

Information effect of entry into credit ratings market: The case of insurers' ratings*

Neil A. Doherty[†] Anastasia V. Kartasheva[‡] Richard D. Phillips[§]

August 21, 2011

Abstract

This paper analyzes the optimal entry strategy of a credit rating agency (CRA) to a market served by the incumbent. We show that a monopolistic CRA pools sellers into multiple rating classes and has partial market coverage. This provides an opportunity for market entry. The entrant targets higher-than-average companies in each rating class of the incumbent's rating scale and employs more stringent rating standards. We use Standard and Poor's entry into the market for insurance ratings previously covered by a monopolist, A.M. Best, to empirically test the impact of entry on the information content of ratings.

Keywords: rating agency, entry, competition, precision and disclosure of information, insurance.

JEL Codes: D8, G22, G28, L1, L43

*The paper previously circulated under the title "Competition among Rating Agencies and Information Disclosure." We are grateful to Gustavo Manso, Gregory Nini, Anjolein Schmeits, Ajay Subramanian, and Bilge Yilmaz for valuable comments, and the participants at CEPR-EC Banking Regulation, NBER Summer Institute, Temple, IMF, UC Berkeley, NBER Insurance Group, FIRS, IIOC, Econometric Society Summer Meeting for useful discussions. We gratefully acknowledge the support of the Rodney L. White Center for Financial Research at Wharton.

[†]Insurance and Risk Management Department, Wharton School of the University of Pennsylvania, doherty@wharton.upenn.edu.

[‡]*Corresponding author.* Insurance and Risk Management Department, Wharton School of the University of Pennsylvania, karta@wharton.upenn.edu.

[§]Department of Risk Management and Insurance, Robinson College of Business, Georgia State University. rphillips@gsu.edu.

Significant debate exists regarding the role that competition should play in the market for credit ratings. Representatives from the credit rating agencies (CRAs) themselves argue that reducing barriers to entry would eventually lead to reduced disclosure of information as the new entrants engage in “race to the bottom” strategies in order to sell their services. Testimony by Mr. Raymond W. McDaniel, chairman and CEO of Moody’s Corporation, before the US Securities and Exchange Commission illustrates this line of reasoning:

Considering the unique dynamics of our market, historically new market entrants and marginal participants have sought to make their products more attractive to issuers by offering higher ratings than do more established market participants. Some new entrants might be inclined to try to compete in this manner because of the ease with which such a strategy could be implemented and the short-term benefits that might accrue to the entrant as a result. Therefore, Moody’s believes that the usefulness of credit ratings in the aggregate for market efficiency, transparency and investor protection would decline in the event that more Nationally Recognized Statistical Rating Organizations (“NRSROs”) are established and rating levels become a more important element of competition within the industry.¹ (McDaniel 2002)

Critics, on the other hand, argue that the lack of competition may be one of the factors that contributed to the inability of the CRAs to provide accurate and timely information in the recent credit crisis and to support the adoption of rules that promote competition (SEC 2008). Recent regulatory changes have favored this point of view—most notably when Congress passed the “Credit Rating Agency Duopoly Relief Act” in

¹Primarily for the purposes of safety and soundness regulation, regulated investors including banks, thrifts, insurance companies, pension funds, and so on are required to follow rules that in part restrict their investments to carry ratings issued only by an NRSRO. Thus, NRSRO status is viewed by many as regulatory approval of the rating agency. See White (2002) for background and discussion. Recent Dodd-Frank regulation is likely to change this approach.

2006 that clarified the process of obtaining NRSRO status. Though market concentration remains very high, several new CRAs have recently obtained NRSRO status.²

Although Congress has already moved in the direction of encouraging entry in the market for credit ratings, little theoretical or empirical research exists that would be helpful in this debate. We seek to address this shortcoming by theoretically and empirically analyzing how the entry of a new CRA affects the informational content of ratings. We begin by examining the information disclosure of a monopoly CRA. Our analysis suggests that it is optimal for the agency to pool information into letter grades, which results in the clustering of companies and issuers into fairly broad rating classes—a result consistent with common practice. Pooling of information, however, generates an opportunity for a new agency to offer additional ratings. Thus, the second objective is to analyze the entry strategy of a new CRA. If a company or an issuer decides to pay for a rating by a new CRA, it suggests that the entrant uses a different rating scale and/or that its rating contains additional information.³ The third objective is to provide evidence of the effect of entry on the information content of the ratings of insurance companies. The insurance industry provides a unique natural experiment as it was served by the monopoly rating agency, the A.M. Best Company, for several decades until it experienced the entry of Standard & Poor's (S&P) at the end of the 1980s. We show that S&P enters by differentiating its rating scale from A.M. Best's scale. Also, we find that, consistent with our theory, A.M. Best reacted to entry by disclosing more information.

Optimal information disclosure was first addressed by Lizzeri (1999) who showed that when all parties are risk neutral a monopoly CRA's optimal disclosure strategy is to pool

²The Herfindahl-Hirschmann Index (HHI) for all NRSRO ratings outstanding is 3,495, which is equivalent to 2.86 equally sized firms. The current total number of NRSROs is 10, including A.M. Best, DBRS, Fitch, Japan Credit Rating Agency, Moody's Rating and Investment Information, Standard & Poor's, Egan-Jones, LACE and Realpoint. (SEC 2008)

³Although the focus for this paper is to investigate the entry strategy of a new CRA, we also developed a simplified extension of our model to analyze the equilibrium rating systems of the incumbent and the entrant CRAs in a Stackelberg setting to account for the incumbent's reaction to entry. The extension shows the results derived in the paper are robust in the Stackelberg game. See supplementary section available online from the authors.

all companies into one rate class. Surprisingly, all sellers pay to be rated in spite of the fact that a seller does not have to obtain a rating, and de facto the CRA discloses no information. At the same time, entry of a new CRA results in full disclosure of information.

Lizzeri's results contrast with practice. For example, as a monopolist CRA in the insurance industry for many decades, A.M. Best had multiple rating categories and did not have complete market coverage. Moreover, in most settings, multiple CRAs are in competition. In 2000, A.M. Best covered 94.7 percent of the insurance companies, while S&P's coverage was 27.5 percent. We present a model that explains these phenomena and addresses the impact of new entry.

We depart from Lizzeri by suggesting that buyers value the precision of information contained in ratings. Because rating-based guidelines are widely used in the conduct of business activities, buyers are ready to pay a higher price for a good with less ambiguous quality.⁴ We consider a model where sellers have private information about the quality of a good, which they cannot communicate credibly to buyers. A CRA can learn what a seller's quality is, but it has discretion in how this information is communicated to buyers. The evaluation of a seller is reached in two stages. In the first stage, a CRA designs a rating system that consists of a disclosure policy and the fee for its services. In the second stage, privately informed sellers decide whether to demand a rating from the CRA. Once the CRA evaluates the sellers who solicited a rating, buyers form a belief about each seller based on that seller's decision to be rated and, possibly, the seller's rating.

We derive four main results. First, when buyers care about the precision of information, Lizzeri's single rating result holds only as a special case. The optimal rating scale derived from the model resembles the interval disclosure rule actually employed by the major CRAs. Pooling the lowest rated seller with better types has two countervailing effects: the expected quality in the eyes of buyers increases, while the precision of infor-

⁴For example, in a survey of 200 plan sponsors and investment managers in the U.S. and Europe (Cantor, Gwilym and Thomas 2007), 60 percent of fund managers and 47 percent of plan sponsors report that CRAs should put more emphasis on the accuracy of ratings.

mation goes down. For a low value of information precision, the first effect dominates, and full pooling is optimal for the CRA. As the value of precision increases, the trade-off between the two effects defines the boundary of the lowest rating. The model results in multiple equilibrium disclosure policies for rated sellers of higher quality. However, in all equilibria, the payoff of rated sellers is non-decreasing in quality. Also, as the value of information precision goes to infinity, the optimal disclosure of a monopoly CRA converges to full disclosure. Hence, the CRA's incentives to provide information cannot be fully attributed to a competitive market structure.

Our second result is that the optimal disclosure policy of a CRA implies partial market coverage. The reason is that the presence of unrated companies widens the gap between the prices rated and unrated sellers receive on their products or securities. As a result, this permits the CRA to charge a higher fee.

The next two results are related to the optimal entry strategy of a new CRA to a market previously served by the incumbent. Our third result deals with the demand for entrant's ratings. For each incumbent's rating grade, sellers of higher-than-average quality are disadvantaged by pooling. We show that an optimal entry strategy for a new CRA is to target these sellers. Hence, the number of ratings each seller obtains depends on its position relative to other sellers in the rating interval of the incumbent. As a result, there is no congruency between the demand for the entrant's rating and the quality of the seller. High and low quality sellers can be rated by both agencies, while the intermediate quality seller obtains only one rating.

