review*, 106, 2552–2581.
- Halac, Marina and Pierre Yared (2016), “Commitment vs. flexibility with costly verification.” Technical report, National Bureau of Economic Research.
- Harris, Milton and Artur Raviv (2005), “Allocation of decision-making authority.” *Review of Finance*, 9, 353–383.
- Harris, Milton and Artur Raviv (2006), “A theory of board control and size.” *The Review of Financial Studies*, 21, 1797–1832.
- Holmstrom, Bengt (1984), “On the theory of delegation.” *Bayesian Models in Economic Theory. Amsterdam: North Holland*.
- Inostroza, Nicolas and Alessandro Pavan (2017), “Persuasion in global games with application to stress testing.” Working paper.
- Julio, Brandon and Youngsuk Yook (2012), “Political uncertainty and corporate investment cycles.” *The Journal of Finance*, 67, 45–83.

- Kamenica, Emir and Matthew Gentzkow (2011), “Bayesian persuasion.” *American Economic Review*, 101, 2590–2615.
- Lazear, Edward P (2006), “Speeding, terrorism, and teaching to the test.” *The Quarterly Journal of Economics*, 121, 1029–1061.
- Levit, Doron (2016), “When words speak louder than actions.” Working paper.
- Milgrom, Paul R (1981), “Good news and bad news: Representation theorems and applications.” *The Bell Journal of Economics*, 380–391.
- Orlov, Dmitry, Pavel Zryumov, and Andrzej Skrzypacz (2017), “Design of macro-prudential stress tests.” Working paper.
- Pástor, L’uboš and Pietro Veronesi (2013), “Political uncertainty and risk premia.” *Journal of Financial Economics*, 110, 520–545.
- Williams, Basil (2015), “Stress tests and bank portfolio choice.” Technical report, Mimeo, New York University.

Appendix

8.1 Optimal Rule, Positive c_s

The regulator will choose to reveal if $V_r(s_r) < V_n(s_n)$. We show this is true for sufficiently low bank cost of failure c_b .

Proposition 3. *If the regulator can optimally choose with commitment the passing thresholds s_n and s_r , then for all $c_s \geq 0$, there exists a $\underline{c}_b > 0$ such that the regulator's optimal policy is to*

$$\begin{cases} \text{Reveal} & \text{if } c_b < \underline{c}_b \text{ and } c_s > 0 \\ \text{Not reveal} & \text{if } c_b > \underline{c}_b \\ \text{Either} & \text{otherwise.} \end{cases}$$

Furthermore, as c_s increases to infinity, $\underline{c}_b$ strictly increases to infinity.

Intuitively, for any positive c_s , not revealing is optimal for high enough c_b , because the optimal cautious threshold ω_b may be induced with a low passing threshold and therefore a low probability of costly test failure.

Although it may seem that model opacity and high capital thresholds are policy substitutes for deterring overinvestment, the following proposition shows that when it is optimal to disclose the regulator's private signal, there are cases in which the optimal threshold under revealing is *lower* than under not revealing.

Proposition 4. *There exists a $c_s > 0$ and $c_b > 0$ such that revealing is strictly optimal and $s_r < s_n$.*

If the regulator does not reveal the model and the bank's cost of test failure is low, the only way to deter bad investments is to set the capital requirement very high, so high that the regulator fails the bank even when the regulator's private signal indicates the asset has positive social value. (Mathematically, $E[v(\omega)|s_n] > 0$). The optimal not revealing threshold s_n balances the marginal

cost of failing good banks with the marginal benefit of deterring bad investment. When c_b is very low, this marginal benefit is very small compared to the marginal cost, so at the optimal threshold s_n , only very little bad investment is deterred.

Under revealing, if c_s is sufficiently high, the optimal threshold s_r is lower than s_n , and the regulator is better off. Why? Under revealing, setting a very high threshold does nothing to deter bad investment, because the bank invests in all states whenever it expects to pass, so the regulator optimally sets the threshold so that all banks with weakly positive expected value pass (that is, $E[v(\omega)|s_r] = 0$). Therefore, the threshold is lower under revealing. The regulator is also better off because the avoided social costs of failure c_s dominate the very small additional bad investment due to gaming.

