

Client Discretion, Switching Costs, and Financial Innovation

Sugato Bhattacharyya

University of Michigan Business School

Vikram Nanda

University of Michigan Business School

We analyze the incentives of investment banks to develop innovative products. We show that client characteristics and market structure affect these incentives significantly. Investment banks with larger market shares have greater incentives to innovate and smaller banks are likely to share their innovations with the largest bank. Innovation incentives increase in volatile environments and regulatory scrutiny actually encourages loophole exploitation activity. Our predictions are consistent with stylized facts and the analysis has broad testable implications for innovative activity in other markets similarly characterized by a lack of comprehensive protection for intellectual property rights, for example, the software industry.

A large number of new financial products have been brought to market in recent years. This profusion of new products has been attributed to a variety of causes: volatility in interest rates, floating exchange rates, regulatory policy and taxes, among others.¹ A leading role in the development and introduction of these innovative products has been played by investment banking firms. In this article we analyze the incentives of investment banks to develop new securities and the impact of competing banks and client characteristics on the innovation process.

It is well recognized that the incentive to innovate is diluted unless an innovator can prevent competitors from freely imitating its innovation.² Such imitation problems may be particularly acute for financial innovations. On the one hand, the costs of security innovation, including those of product development, marketing, and legal expenses, can be substantial. Tufano (1989), for example,

We thank Franklin Allen, Mark Bagnoli, Thomas Chemmanur, Ronen Israel, Francine LaFontaine, Rich Romano, Raj Singh, Anjan Thakor, and seminar participants at the University of Michigan, Georgetown University, University of Miami, Pennsylvania State University, University of Florida, ISI Delhi, IIM Calcutta, Calcutta University, Michigan State University, SFS Meetings at UT Austin 1998, and the Western Finance Association Meetings 1997 for helpful comments. The article has benefited significantly from suggestions provided by the editors, Sheridan Titman, Ravi Jagannathan, and Maureen O'Hara, and two anonymous referees. The authors retain responsibility for errors. Address correspondence to Sugato Bhattacharyya, University of Michigan Business School, Ann Arbor, MI 48109-1234, or e-mail: sugato@umich.edu.

¹ See Finnerty (1992), Miller (1986), Van Horne (1985), Ross (1989), and Allen and Gale (1994), among others, for a discussion of the causes underlying many recent financial innovations.

² Tirole (1988) provides a nice introduction of the literature on research and development.

estimates such costs to be between \$0.5 million and \$5 million. Yet financial products cannot be patented and all details of a new security become publicly available once the offering is filed with the SEC.³ Hence innovating banks may have only a limited “first-mover” advantage before rival banks can offer similar products and compete away any monopoly profits.⁴ As a report in the May 1986 edition of *Institutional Investor* notes, “Quite apart from the proven value of copying a successful product, the time and trouble it takes to get a new security through the regulatory maze provides plenty of incentives for bankers to knock off their competitors’ creations.” Accordingly our model assumes that once a new security is issued by an innovating bank, rival banks can imitate and offer the same product one period later.

In our model, an investment bank decides to develop an innovative financial product if its expected revenues exceed the required expenditure to develop and market the product. The lack of patent protection leaves the bank with only a single period to recover its costs. Revenues are also affected by the market share of the innovating bank among potential clients for the new security. Following the recent literature on relationship banking, we assume that investment banks develop valuable relationships with client firms in the course of providing services.⁵ Hence a firm will switch from its bank to the innovating bank only if the benefit exceeds the cost associated with switching.⁶ Clients are also assumed to have the discretion to postpone the adoption of innovative products. When the cost of delayed adoption is small, the amount a client can be charged for a new product is limited since she can always wait until competition from rival banks drives down the price.

This article analyzes the effect of these and other factors on expected revenues from innovation. Of particular interest is the role of the innovator’s market share in its decision to develop innovative securities. We show that a larger market share allows the innovating bank to derive greater revenues from a given innovation and, hence, gives it greater incentives to engage in innovative activity. This prediction is consistent with the patterns of concentration in financial innovation activity documented in the literature.⁷

The model predicts that banks will tend to pursue innovation opportunities in areas where clients face greater costs of delay. Volatility in the economic environment, for example, may make delay more costly and thereby encourage the demand for products tailored to exploit immediate opportunities.

³ Recent court decisions have, however, extended limited protection to some financial innovations. The seminal case is the patent awarded to Merrill Lynch’s Cash Management Account (CMA).

⁴ See Baldwin and Scott (1987) for a survey of the literature on first-mover benefits.

⁵ For models of bank-client relationships, see, for example, James (1992) and Rajan (1992).

⁶ The impact of switching costs on the nature of product market competition has been explored in the industrial organization literature. See, for example, Farrell and Saloner (1987), Matutes and Regibeau (1988), and David and Greenstein (1990).

⁷ See, for instance, Nanda and Yun (1996). The possibility of a positive association between market share and incentives to innovate was first raised by Schumpeter (1943).

Our argument is quite different from the “standard” explanations based on the derived demand for risk management products by firms. As we discuss, the nonpatentability of financial products also sheds some light on why financial innovation activity—which was vigorous in the late 19th and early 20th centuries—hit such a low from the 1930s to more recent times. We argue that the establishment of the SEC and the introduction of comprehensive disclosure regulations diminished the informational advantage of innovators over imitators. Combined with the lack of demand for investment banking services in the post-Depression years, these changes in the competitive environment were likely to have had a drastic effect on the innovation incentives for investment banks.

Tax considerations, or the avoidance of existing regulations, have often been claimed to be important in the development of innovative financial instruments. Miller (1986) takes the view that most financial innovation is triggered by efforts to skirt regulation or reduce taxes—though new securities may sometimes prove viable even after the regulatory or tax incentive is removed. Our analysis indicates that if the regulatory response is predictable, and is likely to be effective in closing the loophole, it may have the perverse effect of increasing the incentives to seek out such opportunities. This is because the anticipation of a regulatory response increases the costs associated with delay and this, in turn, enhances the demand for loophole-based innovations.

Extending the model to multiple periods, we show that a bank may strategically decide to introduce an innovation in phases. This result suggests that a focus on the number of new financial products may overstate the extent of actual financial innovation. The incentives for phased introduction are shown to increase with imitation pressure and, when clients are heterogeneous in their ability to delay, the market share of the innovating bank. In particular, a smaller bank is shown to have incentives to follow a more aggressive introduction strategy in order to induce switching of clients from larger banks. Thus even though a higher market share is associated with greater innovation incentives, a lower market share promotes a more aggressive introduction strategy.

When cooperative arrangements are permissible, nonpatentability tends to limit the situations in which such arrangements may arise. Consistent with observed patterns, we find that smaller innovators will be more likely to share their innovations with rival banks than the other way around. Consequently, the distribution of market shares across banks is important to the financial innovation process and an asymmetric distribution of market shares supports a higher level of innovation activity.

The literature on financial innovation has several strands. First, there exists a large literature that either describes or analyzes the specific advantages of several innovations. Finnerty (1992) presents a detailed analysis of recent financial innovations in corporate securities, while Allen and Gale (1994)

provide an overview of financial innovation activity spanning bank drafts in ancient Mesopotamia through poison-pill securities in modern times. A second strand studies the forces driving financial innovation activity. Thus Van Horne (1985) focuses on risk-sharing imperatives, while Miller (1986) points to the incentives to work around tax and other regulatory roadblocks. Ross (1989) provides yet another rationale by characterizing such innovations as solutions to moral hazard problems. Merton (1992) further develops the risk-sharing arguments in an international context and refutes the argument that financial innovations only result in private benefits at the expense of social welfare. Allen and Gale (1994) explore the incentives for financial innovation in the context of individual securities and for whole markets. While most of their analysis focuses on efficient risk sharing promoted by the completion of markets, they also briefly discuss the impact of the market power of an innovator and the effects of imitation. Their discussion, in the context of quantity competition, does not address issues of bank-client relationships and optimal introduction strategies that are highlighted in this article. Finally, there exists a "folk explanation" for innovative activity in the financial arena relying on a signaling rationale. Under this explanation, costly innovative activity is valuable because it establishes a reputation for competence that helps attract clients. In our approach, only profitable innovative activity is undertaken.