The final result is that the entrant will design a more stringent rating scale relative to that of the incumbent. A seller will purchase a second rating from the entrant only if this enables it to increase the price charged to buyers. This occurs when the entrant's rating increases the seller's expected quality by pooling it with better quality types or improves information precision by reducing the diversity of the pool. In both cases, an entrant will require that higher standards be met in order to provide a similar rating to that of the incumbent. It also follows from our model that sellers are more likely to demand a second

rating in markets where precision is of greater value.

We test our predictions on the entry strategy of a new CRA using data on the U.S. property-liability insurance market. The insurance industry provides an ideal natural experiment to study entry for two reasons. First, unlike the market for bond ratings, there are no regulatory barriers to entering the market for insurance ratings. Second, until recently, the market for insurance ratings has largely been dominated by a single monopoly agency—the A.M. Best Company. S&P made its initial foray into the insurance ratings market in the late 1980s and dramatically increased the number of ratings it provided to insurers during the 1990s.⁵

Insurance ratings measure an insurer’s financial strength and ability to meet its ongoing insurance policy and contract obligations. Prior research has shown that ratings in this market matter. For example, Epermanis and Harrington (2006) provide evidence that insurance buyers are sensitive to insurers’ ratings with lower rated insurers receiving lower prices for their policies in the marketplace. Since policy liabilities are the primary source of capital for insurers, lower prices imply higher costs of capital. Also the importance of a rating for an insurer depends on the type of buyer. Corporate insurance buyers usually require that insurers are highly rated as their insurance policies are very detailed and tailored to a given company’s profile. These policies are mostly sold through insurance agents and brokers who often will not recommend an insurer with an A.M. Best rating below A- (e.g., see Bradford 2003). Personal automobile insurance and homeowners insurance, on the other hand, are protected by state-guarantee funds and sold to less sophisticated customers. As a result, prices and demand are less sensitive to an insurer’s financial strength.

We employ two methodologies to examine the entry strategy of new CRA. First, we use a hazard model to estimate a one-year probability of insolvency using publicly available

⁵In 1992, S&P issued full rating opinions on only 25 property-liability insurers, and this number increased to over 250 insurers by the end of the decade. By 2000, S&P was the second largest insurance rating agency and now rates over 800 companies, representing more than 45 percent of the industry’s assets.

data for all U.S. property-liability insurers. We show that S&P applied higher standards compared to A.M. Best for an insurer to achieve a similar rating.

The second empirical test is designed to investigate the differences in rating opinions between the incumbent and the entrant using the Heckman-style sample selection methodology (Heckman 1979). It allows us to correct for the strategic decision-making of firms and to decompose the sources of rating differences into two components: standards differences between the two agencies, and the strategic decision of insurers pooled in one rating grade to signal higher financial quality. We find that higher-than-average quality insurers in each rating category chose to receive a second rating from S&P and that S&P required higher standards for an insurer to achieve a similar rating. Both results are consistent with our theory.

The plan of the paper is as follows. The next section discusses related literature. Section 2 presents the model and examines the disclosure policy of the monopoly. Section 3 analyzes the entry strategy of a new CRA. We discuss the institutional details of insurance ratings and describe the data in Section 4. Our empirical analysis is presented in Section 5. Section 6 discusses other aspects of competition and the quality of ratings, and the conclusion follows. All proofs and tables are provided in the Appendix.

1 Related Literature

This paper belongs to the growing literature on the incentives of information intermediaries to manipulate information disclosed to interested parties.⁶ Since Akerlof's (1970) "lemons markets" paper, it has been recognized that information intermediaries may play a crucial role for markets under adverse selection (see Biglaiser 1993). Boot, Milbourn and Schmeits (2006) show that intermediaries can help to coordinate on a desired equilibrium. However, if an intermediary cannot perfectly assess the quality of the good and/or it has discretion about how the results of the assessment are communicated to buyers,

⁶The focus of the paper is on the role of rating agencies as information sellers. In this respect, the paper is related to a wide literature on financial intermediation, e.g. Allen (1990).

incentive problems may reduce the precision of information disclosed to the market.⁷ The central question we address is whether competition between rating agencies improves the information revealed to investors and other stakeholders, or whether (as suggested in the introduction) it leads to a “race to the bottom.”

There is a literature addressing crowding-out of private information due to reputation concerns in a competitive environment. These papers consider cheap-talk models (Crawford and Sobel 1982) in which intermediaries are concerned with establishing a reputation of being well informed. For example, Scharfstein and Stein (1990) and Ottaviani and Sorensen (2006a, 2006b, 2006c) study the impact of reputation concerns on the reports of analysts. In order to signal its ability to provide information with high precision, the intermediary biases its private observation in favor of prior belief. Mariano (2006) applies the cheap talk model in the context of rating agencies.

The complexity of information and participation of naive buyers can lead to biases in information reporting. Skreta and Veldkamp (2008) study how the higher complexity of rated assets affects incentives for ratings shopping. They show that the ability of sellers to compare ratings from different CRAs before the ratings are disclosed to the market leads to ratings shopping and ultimately inflates ratings. Bolton, Freixas and Shapiro (2008) show that a CRA may overstate the seller’s quality when there are more naive investors. Pagano and Volpin (2008) argue that the issuers of structured bonds (sellers) prefer coarse and opaque ratings to expand demand from sophisticated investors (buyers).⁸

⁷Grossman and Hart (1980), Grossman (1981) and Milgrom (1981) analyzed direct communication between the buyer and the seller. In their framework, a seller can disclose any information but must include the true type in its report. Then, the conjecture of the buyer is that the true type is the most pessimistic element of the report, and full disclosure obtains. The key distinction of our approach is that a rating agency provides information about multiple sellers. It permits clustering different sellers into one rating class, and the unraveling need not happen.

⁸There are other explanations of why an information intermediary might manipulate information. Manipulation can occur due to collusion between the intermediary and the seller; for example, Strausz (2005) shows that the threat of collusion makes honest certification a natural monopoly. Peyrache and Quesada (2005) argue that mandatory certification makes intermediaries more prone to collusion by increasing the participation of low quality sellers. Mathis, McAndrews and Rochet (2009) show that reputation is sufficient to discipline CRAs only when a large fraction of their incomes come from rating simple assets. Benabou and Laroque (1992) analyze the incentives of an intermediary to manipulate information when the intermediary also acts as a speculator on the market. Goel and Thakor (2010)

Our theory builds on Lizzeri (1999) who studies the optimal disclosure policies of a monopoly intermediary capable of perfectly ascertaining the quality of the seller and communicating it to the buyer. Lizzeri derives a unique equilibrium in which all sellers pay a positive fee for an uninformative rating. The logic of this result is as follows: the profit of a CRA is a product of market coverage and a uniform rating fee. The fee cannot exceed the willingness to pay of the lowest quality rated seller. By pooling this seller with better quality sellers into one rating, a CRA can charge a higher fee without reducing demand for its services. In contrast, Lizzeri then shows that competition leads to full disclosure and zero fees for certification.⁹

Risk neutrality is essential for Lizzeri's no disclosure result. It implies that buyers are ready to pay the same price regardless of whether the quality is known for sure or is uncertain. In other words, the buyer does not value the precision of information disclosed by the intermediary. We change this assumption and assume that buyers care about the precision of information contained in the rating. In this respect, our analysis is related to the literature on information quality and ambiguity aversion (Veronesi 2000; Epstein and Schneider 2008).

Limited empirical literature exists that investigates the impact of competition in the market for credit ratings. Closest to our paper is Becker and Milbourn (2011) who analyze changes in the informativeness of corporate bond ratings as a third credit rating agency, Fitch, grew in a market previously dominated by Moody's and S&P. Contrary to our work, these authors conclude that the increased competition from Fitch led to a decrease in the overall information content of ratings, i.e., "race to the bottom." To understand the differences between our paper and theirs, it is important to recognize that rating agencies can compete in different dimensions. Becker and Milbourn argue that, since competition

show that CRAs may have incentives to produce coarse ratings in order to reduce the litigation risk.

⁹In spite the fact that most information intermediaries function in oligopolistic markets, there is little research on the impact of competition on the disclosure of information. A few exceptions include Lerner and Tirole (2006) who study competition in standard settings; Farhi, Lerner and Tirole (2008) analyze the interaction between ratings shopping behavior and the transparency of certification; Morrison and White (2005) study banks' decisions to apply to regulators with different perceived abilities.

reduces rents, it thereby undermines the incentive to make costly investments in rating accuracy.

We examine a different mechanism used by rating agencies to compete: differentiation of the rating scales. In this sense, our paper is related to the extensive literature on insurance classification which shows that broadly pooled risk categories can be profitably cherry picked by rivals and that this process will lead to finer, and more informative classes (see Hoy 1982; Bond and Crocker 1991; Thiery and Schoubroeck 2006; and Crocker and Snow 2010). We provide further discussion on different aspects of competition and the quality of ratings in Section 6.

