



Bank Capital Standards for Market Risk: A Welfare Analysis*

DAVID MARSHALL and SUBU VENKATARAMAN

Federal Reserve Bank of Chicago, 230 South LaSalle Street, Chicago, Illinois 60604, USA; E-mail: dmarshall@frbchi.org

Abstract. We propose a simple model that is suitable for evaluating alternative bank capital regulatory proposals for market risk. Our model formalizes the conflict between bank objectives and regulatory goals. Banks' decisions represent a tension between their desire to exploit the deposit-insurance put option and their desire to preserve franchise value. Regulators seek to balance the social value of deposits in mediating transactions against the deadweight costs of failure resolution. Our social welfare criterion is standard: a weighted average agents' utilities.

We demonstrate that banks do not incrementally alter their portfolio risk as the economic environment changes. Rather, banks either choose the minimal feasible risk or the maximal feasible risk. This pattern, in turn, drives regulatory decisions: The first goal of the regulator is to induce banks to choose the minimal risk level. For all nontrivial cases, unregulated banks fail to choose the first-best allocations. Traditional ex-ante capital requirements can induce banks to choose the socially-optimal level of portfolio risk, but the required capital is often inefficiently high. In contrast, variants of the Federal Reserve Board's precommitment proposal imply far smaller efficiency losses, and achieve allocations at or near the first-best for most reasonable model specifications. The ex-post penalties required for the optimal implementation of precommitment are not excessively large. The welfare gains from precommitment are even higher when the precommitment penalty function is precluded from sending banks into default. We conclude that state-contingent regulatory mechanisms, of which the precommitment approach is an example, offer the possibility of substantial gains in regulatory efficiency, relative to traditional state non-contingent regulation.

1. Introduction

Over the last few years, a good deal of attention has been focused on how to set bank capital standards for market risk (the risk that a bank's value may be adversely affected by price movements in financial markets). The risk-based capital standards in the 1988 Basle Accords cover only credit risk. A proposal to extend these accords to cover market risk was published in April 1993. This proposal (known as the 'Standardized Approach') followed the conventional approach of

* The opinions expressed in this paper do not necessarily represent the views of the Federal Reserve Bank of Chicago or the Federal Reserve System. We would like to thank Doug Evanoff, Mark Flannery, George Kaufman, Paul Kupiec, Jeff Lacker, Pierre Mella-Barral, Alistair Milne, and Jim O'Brien for their comments on earlier drafts, and Denise Duffy, Xiongwei Ju, and Glenn McAfee for their valuable research assistance.

state non-contingent regulation: Ex-ante capital requirements would be determined by the regulator's assessment of a bank's market risk. The standardized approach was heavily criticized. Many industry participants claimed that risk assessments by regulators would be highly inaccurate, compared with those computed internally by the banks themselves. In other words, bank portfolio risk fundamentally represents *private information*.

It is noteworthy therefore that subsequent proposals for market-risk capital standards have moved towards *state contingent regulation*,¹ in which regulatory consequences for banks depend on the ex-post performance of the bank's trading portfolio. For example, the standards adopted in December 1996 base regulatory capital on banks' own reports of their trading portfolio risk. State-contingency, in the form ex-post back-testing, is used to induce banks to reveal truthfully their risk level. A more innovative proposal incorporating state-contingent regulation was put forth for comment by the Federal Reserve Board in June 1995.² Known as the 'Precommitment Approach', this proposal would have each bank state the maximum loss that its trading portfolio will sustain over the next period. The capital charge for market risk would equal this pre-committed maximum loss. If the bank's losses exceed the pre-committed level, a penalty would be imposed that is proportional to the amount of the excess loss.

The precommitment approach has attracted substantial attention from both regulatory economists³ and from commercial banks.⁴ However, there is still a good deal to be learned about this proposal. Under what circumstances would precommitment dominate state-noncontingent regulation from a social welfare standpoint? Would the precommitment approach impose undue costs on the banks themselves? What would the optimal precommitment penalty structure look like? In particular, will the optimal precommitment penalty scheme require extremely high (perhaps politically infeasible) ex-post fines?

To address these questions, and to examine the problem of bank capital requirements for market risk more generally, this paper constructs an optimizing model of bank and regulatory behavior in which bank market risk is unobserved. As is typical in welfare economics, our measure of social welfare is an equally-weighted sum of agents' utilities. This criterion makes precise the key conflicts between banks' objectives and the regulatory goals, and formalizes the tradeoffs facing regulators. Unlike some other papers of this type, our focus is not on finding a mechanism that achieves the first-best allocation. Rather, we wish to evaluate regulatory mechanisms that are either currently in use or under consideration by regulatory bodies. In particular, we use the model to study traditional state non-contingent regulation

¹ See, e.g., Laffont and Tirole (1993).

² The precommitment approach was developed in papers by Kupiec and O'Brien (1995a,b). The Board of Governor's proposal is described in Federal Register, Vol. 60, No. 142, pp. 38142-38144.

³ See Kupiec and O'Brien (1997), Bliss (1995), Prescott (1997).

⁴ Most notably, the New York Clearing House has embarked on a pilot study testing the feasibility of precommitment.

(ex-ante capital requirements) and two variants of the precommitment approach. We compute the socially optimal levels of risk and capital, determine how closely the choices of an unregulated bank approximate these socially optimal allocations, and rank the differing regulatory approaches.

Our approach is closely related to models used in papers by John, John, and Senbet (1991), Giammarino, Lewis, and Sappington (1993), Gorton and Winton (1995), and John, Saunders, and Senbet (1996). In all these papers, banks produce socially valuable liquidity as well as making socially valuable investments. The regulator must control the banks' tendency towards excessive risk-taking without unduly suppressing these valuable bank activities. Giammarino, Lewis, and Sappington (1993) and Gorton and Winton (1995) formalize this tradeoff with a social welfare function similar to that used in our model. In particular, Gorton and Winton (1995) formulate a general equilibrium model, in which the regulator chooses allocations to maximize the average utility of agents in the economy. In contrast, we use a partial equilibrium approach, in that the risk-free rate is set outside the model.

In our model, the two bank-specific characteristics are the bank's franchise value and the quality of the bank's loan portfolio. We assume that these bank-specific characteristics are observable, so bank capital regulations can be tailored to the characteristics of the regulated bank. (Alternatively, one can think of this paper as analyzing a simple economy in which all banks are alike.) We do so to focus attention on the issue of moral hazard (unobserved action) in the bank's choice of risk. Of course, there are a host of important questions that arise when both bank actions and bank characteristics are unobservable. To analyze these questions in the context of this model, one must draw in an essential way on the results in this paper.

An implication of our model is that the regulator seeks to eliminate all uncompensated (idiosyncratic) risk in the bank's trading portfolio. In contrast, an unregulated bank either seeks to *eliminate* uncompensated risk or to *maximize* uncompensated risk. In effect, banks endogenously sort into *prudent* banks (those who are most concerned with maximizing expected franchise value) and *risk-seeking* banks (those who are most concerned with maximizing the deposit insurance put option). The most important goal of the regulator is to induce risk-seeking banks to behave prudently. A secondary goal of regulation is to induce the optimal capital structure *given* that the bank is behaving prudently.

For most reasonable parametric specifications, the precommitment approach dominates ex-ante capital regulation, both according to the social welfare criterion and according to the bank's own objective function. The magnitude of the optimal ex-post penalty under precommitment depends on bank characteristics. For example, when banks have high franchise value, the required penalty level is extremely small (less than 5% of the precommitment violation). When franchise value is low, however, the optimal penalty can exceed 30% of the precommitment violation. While this is a substantial penalty, banks still prefer this penalty level to

the optimal capital requirement under state non-contingent regulation, which can be above 80% when franchise value is low.

The remainder of the paper is structured as follows: Section 2 describes our model of bank activity. Section 3 describes the social welfare criterion, and characterizes the first-best allocation. Section 4 characterizes the decisions of a bank in an unregulated economy. In Section 5, we describe the regulatory mechanisms we explore: ex-ante capital requirements and two variants of the precommitment approach. In Section 6, we use numerical methods to compare these regulatory mechanisms. Section 7 summarizes the conclusions we draw from this paper.

2. The Model

2.1. THE ECONOMIC ENVIRONMENT

In this section, we describe in detail a model of bank activity that is suitable for studying capital regulation for market risk. There are two time periods, and two types of agents (households and banks), along with a bank regulator. All agents are risk-neutral price takers. We assume that the risk-free return is set outside the banking sector, and we normalize the gross risk-free rate to unity.

2.2. HOUSEHOLDS

In the first period, the household can invest in bank equity and bank deposits. Deposits are government insured,⁵ and provide liquidity services to households. In particular, let D denote the face value of deposits held by a household. We assume that the liquidity services generated by these deposits have a consumption-good-equivalent of ρD (where ρ is a strictly positive parameter).⁶ The parameter ρ gives the *relative* benefits of deposits versus equity in our model. It gives rise to a wedge between the cost of equity financing to the bank and the cost of deposits. An alternate justification for this wedge would be the 'lemons' cost associated with the sale of equity. (See Gorton and Winton (1996).)⁷

Let p_d denote the price of one unit (face value, in units of the consumption good) of bank deposits, let p_z denote the price per share of bank equity, (with the number of shares normalized to unity) and let $\tilde{v}$ denote the stochastic payoff (to be formalized later) to a share of bank equity. The expected returns to deposits

⁵ We take deposit insurance as given. We do not explicitly model the rationale for deposit insurance, which has been treated elsewhere (e.g., Diamond and Dybvig (1983)).

⁶ The importance of liquidity services provided by deposits has been noted by many authors. See Giammarino, Lewis, and Sappington (1993) and the references cited therein.

⁷ All of our results go through if ρ is interpreted as the relative deadweight cost of issuing equity, rather than the relative social benefit to issuing deposits.

(inclusive of the liquidity value ρ) and to bank equity must be equal to the risk-free rate of unity, implying

$$p_d = 1 + \rho; \quad p_z = E[\tilde{v}]. \quad (1)$$

2.3. BANKS

2.3.1. *The Trading Portfolio*

There are two types of bank assets: loans and marketable securities. The latter compose the bank's trading portfolio. Since the trading portfolio exhibits constant returns to scale, risk neutrality would imply that the optimal portfolio size is indeterminate. For this reason, we fix the size of the trading portfolio exogenously. Without loss of generality, we normalize this size to unity.

The gross return to the trading portfolio is denoted $\tilde{r}$. It is assumed that $\tilde{r}$ is a log-normal random variable with mean $\mu > 1$ (a constant parameter) and standard deviation σ (a choice variable of the bank). Admissible values of σ lie in the interval $[\sigma_0, \bar{\sigma}]$, where parameter $\sigma_0 \geq 0$ and parameter $\bar{\sigma} \leq \infty$. If a bank can perfectly hedge all residual market risk, and if it faces no counterparty risk, then σ_0 would equal zero. Otherwise, σ_0 would exceed zero. If $\bar{\sigma} < \infty$, $\bar{\sigma}$ would represent a finite upper bound $\bar{\sigma}$ on the bank's risk choice. If parameter, $\bar{\sigma} = \infty$ then the bank's choice of portfolio risk is unbounded. Let $f(\tilde{r}|\sigma)$ denote the log normal density function with mean μ and standard deviation σ , and let $F(\tilde{r}|\sigma)$ denote the corresponding cumulative distribution function.