Whether financial innovators succeed in capturing temporary market power has also been empirically explored. Kanemasu, Litzenberger and Rolfo (1986) find that the profits associated with stripping treasury securities declined secularly with the entry of competitors. Their results are in agreement with Van Horne (1985), who points out that we should expect such a pattern of declining profitability with the entry of imitators. Tufano (1989) examines a broader set of innovations and finds, in contrast, no evidence of the innovator charging more prior to the emergence of competition. However, he reports that innovators manage to retain a substantial market share in the new security despite entry by competitors.

The analysis closest to ours is in Boot and Thakor (1997), who explore the differential incentives to innovate between universal and specialized banks. They conclude that universal banks, because of the spillover effects of innovations on their existing lines of business, have less incentive to innovate than specialized banks. Their conclusions are in line with arguments advanced by Arrow (1962) in the context of cost-reducing innovations in a product market. The results in this article, however, do not depend on such spillover effects.

Our analysis has broad implications for innovation activity in all markets characterized by a lack of comprehensive protection for intellectual property rights. A prominent example would be the software industry, where useful features of programs are subject to quick imitation. Our analysis would predict in this case the use of phased introduction strategies resulting in incremental version changes; the persistence of innovation activity among large vendors; the usefulness of selling out innovative shops to the largest vendor;

and the use of “competitive upgrade” prices for users of alternate vendors’ products. Some of our observations also carry over to the case of the toy industry, where successful products also invite rapid imitation by “knock-off” producers; to the fashion industry; and to segments of the popular publishing industry. While the software example shares most of the factors that we focus on, including switching costs, the other examples primarily share the important feature of a costly discretion to delay on the part of consumers.

The article proceeds as follows. Section 1 lays out the basic model and establishes the important role of market share in the provision of innovation incentives. Section 2 discusses the interplay between innovation incentives and regulatory activity, while Section 3 analyzes the optimal introduction strategy of innovations. Section 4 introduces the possibility of cooperative sharing arrangements and examines conditions under which such cooperation is feasible. Section 5 concludes.

1. The Model

1.1 Preliminaries

We assume that several investment banks are engaged in the process of delivering services like security underwriting and provision of advice to client firms. In the process of interacting with clients, or by analysis of market conditions, one of these banks (“the innovator”) realizes that there exist potential benefits from the introduction of an innovative security.⁸ Limiting our analysis to the case of a single innovator allows us to avoid issues associated with competition in the innovation process itself. We argue later that our main results are unaffected by this assumption.

The value of the innovation opportunity to each client is denoted by V and is assumed to be common knowledge.⁹ In order to get the security to market, the innovator has to incur development and marketing costs, denoted by X . The innovator’s alternatives are either to accept this opportunity and proceed to its development, or to reject it and possibly wait for a later one. Under our assumptions, the total value delivered to clients depends on the number of clients who adopt the innovation, while the costs are fixed. This feature can, obviously, be relaxed. While there could, in general, also be economies of scale in the development process, our assumptions allow us to focus on reasons other than such cost effects. For expositional ease, we assume an absence of discounting.

To capture the notion that innovative financial products can be rapidly imitated, we assume that the innovator has a single period in which it is a

⁸ For the sake of convenience, we will sometimes call an innovative product a new security, although it should be understood that the analysis applies to any innovative product or service.

⁹ Our results extend easily to asymmetric and uncertain valuations. However, we do not model sequential adoptions as in Persons and Warther (1997).

monopolist provider. In subsequent periods, rival investment banks enter the market for the new security and compete with the innovator. Further, as we discuss below, the ability of the innovator to extract rents from client firms is constrained by its market share and by the extent to which potential clients can delay their adoption of the new product.

The innovating bank, I , is the primary investment bank for the $\alpha_i \leq 1$ proportion of existing firms.¹⁰ As the primary investment bank, the innovator has a cost advantage, relative to other banks, in the issuance of securities or providing other services to its existing clients. Similarly, it is at a relative cost disadvantage when it seeks to provide services to firms that are clients of other investment banks. We denote the cost of “switching” between investment banks by C_S . This switching cost can be viewed as the cost to the client of making pertinent information available or the cost to the innovator of becoming familiar with the operations of a new client. Consequently, the cost of switching is likely to be dependent on prevailing accounting standards and disclosure mandates. Alternatively, the switching cost can also be viewed as a measure of the strength of the relationship between banks and their clients.

We denote by N the number of potential firms that are in a position to benefit from the particular financial innovation introduced by the innovator in that period. These N firms are assumed to be evenly distributed among existing banks so that, in any period, the innovator’s potential market is composed of $\alpha_i N$ existing clients and $(1 - \alpha_i)N$ firms who may switch banks. While these N firms may be potential clients in the first period, they do not all have to issue the innovative security immediately. Rather, each client may choose to delay either the optimal financing of a profitable venture or the venture itself for one period. The cost associated with delay is denoted by C_D . Delay costs include the loss from using an inferior financing instrument for a period before issuing the innovative security and the possible deterioration in the value of the innovative security over this period. Delay costs may also be bigger or smaller than V , the incremental benefit from the innovation.

To illustrate a case with $V > C_D$, consider a situation in which the benefits from using the innovative security accrue over several periods. Delaying adoption for a period involves a partial or total loss in appropriating the value delivered by the product in the first period, say V_1 . In this case, the cost of delay, C_D , should be at most V_1 . If the entire benefits from the innovative security were to be unavailable one period later, we would have $C_D = V$.

Alternatively, consider a firm that plans to fund a project using long-term financing and has the choice of financing via an innovative security. A possible action for the firm is to delay the project itself in order to postpone

¹⁰ The notion of a single primary investment bank was, until recently, a standard feature of the U.S. corporate scene and is still an important feature in most global markets. In the U.S., however, some of the largest corporations have moved away from having a sole primary investment bank. However, most corporations still have a very small number of investment banks familiar with the details of their operations.

issuance of the new security. In this case, the cost of delay would include any loss in value (e.g., loss of market share) associated with the delayed project.¹¹ The firm's best alternative, however, may be to delay issuing the innovative security and to fund the project for the interim period using some short-term securities. Therefore the cost of delay, C_D , would include underwriting and other costs associated with issuing the short-term securities. It is then possible to have $V < C_D$, if the benefits from the innovation are small.

This (costly) timing discretion is a crucial feature of our model, since the price that the innovator can charge depends on the cost the client faces in waiting for imitators to exert pressure on prices. Alternatively, with the duration of a period not fixed ex ante, this cost could also be interpreted as a measure of the actual time taken for competitors to imitate the innovation. In either interpretation, C_D is a measure of the cost the client faces in waiting for imitators to offer a competitive product.

1.2 Net revenue from innovative products

Since the innovating bank has only a single period in which to derive rents from clients, its revenue is determined by the value V that clients receive from the product as well as the switching costs C_S and the delay costs C_D that they face. We assume that investment banks are Bertrand competitors who compete in prices.