8.2 Proofs

Proof of Proposition 1. Let $J(\omega_b) = \int_{\underline{\omega}}^{\omega_b} [1 - F(s_e|\omega)]v(\omega)dG(\omega)$. By (4), revealing is strictly preferred if $J(\omega_b) > 0$, strictly not preferred if $J(\omega_b) < 0$, and the regulator is indifferent if $J(\omega_b) = 0$. Observe that $J'(\omega_b) = [1 - F(s_e|\omega_b)]v(\omega_b)dG(\omega)$, which has the same sign as $v(\omega_b)$. So by Assumptions 1 and 3, as ω_b increases from $\underline{\omega}$ to ω_r to $\bar{\omega}$, J strictly decreases from $J(\underline{\omega}) = 0$ to $J(\omega_r) < 0$ and then strictly increases to $J(\bar{\omega}) = V_r(s_e) > 0$. Let ω_I be the unique $\omega_b \in (\omega_r, \bar{\omega})$ for which $J(\omega_b) = 0$, and the proposition follows. $\square$

Proof of Corollary 1. Note that $V_r(s_e) > 0$ implies $s_e < \bar{s}$. So given $\omega \in \Omega$, $1 - F(s_e|\omega) \in [0, 1)$, and because $c_b/(u(\omega) + c_b)$ is strictly increasing in c_b , there exists a unique $c_b(\omega) \in \mathbb{R}_+$ satisfying $1 - F(s_e|\omega) = c_b/(u(\omega) + c_b)$. Furthermore, $c_b(\omega)$ is strictly increasing in ω . Let $\bar{c}_b \equiv c_b(\bar{\omega}_I)$, and apply Proposition 1. $\square$

Proof of Lemma 2. If $V_n(s_r) \geq V_r(s_r)$, then $V_n(s_n) \geq V_n(s_r) \geq V_r(s_r)$, so not revealing is weakly optimal. On the other hand, if $V_n(s_r) < V_r(s_r)$, then $\int_{\omega_b(s_r)} [1 - F(s_r|\omega)]v(\omega)dG(\omega) < \int_{\underline{\omega}} [1 - F(s_r|\omega)]v(\omega)dG(\omega)$, which im-

plies that $0 < \int_{\underline{\omega}}^{\omega_b(s_r)} [1 - F(s_r|\omega)]v(\omega)dG(\omega)$, so $\omega_b(s_r) > \omega_r$. If so, there exists $\hat{s} \in S$ such that $\hat{s} < s_r$ and $\omega_b(\hat{s}) = \omega_r$. Then $V_n(s_n) \geq V_n(\hat{s}) = \int_{\omega_b(\hat{s})} [1 - F(\hat{s}|\omega)]v(\omega)dG(\omega) = \int_{\omega_r} [1 - F(\hat{s}|\omega)]v(\omega)dG(\omega) > \int_{\omega_b(s_r)} [1 - F(s_r|\omega)]v(\omega)dG(\omega) = V_r(s_r)$. $\square$

Proof of Proposition 2. For fixed s_e and $F(s|\omega)$, denote

$$J(c_s, \omega_b) = c_s \int_{\omega \geq \omega_b} F(s_e|\omega)dG(\omega) + \int_{\underline{\omega}}^{\omega_b} [1 - F(s_e|\omega)]v(\omega)dG(\omega).$$

By Equation (9), revealing is strictly preferred if $J(c_s, \omega_b(s_e)) > 0$, not revealing is strictly preferred if $J(c_s, \omega_b(s_e)) < 0$, and the regulator is indifferent if $J(c_s, \omega_b(s_e)) = 0$.