Note that there is no 'risk-return' tradeoff. Regardless of the level of risk σ chosen by the bank, the mean return is still μ . Essentially, this is an assumption of mean-variance efficiency: A higher mean return is only justified as compensation for extra risk. In our economy, all agents are risk-neutral, so there is never compensation for additional risk. If the trading portfolio paid the risk-free return on average, one would set $\mu = 1$. However, the trading portfolio would then represent a zero net present value investment, and would provide no net social value. The optimal regulatory strategy would then be to prohibit banks entirely from trading. To avoid this trivial conclusion, we assume that μ exceeds the risk-free rate of unity.⁸ This assumption can be justified if banks have special expertise in performing certain valued trading activities (such as acting as counterparty to swaps transactions for commercial clients), and if there are economies of scale or other natural barriers to entering this activity. (One such barrier might be the informational monopoly a bank possesses with respect to its clients. See Petersen and Rajan (1994)).⁹

⁸ In an earlier version of this paper, we also allowed the bank to include risk-free investments in its trading portfolio. For reasonable specifications of $\mu > 0$, however, we found that the optimal weight on the risk-free asset was always zero.

⁹ Genotte and Pyle (1991) argue that, under some ad hoc specifications of the investment opportunity set, increased capital requirements can induce banks to *increase* risk. Our specification avoids this perverse result. See also Keeley and Furlong (1990).

The bank must fund the trading portfolio by selling a combination of deposits and outside equity to the households. This funding constraint can be written:

$$p_d D + p_z Z = 1, \quad (2)$$

where Z denotes the fraction of bank equity sold to the households. The bank can short-sell neither deposits nor equity, so

$$D \geq 0; \quad 0 \leq Z \leq 1. \quad (3)$$

For regulatory purposes, we measure bank capital as $p_z Z$, the value of outside equity at the time it is issued. Note that Equations (1) and (2) associate each level of D with a particular capital ratio (denoted $K(D)$), defined as the ratio of $p_z Z$ to the size of the trading portfolio:

$$K(D) = 1 - (1 + \rho)D. \quad (4)$$

2.3.2. The Loan Portfolio

While the focus of this paper is on the risk associated with the trading portfolio, we cannot properly evaluate regulation of trading portfolio without explicitly modelling the loan portfolio. We assume that both the composition and the financing of the loan portfolio are predetermined. Let $\tilde{x}$ denote the (random) loan portfolio payoff *in excess* of the deposits used to finance the loan portfolio. Random variable $\tilde{x}$ has the following binomial distribution:

$$\tilde{x} = \begin{cases} \bar{x} & \text{with probability } p \\ 0 & \text{with probability } (1 - p). \end{cases} \quad (5)$$

Under this specification, default can only occur if there are losses to the trading portfolio. In this way, we avoid confounding default induced by market risk with default due to loan portfolio risk. We further assume that default can never occur when $\tilde{x} = \bar{x}$. A sufficient condition for this is

$$\bar{x} > \frac{1}{1 + \rho}. \quad (6)$$

Default can occur when $\tilde{x} = 0$. Parameter p measures the quality of the loan portfolio, in the following sense: The region where $\tilde{x} = \bar{x}$ comprises those realizations of the loan portfolio where default cannot be induced by trading losses of any magnitude. Parameter p gives the probability of this region. As long as Equation (6) is satisfied, the precise value of $\bar{x}$ is unimportant for the results of this paper.

2.3.3. Franchise Value

In the event of bank failure, a portion of the bank's value, denoted ψ , is lost both to the bank shareholders and to society as a whole. We refer to ψ as *franchise value*. Intuitively, the value of the bank is in part determined by the network of relationships it has built up. This relationship-value is highly information-sensitive, so is imperfectly transferrable. We think of franchise value as that portion of this relationship-value that cannot be transferred, sold, or claimed by the deposit insurer in the event of bank failure.¹⁰ We model ψ as a predetermined positive parameter.

2.3.4. Deposit Insurance

The deposit insurer ensures that the depositors of failed banks are paid off in full. Failure resolution generates deadweight costs equal to β times the cash shortfall, where $\beta \geq 0$. The total costs faced by the deposit insurer are paid via a lump-sum tax $\tilde{\tau}$, assessed ex-post on the households:

$$\tilde{\tau} = (1 + \beta)[\max\{D - \tilde{x} - \tilde{r}, 0\}]. \quad (7)$$

Note that we preclude risk-based deposit insurance premiums. We do so to focus on the issue of capital regulation. However, as noted by Berger, Herring, and Szego (1995), deposit insurance premiums in the U.S. are only weakly tied to the underlying levels of bank asset risk.

2.4. FORMAL STATEMENT OF THE BANK'S OBJECTIVE

Under limited liability, the bank's equity payoff, $\tilde{v}$, is given by

$$\tilde{v} = \begin{cases} \tilde{x} + \tilde{r} + \psi - D, & \text{if } \tilde{x} + \tilde{r} \geq D, \\ 0 & \text{otherwise.} \end{cases} \quad (8)$$

The bank seeks to maximize the expected value of inside equity, which is given by:

$$(1 - Z)E[\tilde{v}] = p(\bar{x} + \mu - D + \psi) + (1 - p) \times \int_D^\infty (\tilde{r} + \psi - D)f(\tilde{r}|\sigma) d\tilde{r} - 1 + (1 + \rho)D. \quad (9)$$

Equation (9) follows from Equations (1), (2), (6) and (8). In Equation (9), D and σ must satisfy:

$$0 \leq D \leq \frac{1}{1 + \rho} \equiv D^{\max}; \quad \sigma_0 \leq \sigma \leq \bar{\sigma}. \quad (10)$$

¹⁰ As a formal matter, the portion of this relationship-value that is transferrable can be lumped into the payoff of the loan portfolio.

3. The Regulator's Problem: Computing the First-Best Allocation

In this section, we posit an explicit social welfare criterion. In contrast to many studies on banking regulation that impose ad hoc specifications for the regulator's objective, our model has a natural candidate: We assume that the regulator seeks to maximize an equally-weighted average of the utilities of the agents in the economy. In our economy, there are two types of agents, households and bank insiders, both of whom are risk-neutral. Therefore, the regulator's problem is:

$$\begin{aligned} \max_{D, \sigma} \{ & (1/2)[E(Z\bar{v} - \bar{\tau}) + (1 + \rho)D] + (1/2)[(1 - Z)E(\bar{v})], \\ \text{s.t. Equation (10).} \end{aligned} \quad (11)$$

In Equation (11), the first term in square brackets gives the bank's contribution to household expected wealth (including the liquidity value of deposits). The second term gives the expected value of inside bank equity. Since both types of agents are risk-neutral, equal weighting is the only sensible choice. If the weights were unequal, the social optimum would be attained by maximizing the utility of the higher-weighted type of agent, ignoring the lower-weighted type.

In our model, maximizing this equally-weighted welfare problem is equivalent to maximizing the expected total output of the banking sector, inclusive of the liquidity value of deposits (ρD), the franchise value of the bank (ψ), and taking account of the deadweight costs of failure resolution:

$$\begin{aligned} \max_{D, \sigma} p\bar{x} + \mu + \psi + \rho D - (1 - p) \int_0^D (\beta[D - \bar{r}] + \psi) f(\bar{r}|\sigma) d\bar{r}, \\ \text{s.t. Equation (10).} \end{aligned} \quad (12)$$

The first-order condition with respect to D is as follows:

$$\rho - (1 - p)[\beta F(D|\sigma) + \psi f(D|\sigma)] \begin{cases} \leq 0 & \text{if } D = 0 \\ = 0 & \text{if } 0 < D < D^{\max} \\ \geq 0 & \text{if } D = D^{\max}, \end{cases} \quad (13)$$

where $D^{\max}$ is defined in (10). Note that the optimal default probability is not zero. According to (13), the optimal probability of default when $\psi = 0$ (at an interior solution) equals ρ/β . When $\psi > 0$, the regulator also considers the marginal expected loss in bank franchise value. If the solution for D is interior, the following second-order necessary condition must hold:

$$-(1 - p)[\beta f(D) + \psi f'(D)] < 0. \quad (14)$$

Once the optimal D is determined, the optimal capital ratio follows immediately from Equation (4).

The first-order condition with respect to σ is:

$$(1-p) \int_0^D [\beta(\tilde{r} - D) - \psi] f_\sigma(\tilde{r}|\sigma) d\tilde{r} \begin{cases} \leq 0 & \text{if } \sigma = \sigma_0 \\ = 0 & \text{if } \sigma_0 < \sigma < \bar{\sigma} \\ \geq 0 & \text{if } \sigma = \bar{\sigma}. \end{cases} \quad (15)$$

Not surprisingly, it can be shown that (ignoring the trivial case of $D = 0$, where σ drops out of the regulator's problem entirely) the regulator always chooses $\sigma = \sigma_0$:

PROPOSITION 1. If $D > 0$, the social planning optimum sets $\sigma = \sigma_0$.

Proof. All proofs are in the Appendix.

Let the socially optimal level of D be denoted D^* . The following comparative static results hold:

PROPOSITION 2.

- (i) D^* is strictly positive.
- (ii) If $D^* < D^{\max}$, then D^* is strictly decreasing in ψ and β .
- (iii) If $D^* < D^{\max}$, then D^* is strictly increasing in ρ and p .

Proof. Appendix.

These comparative static results are intuitive. Increasing p reduces the probability of default, and decreasing β reduces the cost of default. Either change reduces the social cost of deposits. Similarly, increasing ρ increases the social benefits to deposits. Decreasing ψ decreases the social cost of default, since reducing ψ reduces the value of the bank to society.

4. The Bank's Optimal Choices in an Unregulated Economy

In the previous section we characterized the first-best allocation in this economy. In this section, we study the bank's optimization problem in the absence of regulation. We assume that the bank's choice of trading portfolio risk, σ , cannot be observed by regulators. This creates the potential for moral hazard between the risk level preferred by the regulator and that selected by the bank. One of the objectives of regulation will be to control this moral hazard problem.