2 Ratings of a Monopoly Credit Rating Agency

2.1 Model

A CRA provides services that can lower the information asymmetry between buyers and sellers. Sellers have private information about their quality, v . The CRA and buyers share a common prior about the quality of a seller. For simplicity, we assume that v is distributed uniformly on $[0, 1]$, and higher v indicates higher quality.

A seller cannot credibly communicate its quality to buyers. A rating agency offers an evaluation service for a fee t and can perfectly observe the type v of a seller.¹⁰ The fee is flat for all sellers purchasing a rating,¹¹ and a CRA cannot screen sellers by demanding a higher fee for a more favorable rating.¹²

Having evaluated a seller, a CRA strategically communicates its results. The disclosure policy of the agency defines how rated sellers' quality types are communicated to buyers.

¹⁰Our results are robust to the specification where the rating agency observes the seller's type with some noise. Since noisy signals do not affect the logic of the results, we focus on the case of perfect signals in the paper. Further results are available from the authors.

¹¹In practice, CRA fees depend on the type of provided service, but do not depend on the assigned rating (Cantor and Parker 1994).

¹²A rated seller does not value an option to withhold its rating. Faure-Grimaud, Peyrache and Quesada (2009) show that firms may have incentives to hide their ratings only if they are sufficiently uncertain about their quality. In our setting, firms have perfect information about their quality, and they will not apply for a rating unless it increases their reservation price.

For example, under full disclosure, a CRA communicates the observed quality v . In general, a disclosure policy is a measurable function from the set of signals $[0, 1]$ into the set of Borel probability distributions on real numbers.

There is a unit mass of identical buyers. A buyer purchases at most one unit of a good from one seller. The buyer's willingness to pay for the good depends on the expected quality and the precision of information about quality.¹³ Precision is measured by the variance of quality conditional on the information available to buyers. For example, under full disclosure, the variance of quality of a rated seller is zero, and the precision is the highest. We model demand for precision by assuming that buyers have mean-variance preferences. Given information I available to buyers, the valuation of a good is equal to

$$u(I) \equiv E[v|I] - a\text{Var}[v|I],$$

where $E[v|I]$ is the expected quality, and $\text{Var}[v|I]$ is the variance of quality. $a > 0$ measures the marginal value of information precision to buyers. Buyers are price takers, and $u(I)$ is the price paid for the good. When $a = 0$, this model is equivalent to Lizzeri.

Obtaining ratings is voluntary to sellers. The expected payoff to a seller of type v depends upon a buyer's valuation and is equal to

$$\begin{aligned} &u_R(v) - t \text{ if it is rated, and} \\ &u_N(v) \text{ if it is not rated,} \end{aligned}$$

where $u_R(v)$ and $u_N(v)$ are the expected payoffs of type v with and without a rating, respectively. Denote δ the mass of sellers demanding a rating. Then the payoff of the rating agency is equal to

$$V = \delta t.$$

The game consists of three stages.

¹³The value of more accurate signals has been extensively analyzed in connection with earnings management. For example, earnings smoothing can increase the informativeness by enhancing the signal/noise ratio. Moreover, more informative earnings can command a stock price premium. See for example, Hunt, Moyer and Shelvin (2000) and Tucker and Zarowin (2006).

- I. Sellers learn their types. A rating agency designs its disclosure policy and sets a fee.¹⁴
- II. Sellers observe the disclosure policy of the rating agency and the fee. They decide whether to purchase a rating. The participating sellers are evaluated, and the results are disclosed to buyers according to the disclosure policy of the CRA.
- III. The buyers observe the disclosure policy and the rating if the seller is rated. They decide whether to purchase the credit sensitive product. Sellers receive a payoff, which depends on the rating status.

We study sequential equilibria of the game. The strategies of all players must be optimal at every stage of the game given the beliefs about other players' information. Beliefs must be consistent with the Bayes rule whenever possible.

2.2 Full Disclosure and the Benefits of Information Pooling

Analysis of full disclosure is useful to highlight the CRA's benefits from pooling information. Under full disclosure, a rating is a perfect signal about a seller's type. If a seller $\hat{v} \in [0, 1)$ decides to purchase a rating, better sellers $v > \hat{v}$ do the same because their payoff when rated is increasing in quality. Then, nonrated sellers must have a quality below $\hat{v}$. Also, if a seller $\hat{v}$ is not rated, a rating does not benefit the sellers with a lower quality $v < \hat{v}$. This intuition implies that, given a fee for ratings, there is a seller type indifferent between purchasing a rating or operating without a rating.¹⁵ The demand for ratings comes from higher seller types. The optimal fee for the rating agency and

¹⁴We assume that the CRA is able to commit to a disclosure policy and a fee. This assumption follows the literature on information disclosure, in particular Lizzeri (1999). It rules out the possibility of collusion between the seller and the intermediary. It would be interesting to extend the model to allow for collusion, but our focus in the paper is on a different question. The model with collusion would be useful to understand how the CRA can transmit information when it has to overcome the problem of credibility. In contrast, our focus in the paper is to explain the CRA incentives to manipulate information by designing the optimal rating scale. However, we believe that our results provide a useful benchmark for analyzing the questions of collusion and credibility.

¹⁵For this to hold, the fee should not exceed the gain of a rating to the highest type $v = 1$ when it is the only rated type, $1 - (\frac{1}{2} - \frac{1}{12}a) = \frac{1}{2} + \frac{1}{12}a$.

the resulting coverage of the market are derived from the trade-off between the marginal benefit of charging a higher fee and the marginal cost of the reduced demand for ratings.

Proposition 1 *Suppose that a monopoly rating agency commits to full disclosure, and the fee for the rating services is less than $\frac{1}{2} + \frac{1}{12}a$. Then the unique sequential equilibrium of the subgame has a threshold structure: there is a type $v_F \in [0, 1]$ such that all types above v_F purchase a rating, and no type below v_F is rated.*

Is full disclosure optimal for the CRA? Suppose that, instead of reporting the type v_F , the rating agency announces that this type is from an interval $[v_F, v_F + \Delta]$, $\Delta > 0$. Thus, v_F is pooled with better types, and the rating agency may be able to charge a higher fee without reducing the demand for ratings. This occurs when the valuation of pooled types is higher than the valuation of the lowest rated type; that is,

$$v_F + \frac{1}{2}\Delta - \frac{1}{12}a\Delta^2 > v_F.$$

If the marginal value of information is zero, $a = 0$, Lizzeri's result holds: all types should be pooled and assigned the same rating grade. When the precision of information matters, $a > 0$, pooling imposes a cost in lost precision, $\frac{1}{12}a\Delta^2$. This intuition suggests that the optimal disclosure policy trades off the benefits of pooling due to higher fees with the cost of pooling due to reduced precision.

2.3 Optimal Disclosure

We now analyze the profit maximizing disclosure policy of a monopoly rating agency. Our objective is to derive the optimal disclosure without putting any restrictions on how the CRA communicates the information about sellers' types to buyers. To do so, we define the CRA disclosure policy as a general mapping from the set of sellers' types into the set of Borel probability distributions on $[0, 1]$. This modeling choice includes the possibility of misrepresentation of sellers type by CRA, randomization, pooling of information, etc. Formally, a disclosure policy is a correspondence $s : [0, 1] \rightarrow [0, 1]$. The expected quality

$\mu(s(v))$ and the variance $\sigma^2(s(v))$ of type v rated $s(v)$ depend on the set of types that obtain the same rating,

$$\begin{aligned}\mu(s(v)) &= E[v' : s(v') = s(v)], \\ \sigma^2(s(v)) &= Var[v' : s(v') = s(v)],\end{aligned}$$

which results in buyers' valuations of seller type v equal to

$$u(s(v)) = \mu(s(v)) - a\sigma^2(s(v)).$$

Denote $V_R(s)$ the set of rated seller types, $V_R(s) \subset [0, 1]$, and $V_N(s)$ the set of nonrated types, $V_N(s) = [0, 1] \setminus V_R(s)$. Sellers purchase a rating only if it has a positive return,

$$u(s(v)) - t \geq \max\{u(V_N(s)), 0\} \text{ for all } v \in V_R(s). \quad (1)$$

The right-hand side of the inequality reflects that a nonrated seller trades only if its valuation without a rating is positive. Given any distribution of types $F(v)$, a disclosure policy $s(\cdot)$ generates demand

$$\delta(s) = \int_{V_R(s)} dF(v).$$

A strategy of the rating agency is a disclosure policy $s(\cdot)$ and a fee for the rating t . A strategy of each seller type is the decision to be rated. We restrict attention to pure strategies and study sequential equilibria of this game. In equilibrium, the following two conditions must be met. First, the disclosure policy is optimal for the rating agency,

$$(s(\cdot), t) \in \arg \max_{\tilde{s}, \tilde{t}} \delta(\tilde{s})\tilde{t}.$$

Second, the decision to obtain a rating is optimal for a seller. That is, for any $(s(\cdot), t)$ and strategies of sellers $[0, 1] \setminus v$, seller type v is rated if and only if (1) holds for this seller.