Bertrand competition ensures that the most a bank can charge an existing client for providing a widely available service is the cost of switching to another supplier. Normalizing the marginal cost of providing the service to zero, C_S then represents the price that can be charged for existing products. The innovator can charge its client firms, at most, an additional amount V for a new product. However, clients with timing discretion will never pay more than C_D over the base level of C_S : if they are charged more, they will find it optimal to delay the adoption of the product until they can reap the benefits of additional competition next period. Hence when $V < C_D$, the innovator can extract the entire benefit V ; when $V \geq C_D$, the innovator can, at most, extract a value of C_D . The incremental revenue can therefore be represented as $\min(V, C_D)$ and the total price charged for the innovative product is: $C_S + \min(V, C_D)$. Therefore we have

Lemma 1. *From an existing client, the innovator receives incremental revenues of $\min(V, C_D)$.*

Consider now the incremental revenue from firms that are currently clients of other investment banks. For these firms, the price charged for the innovative product is also the incremental revenue from the innovation, since firms

¹¹ Such delay costs could be low if there were offsetting "real option" benefits, such as an improved product technology, associated with delaying the project.

would not switch banks in the absence of an innovative product. In comparison with the revenue to the innovator from its own clients, the revenue will be reduced by the extent of the switching costs C_S —to reflect the cost of switching. We have

Lemma 2. *From a firm that is currently a client of another investment bank, the innovator receives an incremental revenue of $\max\{\min(V, C_D) - C_S, 0\}$.*

Proof. Denote by J the potential client's primary investment bank. For a client to switch banks, the price charged by the innovator must make it at least as well off as (i) staying with J and using an existing product and, (ii) staying with J and delaying until the next period. These are the only relevant choices as the other possibilities are strictly dominated.¹²

In the first case, to discourage switching by its client, bank J could offer to reduce the price it charges to zero from C_S , the price it would charge in the absence of the innovation. If the innovator's price for the new product is p , the benefit from switching is given by $V - C_S - p$. Therefore a switch occurs only if

$$V - C_S - p \geq 0.$$

This implies that $p \leq \max(0, V - C_S)$, since the innovator has no incentive to suffer losses in a single-period setting.

The benefit associated with staying with J and delaying adoption for a period can be represented by $V - C_D - C_S + \Delta$, where $V - C_D$ is the value of delayed adoption, C_S is the price that it would ordinarily be charged for an existing product, and Δ is the maximum subsidy that J might offer to discourage switching. If the firm does not switch, J stands to receive a revenue of C_S next period from the firm. Hence the maximum subsidy it would offer is $\Delta = C_S$. Therefore the firm switches only if

$$V - C_S - p \geq \max[V - C_D - C_S + \Delta] = V - C_D.$$

Since $p \geq 0$, this yields

$$p \leq \max(0, C_D - C_S).$$

Combining the two constraints on p gets our required result. ■

The two lemmas establish that the incremental revenue for an innovator is

$$R(V) = \alpha_i N \min(C_D, V) + (1 - \alpha_i) N \max\{\min(V, C_D) - C_S, 0\}. \quad (1)$$

An investment bank faced with an opportunity to develop an innovative product will do so only if the revenue $R(V)$ exceeds the costs X .

¹² For instance, the possibility of delaying after switching is dominated by (ii), since the price charged for the product in the next period is the same irrespective of which investment bank the firm is with. Staying with J and delaying avoids having to incur the switching cost C_S .

1.3 Effect of various parameters on innovation revenues

$R(V)$ is (weakly) increasing in the value, V , the innovation provides potential clients. The nature of the relation between $R(V)$ and V is easily characterized by an examination of Equation (1). The relationship is summarized in the following lemma. The proof of the lemma is simple and is therefore omitted.

Lemma 3. *For $\alpha_i > 0$, $R(V)$ is piecewise linear and (weakly) increasing in V .*

With $C_D > C_S$:

$$\frac{dR(V)}{dV} = \begin{cases} \alpha_i N & \forall C_S > V \\ N & \forall C_D > V \geq C_S \\ 0 & \forall V \geq C_D. \end{cases}$$

With $C_S \geq C_D$:

$$\frac{dR(V)}{dV} = \begin{cases} \alpha_i N & \forall C_D > V \\ 0 & \forall V \geq C_D \end{cases}$$

While the revenue $R(V)$ is increasing in V , the extent to which the innovator can appropriate the value V is determined by its market share, the ability of issuers to delay, and the costs of switching. Note that the revenues appropriated from any innovation are strictly increasing in the market share of the innovator. This is because preferential access to clients due to costs of switching gives rise to market power that can be exploited. Hence investment banks with higher market shares are in a position to accept innovative projects that are not profitable to banks with lower market shares. Such market power disappears *on the margin* for higher valued innovations if delay costs are greater than switching costs. This is because, in the presence of switching, the entire market now belongs to the innovator. These observations form our first proposition.

Proposition 1. *A higher market share is likely to be associated with greater observed innovation activity. Moreover, large innovations will enhance the market share of innovating banks.*

The proposition seemingly contradicts the well-known intuition in Arrow (1962) about competitors in a product market having greater incentives to innovate than monopolists. In the research and development (R&D) literature, the impact of market power on innovation incentives is straightforward: since product market power gives rise to quasi-rents, any innovation is worth less on the margin to someone with existing quasi-rents. Boot and Thakor (1997) rely on broadly similar arguments to show that “universal” banks would be less willing than specialized investment banks to develop innovative products. This intuition fails in the present context because the innovation’s value is

over and above the value delivered by existing products. To put this result into a comparative context would require a product market monopolist to get the same gains over and above any existing quasi-rents that would be available to a competitor. While such a structure strains credulity in the product market context, it is natural in the context of financial innovations, where imitation allows every bank to offer an existing financial product at the same cost.

The result that market shares affect and are, in turn, affected by innovation, finds support in the existing empirical literature on financial innovations. Nanda and Yun (1996) find that investment banks that are market leaders in a particular class of securities are also responsible for many of the innovations in these types of securities. In addition, Tufano (1989) reports that innovating banks are able to maintain a large market share in the new securities they introduce, despite subsequent entry by competing banks.

1.4 Robustness

The association of a higher market share with greater observed innovation activity is a result that is quite robust to perturbing the assumptions of the model. Allowing for economies of scale in the development process, for example, only bolsters the result. In addition, large, well-established banks may face lower switching costs when they market a product to the universe of clients, due to their varied expertise. This change, too, only bolsters our result.

We have modeled competition among banks solely in the form of product imitation rather than competition at the innovation stage itself. There is no reason to expect, however, that such competition would affect the predicted relationship between market share and innovative activity. Consider, for example, the case of banks investing resources in the innovation process. Such resources could be in the form of cash outlays or effort costs. Clearly banks expecting to obtain greater revenues from a given innovation would also engage in the search for innovations with greater intensity. Consequently, an extension of our model to incorporate the effects of competition in the innovation process itself will yield similar results relating market share to innovation intensity. If anything, such a model would predict an even stronger association between market share and the likelihood of innovative activity.

The relationship between innovation activity and market share does not necessarily extend to the sheer size of the innovating bank. As Boot and Thakor (1997) point out, a universal bank will internalize the effect an innovative product may have on the profits from imperfect substitutes in the bank's product portfolio. In our setting, such externalities do not arise since we assume costless imitation and price competition. Consequently, no bank earns profits—other than the switching costs earned from its own clients—for any existing product. If we were to allow for significant spillover effects in the model, as in Boot and Thakor (1997), it would not necessarily yield a predictable relationship between size and innovative activity. It should be

noted, however, that such spillovers can work to promote innovation activity as well. For instance, a bank that is active in several market segments may enhance its overall reputation for problem solving and creativity by the introduction of an innovative product in one segment. The innovation may pay off in terms of attracting clients to activities other than those directly associated with the new product.