Using Equation (9), we can rewrite the bank's problem as:

$$\max_{(D, \sigma)} (p\bar{x} + (\mu - 1) + \rho D + \psi) + (1-p) \int_0^D (D - \tilde{r}) f(\tilde{r}|\sigma) d\tilde{r} - (1-p)\psi F(D|\sigma), \quad \text{s.t. Equation (10).} \quad (16)$$

The first term in parentheses gives the value of the bank under unlimited liability. The second term gives the value of the put option given by the deposit insurer to the

bank. The third term subtracts off the expected lost franchise value due to default. The bank first-order condition with respect to D is

$$\rho + (1-p)[F(D|\sigma) - \psi f(D|\sigma)] \begin{cases} \leq 0 & \text{if } D = 0 \\ = 0 & \text{if } 0 < D < D^{\max} \\ \geq 0 & \text{if } D = D^{\max}. \end{cases} \quad (17)$$

If the solution for D is interior, the following second-order necessary condition must hold:

$$(1-p)[f(D|\sigma) - \psi f'(D|\sigma)] < 0. \quad (18)$$

The first-order condition with respect to σ is:

$$(1-p) \int_0^D (D - \tilde{r} - \psi) f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} \begin{cases} \leq 0 & \text{if } \sigma = \sigma_0 \\ = 0 & \text{if } \sigma_0 < \sigma < \bar{\sigma} \\ \geq 0 & \text{if } \sigma = \bar{\sigma}. \end{cases} \quad (19)$$

If the solution to (19) is interior, the following second-order necessary condition must hold:

$$(1-p) \int_0^D (D - \tilde{r} - \psi) f_{\sigma\sigma}(\tilde{r}|\sigma) d\tilde{r} < 0. \quad (20)$$

Equations (19) and (20) imply an important property of the bank's optimal choice for σ : The bank never chooses an interior value for σ . This is shown in the following proposition:

PROPOSITION 3. The bank's optimal choice of σ is either σ_0 or $\bar{\sigma}$.¹¹

Proof. Appendix.

The intuition behind this result is illustrated in Figure 1. When a bank considers the effects of higher risk, it trades off the benefits of the deposit insurance put option (the dot-dash line in the figure) against the potential loss of franchise value (the dotted lines). Consider the case where franchise value is very low (ψ_{low} in the figure). As risk is increased, more probability mass is placed on the area at the right edge of the bankruptcy region. Here, the (probability weighted) area associated with lost

¹¹ If $\bar{\sigma} = \infty$, Proposition 3 still holds, but the interpretation differs somewhat from the case where $\bar{\sigma} < \infty$. If the bank's optimal choice is $\bar{\sigma} = \infty$, it sets σ arbitrarily high. Note that for any D satisfying Equation (10),

$$\lim_{\sigma \rightarrow \infty} f(D|\sigma) = \lim_{\sigma \rightarrow \infty} F(D|\sigma) = \lim_{\sigma \rightarrow \infty} \int_0^D \tilde{r} f(\tilde{r}|\sigma) d\tilde{r} = 0.$$

It follows from Equation (17) that $D = D^{\max}$, and the value of objective function (16) is its limiting value as $\sigma \rightarrow \infty$, which is $p\bar{x} + (\mu - 1) + \rho D^{\max} + \psi < \infty$.

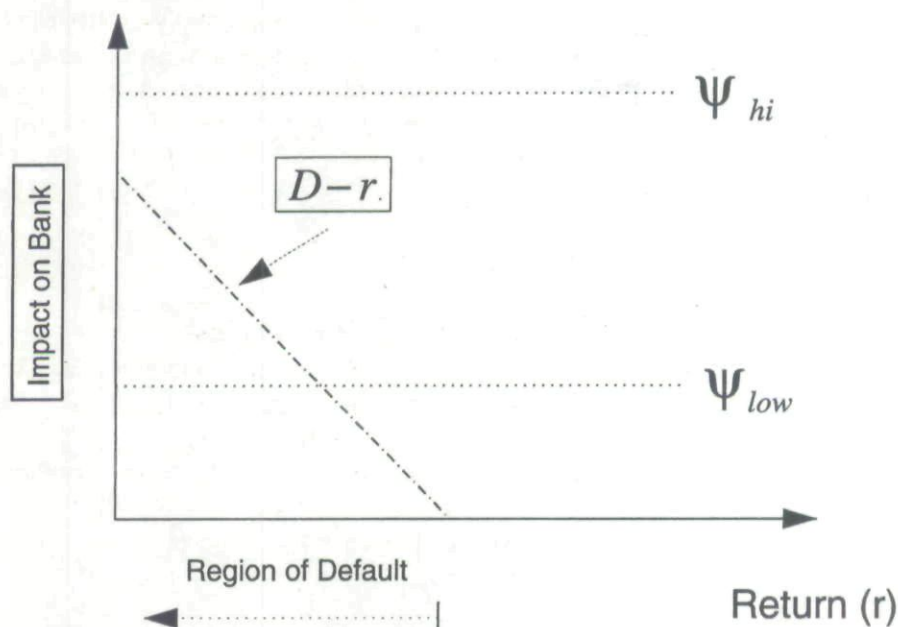


Figure 1. Deposit insurance versus franchise value.

franchise value is larger than the area associated with the deposit insurance put option.¹² As risk is increased, more weight is placed on the area further to the left. As one moves into this region, the area under the put option payoff is increasing much faster than the area under the franchise value line. It is possible that at some point the expected payoff of the put option begins to dominate. This is especially likely when franchise value is low. When franchise value is high (ψ_{hi}), however, the potential loss of franchise value is always the dominant concern to the bank. This analysis implies that the bank's net payoff as a function of risk is either (i) always increasing (when franchise value is zero), (ii) initially decreasing but then uniformly increasing, or (iii) always decreasing. Since these are the only possible outcomes, the bank will always select either minimal or maximal risk.

This proposition implies a striking characteristic of bank behavior in this model: While it is feasible for banks to choose any value of σ in the interval $[\sigma_0, \bar{\sigma}]$, their optimal strategy is to choose either σ_0 , the lowest feasible variance for their trading portfolio, or $\bar{\sigma}$, the highest feasible variance. In other words, banks self-select into one of two types: *prudent* banks (defined as those who choose σ_0) and *risk-seeking* banks (defined as those who chose $\bar{\sigma}$). Prudent banks are those for whom preservation of franchise value is of paramount concern, while risk-seeking banks are those for whom the paramount consideration is exploiting the deposit insurance put option. Of course, this bifurcation into two, and only two, choices for portfolio variance is specific to this model. A more general model that allowed for a limited

¹² This is easiest to see in terms of a mean preserving spread around a uniform distribution.

mean-variance tradeoff (that is, μ increasing in σ) would imply a range of portfolio variances for 'prudent' banks. Still, the key qualitative lesson from Proposition 3 is informative: High franchise value banks eschew uncompensated (idiosyncratic) risk, while low franchise value banks maximally take on uncompensated risk in order to exploit the deposit insurance put option.¹³

The following proposition further characterizes the bank's optimal choices:

PROPOSITION 4.

- (i) The bank's optimal choice of D is strictly positive.
- (ii) The bank's problem has at most one local maximum with respect to D .
- (iii) At $D < D^{\max}$, and for a fixed level of risk σ , the bank's optimal choice of D is strictly decreasing in ψ and strictly increasing in ρ .
- (iv) If $\psi = 0$, the bank's optimal choice of $D = D^{\max}$, and the bank's optimal choice of $\sigma = \bar{\sigma}$.
- (v) Holding other structural parameters constant, there exists $\bar{\psi} > 0$ such that the bank's optimal choice of $\sigma = \sigma_0$, $\forall \psi > \bar{\psi}$, and is $\sigma = \bar{\sigma}$, $\forall \psi < \bar{\psi}$.

Proof. Appendix.

According to Proposition 4(iv) and 4(v), the bank trades off franchise value loss against the deposit insurance put option. If a bank has no franchise value, it has no disincentive to taking maximal risk. On the other hand, a sufficiently high franchise value induces the bank to choose the minimum risk level.

Even if the bank sets $\sigma = \sigma_0$, it does not follow that the bank's choices correspond to the first-best. In general, even prudent banks choose suboptimally low levels of capital relative to the social-planner's optimum. This is formalized in Proposition 5, subject to the following technical restriction that is sufficient to ensure that the regulator's problem is globally concave.¹⁴

$$\frac{1}{1 + \rho} < \frac{\mu^4}{[\mu^2 + \sigma_0^2]^{3/2}}. \quad (21)$$

PROPOSITION 5. Suppose (21) is satisfied and the bank chooses $\sigma = \sigma_0$. Except in the case where both the bank and the regulator set $D = D^{\max}$, the bank's optimal choice of D is strictly above the socially optimal choice of D .

Proof. Appendix.

Recall from Equations (4) and (10) that the $D^{\max}$ corresponds to zero capital. Proposition 5 says that the unregulated bank's choices can only correspond to the

¹³ A similar results was noted in papers by Marcus (1984) and Ritchkin, Thomson, DeGennaro, and Li (1993).

¹⁴ Notice that since the condition is satisfied when σ_0 is sufficiently low.

social optimum in the rather extreme case where deposits are so valuable (high ρ) or default has such a low cost (low β) that the regulatory optimum sets bank capital at zero. In all other cases, the unregulated bank fails to make the socially optimal choices.

5. Regulatory Mechanisms

According to Section 4, above, the choices of banks in an unregulated economy are generally suboptimal. The goals of capital regulation are to induce risk-seeking banks (who would choose $\sigma = \bar{\sigma}$ in the absence of regulation) to behave prudently (that is, to choose $\sigma = \sigma_0$), and to provide them with incentives to hold sufficient capital. In this section we describe three regulatory mechanisms: ex-ante capital requirements, and two variants of the Board of Governors' 1995 precommitment proposal. For each mechanism, we define the optimal implementation of that mechanism as the implementation whose associated allocation (D, σ) achieves the highest value of social welfare function (12).

5.1. STATE NON-CONTINGENT CAPITAL REGULATION: EX-ANTE CAPITAL REQUIREMENTS

Under ex-ante capital requirements, the regulator sets the bank capital ratio, which is equivalent to choosing a value for D . The bank then maximizes its profits subject to this capital constraint by choosing its preferred level of portfolio risk, denoted $\sigma^{\text{ex}}(D)$. Using (16),

$$\sigma^{\text{ex}}(D) = \begin{cases} \sigma_0 & \text{if } \int_0^D (D - \tilde{r} - \psi) f(\tilde{r}|\sigma_0) \geq \int_0^D (D - \tilde{r} - \psi) f(\tilde{r}|\bar{\sigma}) \\ \bar{\sigma} & \text{otherwise.} \end{cases} \quad (22)$$

Let D^{ex} denote the regulator's choice of D under ex-ante capital requirements. As long as franchise value is positive, this mechanism can induce prudent behavior by imposing a sufficiently high capital ratio:

PROPOSITION 6. If $\psi > 0$, there exists D^{ex} such that $\sigma^{\text{ex}}(D^{\text{ex}}) = \sigma_0$.

Proof. Appendix.

Intuitively, the strike price of the deposit insurance put option is D , which is strictly decreasing in the capital ratio. By increasing the required capital ratio, this strike price (and therefore the value of the deposit insurance put option) can be made arbitrarily small. As long as $\psi > 0$, this put-option value can be made sufficiently low that it is dominated by the bank's concern for preserving franchise value. However, for very low ψ 's, the required capital ratio is quite high. In the limit,

when $\psi = 0$, the optimal ex-ante capital regulation is 100% capital, which reduces the put-option value to zero.

5.2. THE BASIC PRECOMMITMENT APPROACH

Under the Board of Governor's precommitment approach, the bank chooses a level K of capital. This is interpreted as a commitment that the loss to the trading portfolio, $(1 - \tilde{r})$, will not exceed K . If the loss exceeds this pre-committed level, a penalty ϕ is imposed that is proportional to the excess loss. Formally, the penalty function $\phi(D, \tilde{r})$ can be written:

$$\phi(D, \tilde{r}) = \begin{cases} 0 & \text{if } (\tilde{r} - 1) + K(D) > 0, \\ \gamma[(1 - \tilde{r}) - K(D)], & \text{otherwise,} \end{cases} \quad (23)$$

where $\gamma \geq 0$ is a constant of proportionality.