To analyze the optimal disclosure policy, we proceed in two steps. In the next proposition, we describe the structure of an optimal disclosure policy for any distribution of types $F(v)$. Then, in Proposition 3, we apply this result to solve for the policy in the context of our model where $F(v)$ is uniform.

Proposition 2 *An optimal disclosure policy of a monopoly rating agency has the following structure. There is a type $v_M \in [0, 1]$ such that all types $v \geq v_M$ are rated, and no type $v < v_M$ is rated. Types $[v_M, v_M + b_M]$, $b_M \geq 0$ and $v_M + b_M \leq 1$ are assigned the same rating. The fee charged for the rating is equal to the value of the rating of the lowest rated type, $t_M = u([v_M, v_M + b_M]) - \max\{u([0, v_M]), 0\}$.*

An optimal disclosure policy is similar to the discrete system of ratings employed by the major rating agencies. Under this system, a CRA partitions the set of realization of v in subintervals and discloses that its estimate of quality belongs to a subinterval.

The rating agency faces demand $\delta_M = 1 - v_M$ and earns profits

$$(1 - v_M)t_M$$

Denote $u(N)$ and $u(L)$ the valuation of nonrated types $N = [0, v_M]$ and the lowest rated types $L = [v_M, v_M + b_M]$, respectively. Then,

$$u(L) = v_M + \frac{1}{2}b_M - \frac{1}{12}ab_M^2, \quad (2)$$

$$u(N) = \max\left(\frac{1}{2}v_M - \frac{1}{12}av_M^2, 0\right),$$

$$t_M = u(L) - \max\{u(N), 0\}. \quad (3)$$

If a seller cannot trade without a rating, $u(N) < 0$, the fee is equal to the valuation of types in the lowest grade L . When $u(N) > 0$, a CRA charges the difference between the valuations of the lowest grade sellers and nonrated sellers, $u(L) - u(N)$.

In equilibrium, nonrated sellers N must be better off without a rating. If a seller $v \in N$ deviates and purchases a rating, the rating agency announces that the seller's quality is from the interval N . Then, the deviation is not profitable and purchasing a rating cannot increase the reservation price charged by these sellers.

An optimal disclosure policy of the rating agency solves

$$\max_{(v_M, b_M)} (1 - v_M)(u(L) - \max\{u(N), 0\}).$$

In the next proposition, we summarize the solution to this problem.

Proposition 3 *The optimal monopoly rating system is summarized in the following table.*

a	v_M	b_M	t_M	π_M	$\max\{u_N, 0\}$
$0 \leq a \leq 2$	0	1	$\frac{1}{2} - \frac{1}{12}a$	$\frac{1}{2} - \frac{1}{12}a$	0
$2 \leq a \leq 6$	$\frac{3}{4} - \frac{3}{2a}$	$\frac{1}{4} + \frac{3}{2a}$	$\frac{1}{4} + \frac{1}{24}a$	$\frac{(a+6)^2}{96a}$	$\frac{3(10-a)(a-2)}{64a}$
$6 \leq a \leq \frac{21}{2}$	$\frac{2}{3} - \frac{1}{a}$	$\frac{3}{a}$	$\frac{(a+3)^2}{27a}$	$\frac{(a+3)^3}{81a^2}$	$\frac{(2a-3)(21-2a)}{108a}$
$\frac{21}{2} \leq a \leq \frac{51}{4}$	$\frac{6}{a}$	$\frac{3}{a}$	$\frac{27}{4a}$	$\frac{27(a-6)}{4a^2}$	0
$a \geq \frac{51}{4}$	$\frac{1}{2} - \frac{3}{8a}$	$\frac{3}{a}$	$\frac{3}{8a} + \frac{1}{2}$	$\frac{(4a+3)^2}{64a^2}$	0

When the marginal value of information is relatively low, $a \leq 2$, all seller types are rated and pooled in the same rating grade. As the value of information increases, the mass of pooled types b_M decreases and the rating becomes more precise. The market coverage decreases in a when $u(N) > 0$, increases when $u(N) = 0$ and decreases when $u(N) < 0$. The profit of the rating agency is non-monotone in the value of information a . It decreases when $u(N) > 0$, increases when $u(N) = 0$ and decreases when $u(N) < 0$. Profit is the highest when the value of information is the lowest, $a = 0$. As $a \rightarrow +\infty$, the profit converges to $\frac{1}{4}$.

What is the mass of types pooled in the lowest rating? From (2), pooling db_M sellers in one rating increases the expected quality of $u(L)$ by $\frac{1}{2}db_M$ and reduces the precision of the rating by $(-\frac{1}{6}ab_M)db_M$. For low values of a , the increase in expected quality from pooling outweighs the precision cost and that leads to extensive pooling. For higher values of a , the interior solution obtained when the marginal increase in expected quality is equal to the marginal cost of reduced precision resulting in $b_M = \frac{3}{a}$. As the value of precision increases, the mass of types pooled in the lowest rating goes to zero.

The marginal value of information entails five disclosure policy regimes. When a is low, $0 \leq a \leq 2$, the optimal disclosure policy of the rating agency is to pool all sellers in the same rating grade. It shows that Lizzeri's "no disclosure" result is more general. It also holds when buyers have a relatively low value for the information precision of the rating.

For moderate information values, $2 \leq a \leq 6$, a monopoly rating agency has partial coverage of the market, $v_M > 0$, but all rated sellers are still pooled in a single rating

grade. This regime resembles a minimum standard setting. Reducing the coverage of the market is beneficial for the rating agency because it widens the difference between the valuation of rated and nonrated sellers, and allows the CRA to charge a higher fee for the rating. At the same time, the value of precision is too low to expand the number of rating categories, so all rated sellers are pooled in order to increase the expected quality of the lowest rated type.

As the value of information increases, $a \geq 6$, providing precision becomes more valuable than increasing expected quality by pooling. The distinction between the last three regimes for $a \geq 6$ is the ability of nonrated sellers to trade. Higher demands for precision imply that rating becomes essential for trade, and the agency expands its coverage of the market for $\frac{21}{2} \leq a \leq \frac{51}{4}$. However, when the payoff of nonrated sellers is negative, $u(N) < 0$, precision becomes secondary to improving the pool of rated companies. Coverage is increasing for $a \geq \frac{51}{4}$.

The profit of the rating agency is non-monotone in the value of information. For relatively low values, the CRA can benefit from its unique ability to screen sellers and selectively disclose the results. However, as the value of information increases, the optimal rating system requires finer information disclosure, and the CRA cannot increase the fee by pooling types in one rating.

Figure 1 shows the boundaries for rating L as a function of a . Types located below the lower curve are not rated. Types located between the lower and the upper curves are pooled in the lowest rating grade L . As with full disclosure, the coverage of the market is non-monotone in a and depends on the ability of nonrated companies to trade. Coverage is decreasing in a for low information values because the rating agency has an incentive to widen the gap between the valuations of rated and nonrated sellers.

The optimal disclosure policy of the monopolist admits multiple equilibrium rating scales for types $[v_M + b_M, 1]$. As long as these sellers are willing to purchase a rating, the CRA is indifferent among all disclosure policies. We focus on one equilibrium type, the interval disclosure, that is employed by the majority of credit rating agencies. In the

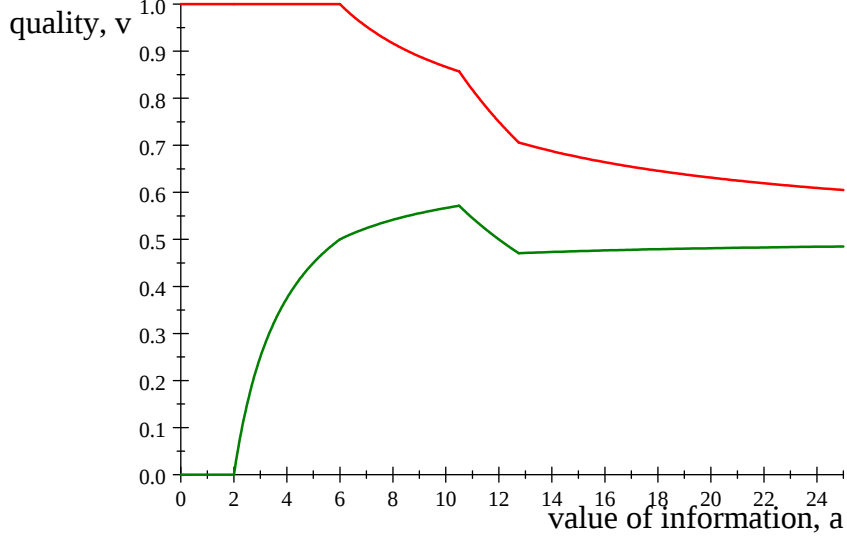


Figure 1: Optimal Rating Scale of a Monopoly CRA

next proposition, we derive a necessary condition for an interval disclosure policy to be an equilibrium, and we show how the width of a rating interval changes with the marginal value of information.