Finally, the result in Proposition 1 should be robust to a dynamic extension of our single-period model. In such an extension, an innovator would also take into account any change in market share resulting from the introduction of an innovative product. Since firms with larger market shares are more likely to innovate, the value of incremental market share increases with size for all nonmonopoly situations. Thus an explicit accounting of market share acquisition should not affect the qualitative nature of our conclusions.

1.5 Implications

The basic model of financial innovation described above makes several predictions. First, revenues earned from innovative activity are directly proportional to the demand for the new product. Therefore innovation activity in a particular market will be more pronounced when demand for products in that market is higher (in terms of the model, a higher N). We should thus expect to see spurts in innovation activity in initial public offerings (IPOs), seasoned equity offerings, and in takeover markets when these markets are “hot.”

Second, the model predicts that more financial innovation will occur in areas where the costs of delay are higher. This is a direct consequence of the nonpatentability of financial innovations. When opportunities are fleeting, the appropriability of revenues from innovation is higher and hence innovation opportunities that were earlier perceived to be unprofitable will appear worthwhile. This prediction is supported by the fact that innovation activity in support of risk management strategies has increased dramatically with deregulation of financial markets. While the “standard” explanation of this phenomenon is demand based, namely, with greater volatility comes a greater need to practice risk management, our results present a different story. In our model, greater volatility in environmental variables like interest rates and exchange rates implies greater costs of delay to a client who must take advantage of prevailing rates. Hence innovators of securities that take advantage of current market conditions are able to appropriate greater revenues from innovations (in terms of the model, a higher C_D). Similar arguments may be made for other areas where delay costs are likely to be high, for example, takeover offers where delays are costly due to information leakage and the arrival of competing bidders.

Under the interpretation that C_D is an effective measure of competitive pressure exerted by imitators, we would expect innovating banks to exert some care to guard against easy imitation. Hu (1989), for example, reports that the originators of the first currency swap went to some lengths to preserve secrecy in order to maintain a competitive advantage. When total

secrecy is hard to achieve, an innovator may try to introduce innovations first to market segments that require less stringent disclosure standards. This would lessen the probability of quick imitation since “reverse engineering” is now harder to do. In general, private placement markets have less stringent disclosure rules than public securities markets and, consequently, we should expect to see innovative products first introduced in such markets, if appropriate. Indeed, there is evidence that privately placed securities often include innovative features that are only later used in publicly offered securities [see Wolf (1988)]. Similarly, the individual components of complex, innovative hedge strategies would not necessarily be exposed to wider market scrutiny, if crossing of orders can be done in-house. This leads to the prediction that hedge books maintained by banks will be guarded well and that only net positions of the plain vanilla variety would be hedged in the public markets. Anecdotal evidence seems to support this prediction.

The discussion above also allows us to throw some new light on why financial innovation activity—which was vigorous in the late 19th and early 20th centuries—took a nosedive in the post-Depression period. The establishment of the SEC and the advent of comprehensive disclosure regulations surely served to diminish the informational advantage enjoyed by innovators over imitators. Also, new registration requirements and passing muster with the SEC approval process imposed additional costs of developing innovations. Combined with the lack of demand for investment activity in the post-Depression years, these factors could plausibly have imposed daunting pressures on the incentives to innovate on the part of investment banks.

Finally, the model highlights an important measurement problem for analysts trying to gauge the profitability of innovation activity in the financial arena. To illustrate, Tufano (1989) finds no significant difference between the average underwriting spreads on new security offerings before and after rival banks enter. As a result, he concludes that innovators do not engage in monopolistic pricing. Our model suggests potential problems with such an interpretation. While it does predict a decline in spreads charged to the innovating bank’s own customers with time,¹³ it also points out that larger innovations will trigger switches. When switching costs are primarily borne by clients firms, the average spread over periods is not a good measure of market power. This is because the spreads offered to potential switchers would have to be less than those charged to the bank’s own clients and could, in fact, be lower than the spreads charged after the emergence of imitators. Consequently, average spreads in the period of monopoly could be higher or lower than those charged after competition emerges.

We next turn to the issue of what systematic patterns we should expect to observe in markets where rapid imitation limits appropriability. From the literature on patent protection, we know that such markets will underprovide

¹³ This prediction is consistent with evidence in Kanemasu, Litzenberger, and Rolfo (1986).

innovations that require major outlays, since such investments can only be recouped if monopoly power is granted to augment appropriability. Consequently, we should expect to see mostly innovations that do not require huge outlays to develop and market. However, as we show below, we should also expect innovating firms to exercise significant ingenuity in exploring opportunities to enhance appropriability in the face of imitation. Such adaptability could exhibit itself in the choice of areas in which to concentrate innovation activity, in altering the introduction strategy of larger innovations, and in exploring sharing opportunities with competitors with their captive client bases. These avenues are explored in the next three sections.

2. Regulation and Innovation Incentives

Tax considerations and the avoidance of existing regulations have often been claimed to be major driving forces in the development of innovative financial instruments.¹⁴ Miller (1986) takes the view that most financial innovations are triggered by efforts to skirt regulation or reduce taxes. On the other hand, strategies designed to exploit loopholes are not usually looked upon favorably by regulators and tax authorities, and typically trigger regulatory responses that try to make such skirting harder or impossible. In this section we argue that anticipated regulatory responses seeking to limit loophole exploitation activity can actually have perverse effects.

Consider the introduction of an innovation that is profitable for the client but that may be deemed undesirable by the regulator. Accordingly, the regulator may close the tax or regulatory loophole in the next period with a probability π . In such a situation, products already introduced are assumed to be grandfathered, but any new introductions are banned. We also assume that the introduction of an innovative security does not require *ex ante* approval from the regulator. Both assumptions are consistent with the general nature of regulatory interventions in the United States.¹⁵ The differential impact of alternative regulatory regimes is discussed more fully at the end of the section.

When an innovative product may effectively disappear from the scene, it becomes necessary to distinguish between distinct sources of the cost of delay. Following our earlier discussion, delay costs can be viewed as having

¹⁴ For example, tax considerations were paramount in the emergence of zero-coupon bonds.

¹⁵ A canonical example of such an intervention occurred during the introduction of *Primes and Scores*. To quote from Allen and Gale (1994, p. 21):

Primes and Scores: The first example of primes and scores was the Americus Trust offered to owners of common stock in AT & T in October 1983. This consisted of a trust where shares in the company were deposited, and for each underlying share the trust issues a prime and score security which are listed separately on the American Stock Exchange. . . In March 1986 the IRS implemented changes in the tax code, which made trusts created after that date subject to a separate corporate income tax. This effectively prevented the creation of further trusts. However, before that date Americus trusts for twenty-seven U.S. corporations including firms such as Exxon, DuPont, American Home Products and Bristol-Myers were created.

two components: (1) C_1 , the extent to which the benefits from the product are attenuated by delaying adoption by a period and (2) C_2 , the costs associated with alternative actions undertaken to arrange for the delay itself. The first cost is incurred only if the innovative product is actually adopted in the second period. The second cost, however, is incurred irrespective of actual adoption. While the second category of costs may be quite large, the loss of benefits from waiting alone should not exceed the first period benefits associated with the product.

Consider a current client of an innovator with a product that delivers a value of V_1 in the first period and V_2 in the periods thereafter. The maximum incremental revenue that can be extracted from such a client is given by p , where

$$V_1 + V_2 - p = \max[(1 - \pi)(V_1 + V_2 - C_1) - C_2, 0].$$

The first two terms on the left-hand side constitute the value of the innovation to the client and the third term represents the price paid for the product. On the right-hand side is reflected the value of waiting for competition to emerge. In this case, the client gets to avail of the product with a probability $1 - \pi$ if there is no regulatory action. However, the value of the product is reduced by the amount C_1 . Other costs associated with delay, C_2 , are incurred whether or not the regulatory action takes place. Of course, the client cannot be charged more than the product is worth to her, and this is reflected in the use of the $\max[\cdot]$ operator.