We now modify the bank's objective function (16) to incorporate the penalty function $\phi(D, \tilde{r})$. First, we must fix some additional notation. Denote the portfolio return that triggers the penalty by r^ϕ :

$$r^\phi(D) \equiv (1 + \rho)D, \quad (24)$$

where we use Equation (4). We must also determine when the bank is in default under precommitment. In the basic precommitment approach, we cannot rule out the possibility of default even when the loan portfolio pays off $\bar{x}$. Let $\bar{r}^*$ and $\underline{r}^*$ denote the portfolio returns that trigger bankruptcy under precommitment when $\tilde{x} = \bar{x}$, and $\tilde{x} = 0$, respectively:

$$\bar{r}^*(D) \equiv \max \left[\frac{(1 + \gamma(1 + \rho))D - \bar{x}}{(1 + \gamma)}, 0 \right], \quad \underline{r}^*(D) \equiv \frac{(1 + \gamma(1 + \rho))D}{(1 + \gamma)}, \quad (25)$$

Note that if both $\rho > 0$ and $\gamma > 0$, then $r^\phi(D) > \underline{r}^*(D) > \bar{r}^*(D)$ for all $D > 0$.

The bank's problem is analogous to Equation (16), except that the expected value of $\phi(D, \tilde{r})$ is subtracted off, and bankruptcy when $\tilde{x} = \bar{x}$ is allowed to have non-zero probability:

$$\begin{aligned} \max_{D, \sigma} p & \left(\int_{\bar{r}(D)}^{\infty} [\bar{x} + \tilde{r} - D + \psi] f(\tilde{r}|\sigma) d\tilde{r} - \gamma \right. \\ & \times \left. \int_{\bar{r}^*(D)}^{r^\phi(D)} [(1 + \rho)D - \tilde{r}] f(\tilde{r}|\sigma) d\tilde{r} \right) \\ & + (1 - p) \left(\int_{\underline{r}^*(D)}^{\infty} [\tilde{r} + \psi - D] f(\tilde{r}|\sigma) d\tilde{r} - \gamma \right) \end{aligned}$$

$$\begin{aligned}
& \times \int_{r^*(D)}^{r^{\phi}(D)} [(1 + \rho)D - \tilde{r}] f(\tilde{r}|\sigma) d\tilde{r} \Big) \\
& -1 + (1 + \rho)D, \quad \text{s.t. Equation (10)}. \quad (26)
\end{aligned}$$

5.3. RE-NEGOTIATION PROOF PRECOMMITMENT

In the basic precommitment mechanism, the penalty is imposed even if the penalty itself drives the bank into default. This possibility ('hitting them when they're down') has received a fair amount of comment. Some argue that to modify the penalty whenever it induces financial distress would constitute forbearance, with all the associated negative consequences for regulatory credibility. On the other hand, a penalty that induces default imposes deadweight societal costs *ex post*. This arrangement is clearly not re-negotiation proof: When there is imminent danger of a penalty-induced default, both the bank and the regulator can be made better off if the regulator modifies the penalty to keep the bank solvent. Therefore, in this section we consider an alternative specification for precommitment in which the penalty function is modified to avoid triggering default.

The renegotiation-proof penalty structure analogous to (23) has the following form:

$$\phi(D, \tilde{r}, \tilde{x}) = \begin{cases} 0 & \text{if } (\tilde{r} - 1) + K(D) > 0 \\ \min\{\gamma[(1 - \tilde{r}) - K(D)], \max[\tilde{r} + \tilde{x} - D, 0]\} & \text{otherwise.} \end{cases} \quad (27)$$

The penalty function in (27) is the same as in (23) *unless* the penalty would induce default, in which case the regulator imposes the maximum penalty that does not trigger default. From the bank shareholders' perspective, the only change is that they keep the bank franchise value unless $\tilde{r}$ falls below $D - \tilde{x}$. (Here, we are assuming that if $\tilde{r} + \tilde{x} - \phi$ exactly equals D , the bank is not technically in default.) Equation (26) must be changed to:

$$\begin{aligned}
\max_{D, \sigma} p & \left(\int_{\tilde{r}(D)}^{\infty} [\tilde{x} + \tilde{r} - D] f(\tilde{r}|\sigma) d\tilde{r} - \gamma \int_{\tilde{r}(D)}^{r^{\phi}(D)} [(1 + \rho)D - \tilde{r}] f(\tilde{r}|\sigma) d\tilde{r} \right) \\
& + (1 - p) \left(\int_{r^*(D)}^{\infty} [\tilde{r} - D] f(\tilde{r}|\sigma) d\tilde{r} - \gamma \int_{r^*(D)}^{r^{\phi}(D)} [(1 + \rho)D - \tilde{r}] f(\tilde{r}|\sigma) d\tilde{r} \right) \\
& + \psi(p + (1 - p)(1 - F(D|\sigma)) - 1 + (1 + \rho)D, \quad \text{s.t. Equation (10)}. \quad (28)
\end{aligned}$$

6. Comparison of Regulatory Mechanisms: Numerical Results

In this section, we examine the allocations of risk and capital that can be implemented by the alternative mechanisms. We approach this inquiry in two ways. First,

we choose parameters randomly, numerically solve the model at each parameterization, and analyze the performance of each type of regulation. The purpose of this exercise is to seek fairly robust general results. Second, we examine in depth a plausible 'baseline' parameterization, along with perturbations away from the baseline.

6.1. GENERAL RESULTS FROM RANDOM NUMERICAL EXERCISES

We select randomly 250 parameter combinations from the following uniform distributions: $\sigma_0 \in [0.001, 1]$; $\bar{\sigma} \in [10, 15]$; $\mu \in [1.001, 2]$; $\rho \in [0.001, 1]$; $\beta \in [0.001, 3]$; $p \in [0.01, 0.99]$; $\psi \in [0, 3]$. Using the results of these experiments, we determine how the various regulatory mechanisms operate, and we evaluate the approaches according to both the social welfare criterion and the bank's objective function.

In all our parameterizations, $\bar{\sigma}$ is big relative to σ_0 . As a result, the portfolio risk chosen by the bank is the critical determinant of social welfare. Stated more formally,

Numerical Result 1. For all of the random parameter draws, an allocation (D_1, σ_0) attains a higher value of the social welfare criterion (12) than an allocation $(D_2, \bar{\sigma})$, for any feasible D_1, D_2 .

That is, for the parameter ranges we consider, the regulator always seeks to deter risk-seeking behavior, even at the cost of excessive capital.

6.1.1. Ex-ante Capital Regulation

First, let us consider the behavior of the bank under the ex-ante capital regulation described in Section 5.1, above. Numerical Result 2 characterizes the function $\sigma^{\text{ex}}(D)$, defined in (22):

Numerical Result 2. For all 250 random parameter draws, we find that

- (i) If $\psi > 0$, there exists a critical level $\bar{D} \geq 0$ such that

$$\sigma^{\text{ex}}(D) = \sigma_0, \quad \forall D \leq \bar{D}, \text{ and}$$

$$\sigma^{\text{ex}}(D) = \bar{\sigma}, \quad \forall D > \bar{D};$$

- (ii) $\bar{D} = 0$ if and only if $\psi = 0$.

- (iii) $\bar{D}$ is strictly increasing in ψ .

Figure 2 illustrates Numerical Result 2. If the regulator imposes stringent capital requirements (i.e., D is 'sufficiently low') the bank behaves prudently, setting $\sigma =$

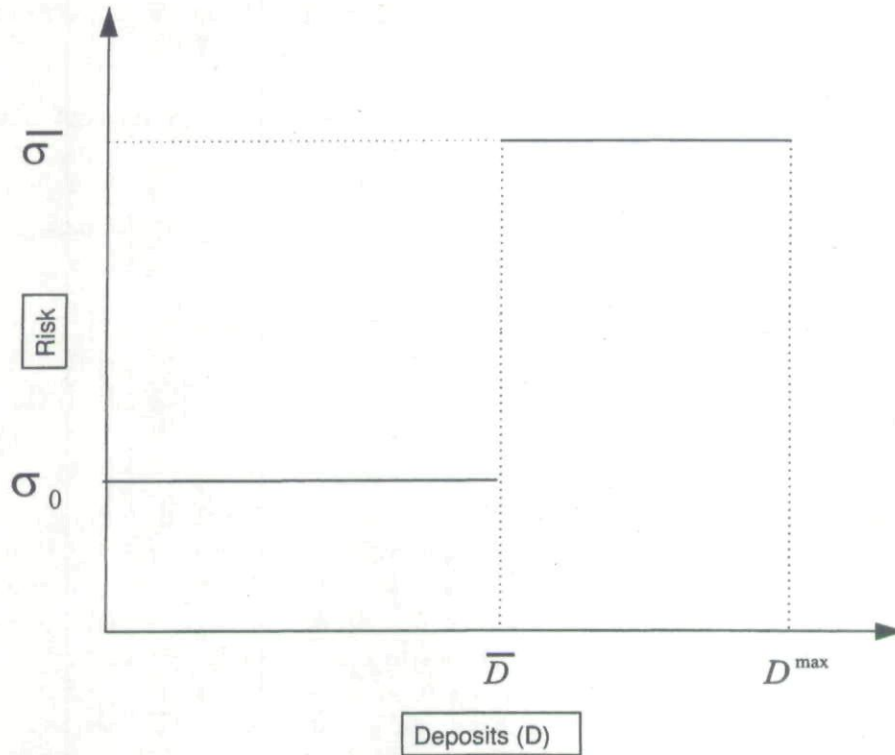


Figure 2. Ex-ante capital regulation.

σ_0 . As these capital requirements are lowered (D is increased), a point is reached beyond which the bank switches from prudent behavior to risk-seeking behavior (i.e., it selects $\bar{\sigma}$). We denote this switching point by $\bar{D}$.¹⁵

Numerical Result 2 enables us to characterize when ex-ante capital requirements can achieve the first-best allocation:

Numerical Result 3. For all 250 random parameter draws, the first-best allocation (D^* , σ_0) is attained by ex-ante capital requirements if and only if $D^* \leq \bar{D}$.

That is, if mandating D^* induces the bank to behave prudently, ex-ante capital requirements can implement the first-best allocation.