Proposition 4 *Any system of intervals $(R_1, \dots, R_N)$, $N \leq +\infty$ that satisfies $b_{k+1} \leq b_k + \frac{6}{a}$ is consistent with the optimal disclosure policy. As the value of precision increases, the width of the rating interval decreases.*

Interval disclosure policies are not equivalent from the seller's perspective. In each pooling interval, the types at the bottom of the interval benefit from pooling at a cost of types on the top of an interval. It is immediate to show the following.

Proposition 5 *In each pooling interval, the mass of types that prefer full disclosure is greater than the mass of types that prefer pooling, and the difference is increasing in the value of precision a .*

This result implies that the number of rated sellers that are willing to pay a positive price for a second rating to improve the precision of the signal about their quality is

increasing in the value of precision to buyers. In the next section, we study entry under the assumption that the incumbent CRA uses interval disclosure. Our motivation is twofold. First, we show that there is an equilibrium that is consistent with industry practice. Second, we show that interval disclosure allows entry in multiple segments of the market. Indeed, if there are segments of the market where the incumbent makes sellers' types perfectly known to the buyers, a new rating agency gains no benefit from entering these segments.

The analysis of the optimal disclosure policy has been conducted under the assumption that the sellers and the CRA can perfectly observe the type. In this setting, the sellers do not value the option to withhold the rating. Indeed, when the seller solicits a rating, he can perfectly infer the rating that he will be assigned from the CRA's rating system. Thus the value of a rating to a seller does not change after the CRA observes the seller's type.

Withholding a rating may become valuable when the CRA receives an imperfect signal about seller's type. A higher quality seller may choose to withhold the unfavorable rating when its value is sufficiently below seller's type. Compared to the situation when CRA's signals are perfect, this will produce two effects. First, it will increase the average quality of the pool of unrated sellers. The reason is that higher quality sellers with a low rating are more prone to withhold the rating than lower quality sellers with a high rating. Second, the rating system will have finer rating intervals. Noisy signals reduce the precision of information contained in ratings, and will make pooling more costly in term of precision. However, qualitatively the rating system will remain the same. The reason is that a noisy signal about the seller's type does not change the trade-off that drives the interval disclosure. It would be determined by the balance between the benefit of pooling due to higher willingness-to-pay of the lowest rated type and the cost of pooling due to less precision under pooling.

3 Entry of a New Credit Rating Agency

We analyze the entry strategy of a new agency on the following time line. After the ratings have been purchased from the incumbent, but before the transactions between buyers and sellers, a new agency offers an additional rating for a fee. If a new rating agency attracts any sellers, these are rated by the entrant. Then, buyers form their valuations based on all available sellers' ratings (i.e., from the incumbent and the entrant) and trade takes place.

The main focus of this section is the transition from a monopolistic to a duopolistic structure. Our objective is to find which segments of the market served by the incumbent CRA can be attractive for the entrant. In this setup, the incumbent does not adjust its disclosure policy. Although understanding the long-term impact of competition is impossible without accounting for the incumbent's reaction to entry, our approach provides a clear intuition about the entry strategy of a new CRA.¹⁶

A seller will pay for an additional rating only if it increases its value in the eyes of the buyers. This occurs when the second rating allows the seller to signal that it has higher quality and/or when it improves the precision of information. If a seller is rated by the incumbent and the entrant, it must be that two ratings are better than one,

$$u(R_m, R_e; v) - t_m - t_e \geq u(R_m; v) - t_m,$$

where $u(R_m, R_e; v)$ and $u(R_m; v)$ are the payoffs of seller v when rated by both agencies and when rated only by the incumbent, respectively, and t_m and t_e are the fees of the two CRAs. If a seller is rated only by the entrant, then

$$u(R_e; v) - t_e \geq \max\{u(V_N), 0\}.$$

In the next proposition we characterize the demand for the entrant's ratings.

¹⁶In a supplementary section available online from the authors we develop a simplified version of the model with discrete types and derive the equilibrium rating systems of the incumbent and the entrant in a Stackelberg competition setting. We show that the properties of the equilibrium of the entry game are consistent with the results presented in this section.

Proposition 6 *Suppose that the incumbent CRA employs an interval disclosure rating system. The demand for ratings of the entrant comes from the sellers at the top of each rating interval of the incumbent. The interval disclosure rating system of the entrant is such that the average quality of a seller with an n th highest rating of the entrant is higher than that of a seller with the n th highest rating of the incumbent. Purchasing a second rating increases the precision of information about the seller.*

The first result implies that the decision to obtain a second rating need not be congruent with the quality of a seller. Higher-than-average sellers in each rating category of the incumbent CRA are the ones who can benefit from obtaining a second rating. The second result is that the rating scale of the entrant must be more stringent in the sense that a seller needs to meet higher standards to obtain a comparable rating. The intuition for this result is illustrated in Figure 2. In the figure, sellers $v \in [v_M, 1]$ are rated by the incumbent, while sellers $v \in [0, v_M]$ are not rated. The rating scale of the incumbent consists of two ratings, A and B . Lower quality sellers located in the interval N are not rated. In this example, there are three potential entry segments, each of which is located on the top of intervals A , B and N . It implies that a rating scale of the entrant consists of intervals $A_e = [x, 1]$, $B_e = [y, v_M + b_M]$ and $C_e = [z, v_M]$. The relationship between the two scales is such that $A_e \subset A$ and $B_e \subset A \cup B$. Hence, the average quality of sellers in each of these intervals is higher than the average quality under the incumbent's rating scale. As for the precision of the entrant's ratings, it is determined by the width of the rating interval. Interestingly, though the precision is always higher for the best entrant's rating A_e , it may be higher or lower for the entrant's lower ratings B_e and C_e . However, the ultimate precision of information about sellers with two ratings is always higher than the precision with a single rating prior to entry.

The mass of types that are disadvantaged by pooling is increasing in the marginal value of information (Proposition 5), so does the mass of potential entry segments for a new CRA. It implies that entry is more profitable in certification markets where the value of information is high.

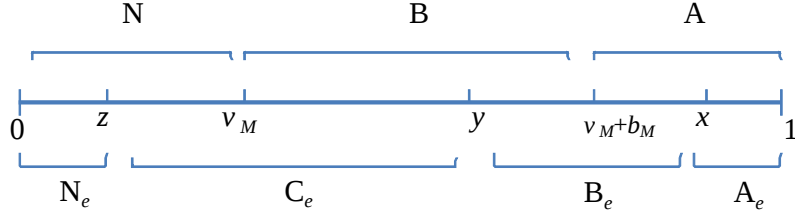


Figure 2: Demand for entrant's ratings

The exact boundaries (x, y, z) of the entrant's rating scale depends on the rating scale of the incumbent and on the marginal value of information to sellers. The entrant can target multiple groups, including unrated sellers. An optimal entry strategy is a trade-off between the fee the entrant can charge and its coverage of the market. The highest quality rated sellers are willing to pay the highest price to refine the information about their quality. However, charging high fees to this group reduces the demand from other sellers.

Proposition 6 does not rest on any distributional assumption over seller types. The particular incumbent rating categories in which the entrant chooses to enter depends on how the incumbent formed its rating categories in the first place. The size, number and thresholds of rating categories depend on the distributional assumption and the value of information a . For example, when the market coverage of the incumbent is low, rating nonrated companies may provide the highest value to the entrant. However, when the incumbent has substantial market coverage, but pools rated companies in wide rate grades, the entrant's profitable strategy is to offer the second rating to disadvantaged rated companies.

4 Institutional Setting and Data

We test the primary predictions of our theoretical model using the data on ratings of U.S. property-liability insurance companies in the remainder of the paper. Proposition 6 implies that (i) the entrant CRA will have higher standards, on average, in order for a

firm to receive a rating similar to the one received from the incumbent CRA; and (ii) the entrant CRA will obtain the greatest demand for its services from higher-than-average quality insurers within a rating class of the incumbent CRA. Proposition 5 means that (iii) insurers for which market participants have a more difficult time assessing the true financial strength of the firm will be more likely to seek an additional rating. We take advantage of a natural experiment, which began in the late 1980's and continued through the 1990's, when the well-known bond rating agency S&P entered the market for insurance ratings.¹⁷

In this section, we present the institutional setting for our tests and describe our data sources. In the following section we discuss our empirical tests and results from (i) univariate comparisons of the stringency standards employed by S&P and A.M. Best; and (ii) the determinants of the differences in the ratings assigned by the incumbent and the entrant CRAs while controlling for the strategic decision of an insurer to obtain a second rating. We conclude the empirical section by presenting an empirical analysis of the A.M. Best's reaction to S&P entry.