The equation above can be rewritten in the more familiar form

$$V - p = \max[(1 - \pi)V - C_D, 0],$$

where the cost of delay is given by $C_D = (1 - \pi)C_1 + C_2$, and the value of the product, if adopted this period, is given by $V = V_1 + V_2$. Note that with the possibility of disappearance of the product, the magnitude of the costs of delay will generally depend on the probability of such a disappearance.

In all cases where the client expects to get a surplus, that is, so long as $C_D < (1 - \pi)V$, the above equation reduces to

$$p = C_D + \pi V = C_1 + C_2 + \pi(V - C_1).$$

Hence the client faces an effective delay cost of C'_D , where $C'_D = C_D + \pi V$. Since $C_1 \leq V_1 < V$, C'_D is higher than C_D and is increasing in π . Thus even though the expected delay cost C_D declines with the probability of the regulatory intervention, π , the revenues extracted by the innovator actually increase. A similar argument holds for firms that are not clients of the innovator, since the revenues available from them also increase with the cost of delay. Hence the incentives to innovate are enhanced in the presence of anticipated regulatory action.

The exact magnitude of the effective cost of delay depends on the nature of the regulatory regime. If, unlike the U.S. case, regulatory action does not involve grandfathering, delaying adoption may not be as costly. To see this, consider the polar case when regulatory action eliminates all future benefits associated with an already introduced product. In this case, p will be given by $\max[C_D + \pi V_1, (V_1 + (1 - \pi)V_2)]$. It is easy to see that the effective cost of delay still increases with π , so long as $C_1 < V_1$. For the extreme case of $C_1 = V_1$, however, the price charged is independent of the probability of the regulatory action.

We summarize these arguments in the following proposition:

Proposition 2.

- (i) *When a regulatory regime grandfathers introduced products, a higher probability of adverse regulatory action actually increases the incentives to develop financial innovations which may attract regulatory scrutiny, as long as $C_D < (1 - \pi)V$.*
- (ii) *Without the grandfathering of products, the incentives to innovate still increase with the probability of adverse regulatory action if, in addition, $C_1 < V_1$.*

When an innovation is not patentable, the appropriability of its benefits depends critically on the emergence horizon of competition. By banning the new product, the regulatory agency not only removes it in the second period, but also casts a shadow over the economics of its introduction. Since the possibility of a ban works against the incentive of clients to delay adoption, their bargaining power is weakened. This leads to greater revenues for the innovator.¹⁶ Examples of successful innovations subsequently banned by the government are not hard to find. Leveraged preferreds (eliminated in 1984), tax-benefit transfers (eliminated in 1982), evasion of taxable income on the buying back of discounted debt (eliminated in 1984), adjustable rate convertible notes redeemable at a discount (eliminated in 1983) are all examples of very successful products which were later eliminated by changes in regulations. The structuring of these products makes it clear that the innovators who peddled them and the adopting clients were all reasonably certain that they would be viewed as tax dodges and would therefore be eliminated.

The general point to note is that the threat of efficient ex post regulatory action may only diminish competition and thereby benefit appropriability. While lax regulatory response allows more clients to take advantage of innovations exploiting loopholes *once they are introduced*, this effect can be overwhelmed by the dilution in the innovator's incentives to introduce such innovations in the first place. Thus an innovator actually has incentives

¹⁶ We implicitly assume that the time frame in which the regulator acts is large enough for significant benefits to accrue before the rules are changed. When this is not the case, the economic significance of our result is reduced.

to exaggerate the likelihood of an adverse regulatory response in order to stimulate demand for its product.

In examining the impact of regulation on financial innovation, one must distinguish between *ex ante* regimes, where regulators approve new products before they can be introduced, and *ex post* regimes, where the introduction process is relatively unrestricted. Our analysis predicts that countries with mostly *ex ante* securities regulation regimes will exhibit less evidence of innovations that exploit loopholes. A similar conclusion holds for countries with lax regulations of the *ex post* variety. In this case, however, one would also see widespread imitation of innovations introduced elsewhere, which may not happen in *ex ante* regimes. Finally, more financial innovations should occur in regimes characterized by relatively efficient *ex post* regulatory structures. Accordingly, we should expect a greater focus on innovations seeking to exploit regulatory loopholes in countries like the United States which have such a regulatory regime.

This close connection between the nature of the regulatory regime and the incentives to innovate around loopholes has not entirely escaped the eyes of savvy regulators. In recent times, both the U.S. administration and Congress have considered retroactive changes in legislation to curb excessive exploitation of loopholes. Thus a recent debate in Congress focused on requiring investment banks to register new tax-oriented products in order to give the government a chance to preempt their introduction. Perhaps the most spectacular change in attitudes was witnessed in the recent case of step-down preferreds introduced by J. P. Morgan and marketed aggressively by Bear, Sterns and Co. and by Morgan Stanley and Co. In a matter of a few weeks in January and February 1997, the trio of firms sold the product to clients to the tune of \$10 billion in issuances. In March of the same year, the government announced that the transactions were taxable under current law and *just in case they were not, the law would be changed retroactively*.¹⁷

3. Optimal Introduction of Innovations

Without patent protection, investment banks have incentives to spread out the introduction of innovations over time.¹⁸ While a sequential introduction strategy may be possible only with certain innovative products, such a strategy can contribute to the image of hectic innovation activity. Since a focus

¹⁷ See "Crackdown on Creativity," by Ann Monroe (*Institutional Investor*, June 1997, pp. 77–80).

¹⁸ A well-noted example of this phenomenon in another market is that of Kodak and Polaroid in their battle over the instamatic market. Over several rounds, in response to Kodak's introduction of a new camera, Polaroid introduced a better camera within a matter of weeks. It was clear that Polaroid had the technology well before their introduction of the new products and, presumably, would have introduced the products over time anyway.

on the number of innovations misses the actual magnitude of each innovation, observers may very well have exaggerated the degree of actual innovative activity in the financial sector. Note that the sequential introduction of improved versions of financial products is quite common. Rosenthal and Ocampo (1988) describe the example of Salomon Brothers who, in their development of credit card-backed securities, continued to modify the design features despite the moderate success of their initial (Bank One) offering. After the design of these securities was modified—a change in the form of credit enhancement leading to AAA ratings—the product proved even more successful and served as a prototype for future offerings.¹⁹

To be specific, suppose an innovator has developed a new security design that, on account of taxes or other savings, can provide clients with an incremental benefit of V . Instead of introducing this security, she may initially deliberately introduce a less efficient version that delivers benefits of only $V_1 < V$ to clients. At a later stage, the innovator can then introduce the more efficient security. The incremental benefit associated with the second introduction is, then, $V_2 = V - V_1$.

For simplicity, we allow the innovator full flexibility in her choice of introduction strategy. Thus the innovator chooses the number of stages, say T , and the various pieces V_t , such that $\sum_{t=1}^T V_t = V$. We assume that a marketed product is imitated in the following period. Given the significant employee mobility seen in the investment banking industry, we also allow for the possibility of information about unintroduced products leaking out. Thus we assume that in each period there is a probability $(1 - \beta)$ of the innovation leaking to competitors. Given the possibility of leakage, the innovator prefers a single-stage introduction strategy if the entire benefits were appropriable.

Each period, N firms realize the need to issue securities that can benefit from this innovation. While a firm can delay issuance at a cost of C_D per period, we assume, for simplicity, that it needs to issue new securities only once. This allows us to directly extend the one-period model to multiple periods. We discuss the effect of relaxing this assumption later. For convenience, we assume that the size of the total innovation, V , is known to all.