6.1.2. State Contingent Regulation: The Precommitment Approach

We now turn to bank behavior under the precommitment mechanism. Let $D^{pc}(\gamma)$ and $\sigma^{pc}(\gamma)$ denote the bank's optimal choices of D and σ under precommitment

¹⁵ If $D^{ex} = \bar{D}$, the bank is indifferent between σ_0 and $\bar{\sigma}$. We assume that the bank chooses the prudent risk level σ_0 in this case.

when the penalty function parameter equals γ . These functions are characterized in Numerical Result 4:

Numerical Result 4. For all random parameter draws, and for both variants of precommitment,

- (i) For fixed $\sigma = \sigma_0$, if $D^{pc}(\gamma) < D^{\max}$, $D^{pc}(\gamma)$ is strictly decreasing in γ .
- (ii) For very low values of p and ψ , $\sigma^{pc}(\gamma) = \bar{\sigma}$, $\forall \gamma$.
- (iii) For p and ψ sufficiently high, there exists $\bar{\gamma} < \infty$ s.t.

$$\sigma^{pc}(\gamma) = \sigma_0, \forall \gamma \geq \bar{\gamma}, \text{ and}$$

$$\sigma^{pc}(\gamma) = \bar{\sigma}, \forall \gamma < \bar{\gamma}.$$

Part (i) of Numerical Result 4 is intuitive: If the bank is already acting prudently, increasing the penalty induces the bank to increase its capital level. (If the bank sets $\sigma = \bar{\sigma}$, we find that the bank also sets $D = D^{\max}$, so D is unresponsive to small changes in γ .) According to part (ii) of Numerical Result 4, there may not exist a finite γ that induces banks to behave prudently. Intuitively, $\gamma = \infty$ would mean that *any* loss to the trading portfolio results in immediate default, with consequent closure of the bank. If the costs of default to the bank insiders are sufficiently small, (because p and ψ are sufficiently small) and the benefits from exploiting the deposit insurance put option are sufficiently large, even this draconian penalty might be insufficient to deter risk-seeking behavior. Barring this case, there exists a critical value of the penalty parameter (denoted $\bar{\gamma}$) above which the bank will set $\sigma = \sigma_0$. That is, the penalty can be set sufficiently high so that an otherwise risk-seeking bank behaves prudently. The only cases we found where did not exist were when $p < 0.4$ and $\psi < 0.3$. These parameterizations represent very poor quality banks: A value of $p < 0.4$ implies that with probability exceeding 60% the value of the loan portfolio would fall below the face value of deposits, leaving the bank vulnerable to trading-induced default. Such a bank would clearly be considered extremely troubled, and would be subject to prompt corrective action under FDICIA.

For all other parameterizations, $\bar{\gamma}$ exists. Figures 3a and 3b illustrate the way precommitment functions. When $\gamma > \bar{\gamma}$, the bank selects the low risk σ_0 . In this region, increasing γ induces the bank to increase capital (reduce D). The location of determines whether precommitment can implement the optimal allocation:

Numerical Result 5. Let γ^* denote the regulator's optimal choice of the precommitment penalty parameter. For all random parameter draws

- (i) If $\bar{\gamma} < \infty$ exists and $D^* \leq D^{pc}(\bar{\gamma})$, precommitment can implement the first-best allocation (D^*, σ_0) , and γ^* is defined by $D^* = D^{pc}(\gamma^*)$.

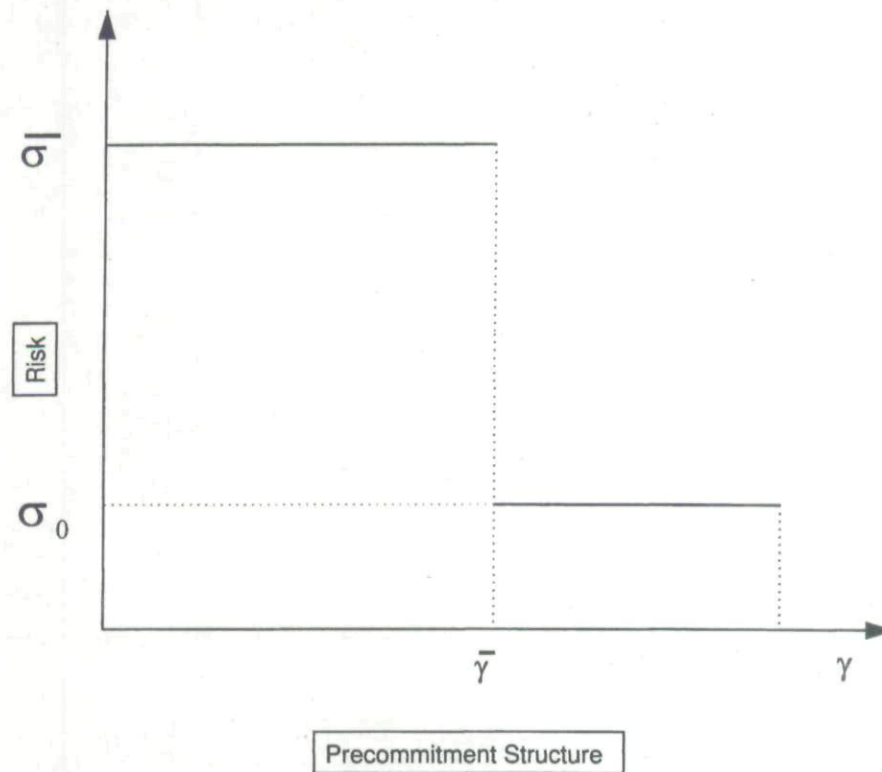


Figure 3a. Precommitment and risk.

- (ii) If $\bar{\gamma} < \infty$ exists and $D^* > D^{pc}(\bar{\gamma})$, precommitment cannot implement the first-best allocation, and $\gamma^* = \bar{\gamma}$.

Numerical Result 6 describes the way γ^* varies with bank characteristics.

Numerical Result 6. For all of the random parameter draws where $\bar{\gamma} < \infty$ exists and where D^* is interior (that is, $D^* < D^{\max}$), γ^* is strictly decreasing in ψ and p for both basic precommitment and re-negotiation proof precommitment.

6.1.3. Comparison of the Regulatory Mechanisms

In this section, we compare the various regulatory approaches according to the social welfare criterion. According to Numerical Result 1, the critical task of bank regulation is to induce banks to behave prudently. According to Proposition 6, ex-ante capital requirements can always do this (although, as we shall see, the requisite capital ratios can be extremely high for low ψ 's). According to Numerical Result 4(ii), precommitment cannot always do this. It follows that, when p and ψ are very low, ex-ante capital requirements dominate precommitment from a social welfare standpoint.

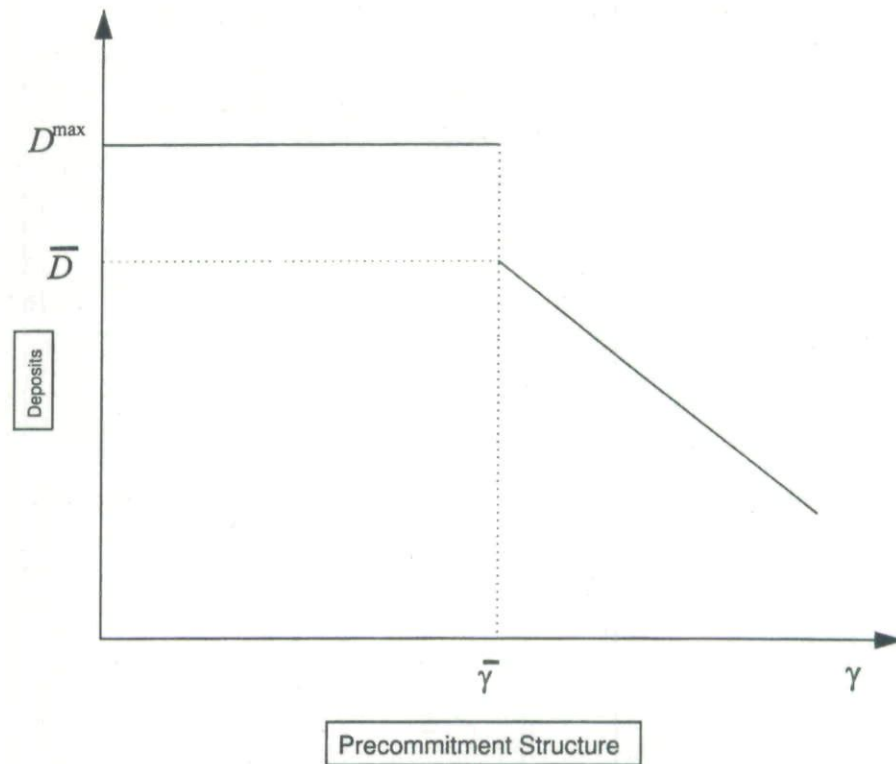


Figure 3b. Precommitment and deposits.

When there exists a finite $\bar{\gamma}$ (in Numerical Result 4(iii)), precommitment can induce prudent behavior. For most such parameterizations, basic precommitment weakly dominates ex-ante capital requirements in the following sense:

Numerical Result 7. For most cases where $\bar{\gamma}$ exists,

- (i) When ex-ante capital regulation can achieve the first-best allocation, the basic precommitment approach can also achieve the first-best allocation.
- (ii) For parameterizations where ex-ante capital regulation fails to achieve the first-best allocation, the optimal implementation of the basic precommitment approach attains a strictly higher value of both the social welfare criterion and the bank's objective function than ex-ante capital regulation.

In particular, of the 250 random parameter draws, we found only two parameterizations where a value $\bar{\gamma}$ exists (so precommitment could induce prudent behavior) but where ex-ante capital requirements attains a strictly higher value of the so-

cial welfare criterion than precommitment.¹⁶ For the remaining parameterizations where $\bar{\gamma}$ exists, precommitment dominates ex-ante capital requirements according to both the social welfare criterion and the bank's objective function. We conclude that precommitment provides a superior regulatory environment unless franchise value or loan portfolio quality are extremely low.

In Numerical Result 8, we find a similar weak domination of basic precommitment by re-negotiation proof precommitment.

Numerical Result 8. For all random parameter draws, where $\bar{\gamma}$ exists (in the sense of Numerical Result 4(iii)) for renegotiation-proof precommitment,

- (i) Whenever basic precommitment can achieve the first-best allocation, re-negotiation proof precommitment approach can also achieve the first-best allocation.
- (ii) For all parameterizations where basic precommitment fails to achieve the first-best allocation, the optimal implementation of re-negotiation proof precommitment attains a strictly higher value of both the social welfare criterion and the bank's objective function than basic precommitment.

According to Numerical Result 8, it is not necessarily bad to reduce the severity of regulatory sanctions for a weak bank. (Gorton and Winton (1995) arrive at a similar conclusion.) However, one should build that contingency explicitly into the regulation. This does not constitute 'forbearance', since the regulator is in no way failing to enforce the regulation as written.

6.2. NUMERICAL RESULTS FOR SPECIFIC PARAMETERIZATIONS

In this section, we display detailed results for a 'baseline' parameterization, and we discuss perturbations away from the baseline. We think of the length of each period as one year. We set $\rho = 0.05$, implying a 5% annual return differential between monetary and non-monetary assets. We set $\mu = 1.04$, which gives the bank's trading portfolio a risk-adjusted premium of 4% over the risk-free rate. To calibrate β , we note that if $\psi = 0$, the socially-optimal bank failure rate equals ρ/β . It would be difficult to justify the optimality of a failure rate above 5% per year, even with zero franchise value, so we set $\beta = \rho/0.05 = 1$. Literally interpreted, this value of β implies that the deadweight loss due to failure resolution equals 100% of the failed bank's shortfall. This high value can be justified if we think of β as a parameter summarizing all of the factors that make bank failure costly from a social welfare standpoint. These factors might include costs of systemic risk, loss of confidence in the banking system, and losses due to the portion of bank deposits that is uninsured.