Note that

$$\frac{\partial J(c_s, \omega_b)}{\partial \omega_b} = \left([1 - F(s_e|\omega_b)]v(\omega_b) - c_s F(s_e|\omega_b) \right) dG(\omega_b),$$

which has the same sign as

$$[1 - F(s_e|\omega_b)]v(\omega_b) - c_s F(s_e|\omega_b), \quad (10)$$

(Part 1): We first show that $J(c_s, \cdot)$ is strictly quasiconvex and, therefore, has a unique minimizer in $[\underline{\omega}, \bar{\omega}]$. If $c_s = 0$, then by Assumption 1, (10) is strictly decreasing when $\omega < \omega_r$ and strictly increasing when $\omega > \omega_r$. (By Assumption 5, $F(s_e|\omega) > 0$ for every $\omega > \underline{\omega}$.) If $c_s > 0$, then if $\omega_b \leq \omega_r$, (10) is strictly negative, and if $\omega_b > \omega_r$, (10) is strictly increasing in ω_b , crossing zero at most once. So $J(c_s, \cdot)$ is strictly quasiconvex. We denote the unique minimizer in $[\underline{\omega}, \bar{\omega}]$, as $\omega_m(c_s)$. If (10) is negative for all $\omega \in \Omega$, then $\omega_m(c_s) = \bar{\omega}$. Otherwise, $\omega_m(c_s)$ is the unique zero of (10), and so, $\omega_m(c_s) \geq \omega_r$.

Next, we show that the minimum $J(c_s, \omega_m(c_s))$ is increasing in c_s , and that there is a unique $c_s > 0$ such that $J(c_s, \omega_m(c_s)) = 0$. By the envelope theorem,

$$\frac{dJ(c_s, \omega_m(c_s))}{dc_s} = \frac{\partial J(c_s, \omega_m(c_s))}{\partial c_s} = \int_{\omega \geq \omega_m(c_s)} F(s_e|\omega)dG(\omega) \geq 0,$$

and the inequality is strict whenever $\omega_m(c_s) < \bar{\omega}$

Consider the value of $J(c_s, \omega_m(c_s))$ at the extreme point $c_s = 0$. If $c_s = 0$, (10) reduces to $[1 - F(s_e|\omega)]v(\omega)$ and therefore $\omega_m(0) = \omega_r$, so $J(0, \omega_m(0)) = J(0, \omega_r) < 0$. Now consider the value of $J(c_s, \omega_m(c_s))$ as $c_s \rightarrow \infty$. For $\omega_m(c_s) \in [\omega_r, \bar{\omega}]$, applying the implicit function theorem to (10) gives

$$\omega'_m(c_s) = \frac{F(s_e|\omega)}{[1 - F(s_e|\omega)]v'(\omega) + (c_s + v(\omega))(\partial[1 - F(s_e|\omega)]/\partial\omega)} \Big|_{\omega=\omega_m(c_s)} \geq 0,$$

so $\omega_b(c_s)$ is weakly increasing in c_s , and must therefore converge to some limit $L \in [\underline{\omega}, \bar{\omega}]$ as $c_s \rightarrow \infty$. There are two cases. Case 1: $L = \bar{\omega}$. Then since $V_r(s_e) > 0$, it must be that $\lim_{c_s \rightarrow \infty} J(c_s, \omega_m(c_s)) \geq V_r(s_e) > 0$. Case 2: $L < \bar{\omega}$. Then $\lim_{c_s \rightarrow \infty} \frac{dJ(c_s, \omega_m(c_s))}{dc_s} = \int_{\omega \geq L} F(s_e|\omega) dG(\omega) > 0$, so $\lim_{c_s \rightarrow \infty} J(c_s, \omega_m(c_s)) = \infty$.¹⁸ In either case, $J(c_s, \omega_m(c_s)) > 0$ for high enough c_s .

By the intermediate value theorem, there exists a $c_s > 0$ for which $J(c_s, \omega_m(c_s)) = 0$. To show uniqueness, suppose there exist two zeros c'_s and c''_s , with $c'_s < c''_s$. That $J(c_s, \omega_m(c_s))$ is weakly increasing gives $dJ(c_s, \omega_m(c_s))/dc_s = 0$ for all $c_s \in [c'_s, c''_s]$, which implies $\omega_m(c_s) = \bar{\omega}$, and therefore $J(c_s, \omega_m(c_s)) = J(c_s, \bar{\omega}) = V_r(s_e) > 0$, a contradiction.