Proposition 3. *For an integer M , such that $(M + 1)C_D \geq V \geq MC_D$, the optimal introduction strategy is to sequentially introduce, over the first M stages, innovative products each with a value of C_D . The residual value*

¹⁹ As another example, Smith Barney first introduced Adjusted Tender Preferred Stock (ATPS) to compete with Dutch Auction Rate Preferred Stock. In lieu of an auction, this product used a remarketing agent to reprice the issue every 49 days. A modified version introduced later, Share-Adjusted Broker-Remarketed Equity Securities (SABRES), allowed the remarketing agent to vary the dividend period.

Unfortunately, the staged introduction of innovations is observationally indistinguishable from situations where an innovator experiments with product design before settling on a final version.

$(V - MV)$ is introduced at the last stage. The total expected revenue to the innovator from an innovation of size V is

$$E(R(V)) = \left\{ \frac{1 - \beta^{M-1}}{1 - \beta} \right\} N[\alpha_i C_D + (1 - \alpha_i) \max\{0, C_D - C_S\}] \\ + \beta^M N\{\alpha_i \Gamma + (1 - \alpha_i) \max\{0, \Gamma - C_S\}\}$$

where Γ denotes $V - MC_D$.

Proof. Consider, first, the introduction of a stage security of value $V_t = C_D + \delta$, where $\delta > 0$. This strategy has no effect on revenues at stage t , but can only reduce later revenues. Alternatively, choosing $V_t = C_D - \delta$ reduces revenues at stage t . The lost revenues may not be recoverable later if information leaks to competitors. Therefore it is optimal to introduce amounts $V_t = C_D$ for M periods, followed by the residual at the last stage.

The derivation of the expected revenues is straightforward and is therefore omitted. ■

This simple form of the optimal introduction strategy survives the relaxation of some of our simplifying assumptions. For example, while the level of the expected revenues depends on the magnitude of β , the form of the optimal introduction strategy does not. Intuitively, it is always better to take a chance on leakage rather than forego added profits for sure. Introduction of marketing and development costs, too, does not affect the form of the introduction strategy. Even when the staging of introductions requires additional costs each period, the optimal introduction strategy remains unchanged. Again, revenues change, but the form of the introduction strategy does not.

Up to now, we have assumed delay costs to be independent of the value of the introduced innovation. It may be more realistic to assume, instead, that each period's cost of delay, say $C_D(t)$, depends on the value of the innovation introduced during the period, V_t . As long as the delay cost exceeds the value for a minor innovation (e.g., if there is a fixed component to C_D) and is lower than the value of a major innovation, it still pays for the innovator to phase in introduction for large enough innovations. Assuming a unique solution to the equation $C_D(V) = V$, the optimal introduction strategy can be shown to be one of staged introduction. Given the possibility of leakage, however, it never pays to introduce an innovation strictly less than its associated cost of delay. Hence the optimal introduction strategy generalizes to introducing the innovation in pieces such that $C_D(t) = V_t$ for all t . For simplicity, for the rest of the discussion, we stay with our original assumption that C_D is independent of V .

As pointed out earlier, it is possible to interpret C_D as a measure of the imitation pressure faced by the innovator. Under this interpretation, the proposition shows that a greater threat of imitation—a lower value of C_D —

makes the innovator resort to staging introduction over a longer time. Thus we have

Lemma 4. *A greater pressure from imitators results in the introduction of smaller innovations each period.*

We have assumed that a candidate client firm issues securities only once over the period of introduction of the innovation. This assumption considerably simplifies the analysis and allows us to directly extend the single-period model to a multiperiod setting. Allowing firms to issue multiple securities over time substantially complicates the analysis in the presence of switching. Analyzing switches requires us to characterize the evolution of the firm's relationship to the innovator as well as to its original bank. For instance, it can be shown that the optimal introduction strategy above is unaffected if, despite using the services of the innovator for a single offering, a firm remains the client of its original bank and still faces switching costs C_S when it employs the services of the innovator again. If, on the other hand, switching to the innovator results in the firm becoming a client of the innovator, the optimal introduction strategy has to take into account dynamic considerations. A complete dynamic analysis of this problem with endogenous market shares is quite complicated and beyond the scope of the current article. However, our conjecture is that, in such a dynamic setting, firms with relatively small market shares may find it optimal to initially raise their market shares by aggressively introducing innovative products of value greater than C_D , in order to induce switches. The notion is that, while the innovator would lose revenue from existing clients, this would be offset by the revenue increase from greater market share in subsequent periods. The idea that innovators with smaller market shares may tend to introduce innovative products of greater value to induce switches can, however, be derived in a simpler setting, as shown below.

Our strong result that the innovation will be introduced in pieces of similar incremental value, irrespective of innovator market share, is very much dependent on our assumption that all clients face similar delay costs. Indeed, if this assumption were relaxed, one need not get the implication that the optimal introduction strategy is invariant across innovators with different market shares. For example, consider the case when each bank's clients have a distribution of delay costs: a fraction γ have low delay costs (less than C_S) while the rest, fraction $(1 - \gamma)$, have higher delay costs (greater than C_S). Clearly a bank with a high-enough market share will derive most of its revenues from its own clients and thus will choose to divide up its innovation in chunks corresponding to the lower delay cost if γ were large enough. However, this would not necessarily be the optimal strategy for a bank with a small market share. The reason is that such a bank may prefer to introduce its innovations in chunks corresponding to the large delay costs in order to induce

switches by clients of other banks. Given the large potential magnitude of such switchers, the revenues earned from them may very well overwhelm the revenues earned from catering to their own clients by way of successive introduction of smaller innovations over a larger period in time. This strategy is more attractive the higher the probability of leakage. We state this result below without proof.

Proposition 4. *With heterogeneous delay costs, a bank with a smaller market share is more likely to introduce its innovations aggressively and induce switches. Banks with higher market shares will appear to be introducing a greater number of innovations without inducing switches.*

This discussion on the optimal introduction strategy strongly suggests that the possibility of leakage may play an important role in determining which innovation opportunities a bank pursues. In particular, the value of a large innovation to an innovating bank is concave in the full value V of the innovation. The degree of the concavity is decreasing in β . This suggests that in environments characterized by significant probability of leakage, a bank may very well pass up large innovative projects requiring greater development outlays and focus more on smaller innovations in the first place. This only bolsters the possibility of smaller innovations being introduced in the financial marketplace.

4. Sharing the Innovation with Competitors

The process of changing banks involves incurring dissipative switching costs. In addition, if a client of a noninnovating bank chooses to wait for competition to emerge, she incurs dissipative delay costs. Alleviating the impact of these costs through cooperation between investment banks may make everyone better off.²⁰ Such cooperation could be structured in a variety of ways, although the typical cooperative arrangement in the securities industry takes the form of jointly managed offerings. Abstracting from the issue of the precise form of cooperation, we focus, instead, on the conditions under which cooperation between banks may emerge.

When innovations are patentable, licensing can enhance effective market size through profitable sharing with others. In the absence of patents, however, sharing an innovation immediately creates competition, since partner banks are free to introduce their own versions of the innovative product. Thus, for cooperation to be feasible, the gains from sharing must exceed the losses from enhanced competition.

To capture the trade-offs involved with sharing of innovations, we modify the model to allow for cooperation between bankers. The initial period is assumed to consist of three stages. In the first stage, the innovator approaches

²⁰ Of course, costs could also be reduced by taking over the innovator. On the other hand, takeovers may involve costs of their own. We do not address the issue of takeovers in this article.

its competitors with take-it-or-leave-it offers to familiarize them with the new product.²¹ Such offers are accepted or rejected. For simplicity, we assume that acceptance or rejection is not publicly observable. The acceptance of an offer is followed by a quick familiarization with the innovation technology. The familiarization process is taken to be costless, although small costs would not change our results. In the second stage, every bank with access to the innovation can approach other banks and propose sharing arrangements of their own. At this stage, they are assumed to compete as Bertrand competitors. After all such arrangements have been completed, we proceed to the third stage, when the innovative product is marketed to clients. If no offers are accepted at the first stage, the innovator markets its product directly to all potential clients.