¹⁶ Both cases involved a low value of p , a low value of ψ , and a very high value of σ_0 (above 0.70).

The remaining parameters in our baseline case are as follows: We set $\sigma_0 = 0.10$ (so the minimum standard deviation for the trading portfolio return is 10%). The value of $\bar{\sigma}$ has little effect on the results, provided it is high enough to approximate infinite risk. We set $\bar{\sigma} = 10$. The exact value of $\bar{x}$ does not affect bank or regulatory choices as long as it is sufficiently large to preclude trading-induced default, including defaults triggered by the penalty in the basic precommitment approach. We set $\bar{x} = 4$. The bank-specific parameters are ψ (franchise value, as a fraction of the ex-ante value of the trading portfolio) and p (the probability of loan portfolio outcomes where default has zero probability). We explore the following ranges: $\psi \in [0, 1.5]$; $p \in [0.2, 0.95]$.

In Figure 4, we fix $p = 0.8$ and we vary ψ , while Figure 5 fixes ψ at 0.5 and varies p . For both figures, Panel A displays the capital ratios $K(D)$ for the regulatory optimum (solid line), the unregulated bank's optimum (dotted line), and the optimal implementation of ex-ante capital regulation (dashed line), while Panel B gives capital ratios for the optimal implementation of the basic precommitment approach (dashed line), and the re-negotiation proof precommitment (dotted line). The regulatory optimum is also displayed as the solid line in Panel B. (Note that the vertical scale differs from the vertical scale in Panel A.) Consider first Figure 4. For all values of ψ below 0.84 the unregulated bank is risk-seeking, choosing zero capital and setting $\sigma = \bar{\sigma}$. For $\psi \geq 0.84$, the unregulated bank sets $\sigma = \sigma_0$ and chooses capital ratios only slightly below the regulatory optimum. In this sense, high franchise-value banks are almost self-regulating. For low values of ψ , ex-ante capital regulation induces prudent behavior, but only by imposing very high capital requirements. For example, when $\psi = 0.1$, the optimal ex-ante capital requirement equals 85% of the trading portfolio. As ψ increases, the optimal capital requirement falls dramatically. For $\psi = 0.69$, ex-ante capital regulation achieves the first-best allocation. For example, with $\psi = 1$, ex-ante capital regulation achieves the first-best allocation with a capital requirement of 14.2%.

Panel B of Figure 4 shows that, for all values of ψ , the basic precommitment specification induces banks to choose $\sigma = \sigma_0$. However, for positive values of ψ less than 0.69, the optimal precommitment specification induces banks to overcapitalize slightly, relative to the regulatory optimum. When $\psi = 0$, precommitment attains the socially optimal capital ratio of zero. Intuitively, with $\psi = 0$ the bank has little value to society. The marginal social cost of mandating positive capital exceeds the marginal social value of reducing the default risk for such a low-value bank. Panel B of Figure 4 also illustrates the weak domination of basic precommitment by re-negotiation proof precommitment. Notice that when the capital level induced by the latter differs from the regulatory optimum, it is always closer to the regulatory optimum than the capital level induced by the basic precommitment approach.

Panels A and B of Figure 5 show that the first-best capital ratio and the capital ratios under the optimal implementations of precommitment fall as p increases. In this example the optimal ex-ante capital requirement (dashed line in Panel A) is unaffected by p . The reason for this is that $\bar{D}$ is invariant to p . (This follows

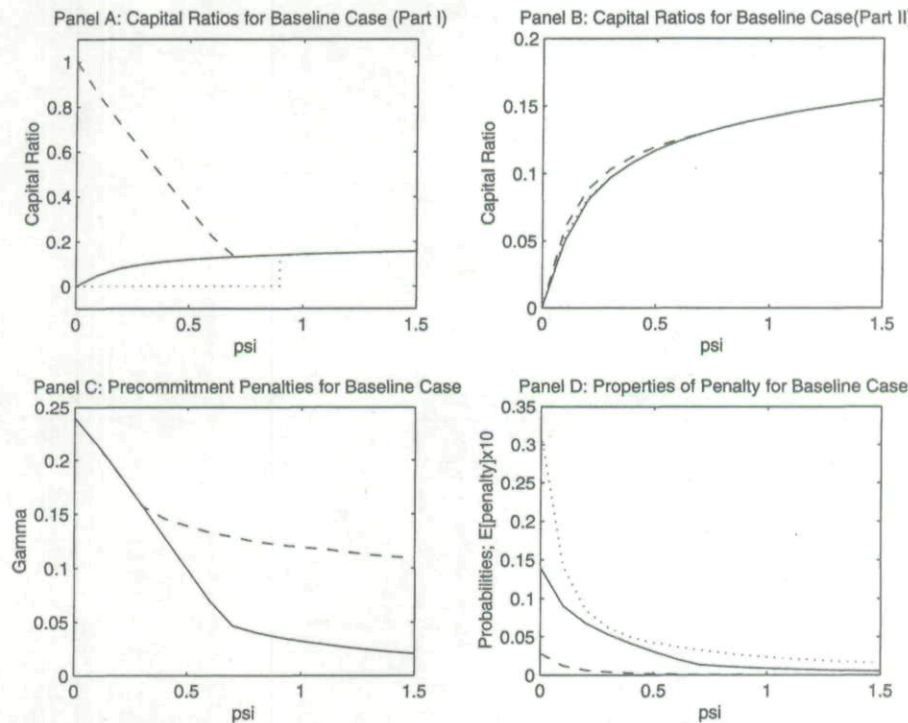


Figure 4. The baseline parameterization. This figure displays the implications of the model under the following parameterization: $\sigma_0 = 0.1$; $\mu = 1.04$; $\rho = 0.05$; $\beta = 1.0$; $\bar{\sigma} = 10$; $\bar{x} = 4$; $p = 0.8$. In each panel, the horizontal axis varies ψ from 0 to 1.5. Panel A plots the capital ratio (defined as $1 - D/D^{\max}$) under the first-best allocation (solid line), unregulated bank (dotted line), and optimal ex-ante capital regulation (dashed line). Panel B plots the capital ratio under the optimal basic precommitment regime (dashed line), the optimal re-negotiation proof precommitment regime (dotted line), along with the capital ratio for the first-best allocation (solid line, identical to the solid line in Panel A). Panel C gives the values of the optimal choice of γ in the basic precommitment regime (solid line) and under the re-negotiation proof precommitment regime (dashed line). Panel D gives the following properties of the penalty function under the basic precommitment regime (with γ chosen optimally, as in Panel C): probability that the penalty is imposed (dotted line); probability that the penalty triggers default, conditional on $x = 0$ (dashed line); and the expected value of the penalty, conditional on the penalty being imposed, scaled by a factor of 10 (solid line). The last statistic is scaled to make the units comparable to the other two plots in this panel.

from Equation (19).) For the parameterization of Figure 5, D^{ex} always equals $\bar{D}$, so ex-ante capital requirements fail to achieve the first-best allocation. In contrast, basic precommitment achieves the first-best for $p \geq 0.9$, and re-negotiation proof precommitment does so for $p \geq 0.5$.

An important question is whether precommitment requires extremely harsh penalties. Panels C and D of Figures 4 and 5 provide some answers. Panel C plots γ^* under the basic precommitment mechanism (solid line) and re-negotiation proof precommitment (dashed line). Consider first Figure 4. When $\psi = 0$, $\gamma^* =$

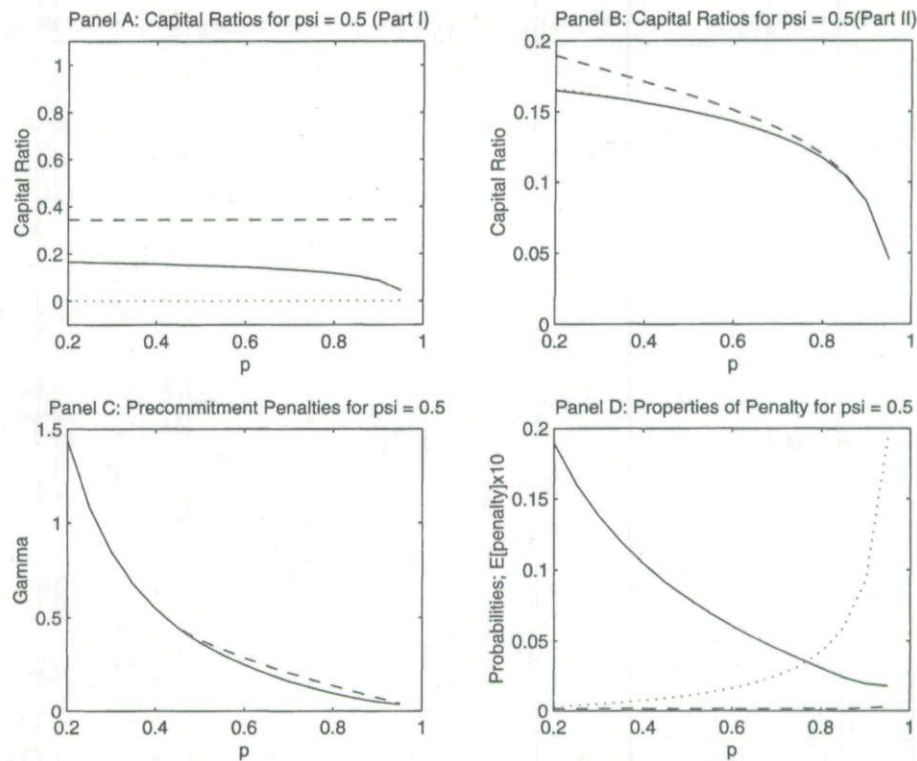


Figure 5. Parameter p varies from 0.2 to 0.95. This figure is analogous to Figure 4 except that ψ is fixed at 0.5, and p is varied (along the horizontal axis) from 0.2 to 0.95. The parameterization is as follows: $\sigma_0 = 0.1$; $\mu = 1.04$; $\rho = 0.05$; $\beta = 1.5$; $\bar{\sigma} = 10$; $\bar{x} = 4$; $\psi = 0.5$; $p \in [0.2, 0.95]$.

0.241, implying a fine slightly less than 25% of the portfolio loss in excess of the pre-committed amount. As franchise value increases, γ^* decreases rapidly. The magnitude of these fines are non-trivial, but they are not unreasonably large. Furthermore, when we perturb parameters $\{\sigma, \rho, \mu, \beta\}$ away from our baseline (keeping $p = 0.8$, as in Figure 4), γ^* changes only modestly. For example, when we double portfolio risk by increasing σ_0 from 0.10 to 0.20, the value of γ^* for $\psi = 0$ rises from 0.24 to 0.35 under both precommitment mechanisms. Similarly, a 50% increase in the cost of default (from $\beta = 1.0$ to $\beta = 1.5$) increases γ^* from 0.24 to 0.30 when $\psi = 0$.

According to Figure 4, Panel C, γ^* for re-negotiation proof precommitment is considerably higher than that for the basic precommitment approach, except for the lowest levels of ψ . For example, when $\psi = 1$, the optimal implementation of basic precommitment uses a γ^* of only 3.2%, while re-negotiation proof precommitment sets $\gamma^* = 11.9\%$. This is no surprise: For a given γ , re-negotiation proof precommitment imposes a less onerous penalty than basic precommitment, so a higher value of γ is needed to have the same effect on bank incentives.