4.1 Insurance Ratings and Standard & Poor's Entry

Before the end of 1980s, the market for insurance ratings was largely dominated by the A.M. Best Company. Incorporated in 1899, A.M. Best publishes "financial strength ratings" for the majority of U.S. insurers, and, for most of its history, it was the only agency

¹⁷Although our theoretical results could be useful to analyze entry into the bond rating industry, none of the new NRSROs has obtained a significant market share to perform the tests. As noted by White (2002): "A striking fact about the structure of the industry in the U.S. is its persistent fewness of incumbents. There have never been more than five general-purpose bond rating firms; currently there are only three. Network effects—users' desires for consistency of rating categories across issuers—are surely part of the explanation. But for the past 25 years, regulatory restrictions (by the Securities and Exchange Commission) on who can be a "nationally recognized statistical rating organization" (NRSRO) have surely also played a role." Becker and Milbourn (2011) argue the significant growth of Fitch Ratings during the 1990s through 2006 can be thought of an increasing competition in the market for corporate debt ratings and they conduct an analysis that attempts to document the effects thereof. The market for insurance company ratings differs quite substantially as S&P has been providing rating opinions on corporate debt since the early 1900s. S&P's decision to begin offering insurance ratings truly represented a new entrant and thus is a more direct test of our model.

doing so. A.M. Best’s monopoly position began to erode after it was criticized following the liability insurance crisis of the mid-1980s and several natural catastrophes in the early 1990s bankrupted numerous insurers.¹⁸

The most aggressive agency to enter the market was S&P who entered in three phases.¹⁹ As shown in Table 1, phase 1 occurred in the late 1980’s when S&P announced it would publish “claims paying ability” ratings on property-liability insurers. Phase two began in 1991 when S&P introduced its “qualified solvency rating” service to complement its traditional ratings. Qualified ratings were unsolicited ratings based solely upon publicly available data. An important feature of the service was that no insurer could receive a rating above BBB. As a result, qualified ratings did not provide much new information to insurance buyers about individual companies. However, they advertised a new source of ratings to the insurance industry that used to view A.M. Best as the sole ratings supplier for several decades. The final entry phase occurred in late 1994 when S&P relaxed the BBB ceiling on qualified ratings. As shown in Table 1, this decision led to an increase in the demand for its services over the next couple of years. By the end of the 1990s, S&P provided full rating opinions to about 30–50 percent of insurers measured by the asset size.²⁰

The key feature of the entry strategy discussed in Section 3 is that the entrant offers ratings to companies that have already decided whether to be rated by the incumbent CRA. Consistent with the model assumption, almost all insurers in the sample obtained S&P rating as their second rating. The average percentage of firm-year observations where the insurer had an S&P rating and no Best rating in the prior year and then had both an S&P and Best rating in the current year represents only 0.14% of our sample. At the

¹⁸Winter (1991) overviews the liability insurance crisis in the U.S. that occurred between 1984–1986 while Lewis and Murdock (1996) describe the state of U.S. property insurance markets following several large natural disasters that occurred in the early 1990’s.

¹⁹The three phases have been documented in the financial press. See PR Newswire, 27 August 1987, “S&P launches insurance rating service”; Risk Management, June 1991, “S&P launches service to rate solvency”; National Underwriter, 25 September 1995, “S&P expanding P-C range.”

²⁰We focus on S&P’s solicited ratings in the paper as we assume the agencies learn private information that is partially revealed to buyers through the ratings assignment evaluation.

same time, the percentage of firms that requested a new rating from S&P that already had an existing A.M. Best rating is 83%.

4.2 Data

We gathered two sets of data for this study. The first documents the financial quality, business strategy and organizational characteristics of U.S. insurance companies, while the second reports the financial strength/claims paying ability ratings assigned to insurance companies by A.M. Best and S&P, respectively. Information on insurers' financial quality and other relevant characteristics comes from the annual regulatory statements of all property-liability insurers maintained in electronic form by the National Association of Insurance Commissioners (NAIC). We include all firms that meet our data requirements (discussed below). The ratings information for A.M. Best comes from Best's annual *Key Ratings Guide* (various years). We obtained the S&P ratings directly from S&P via a custom data order. Our data spans the years of S&P's entry into this market: 1989–2000.

5 Econometric Analysis and Results

5.1 Comparing Rating Stringency

Empirical strategy. We first seek to compare the stringency of the ratings assigned by the incumbent firm (A.M. Best) relative to the entrant (S&P). To do so, we need a statistic that summarizes the financial quality of each insurer in our data. Fortunately a single metric is sufficient to make this comparison since both A.M. Best and S&P have similar objectives for their rating systems.²¹ Thus, consistent with the ratings literature and with the agencies own stated objectives, our proxy for financial quality is the insurer's one-year

²¹A.M. Best describes its Financial Strength Rating as an "independent opinion of an insurer's financial strength and ability to meet its ongoing insurance policy and contract obligations." S&P describes their claims-paying ability rating as "an assessment of an operating insurance company's financial capacity to meet its policyholder obligations in accordance with their terms." Thus, both definitions suggest each agency's primary concern is to estimate the probability of insolvency. See A.M. Best (2004) and Standard & Poor's (1995).

probability of default. We estimate this probability using the discrete-time hazard model of Shumway (2001).

Following Shumway, the dependent variable for the hazard model, y_{it} , is a binary indicator that is set equal to 1 if firm i is declared insolvent in year $t + 1$ and equals 0 otherwise. Following the literature, we classify the year of insolvency when the first formal regulatory action is taken against a troubled insurer (Cummins, Harrington and Klein 1995). We use various industry and regulatory sources to identify the 300 plus property-liability insurers that failed between 1990 and 2001.²²

The explanatory variables used to estimate the model are the 19 balance sheet and income statement ratios that constitute the NAIC's Financial Analysis and Surveillance Tracking (FAST) system plus two additional control variables: one for firm size and the second an indicator variable for firms organized as a mutual or reciprocal insurers. These variables are similar to those used to model corporate debt ratings (e.g., Altman 1968 or Shumway 2001) but there are more of them and they are tailored to the specifics of this industry.²³

We estimate the probability of default for all insurers for which we have data to calculate the dependent and independent variables. Thus, not only do we include insurers rated by A.M. Best or S&P, but we also include insurer firm-year observations that do not receive ratings from either of these two agencies. In an effort to include as many insolvent observations in the analysis as possible, we include insurers who report data two years prior to their first event year, but who do not report in the year prior to their first event year. We delete any bankrupt firms for which we were unable to locate data within two years of their first event year. The final data set contains 24,236 solvent firm-year observations and 217 insolvent firm-year observations.

²²We use the NAIC's Report on Receiverships (various years), the Status of Single-State and Multi-State Insolvencies (various years) and the list of insolvent insurers provided by the A.M. Best Company (A.M. Best 2004).

²³Grace, Harrington and Klein (1995) suggest there are diminishing marginal returns to incorporating additional balance sheet and income statement ratios not already included in the FAST system except to include controls for firm size and organizational form.

To compare the strategies of the two CRAs, we need to define a mapping between their rating systems. For this study, we reviewed the verbal descriptions each agency ascribed to their individual rating categories at the time S&P began to enter the market for insurance ratings (A.M. Best 1991; Standard & Poor’s 1991). Identifying the upper-end of each agency’s ratings scale is straightforward, so we map S&P’s AAA rating into A.M. Best’s A++/A+ category.²⁴ For the other categories, we use the A.M. Best and S&P descriptions to map the ratings. Table 2 summarizes this mapping along with the verbal descriptions each agency uses to describe their system.²⁵ Also shown in Table 2 are numerical values assigned to each rating category.

Clearly, while A.M. Best and S&P may use similar descriptions of rating classes, the fact that they adopt different criteria, as predicted by theory, implies there will be an imperfect match between the two systems. We explore those differences next.

Results. Table 3 presents summary statistics of the data set used to estimate the hazard model comparing the solvent and insolvent samples while Table 4 Panel A displays the results of the discrete-time hazard regression model. The regression results are consistent with the summary statistics and with intuition. For example, highly leveraged firms, rapidly growing firms and firms that rely more heavily upon reinsurance to support their capital positions are associated with higher failure rates. Larger firms and insurers part of mutual organizations are relatively less likely to default. Overall the explanatory power of the model is reasonable as the pseudo R^2 statistic is 26 percent.

Panels B and C in Table 4 demonstrate the hazard model does a reasonable job identifying troubled insurers. For example, Panel B shows the average/median esti-

²⁴As discussed in more detail later in the paper, A.M. Best’s A++ and B++ ratings did not exist when S&P first entered the market for insurance ratings. Midway through our sample period, in 1992, A.M. Best announced the expansion of its rating scale to add finer distinctions in the A+ and B+ ratings by splitting each into two categories: A+ became A++ and A+ and B+ became B++ and B+. See A.M. Best (1991,1992).

²⁵The literature contains several additional ways to define mappings between ratings across CRAs including an alternative four-category system proposed by Pottier and Sommer (1999) and a five-level system used by the Government Accountability Office (1994) and Doherty and Phillips (2002). Although some of the details are different, the overall results discussed in this paper are robust to the choice of mapping across the two rating systems.

mated probability of default for the healthy firm-year observations is 0.8/0.2 percent while the average/median statistics for the firms in the year before they become bankrupt is ten/twenty times larger at 9.2/4.4 percent. The classification table in Panel C shows the model correctly predicts 82(83) percent of the insolvent(solvent) firm-year observations.