Denote by $\bar{c}_s$ the unique zero of $J(c_s, \omega_m(c_s))$. We have that $c_s > \bar{c}_s$ if and only if $J(c_s, \omega_m(c_s)) > 0$. So $c_s > \bar{c}_s$ if and only if for all $\omega_b(s_e) \in \Omega$, $J(c_s, \omega_b(s_e)) \geq \min_{\omega \in \Omega} J(c_s, \omega) = J(c_s, \omega_m(c_s)) > 0$, and Part 1 is proved.

(Part 2): If $c_s < \bar{c}_s$, then by the proof of Part 1, $J(c_s, \omega_m(c_s)) < 0$. As ω increases from $\underline{\omega}$ to $\bar{\omega}$, the strict quasiconvexity of $J(c_s, \cdot)$ implies that $J(c_s, \cdot)$ strictly decreases from $J(c_s, \underline{\omega}) \geq 0$ to $J(c_s, \omega_m(c_s)) < 0$ and then strictly increases to $J(c_s, \bar{\omega}) > 0$. So there exist exactly two zeros of $J(c_s, \cdot)$ in Ω : a

¹⁸If $\lim_{x \rightarrow \infty} h'(x) > 0$, then $\lim_{x \rightarrow \infty} h(x) = \infty$.

unique $\underline{\omega}_I \in [\underline{\omega}, \omega_m(c_s))$ and a unique $\bar{\omega}_I \in (\omega_m(c_s), \bar{\omega})$ which satisfy

$$J(c_s, \omega_b(s_e)) \begin{cases} < 0 & \omega_b(s_e) \in (\underline{\omega}_I, \bar{\omega}_I) \\ > 0 & \omega_b(s_e) \in [\underline{\omega}, \underline{\omega}_I) \cup (\bar{\omega}_I, \bar{\omega}] \\ = 0 & \omega_b(s_e) \in \{\underline{\omega}_I, \bar{\omega}_I\}. \end{cases}$$

Finally, if $c_s = \bar{c}_s$, then $J(c_s, \omega_m(c_s)) = 0$. The value $\omega_m(c_s)$ is the unique minimizer of $J(c_s, \cdot)$, so $\omega_m(c_s)$ is the only zero of $J(c_s, \cdot)$, and for all $\omega_b(s_e) \neq \omega_m(c_s)$, $J(c_s, \omega_b(s_e)) > 0$.

(Part 3): From above, $\underline{\omega}_I$ and $\bar{\omega}_I$ satisfy $0 = J(c_s, \underline{\omega}_I(c_s)) = J(c_s, \bar{\omega}_I(c_s))$. Applying the implicit function theorem gives $\underline{\omega}'_I(c_s) = -\frac{\partial J}{\partial c_s} / \frac{\partial J}{\partial \omega_b} \Big|_{(c_s, \underline{\omega}_I(c_s))}$. Because $\omega_m(c_s)$ is the unique zero of (10) and $\underline{\omega}_I(c_s) < \omega_m(c_s)$, it must be that $\frac{\partial J}{\partial c_s} < 0$ at $(c_s, \underline{\omega}_I(c_s))$. Furthermore, $\frac{\partial J}{\partial \omega_b}$ is strictly positive unless $\underline{\omega}_I = \underline{\omega}$, which occurs only if $c_s = 0$. So $\underline{\omega}'_I(c_s) < 0$ for all $c_s \in (0, \bar{c}_s)$. Similarly, $\bar{\omega}'_I(c_s) = -\frac{\partial J}{\partial c_s} / \frac{\partial J}{\partial \omega_b} \Big|_{(c_s, \bar{\omega}_I(c_s))}$. Because $\bar{\omega}_I(c_s) > \omega_m(c_s) \geq \omega_r > \underline{\omega}$, it must be that $\frac{\partial J}{\partial c_s} > 0$ and $\frac{\partial J}{\partial \omega_b} > 0$ at $(c_s, \bar{\omega}_I(c_s))$, so $\bar{\omega}'_I(c_s) > 0$ for all $c_s \in [0, \bar{c}_s)$. Finally, if $c_s = 0$ and $\bar{\omega}_I = \bar{\omega}$, then $J = V_r(s_e) > 0$, a contradiction. So if $c_s = 0$, $\bar{\omega}_I < \bar{\omega}$. $\square$