The acceptance of a first-stage offer ensures the free availability of the innovation to all remaining banks in the second stage. This is because of the Bertrand competition between the innovating bank and its partner(s) in the second stage. Thus our assumption of Bertrand competition is a rather strong one. However, it easily captures the idea that the lack of patentability prevents the innovator from effectively asserting control over the speed (and price) of dissemination of the innovation. In turn, our assumptions imply that if the innovator is to derive a benefit from sharing its knowledge with a competitor, it comes only from the bank(s) that accept(s) its offer in the first stage. This is in sharp contrast to the situation with patent protection when the innovator can license its technology to competitors sequentially. Our qualitative results do not require the innovator to lose all control over the dissemination process—any significant reduction in power due to competition in the second stage would suffice.

The j th bank's market share is represented by α_j , $j = 1, \dots, K$ and, as before, the innovating bank's market share is given by α_I . The index M is used to represent the bank with the largest market share. We first derive the incremental revenues to the innovator if the j th bank alone were to accept the first-stage offer from the innovator, I .

Lemma 5. *The maximum total revenue earned by the innovator if the j th bank alone accepts a take-it-or-leave-it offer is given by*

$$R(V) = \begin{cases} \alpha_j N \min(C_D, V) & \forall C_S > \min(C_D, V) \\ \alpha_j N C_S & \forall C_S \leq \min(C_D, V). \end{cases}$$

Proof. As discussed above, Bertrand competition between I and j ensures that all other banks obtain the innovation free of cost in the second stage. Also, since banks compete in offering the innovative product, I can no longer

²¹ The qualitative nature of our results is unchanged under alternate bargaining structures.

charge its own clients an incremental revenue over above their cost of switching. Hence, the sole incremental revenue received for the innovation comes from fees received from bank j . There are two cases to examine:

- (i) When $C_S \leq \min(C_D, V)$: In this case, the innovator can induce j 's clients to switch if he were to market to them directly. However, this involves dissipation of switching costs. By entering into a sharing arrangement, j can still offer the innovation to its own clients and charge them C_S . Hence the innovator can extract a maximum of $\alpha_j N C_S$ from bank j for entering into the sharing arrangement.
- (ii) When $C_S \geq \min(C_D, V)$: In this case, the innovator cannot bring about switches by marketing the product directly to the clients of the j th bank. But the innovation still reduces what j can charge its own clients for its existing product.

When $V < C_D$, a switch nets the client a net benefit of $V - C_S - p$, where p is the price charged by the innovator. Bank j will choose to prevent a switch by charging a price of $C_S - V + p$. Bertrand competition implies that the price charged by j for its existing product will be $C_S - V$.

When $V \geq C_D$, a client of j has two choices when I innovates. Switching right away yields a maximum benefit of $V - C_S$. Alternatively, delaying until j can imitate the innovation gives it a benefit of $V - C_D - C_S + \Delta$, where Δ is any subsidy that j might offer. To prevent a switch, j is forced to offer a subsidy of C_D , which reduces the price it can charge for its existing product to $C_S - C_D$. A sharing arrangement, however, allows j to offer the innovative product immediately and still charge its clients C_S . Hence, j would be willing to $\alpha_j N \min(C_D, V)$ to enter into a sharing arrangement. ■

The lemma implies that if I were to enter into a sharing arrangement with a single bank in the first stage, it prefers to share with M . We therefore present an equilibrium in which I approaches M with a take-it-or-leave-it offer in the first stage. While other Nash equilibria exist, the equilibrium we focus on is Pareto superior. We denote this equilibrium as equilibrium A.

Proposition 5. *There exists an equilibrium A in which the innovator enters into a cooperative arrangement with the competitor with the largest market share, M , if and only if: $\alpha_M > \alpha_i$ and $C_S \geq (\min(C_D, V))/(1 - \alpha_i + \alpha_M)$.*

Proof. From the innovator's perspective, an offer to bank M is optimal only if the revenue from a cooperative arrangement exceeds that achievable by marketing directly to clients. We consider two cases:

- (i) When $\min(C_D, V) < C_S$: Since inducing switches is too costly, the loss of revenue from own customers due to the creation of competition can only be offset by cooperative arrangements with a larger

competitor. This follows from the fact that the revenue under the cooperative arrangement is larger when

$$\alpha_M N \min(C_D, V) > \alpha_i N \min(C_D, V)$$

or when $\alpha_M > \alpha_i$.

- (ii) When $\min(C_D, V) \geq C_S$: In this case, it is possible to induce clients of all banks to switch. Thus cooperation with one competitor has to overcome the revenue losses from all potential switchers. This is the case when

$$\alpha_M N C_S > \alpha_i N \min(C_D, V) + (1 - \alpha_i) N \{\min(C_D, V) - C_S\}$$

or when

$$C_S \geq \frac{\min(C_D, V)}{1 - \alpha_i + \alpha_M}. \quad (2)$$

Since $\alpha_M > \alpha_i$ and $\min(C_D, V) < C_S$ together satisfy Equation (2), we have demonstrated that the innovator has the incentives to follow the take-it-or-leave-it strategy of equilibrium A.

To complete the proof, we need to verify that bank M will indeed accept such an offer. In the absence of other banks, M is indifferent between acceptance and nonacceptance. With other banks present, a deviation from the candidate equilibrium may be in M 's interest if one of the other banks would receive and accept a similar offer from the innovator. This is because Bertrand competition ensures that the innovation is available at no cost to other banks in the second stage. But acceptance and rejection of an offer is not observable, by assumption. In the candidate equilibrium, therefore, other banks would reject the opportunity to share knowledge of the innovation at any positive price. Hence, M , on receiving an offer from the innovator that leaves it no worse off than under noncooperation, will accept. ■

If sharing does take place, our results are not affected by the exact form of the bargaining protocol used. Equilibrium A is, however, not Pareto dominant when forcing offers whose outcomes are contingent on the responses of all potential collaborators are permitted. In such a case, the innovator could offer the innovation to collaborators conditional on all of them adopting it. Otherwise the innovative product would not be introduced. Such forcing strategies allow the innovator to share with all banks and to appropriate the entire savings on dissipative costs marketwide. In the present context, with innovation opportunities arising randomly, it seems unreasonable to endow the innovator with such ultimatum powers. Nevertheless, in contexts where material renegotiation possibilities do not exist, such industry-wide sharing may be the Pareto superior outcome.

Proposition 5 yields a number of predictions. First, cooperative arrangements are more likely to be observed when the innovating bank is not the market leader in the class of securities affected by the innovative product. When cooperative arrangements like joint management of initial offerings are observed, the proposition predicts the leading bank(s) to be involved as comanagers. These predictions seem broadly consistent with patterns documented in Nanda and Yun (1996): initial offerings of innovative securities are less likely to be comanaged when the lead underwriter is a market leader; for jointly managed initial offerings of innovative securities, comanagers are typically banks with a substantially larger market share than the lead underwriter. Of course, the evidence is also consistent with the notion that market leaders are brought in as comanagers for their distribution skills and reputation alone.