Table 5 displays average and median estimated probability of default statistics across the two agencies by year. The column labeled "t-test" reports the results of the null hypothesis of equal means for the probability of default for A.M. Best rated insurers versus S&P rated insurers. Figure T5, which displays the average probability of default time series for each agency, clearly shows that the average a probability of default for firms that requested an S&P rating were always lower than for the A.M. Best rated firms through 1998 (except for 1989 when S&P only rated three insurers). These difference are always statistically significant at a 1 percent p-value based on a one-sided t-test (see accompanying table).²⁶ To give economic significance, the analysis suggests that from 1989 through 1998, the one-year default probability for the average insurer that requested a rating from S&P was 56 percent lower than the probability of default for the average firm rated by A.M. Best. We also note that differences in average or median quality appear to converge over time as S&P's presence in this market grew more substantial. This result is most easily seen in the t-statistics, which generally decline over time and, in 2000, is insignificant.

Table 6 and the accompanying chart, Figure T6, show the average rating issued by each agency where the numerical scores assigned to each category shown in Table 2. The results stand in stark contrast to those reported in Table 5 as the average ratings assigned by the two agencies were very similar even though the probability of default for the average insurer that requested a rating from S&P was significantly lower. This result is consistent with our hypothesis that the entrant CRA has higher standards in order for a firm to receive a rating similar to the one it received from the incumbent CRA.

²⁶The statistics for the null hypothesis of differences in medians provide similar results and thus are not shown.

The final univariate comparison is shown in Table 7 where we compare all firm-year observations over the years 1989–2000 that received a rating from both A.M. Best and S&P. Each cell of the matrix, c_{ij} , equals the number of firm-year observations that receive rating i from A.M. Best and rating j from S&P, where $i, j \in \{\text{Superior, Excellent, Good, Marginal}\}$. The results suggest that S&P generally agreed with the rating A.M. Best assigned to the same insurer as the two agencies issued the identical rating to 57.7 percent of the firm-year observations. However, when S&P differed it was generally more pessimistic as 38.3 percent of the firm-year observations represent cases where S&P assigned a lower rating to the insurer than did A.M. Best. S&P assigned a higher rating to an insurer than did A.M. Best in only 4.0 percent of the firm-year observations. Again, this result is consistent with our theory.

5.2 Investigating the Differences in Ratings

The univariate results suggest that insurers who opted to receive an S&P rating were, on average, of higher financial quality, yet the average rating they received was either the same or lower than the rating they received from A.M. Best. Although consistent with our theory, this analysis suffers from two shortcomings. First, we have only been able to compare the two rating systems for firms that received a rating from both agencies. Thus, we have not ruled out the possibility that differences in the assigned ratings may be distorted by a potential selection bias between insurers that chose to be rated by the new entrant and those that did not. Second, the hazard model used to calculate the one-year probabilities of default was estimated using only publicly available information. Presumably, one of the advantages of a rating system is the ability of a CRA to learn private information. In this section, we investigate the determinants of differences in the ratings assigned by the incumbent versus the entrant CRA while controlling for the private information the agencies learn through the rating process and the strategic incentives of the insurer.

Empirical Strategy. Assume the following model is used by the incumbent rating

agency to determine the rating for a particular firm:

$$r_{if} = \alpha_i + \beta_i' X_f + \varepsilon_{if} \quad (4)$$

where

r_{if} = rating issued to firm f by the incumbent agency

α_i = constant term for the incumbent agency

β_i = vector of coefficients summarizing the incumbent agency's rating technology

X_f = vector of observable information for firm f

ε_{if} = error term of the incumbent agency's rating of firm f

In addition, assume the new entrant has a model of similar structure. We want to explain differences between the new entrant's ratings and the incumbent's; that is,

$$r_{ef} - r_{if} = (\alpha_e - \alpha_i) + (\beta_e - \beta_i)' X_f + (\varepsilon_{ef} - \varepsilon_{if}), \quad (5)$$

where all variables subscribed with an e denote the new entrant agency. As discussed by Cantor and Packer (1997), estimating equation (5) directly using OLS leads to biased results if the decision to seek a second rating is correlated with the ratings assigned by that agency. This selection bias makes it impossible to know if the estimated differences between the two rating systems are due to differences between the rating scales or whether the sample of firms that choose to get a second rating have a common set of characteristics. The theory presented in this paper suggests that firms which elect to receive a second rating are those with higher than average financial quality in their rating category and have some belief that they are likely to obtain a favorable outcome from the entrant. Thus, the average rating difference based upon univariate comparisons may underestimate the true difference in standards across the two rating systems.

We employ a standard Heckman sample selection methodology to estimate the true difference in standards across the two agencies. The methodology is ideal in this setting not only because it allows us to control for possible strategic behavior by the insurers but it also allows us to incorporate private information garnered in the ratings process. The

empirical procedure is as follows: we first estimate a Probit regression that models the insurer’s decision to request a second rating by S&P; second, we use the results of the Probit regression to estimate an inverse Mill’s ratio, which, when included in the ratings difference model, controls for the selection bias. Thus, in the second stage, we estimate

$$r_{ef} - r_{if} = \alpha + \gamma IMR_f + n_f,$$

where the estimated constant term measures the mean unbiased difference in ratings standards across the two agencies, and the coefficient on the inverse Mill’s ratio (IMR) captures the sample selection effect.²⁷ We hypothesize α will be negative consistent with our theory that the new entrant, on average, will employ higher rating standards than the incumbent firm. In addition, we predict the estimated coefficient γ will be positive consistent with the hypothesis that insurers who believe that they will receive a favorable rating from S&P will self-select to receive that rating.

The following variables are used to explain the demand for a second rating. First, our theory suggests that higher quality insurers within each rating category of the incumbent will have a stronger demand for a second rating from the entrant. To test this hypothesis, we construct, for each insurer, a variable that equals the median probability of default of all insurers within the same A.M. Best’s rating category as the insurer minus the estimated default probability for that insurer. Insurers for which this variable is positive have lower expected default rates than the median insurer in their same A.M. Best rating class. We expect a positive estimated coefficient on this variable.

Second, we include a variable that interacts the “higher quality than the median insurer” variable with a time trend to control for differences in demand over time. We expect insurers with the highest demand will choose a rating from the new entrant first

²⁷Note that differences in rating standards can be due to a shift in the cardinal ranking across the two systems (i.e., differences in the intercept terms) or due to different weightings employed by the two agencies (i.e., differences in the beta coefficients). We are unaware of any theory that can guide us in selecting exogenous variables that might explain why two agencies might place different weighting on rating factors. Therefore, like Cantor and Packer (1997), we only include an intercept term and a control for sample selection bias in the second-stage rating difference regressions.

and this effect will dissipate over time. Thus, we expect a negative coefficient on the interacted variable.

Our model predicts that rating agencies have incentives to reveal more information when the buyers value information more highly. We test this hypothesis by including a variable equal to the percentage of the insurer's premiums in retail lines of insurance.²⁸ as we expect these policyholders place less value on information for two reasons. First, state-guarantee funds provide more complete protection to the retail policyholders than to commercial buyers. And second, because most consumer lines of insurance are amenable to objective data and modeling, and thus the insurer's underwriting portfolio is relatively less opaque. For both reasons we expect that retail consumers value the additional information revealed by obtaining the second rating less relative to commercial buyers of insurance.

We include three variables to control for the insurer's previous experience with S&P. The first is an indicator variable equal to one for any insurer that is part of a group that has corporate debt already rated by S&P. Second, we include an indicator equal to one if the insurer requested a claims paying ability rating from S&P in the previous year. We expect that insurers with an existing relationship with S&P are more likely to request a second rating from the entrant because the managers are already aware of the methodology employed by S&P and therefore can forecast better the outcome of the review process.

Our third variable on previous S&P experience equals one if the insurer was assigned a qualified rating by S&P in the previous year and zero otherwise. Recall, in 1991, S&P introduced its qualified rating service where they assigned ratings to the insurer, without the insurer's consent, based upon quantitative analysis only. From 1991–1994, the highest qualified rating S&P would assign was BBB. In 1994, they relaxed this constraint and began issuing A, AA and AAA ratings on a qualified basis. Thus, we interact the qualified

²⁸The lines of insurance we considered to be retail included personal automobile insurance (both liability and property damage), homeowners insurance and farmowners insurance.

rating indicator with a dummy variable for the years 1990–1994 and another dummy for year 1995–1999 to control for the differences in the qualified rating scheme across time. We do not have strong priors regarding the expected signs on these indicator variables because it is not clear if consumers viewed qualified ratings as informative since S&P only used publicly available information to determine the rating.