Proof of Corollary 2. Note that $V_r(s_e) > 0$ implies $s_e < \bar{s}$. So given $\omega \in \Omega$, $1 - F(s_e|\omega) \in [0, 1)$, and because $c_b/(u(\omega) + c_b)$ is strictly increasing in c_b , there exists a unique $c_b(\omega) \in \mathbb{R}_+$ satisfying $1 - F(s_e|\omega) = c_b/(u(\omega) + c_b)$. Furthermore, $c_b(\omega)$ is strictly increasing in ω . Let $\underline{c}_b \equiv c_b(\underline{\omega}_I)$ and $\bar{c}_b \equiv c_b(\bar{\omega}_I)$, and apply Proposition 2. $\square$

Proof of Proposition 3. Since $s_r \in (\underline{s}, \bar{s})$, it follows that $V_r(s_r) > V_r(\bar{s}) = 0$.

Fix some $c_s > 0$. Suppose $c_b = 0$. Then for all $s \in S$, $\omega_b(s, 0) = \underline{\omega}$. If $s = \underline{s}$, then $F(s|\omega) = 0$, so $V_n(c_s, 0, s) = \int_{\underline{\omega}} v(\omega) dG(\omega) = V_r(\underline{s}) < V_r(s_r)$. If $s > \underline{s}$, then $F(s|\omega) > 0$, so $V_n(c_s, 0, s) < V_r(s) \leq V(s_r)$. So for all $s \in S$, $V_n(c_s, 0, s) < V_r(s_r)$, and therefore $V_n(c_s, 0, s_n) < V_r(s_r)$.

Next, we show that there exists a $c_b > 0$ such that V_n strictly dominates V_r . Note that $\int_{\omega_r} v(\omega) dG(\omega) > V_r(s_r)$. By the continuity of $F(s|\omega)$, there exists an

$s \in (\underline{s}, \bar{s})$ such that $\int_{\omega_r} ([1 - F(s|\omega)]v(\omega) - F(s|\omega)c_s)dG(\omega) > V_r(s_r)$. Because s is interior, the image of $(0, \infty)$ under $\omega_b(s, \cdot)$ is $[\underline{\omega}, \bar{\omega}]$, so there exists a $c_b > 0$ such that $\omega_b(s, c_b) = \omega_r$. Therefore, $V_n(c_s, c_b, s_n) \geq V_n(c_s, c_b, s) > V_r(s_r)$.

By the continuity of $V_n(c_s, \cdot, s_n(c_s, \cdot))$, there exists at least one $\underline{c}_b > 0$ such that $V_n(c_s, \underline{c}_b, s_n(c_s, \underline{c}_b)) = V_r(s_r)$. We now show that $V_n(c_s, \cdot, s_n(c_s, \cdot))$ is strictly increasing at any $\underline{c}_b$, which implies $\underline{c}_b$ is unique. By the envelope theorem,

$$\begin{aligned} \left. \frac{dV_n(c_s, c_b, s_n(c_s, c_b))}{dc_b} \right|_{\underline{c}_b} &= \left. \frac{\partial V_n(c_s, c_b, s_n(c_s, c_b))}{\partial c_b} \right|_{\underline{c}_b} \\ &= - \frac{\partial \omega_b(s_n, c_b)}{\partial c_b} \left([1 - F(s_n|\omega_b)]v(\omega_b) - F(s_n|\omega_b)c_s \right) dG(\omega_b) \Big|_{\underline{c}_b}, \end{aligned} \quad (11)$$

where we have abbreviated $s_n(c_s, c_b)$ and $\omega_b(s_n(c_s, c_b), c_b)$ with s_n and ω_b , respectively. If $\omega_b \in \{\underline{\omega}, \bar{\omega}\}$, then $V_n(c_s, \underline{c}_b, s_n(c_s, \underline{c}_b)) < V_r(s_r)$, a contradiction, so $\omega_b \in (\underline{\omega}, \bar{\omega})$, and therefore $\left. \frac{\partial \omega_b(s_n, c_b)}{\partial c_b} \right|_{\underline{c}_b} > 0$. To sign the second factor of (11), consider

$$\begin{aligned} \left. \frac{\partial V_n(c_s, c_b, s)}{\partial s} \right|_{(\underline{c}_b, s_n)} &= - \int_{\omega_b(s, c_b)} f(s|\omega)(v(\omega) + c_s)dG(\omega) \\ &\quad - \frac{\partial \omega_b(s, c_b)}{\partial s} \left([1 - F(s_n|\omega_b)]v(\omega_b) - F(s_n|\omega_b)c_s \right) dG(\omega_b) \Big|_{(\underline{c}_b, s_n)}. \end{aligned} \quad (12)$$