Second, Proposition 5 shows that smaller banks' incentives to innovate depend crucially on the distribution of market shares across banks. When banks have roughly equal market shares, switching costs are unlikely to satisfy the condition in the Proposition. In such a case, marketing directly to clients dominates sharing of the innovation. Since sharing arrangements tend to reduce dissipative costs, the absence of such opportunities is bad for appropriation of revenues from innovations. In other words, asymmetry in market shares raises the innovation incentives for both small and large banks. In the former case, the advantage arises from the possibility of cooperation with the large bank; in the latter case, the advantage comes from the size of the client base that can be exploited. Thus, when cooperation is not prohibited, we would expect to see more innovation activity when market shares are asymmetric.

Third, symmetry in market shares is likely to be associated with more switching activity, since cooperation possibilities are reduced. Symmetric market shares thus should be accompanied by both fewer innovations and diminished client loyalty.

Finally, our analysis in the section shows clearly the twofold importance of market share in the presence of switching costs. That a higher market share enhances appropriability has been already pointed out in an earlier section.²² A second benefit to a high market share is that it enhances one's attractiveness as a partner in a sharing arrangement. In the context of our model, this is exploited by joint management of new issues. In other contexts, this might be manifested in takeovers by larger companies (e.g., software), the importance of brand names and distribution channels established by marketing networks (e.g., acquisitions of consumer product lines with substantial potential by marketing companies like American Home Products), and in the establishment of joint venture partnerships. We believe that this facet of market share has not received adequate attention in the literature to date.

²² On a related note, Klempere (1987) points out the first-mover advantages that arise with switching costs.

5. Conclusion

We have analyzed aspects of the financial innovation process when bank-client relationships make it costly to switch between investment banks, but in which client firms have the discretion to delay utilizing bank services. In the absence of patent protection, innovating investment banks are severely limited in their ability to appropriate the incremental value of a new product to clients. The anticipated revenue from an innovative product and, therefore, the decision to proceed with its development and marketing, is shown to be critically affected by the market share of the innovating bank, the cost of switching clients from other banks, and the discretion of clients to delay.

In such an environment, banks will innovate in areas where clients face effectively greater costs of delay. As a result, a predictable and efficient ex post regulatory system may, somewhat paradoxically, tend to encourage the search for products designed to exploit tax or regulatory loopholes. Innovators will also tend to focus on the development of relatively smaller innovations and prefer to introduce larger innovations in stages. This suggests that a focus on the number of new financial products may overstate the extent of actual financial innovation. However, in an effort to attract clients of other banks, innovators with smaller market shares may be motivated to introduce larger innovations at one go.

The absence of patents is also shown to limit the situations in which socially beneficial cooperative arrangements may arise. Innovators with smaller market shares are more likely to share their innovations with rival banks. In addition, asymmetric market shares are shown to provide a fillip to innovation incentives when sharing arrangements are feasible.

We now consider some directions in which our model could be extended. Clearly, explicit modeling of competition in the innovation process may provide additional insights and allow a welfare analysis of the question of granting patents to financial innovators. The strength of bank-client relationships could also be modeled as a function of the length or intensity of the relationship. In this case, switching costs would depend on whether the firm has switched banks recently and may cause spillover effects between unrelated innovations. For example, a major financial innovation that causes switching may trigger subsequent innovations purely on account of the lowering of average switching costs. This, in turn, could make more innovation likely. Also, our analysis could be extended to analyze competition in contract introductions between financial exchanges.

Our analysis also has implications for other industries characterized by the lack of comprehensive protection from imitation. For example, in the software industry, although products may be patented, innovative features of programs are widely copied by competitors. A familiar example is the case of spell-check programs which, when introduced in the first word processor, were quickly incorporated by competing word processors. Second, the

software industry is also characterized by switching costs due to familiarity and compatibility issues. Third, when features incorporated have productivity benefits, any delays in adoption have associated delay costs. Consequently, we would expect the economics of this industry to exhibit strong commonalities with that of the investment banking industry. Clearly market share acquisition is considered important in this industry and periodic innovations in the form of incremental changes are commonplace. Anecdotal evidence also suggests that smaller firms generating innovations frequently sell out to large players who then incorporate their innovations into existing products. A formal analysis of such similarities is beyond the scope of this article and is a topic we would like to throw up for further research.

References

- Allen, F., and D. Gale, 1994, "Financial Innovation and Risk Sharing," *MIT Press*, Cambridge, MA.
- Arrow, K., 1962, Economic Welfare and the Allocation of Resources for Inventions, in R. Nelson (ed.), *The Rate and Direction of Inventive Activity*, Princeton University Press, Princeton, N.J.
- Baldwin, W., and J. T. Scott, 1987, *Market Structure and Technological Change*, Harwood Publishers, Chur.
- Boot, A., and A. Thakor, 1997, "Banking Scope and Financial Innovation," *Review of Financial Studies*, 10, 1099–1131.
- David, P., and S. Greenstein, 1990, "The Economics of Compatibility Standards: An Introduction to Recent Research," *Economics of Innovation and New Technology*, 1, 3–42.
- Farrell, J., and G. Saloner, 1987, "The Economics of Horses, Penguins and Lemmings," in L. G. Gable (ed.), *Production Standardization and Competitive Strategies*, North-Holland, Amsterdam.
- Finnerty, J. D., 1992, "An Overview of Corporate Securities Innovation," *Journal of Applied Corporate Finance*, 4, 23–39.
- Hu, H., 1989, "Swaps, the Modern Process of Financial Innovation and the Vulnerability of a Regulatory Paradigm," *University of Pennsylvania Law Review*, December, 333–435.
- James, C., 1992, "Relationship-Specific Assets and the Pricing of Underwriter Services," *Journal of Finance*, 47, 1865–1885.
- Kanemasu, H., R. Litzenberger, and J. Rolfo, 1986, "Financial Innovation in an Incomplete Market: An Empirical Study of Stripped Government Securities," working paper, University of Pennsylvania.
- Klemperer, P., 1987, "The Competitiveness of Markets with Switching Costs," *RAND Journal of Economics*, 18, 138–150.
- Matutes, C., and P. Regibeau, 1988, "Mix and Match: Product Compatibility without Network Externalities," *RAND Journal of Economics*, 19, 221–234.
- Merton, R. C., 1992, "Financial Innovation and Economic Performance," *Journal of Applied Corporate Finance*, 4, 12–22.
- Miller, M. H., 1986, "Financial Innovation: The Last Twenty Years and the Next," *Journal of Financial and Quantitative Analysis*, 21, 459–471.
- Persons, J. C. and V. A. Warther, 1997, "Boom and Bust Patterns in the Adoption of Financial Innovations," *Review of Financial Studies*, 10, 939–967.
- Nanda, V., and Y. Yun, 1996, "Sharing the Limelight: On Why Investment Banks Co-manage Initial Offerings of Innovative Securities with Competitors," working paper, University of Michigan.

- Rajan, R. G., 1992, "Insiders and Outsiders: The Choice Between Informed and Arm's-Length Debt," *Journal of Finance*, 47, 1367–1400.
- Rosenthal, J. A., and J. M. Ocampo, 1988, *Securitization of Credit*, Wiley, New York.
- Ross, S. A., 1989, "Institutional Markets, Financial Marketing, and Financial Innovation," *Journal of Finance*, 44, 541–556.
- Schumpeter, J., 1943, *Capitalism, Socialism and Democracy*, Unwin University Books, London.
- Tirole, J., 1988, *The Theory of Industrial Organization*, MIT Press, Cambridge, MA.
- Tufano, P., 1989, "Financial Innovation and First-Mover Advantages," *Journal of Financial Economics*, 25, 213–240.
- Van Horne, J. C., 1985, "Of Financial Innovations and Excesses," *Journal of Finance*, 40, 621–631.
- Wolf, C. R., 1988, "Private Placements" in J. P. Williamson (ed.), *The Investment Banking Handbook*, Wiley, New York.