Figure 5, Panel C, shows how sensitive the optimal precommitment penalty is to the quality of the loan portfolio. When $p = 0.2$ (the poorest quality loan portfolio we examine), the optimal precommitment penalty requires $\gamma \approx 150\%$. However, we regard this value of p as unrealistically low. When p is increased to 0.95 (implying that, even with a total loss to the trading portfolio, the bank will be solvent 95% of the time), the optimal values of γ are only 3.8% with basic precommitment and 4.4% with re-negotiation proof precommitment.

Panel D of Figures 4 and 5 illustrates the impact of the penalty function. In Figure 4, the probability that the penalty is imposed (dotted line) ranges from 31.5% (for $\psi = 0$) down to 2% (for the higher values of ψ). Conditional on the penalty being imposed, the average penalty (solid line) ranges from 0.014 for $\psi = 0$ (that is, 1.4% of the value of the trading portfolio) to less than 0.0009 (0.09% of the value of the trading portfolio) for ψ 's above 1.0. As with Panel C, these results do not imply a particularly burdensome penalty, except possibly for the lowest franchise values. Figure 5 shows that the probability that the penalty is imposed increases sharply for the highest values of p , even as the expected magnitude of the penalty falls. That is, a very high quality of the loan portfolio reduces the expected social costs associated with poor performance in the trading portfolio, so the regulator views precommitment violations as less dangerous.

Finally, we have chosen a value of $\bar{x}$ sufficiently large that the precommitment penalty never triggers default when $\bar{x} = \bar{x}$. The dashed line in Panel D gives the probability that the penalty triggers default when $\bar{x} = 0$. According to Figure 4, this happens 2.9% of the time for $\psi = 0$. However, this probability falls off rapidly as ψ increases, becoming negligible (less than 0.5%) for ψ 's over 0.25. According to Figure 5, this probability is unaffected by changes in p .

7. Conclusions

Our model provides a precise characterization of the conflicts between regulators, who seek to enhance public welfare, and banks who act as private value maximizers. It implies that banks either behave prudently or seek maximal uncompensated risk, but do not choose intermediate risk levels. In this environment, the most important role of the regulator is to induce the risk-seeking banks to behave prudently, by switching to a low level of risk. While traditional ex-ante capital requirements can induce prudent behavior even for the most risk-seeking of banks, the needed capital ratios are often inefficiently high. The precommitment approach cannot always induce prudent behavior, especially with the poorest quality banks. This result supports the provision in the Board of Governors' 1995 proposal that prohibits weak banks from using precommitment to set capital requirements. For most other bank types, precommitment is generally preferable to ex-ante capital regulation, both according to the social welfare criterion and according to the bank's objective function. Furthermore, the optimal precommitment penalty levels, while non-trivial, are not excessively large.

Given the simplicity of the precommitment mechanism, we find these results striking. They support the work currently being done to explore precommitment as an alternative to the current structure of capital regulation. More generally, they suggest that substantial efficiency gains can be obtained by moving away from ex-ante capital requirements towards structures that seek to modify the ex-ante behavior of regulated firms via state contingent rewards and penalties.

Appendix. Proofs and Lemmas

The following Lemma is used in the proof of Proposition 1.

LEMMA A.1.

- (i) $F_\sigma(r|\sigma) > 0, \forall r < \mu$.
- (ii) $\int_0^r F_\sigma(\tilde{r}|\sigma) d\tilde{r} > 0$.

Proof of (i): The cumulative log-normal distribution function can be written:

$$F(r|\sigma) = \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{1}{\sqrt{2}} \frac{\ln(r) - \mu_n}{\sigma_n} \right) \right], \quad (\text{A.1})$$

where

$$\begin{aligned} \operatorname{erf}(x) &\equiv \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2/2} dt, & \mu_n &\equiv \frac{1}{2} \log \left(\frac{\mu^4}{\mu^2 + \sigma^2} \right), \\ \text{and } \sigma_n &\equiv \sqrt{\log \left(\frac{\mu^2 + \sigma^2}{\mu^2} \right)}. \end{aligned} \quad (\text{A.2})$$

In (A.1), μ_n and σ_n are the mean and variance of the normal distribution to which the log-normal corresponds. Differentiating $F(r|\sigma)$ with respect to σ , one obtains:

$$\begin{aligned} F_\sigma(r|\sigma) &= \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{\log(r) - \mu_n}{\sigma_n} \right)^2} \left(\frac{\sigma}{\sigma_n(\mu^2 + \sigma^2)} \right) \\ &\times \left[1 - \frac{\log(r) - \mu_n}{\sigma_n^2} \right]. \end{aligned} \quad (\text{A.3})$$

Equation (A.3) implies that

$$\begin{aligned} \operatorname{sign}(F_\sigma(r|\sigma)) &= \operatorname{sign} \left[1 - \frac{\log(r) - \mu_n}{\sigma_n^2} \right] = \operatorname{sign}[\sigma_n^2 - \log(r) + \mu_n] \\ &= \operatorname{sign} \left[\log \left(\frac{\mu^2 + \sigma^2}{\mu^2} \right) - \log(r) + \frac{1}{2} \log \left(\frac{\mu^4}{\mu^2 + \sigma^2} \right) \right] \\ &= \operatorname{sign} \left[\log \left(\frac{\sqrt{\mu^2 + \sigma^2}}{r} \right) \right] > 0, \end{aligned} \quad (\text{A.4})$$

where the final inequality is implied by $r < \mu$.

Proof of (ii): With the log-normal distribution, an increase in σ without changing μ represents a mean-preserving spread. This immediately implies the result. (See Rothschild and Stiglitz, 1970.) ■

Proof of Proposition 1. To prove the proposition, it is sufficient to show that the left-hand side of (15) is strictly negative for all $D > 0$. To that end, we use integration by parts to write the left-hand side of (15) as

$$-(1-p) \left[\psi F_\sigma(D|\sigma) + \beta \int_0^D F_\sigma(\tilde{r}|\sigma) d\tilde{r} \right]. \quad (\text{A.5})$$

Since $D \leq D^{\max} \equiv 1/(1+\rho) < 1 \leq \mu$, Lemma A.1 implies that both terms in the bracketed expression in (A.5) are strictly positive. This implies the conclusion of the Proposition.¹⁷ ■

Proof of Proposition 2. Proof of (i): Note that $F(0|\sigma) = f(0|\sigma) = 0$, $\forall \sigma$. Since $\rho > 0$, the left-hand side of (13) is strictly positive at $D = 0$, implying that $D = 0$ is never an optimum. Proof of (ii) and (iii): D^* is interior, so we can totally differentiate the first-order condition (13), first with respect to ψ , second with respect to β , third with respect to ρ , and fourth with respect to p , to obtain

$$\begin{aligned} -(1-p) \left[\beta f(D|\sigma) \frac{dD}{d\psi} + f(D|\sigma) + \psi f'(D|\sigma) \frac{dD}{d\psi} \right] &= 0 \\ -(1-p) \left[\beta f(D|\sigma) \frac{dD}{d\beta} + F(D|\sigma) + \psi f'(D|\sigma) \frac{dD}{d\beta} \right] &= 0 \\ 1 - (1-p) \left[\beta f(D|\sigma) \frac{dD}{d\rho} + \psi f'(D|\sigma) \frac{dD}{d\rho} \right] &= 0 \\ \beta F(D|\sigma) + \psi f(D|\sigma) - (1-p) \left[\beta f(D|\sigma) \frac{dD}{dp} + \frac{\psi f'(D|\sigma) dD}{dp} \right] &= 0, \end{aligned} \quad (\text{A.6})$$

implying

$$\begin{aligned} \frac{dD}{d\psi} &= -\frac{f(D|\sigma)}{\beta f(D|\sigma) + \psi f'(D|\sigma)} < 0, \\ \frac{dD}{d\beta} &= -\frac{F(D|\sigma)}{\beta f(D|\sigma) + \psi f'(D|\sigma)} < 0, \end{aligned} \quad (\text{A.7})$$

¹⁷ Note that the only property of the log-normal distribution used in this proof is that $F_\sigma(r) > 0$, $\forall r < \mu$. This property holds for most distribution functions typically used to model portfolio returns, so the proposition is more general than for the log-normal returns assumed here. Indeed, if the portfolio return distribution did not satisfy this property, then value-at-risk might be *decreasing* in σ , in which case portfolio variance would arguably be a poor measure of portfolio risk. Also note that, in Proposition 1, we exclude the case of $D = 0$ because, in that case, there is zero probability of default, so the social welfare function is unaffected by σ .

$$\frac{dD}{d\rho} = \frac{1}{\beta f(D|\sigma) + \psi f'(D|\sigma)} > 0,$$

$$\frac{dD}{dp} = \frac{\beta F(D|\sigma) + \psi f(D|\sigma)}{(1-p)[\beta f(D|\sigma) + \psi f'(D|\sigma)]} > 0,$$

where the final inequalities are implied by second-order condition (14). ■

The following Lemma is required for the proof of Proposition 3.

LEMMA A.2. For any fixed $r < \mu$,

$$\frac{\partial}{\partial \sigma} \left[\frac{\int_0^r \tilde{r} f_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} \right] < 0. \quad (\text{A.8})$$

Proof. The numerator of the object in brackets in (A.8) is

$$\begin{aligned} \int_0^r \tilde{r} f_\sigma(\tilde{r}|\sigma) d\tilde{r} &= -\frac{\mu}{\sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{-\log(r) + \sigma_n^2 + \mu_n}{\sigma_n} \right)^2} \left(\frac{\sigma}{\sigma_n(\mu^2 + \sigma^2)} \right) \\ &\times \left[1 - \frac{-\log(r) + \sigma_n^2 + \mu_n}{\sigma_n^2} \right]. \end{aligned} \quad (\text{A.9})$$

The denominator of the object in brackets in (A.8) is given by equation (A.3). Taking the ratio of (A.9) to (A.3), substituting definitions from (A.2) and doing a bit of algebra, one obtains

$$\begin{aligned} &\frac{\int_0^r \tilde{r} f_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} \\ &= r \left[\frac{\log(\mu^2) - \frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)}{\frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)} \right]. \end{aligned} \quad (\text{A.10})$$

The derivative of the right-hand side of (A.10) has the same sign as

$$\begin{aligned} &\frac{\partial}{\partial \sigma} \left[\frac{\log(\mu^2) - \frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)}{\frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)} \right] \\ &= \frac{-2\sigma(\log(\mu) - \log(r))}{\left(\frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r) \right)^2 (\mu^2 + \sigma^2)} < 0, \end{aligned} \quad (\text{A.11})$$

where the final inequality follows from $r < \mu$. ■

Proof of Proposition 3. To prove the proposition, it is sufficient to show that (20) cannot hold if (19) holds at equality. Equation (19) holding at equality can be written:

$$(1 - p) \left[(D - \psi) F_{\sigma}(D|\sigma) - \int_0^D \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} \right] = 0, \quad (\text{A.12})$$

where we use

$$\int_r^{\infty} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} = - \int_0^r f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}$$

and

$$\int_r^{\infty} \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} = - \int_0^r \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}.$$