Note that because $\omega_b(s_n(c_s, \underline{c}_b), \underline{c}_b) \in (\underline{\omega}, \bar{\omega})$, it must be that $\left. \frac{\partial \omega_b(s, c_b)}{\partial s} \right|_{(s_n, \underline{c}_b)} > 0$ and $s_n(c_s, \underline{c}_b) \in (\underline{s}, \bar{s})$. If $[1 - F(s|\omega_b)]v(\omega_b) - F(s|\omega_b)c_s \Big|_{(\underline{c}_b, s_n)} \geq 0$, then $v(\omega_b) \geq 0$, so $\omega_b(s_n(c_s, \underline{c}_b), \underline{c}_b) \in [\omega_r, \bar{\omega})$, which implies the first term of (12) is strictly negative, and the second term is weakly negative. But then $\left. \frac{\partial V_n(c_s, c_b, s)}{\partial s} \right|_{(\underline{c}_b, s_n)} < 0$, contradicting the optimality of $s_n \in (\underline{s}, \bar{s})$. So $([1 - F(s|\omega_b)]v(\omega_b) - F(s|\omega_b)c_s) \Big|_{(\underline{c}_b, s_n)} < 0$, and therefore from (11) we have that $\left. \frac{dV_n(c_s, c_b, s_n(c_s, c_b))}{dc_b} \right|_{\underline{c}_b} > 0$. So wherever $V_n(c_s, c_b, s_n) = V_r(s_r)$, V_n must be

strictly increasing in c_b ; together with the continuity of $V_n(c_s, \cdot, s_n)$, this implies that $V_n(c_s, \cdot, s_n)$ must cross $V_r(s_r)$ exactly once, namely at $\underline{c}_b$.

Suppose $c_s = 0$. If $c_b = 0$, then for all $s \in S$, $\omega_b(s, 0) = \underline{\omega}$, $V(0, 0, s) = V_r(s)$, so $s_n = s_r$, and therefore $V_n(s_n) = V_r(s_r)$. Using an identical argument to the case of $c_s > 0$, there exists a $c_b > 0$ such that V_n strictly dominates V_r . Now consider (11) for $c_s = 0$ and $c_b > 0$. First note that $\omega_b < \bar{\omega}$ and therefore $s_n < \bar{s}$, as otherwise $V_n = 0$, which is strictly dominated by selecting some $s < \bar{s}$ such that $\omega_b(s, c_b) = \omega_r$. Next, note that if $[1 - F(s_n|\omega_b)]v(\omega_b) \geq 0$, then $\omega_b \in [\omega_r, \bar{\omega})$ and $s_n \in (\underline{s}, \bar{s})$, which implies the first term of (12) is strictly negative and the second term is weakly negative, contradicting the optimality of $s_n > \underline{s}$. So $([1 - F(s_n)]|v(\omega_b) < 0$, and since $\frac{\partial \omega_b(s, c_b)}{\partial c_b} \geq 0$, (11) implies $V_n(0, c_b, s_n(0, c_b))$ is weakly increasing in c_b . So for $c_s = 0$, let $\underline{c}_b$ be the smallest c_b such that $V_n(0, c_b, s_n(0, c_b)) = V_r(s_r)$.