(The latter result uses the fact that an increase in σ does not affect the mean of the distribution for a mean-preserving spread.) Similarly, second-order condition (20) can be written

$$(1 - p) \left[(D - \psi) F_{\sigma\sigma}(D|\sigma) - \int_0^D f_{\sigma\sigma}(\tilde{r}|\sigma) d\tilde{r} \right] < 0. \quad (\text{A.13})$$

Let us assume towards a contradiction that, for some (D, σ) combination, (A.12) and (A.13) both hold. Lemma A.1 (along with $D \leq D^{\max} < 1 \leq \mu$, implied by (12)) implies that $F_{\sigma}(D) > 0$, so (A.12) implies

$$D - \psi = \frac{\int_0^D \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{F_{\sigma}(D|\sigma)}. \quad (\text{A.14})$$

Substituting (A.14) into (A.13), and rearranging, we obtain

$$\int_0^D \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} \cdot F_{\sigma\sigma}(D|\sigma) - F_{\sigma}(D|\sigma) \cdot \int_0^D \tilde{r} f_{\sigma\sigma}(\tilde{r}|\sigma) d\tilde{r} < 0. \quad (\text{A.15})$$

Equation (A.15) can only hold if

$$\frac{\partial}{\partial \sigma} \left[\frac{\int_0^D \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{F_{\sigma}(D|\sigma)} \right] > 0. \quad (\text{A.16})$$

However, in Lemma A.2 (proved earlier), it is demonstrated that, for any fixed $r < \mu$,

$$\frac{\partial}{\partial \sigma} \left[\frac{\int_0^r \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{F_{\sigma}(r|\sigma)} \right] < 0. \quad (\text{A.17})$$

Equations (A.16) and (A.17) provide the needed contradiction. ■
The following Lemma is required for the proof of Proposition 4.

LEMMA A.3. For all $r < \mu$,

$$\frac{\partial}{\partial r} \left[\frac{\int_0^r F_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} \right] > 0.$$

Proof.

$$\begin{aligned} \frac{\int_0^r F_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} &= r - \frac{\int_0^r \tilde{r} f_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} \\ &= r \left[1 - \frac{\log(\mu^2) - \frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)}{\frac{1}{2} \log(\mu^2 + \sigma^2) - \log(r)} \right], \end{aligned} \quad (\text{A.18})$$

where the first equality in (A.18) uses integration by parts, and the second equality uses Equation (A.11). Differentiating the right-hand side of (A.18) with respect to r , one obtains

$$\begin{aligned} \frac{\partial}{\partial r} \left[\frac{\int_0^r F_\sigma(\tilde{r}|\sigma) d\tilde{r}}{F_\sigma(r|\sigma)} \right] &= \left[1 + \frac{1}{-\log(r) + \frac{1}{2} \log(\mu^2 + \sigma^2)} \right] \\ &\times \left[\frac{1}{-\log(r) + \frac{1}{2} \log(\mu^2 + \sigma^2)} \right] \left[\frac{1}{2} \log \left(\frac{(\mu^2 + \sigma^2)^2}{\mu^4} \right) \right] > 0, \end{aligned}$$

where the inequality follows from $r < \mu$. ■

Proof of Proposition 4. Proof of (i): Note that $F(0|\sigma) = f(0|\sigma) = 0, \forall \sigma$. Since $\rho > 0$, the left-hand side of (17) evaluated at $D = 0$ equals $\rho > 0$, implying that $D = 0$ is never an optimum.

Proof of (ii): We show that there exists $\bar{D}$ such that the second derivative of the bank's objective function with respect to D (which is the left-hand side of (18)) is nonpositive for $D \leq \bar{D}$ and is positive for $D > \bar{D}$. Using

$$f(r|\sigma) = \frac{1}{\sqrt{2\pi}\sigma_n r} e^{-\frac{1}{2} \left(\frac{\log(r) - \mu_n}{\sigma_n} \right)^2}$$

and

$$f'(r|\sigma) = -f(r|\sigma) \frac{1}{\sigma_n^2 r} (\log(r) - \mu_n + \sigma_n^2),$$

one can evaluate the left-hand side of (18) as:

$$(1-p)f(D|\sigma) \frac{1}{\sigma_n^2 r} [\psi \log(D) + \sigma_n^2 D + \psi(\sigma_n^2 - \mu_n)].$$

The object in brackets is strictly increasing in D . In particular, there exists a unique $\bar{D} > 0$ such that $[\psi \log(\bar{D}) + \sigma_n^2 \bar{D} + \psi(\sigma_n^2 - \mu_n)] = 0$. For $D < \bar{D}$ the left-hand side of (18) is nonpositive; for $D > \bar{D}$ the left-hand side of (18) is positive.

Proof of (iii): According to part (ii) of this proposition, there exists at most one interior solution with respect to D . Therefore, part (iii) of the proposition can be proved by totally differentiating the interior solution case in equation (17), first with respect to ψ , and second with respect to ρ . The derivatives $\partial D / \partial \psi$ and $\partial D / \partial \rho$ can then be signed by using the second-order condition (17).

Proof of (iv): If $\psi = 0$, the left-hand side of (17) equals $\rho + F(D|\sigma) > 0$, so $D = D^{\max}$. To prove that the banks optimal $\sigma = \bar{\sigma}$ in this case, note that the left-hand side of (19) can be written:

$$\begin{aligned} (1-p) \left[(D-\psi) F_{\sigma}(D|\sigma) - \int_0^D \tilde{r} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r} \right] \\ = (1-p) \left[\int_0^D F_{\sigma}(\tilde{r}|\sigma) d\tilde{r} - \psi F_{\sigma}(D|\sigma) \right], \end{aligned} \quad (\text{A.19})$$

where the equality in (A.19) uses integration by parts. If $\psi = 0$, the right-hand side of (A.19) is positive.

Proof of (v): According to Lemma A.3, for any given σ ,

$$\frac{\partial}{\partial r} \left[\frac{\int_0^r F_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{F_{\sigma}(r|\sigma)} \right] > 0, \quad \forall r < \mu.$$

Therefore,

$$\max_{D \in [0, D^{\max}], \sigma \in (\sigma_0, \bar{\sigma})} \left\{ \frac{\int_0^D F_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{F_{\sigma}(D|\sigma)} \right\} < \infty. \quad (\text{A.20})$$

For any ψ greater than the object in brackets on the left-hand side of (A.20), the left-hand side of (19) is strictly negative (see Equation (A.19)), implying that the optimal $\sigma = \sigma_0$. ■

Proof of Proposition 5. Under the conditions of this Proposition, (and using Propositions 2(i) and 4(i)) both the regulator's optimal deposit level, D^* , and the bank's optimal deposit level (denoted D^B) represent interior solutions to the respective agent's first-order conditions (13) and (17). Condition (21) ensures that this solution for the regulator's problem is unique. Note that, for any given D , the left-hand side of (17) lies strictly above the left-hand side of (13). Note also that, according to Propositions 2(i) and 4(i), the left-hand side of both (13) and (17) are strictly positive. Finally, Proposition 4(ii) states that the left-hand side of (17) attains a value of zero at no more than one value of D . These facts together imply that if both D^* and D^B are interior, then $D^B > D^*$. ■

Proof of Proposition 6. According to Equation (19), the bank chooses $\sigma^{\text{ex}}(D^{\text{ex}}) = \sigma_0$ if

$$\frac{\int_0^{D^{\text{ex}}} (D^{\text{ex}} - \tilde{r}) f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{\psi \int_0^{D^{\text{ex}}} f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}} < 1,$$

for $\sigma = \sigma_0, \bar{\sigma}$. It is sufficient to demonstrate that

$$\lim_{D \downarrow 0} \frac{\int_0^D (D - \tilde{r}) f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}}{\psi \int_0^D f_{\sigma}(\tilde{r}|\sigma) d\tilde{r}} = 0, \quad \forall \sigma. \quad (\text{A.21})$$

For the log normal distribution, $f_{\sigma}(0|\sigma) = 0$, so the ratio on the left-hand side of Equation (A.21) is an indeterminate form. Repeated application of l'Hôpital's rule verifies Equation (A.21). ■

References

- Berger, A., Herring, R., and Szego, G. (1995) The role of capital in financial institutions, *Journal of Banking and Finance* **19**, 393–430.
- Bliss, R. (1995) Risk-based bank capital: issues and solutions, *Economic Review*, Federal Reserve Bank of Atlanta **80**(5), 32–40.
- Diamond, D. and Dybvig, P. (1983) Bank runs, deposit insurance, and liquidity, *Journal of Political Economy* **91**, 401–419.
- Gennotte, G. and Pyle, D. (1991) Capital controls and bank risk, *Journal of Banking and Finance* **15**, 805–824.
- Giammarino, R., Lewis, T., and Sappington, D. (1993) An incentive approach to banking regulation, *Journal of Finance* **48**, 1523–1542.
- Gorton, G. and Winton, A. (1995) *Bank Capital Regulation in General Equilibrium*, Northwestern University Working Paper.
- John, K., John, T. A., and Senbet, L. W. (1991) Risk-shifting incentives of depository institutions: A new perspective on federal deposit insurance reform, *Journal of Banking and Finance* **15**, 895–915.
- John, K., Saunders, A., and Senbet, L. W. (1996) *A Theory of Bank Regulation and Management Compensation*, manuscript.
- Laffont, J. and Tirole, J. (1993) *A Theory of Incentives in Procurement and Regulation*, The MIT Press, Cambridge, MA.
- Keeley, M. and Furlong, F. (1990) A re-examination of mean-variance analysis of bank capital regulation, *Journal of Banking and Finance* **14**, 69–84.
- Kupiec, P. and O'Brien, J. (1995a) Model alternative, *Risk* **8**, 37–40.
- Kupiec, P. and O'Brien, J. (1995b) *A Pre-Commitment Approach to Capital Market Requirements for Market Risk*, Working Paper 95-36, Federal Reserve Board of Governors.
- Kupiec, P. and O'Brien, J. (1997) *Regulatory Capital Requirements for Market Risk and the Pre-Commitment Approach*, manuscript, Federal Reserve Board of Governors.
- Marcus, A. (1984) Deregulation and bank financial policy, *Journal of Banking and Finance* **8**, 557–565.
- Petersen, M. and Rajan, R. (1994) The benefits of lending relationships: Evidence from small business data, *Journal of Finance* **49**(1), 3–37.

- Prescott, E. S. (1997), The pre-commitment approach in a model of regulatory banking capital, *Economic Quarterly*, Federal Reserve Bank of Richmond **83**(1), 23–50.
- Ritchkin, P., Thomson, J., DeGennaro, R., and Li, A. (1993) On flexibility, capital structure and investment decisions for the insured bank, *Journal of Banking and Finance* **17**, 1133–1146.
- Rothschild, M. and Stiglitz, J. (1970) Increasing risk: a definition, *Journal of Economic Theory* **2**, 225–243.

Copyright of European Finance Review is the property of Springer Science & Business Media B.V.. The copyright in an individual article may be maintained by the author in certain cases. Content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.