To show that for $c_s = 0$, $\underline{c}_b > 0$, we show that there exists a $c_b > 0$ such that $V_n(0, c_b, s_n) = V_r(s_r)$. Let $\hat{s}(c_b)$ be the highest $s \in S$ such that $\omega_b(s, c_b) = \underline{\omega}$. Then given $c_b > 0$, $s \leq \hat{s}(c_b)$ implies $\omega_b(s, c_b) = \underline{\omega}$, so $V_n(0, c_b, s) = V_r(s)$. Note that $\lim_{c_b \rightarrow 0} \hat{s}(c_b) = \bar{s}$, and $\lim_{\hat{s} \rightarrow \bar{s}} V_r(\hat{s}) - \int_{\underline{\omega}}^{\omega_r} [1 - F(\hat{s}|\omega)]v(\omega)dG(\omega) = 0$. So there exists a $\delta > 0$ such that $c_b < \delta$ implies $\hat{s} > s_r$ and $V_r(\hat{s}) - \int_{\underline{\omega}}^{\omega_r} [1 - F(\hat{s}|\omega)]v(\omega)dG(\omega) < V_r(s_r)$. Given $c_b < \delta$, $\max_{s \leq \hat{s}(c_b)} V_n(0, c_b, s) = \max_{s \leq \hat{s}(c_b)} V_r(s) = V_r(s_r)$, whereas $\max_{s > \hat{s}(c_b)} V_n(0, c_b, s) = \max_{s > \hat{s}(c_b)} V_r(s) - \int_{\underline{\omega}}^{\omega_b(s, c_b)} [1 - F(s|\omega)]v(\omega)dG(\omega) \leq V_r(\hat{s}) - \int_{\underline{\omega}}^{\omega_r} [1 - F(\hat{s}|\omega)]v(\omega)dG(\omega) < V_r(s_r)$. So for all $c_b < \delta$, $V_n(0, c_b, s_n(0, c_b)) = \max_{s \in S} V_n(0, c_b, s) = V_r(s_r)$. Therefore, $\underline{c}_b \geq \delta > 0$.

To show that $\underline{c}_b$ is a strictly increasing function of c_s , apply the implicit function theorem to $V_r(s_r) = V_n(c_s, \underline{c}_b(c_s), s_n(c_s, \underline{c}_b(c_s)))$ and the envelope theorem to get

$$0 = \frac{dV_n(c_s, \underline{c}_b(c_s), s_n(c_s, \underline{c}_b(c_s)))}{dc_s} = \frac{\partial V_n}{\partial c_s} + \frac{\partial V_n}{\partial c_b} \underline{c}_b'(c_s) \Big|_{(\underline{c}_b, s_n)}.$$

At $c_b = \underline{c}_b$, ω_b and s_n are interior, so $\frac{\partial V_n}{\partial c_s} = - \int_{\omega_b} F(s_n|\omega)dG(\omega) < 0$, and from

above, $\frac{\partial V_n}{\partial c_b} > 0$, which gives $\underline{c}_b'(c_s) > 0$.

Finally, to show that $\lim_{c_s \rightarrow \infty} \underline{c}_b(c_s) = \infty$, We first show that $\lim_{c_s \rightarrow \infty} V_n(c_s, c_b, s_n(c_s, c_b))$ exists and is strictly less than $V_r(s_r)$. Note that $\frac{dV_n(c_s, c_b, s_n(c_s, c_b))}{dc_s} = \frac{\partial V_n}{\partial c_s} = - \int_{\omega_b} F(s_n|\omega) dG(\omega) \leq 0$, so the optimized V_n is weakly decreasing in c_s and therefore the limit exists. If $\liminf_{c_s \rightarrow \infty} s_n(c_s, c_b) = \underline{s}$, then $\liminf_{c_s \rightarrow \infty} \omega_b(s_n, c_b) = \underline{\omega}$, and so $\lim_{c_s \rightarrow \infty} V_n(c_s, c_b, s_n(c_s, c_b)) = V_r(\underline{s}) < V_r(s_r)$. If $\liminf_{c_s \rightarrow \infty} \omega_b(s_n, c_b) = \bar{\omega}$, then $\lim_{c_s \rightarrow \infty} V_n(c_s, c_b, s_n(c_s, c_b)) = 0 < V_r(s_r)$. If $\liminf_{c_s \rightarrow \infty} s_n(c_s, c_b) > \underline{s}$ and $\liminf_{c_s \rightarrow \infty} \omega_b(s_n, c_b) < \bar{\omega}$, then $\lim_{c_s \rightarrow \infty} V_n(c_s, c_b, s_n(c_s, c_b)) = -\infty < V_r(s_r)$. So regardless of the limit behavior of s_n and w_b , $\lim_{c_s \rightarrow \infty} V_n(c_s, c_b, s_n(c_s, c_b)) < V_r(s_r)$. Therefore, given $\epsilon > 0$, there exists a $\delta > 0$ such that $V_n(\delta, \epsilon, s_n(\delta, \epsilon)) < V_r(s_r)$, which implies $\underline{c}_b(\delta) > \epsilon$. Because $\underline{c}_b$ is strictly increasing in c_s , for all $c_s > \delta$ we must have $\underline{c}_b(c_s) > \underline{c}_b(\delta) > \epsilon$, and therefore $\lim_{c_s \rightarrow \infty} \underline{c}_b(c_s) = \infty$. $\square$

Proof of Proposition 4. Given $c_b \geq 0$, let $\hat{s}(c_b)$ be the highest $s \in S$ such that $\omega_b(s, c_b) = \underline{\omega}$. As c_b increases from zero to infinity, $\hat{s}$ decreases from $\bar{s}$ to $\underline{s}$, so there exists $\hat{c}_b$ such that $\hat{s}(\hat{c}_b) = s_r \in (\underline{s}, \bar{s})$. Denote by $s_n^-(c_s)$ a maximizer of $V_n(c_s, \hat{c}_b, s)$ over $[\underline{s}, s_r]$ and $s_n^+(c_s)$ a maximizer of $V_n(c_s, \hat{c}_b, s)$ over $[s_r, \bar{s}]$.

Because $E[v(\omega)|s_r] = 0$, taking the right derivative of V_n with respect to s and evaluating it at $(c_s, c_b, s) = (0, \hat{c}_b, s_r)$ gives

$$\frac{\partial V_n}{\partial s} = -\frac{\partial \omega_b}{\partial s} [1 - F(s_r|\underline{\omega})] v(\underline{\omega}) dG(\underline{\omega}) > 0,$$

so there exists $s > s_r$ such that $V_n(0, \hat{c}_b, s) > V_n(0, \hat{c}_b, s_r)$. Therefore, $V_n(0, \hat{c}_b, s_n^+(0)) > V_n(0, \hat{c}_b, s_r) = V_r(s_r)$.

By the maximum theorem and the envelope theorem, as c_s increases from zero to infinity, $V_n(\cdot, \hat{c}_b, s_n^+(\cdot))$ decreases continuously from $V_n(0, \hat{c}_b, s_n^+(0)) > V_r(s_r)$ to zero. So there exists $\hat{c}_s > 0$ such that $V_n(\hat{c}_s, \hat{c}_b, s_n^+(\hat{c}_s)) = V_r(s_r)$. Also note that for all $s \in (\underline{s}, s_r]$, $V_n(\hat{c}_s, \hat{c}_b, s) < V_r(s) \leq V_r(s_r)$, and for $s = \underline{s}$, $V_n(\hat{c}_s, \hat{c}_b, \underline{s}) = V_r(\underline{s}) < V_r(s_r)$. Therefore, $V_n(\hat{c}_s, \hat{c}_b, s_n^-(\hat{c}_s)) < V_r(s_r) = V_n(\hat{c}_s, \hat{c}_b, s_n^+(\hat{c}_s))$.

Because $V_n(\cdot, \hat{c}_b, s_n^-(\cdot))$ and $V_n(\cdot, \hat{c}_b, s_n^+(\cdot))$ are continuous and strictly decreasing at $\hat{c}_s$, there exists $c_s > \hat{c}_s$ such that $V_n(c_s, \hat{c}_b, s_n^-(c_s)) < V_n(c_s, \hat{c}_b, s_n^+(c_s)) < V_r(s_r)$. Therefore, $s_n(c_s, \hat{c}_b) = s_n^+(c_s) > s_r$ and $V_n(c_s, \hat{c}_b, s_n(c_s, \hat{c}_b)) < V_r(s_r)$. $\square$