We include two variables to control for differences in insurer complexity: a size variable (equal to the natural logarithm of the firm’s real assets deflated using the CPI) and a variable measuring the geographical concentration of the firm’s premium writings (a Herfindahl index of the premiums written across each state in which the insurer operates). We expect a positive coefficient on the firm size variable and a negative coefficient on the Herfindahl index consistent with the hypothesis that larger and more geographically dispersed insurers are more complex.

The final insurer characteristic variable we include is an indicator variable for whether an insurer is organized as a mutual or reciprocal. We have two competing hypotheses for this variable. First, the managerial discretion literature predicts that mutual insurers underwrite less risky lines of insurance and have more transparent business models than stock insurers due to the reduced ability of a diffuse set of owner/policyholders to monitor management (Mayers and Smith 1987). This logic suggests that mutuals should have a lower demand for a second rating. An alternative is that stock insurers have lower demand for an additional rating because their financial quality is already conveyed in share price.

Finally, we include indicator variables for each rating category assigned by A. M. Best to the insurer for two reasons. First, we wish to control for the private information that A.M. Best learns during the rating process. Second, we wish to control for any differences in the demand for a second rating based upon the insurer’s overall credit quality as judged by A.M. Best. We do not have any prior hypothesis for these variables.

The data for the rating difference tests includes any insurer that received a rating from A.M. Best over the time period 1990–2000 (we lose one year of data due to our inclusion of two lagged variables). We estimate the first-stage Probit regression using all

firm-year observations and the second-stage OLS regressions only for insurers that receive full ratings from both A.M. Best and S&P. There are 13,353 insurer-year observations in the A.M. Best sample of which 1,503 also obtained an S&P rating.

Results. Table 8 displays the summary statistics for all the variables used in the ratings differences tests. The columns labeled “A.M. Best and Standard & Poor’s” display summary statistics for just those insurer-year observations that were assigned ratings by both agencies. The columns labeled “A.M. Best Only” are insurer-year observations for firms that requested a rating only from A.M. Best. The most important statistic in this table is the average difference in rating assigned to an insurer by S&P relative to A.M. Best, which, over the time period of this study, was 0.35 notches lower. The results in Table 8 also suggest that insurers who requested an S&P rating were (i) less likely to be a members of a mutual group of insurers than the average A.M. Best insurer, (ii) were much larger than the average A.M. Best insurer, and (iii) conducted less business in retail lines of insurance. All of these latter results are consistent with our prior hypotheses.

Table 8 also shows that over 70 percent of the insurers that requested a rating from S&P in the current year already had a corporate debt rating from S&P. Likewise, close to 70 percent of S&P rated insurers already had an S&P claims paying ability rating from the previous year. The corresponding percentages for insurers that only had a rating from A.M. Best were 19.6 and 1.7 percent, respectively.

The first-stage Probit regression results are shown in Table 9. Panel A displays the estimated coefficients on the independent variables. Panel B displays the marginal effects.²⁹ The most important results concern the test that higher-than-median quality insurers in each rating category of the incumbent are more likely to demand a second rating. The estimated coefficient on variable measuring the difference between the median insurer’s probability of default in the A.M. Best rating category and the estimated default probability for the insurer is both positive and significant as predicted. At the same time,

²⁹The model also contains year indicator variables but the results are suppressed to save space. The full results with the estimated coefficients for the year indicator variables are available upon request.

the interaction of this variable with the time indicator is negative. This suggests that the effect of S&P's entry strategy to target the best companies in each rating category declined over time.

The insurer characteristic variables suggest that larger insurers, those writing across geographical locales, firms organized as stock insurers and those writing a greater percentage of their business in commercial lines of insurance are more likely to request a second rating from S&P. These results are all consistent with our theory that more complex insurers and insurers whose customers place a higher value on information have a greater demand to communicate their financial strength to the market place.

The previous S&P relationship variables are also consistent with our prior hypotheses as the estimated coefficients on both the corporate debt rating indicator and the previous claims paying ability rating indicator are positive and significant. The indicator variables for S&P assigning a qualified rating to an insurer are largely insignificant suggesting that the insurers' management viewed these ratings as having little ability to reveal new information to market participants.

The results for the indicator variables controlling for A.M. Best rating classes reveal the likelihood of requesting a second rating from S&P increases as the quality of the insurer (as rated by A.M. Best) also increases. Note, however, the economic significance of these effects are relatively small as the estimated marginal increase in probability for an Excellent or Superior rated A.M. Best insurer (over a Marginal insurer) is only 1.5 or 1.7 percent, respectively.

The results from the second stage OLS regressions where the dependent variable is the difference between the rating assigned by S&P minus the rating assigned by A.M. Best in year t are shown in Table 10. The inverse Mill's terms in each model were calculated using the results of each Probit regression shown in Table 9, respectively.

Our first conclusion is that we find evidence of a selection effect because the estimated coefficient on the inverse Mill's ratio is always positive and significantly different from zero. Thus, we find strong evidence that insurers who seek a second rating are acting

strategically and expect, on average, to receive a favorable rating from S&P even after we control for publicly available information to market participants and for the private information revealed during the ratings process. For example, based on the results shown in Model 4, the insurers strategically choosing a second rating expect to receive, on average, a 0.08 higher rating from S&P than they received from A. M. Best. These results support the contention that better than average insurers within any pre-existing A.M. Best rating category sought to differentiate themselves through a second rating.

The second conclusion we draw is that S&P maintained significantly higher standards relative to A.M. Best because the intercept term in each model is negative and significantly different than zero. Focusing again on Model 4, we see that the estimated mean difference in rating standards is 0.43 grades lower on the S&P scale than under the A.M. Best rating system. This result supports our theory that, conditional upon the rating provided by the incumbent agency, the insurers sought to differentiate themselves by seeking an additional rating from a new entrant agency that required higher standards in order to maintain the same rating.

5.3 A.M. Best’s reaction to S&P’s entry

A prediction from our theoretical model is that entry by a new CRA will provide higher than average quality firms the opportunity to differentiate themselves from firms pooled in the incumbent’s broad rating category but who are of lower quality. One possible strategy the incumbent can undertake that might either thwart entry and/or deter defection to the new entrant is to release voluntarily additional information on these higher than average quality customers so they can differentiate themselves without having to use the services of the entrant. Although by itself not conclusive evidence, it appears A.M. Best may have attempted to do exactly this.

As discussed earlier, prior to 1992, the highest category in the A.M. Best rating system was A+ and the highest of the “Secure” ratings was B+ (see Table 2 and accompanying discussion in Section 5.1). In 1992, just a few years after S&P’s entry into this market,

Best announced it would expand its rating scale to add finer distinctions among the A+ and B+ rating categories. It did so by splitting its A+ rating into A++ and A+ and its B+ category into B++ and B+ along with the pronouncement that the standards necessary to be assigned the new ratings would be higher than what previously existed.

Table 11 Panel A displays the average probability of default for all firms judged to be either A+ by A++ A.M. Best and for all firms rated AAA by Standard & Poor's over the time period 1989-2000.³⁰ The average probability of default for an insurer judged to be AAA by S&P over the years 1989-2000 was 0.16 percent which compares quite favorably to firms rated A++ by A.M. Best which had an average probability of default from 1991 – 2000 of 0.18 percent. Consistent with intuition, the average probability of default for an A.M. Best A+ rated insurer was much higher at 0.27 percent. Similarly, the average ratio (across years) of the average probability of default for a S&P AAA rated firm relative to an A.M. Best A++ rated firm was 0.993 suggesting Best and S&P required similar standards to garner their top ratings. Also consistent with intuition, the average (across years) of the average probability of default for a S&P AAA rated firm relative to an A.M. Best A+ rated firm was much lower at 0.653.

Panel B shows similar data but this time compares the average probability of default for firms rated BBB by S&P relative to firms assigned B+ and B++ ratings by A.M. Best. Overall the conclusions to be drawn are similar to those discussed above. That is, in the important rating category that separates “Secure” insurers from ones judged to be “Vulnerable”, Best voluntarily created a finer rate classification scheme and provided additional information to market participants relative to what it revealed in previous years.

³⁰There are no A++ or B++ observations for 1989 and 1990 as these years precede Best's decision to create the finer distinction in the A+ and B+ rating categories.

6 Competition and the quality of ratings

Our results provide strong evidence that an entrant CRA competes with the incumbent by differentiating the rating scale. In the context of our model, the results suggest that the entry makes ratings more informative. However, it is important to recognize that competition among CRAs may involve other dimensions that are outside the scope of the paper. In particular, in this section we discuss the recent literature that analyzes the interplay of competition and incentives to invest in ratings accuracy.

Reputation can provide a monopolist with incentives to produce high quality goods in markets where buyers have limited ex-ante information about quality.³¹ The basic mechanism is that higher reputation leads to higher future monopoly rents that exceed the short-term benefits of producing low quality goods. Since competition reduces future rents, it also weakens the incentives to invest in quality and leads to race to the bottom. Becker and Milbourn (2011) test this hypothesis on corporate bond ratings and report cross industry evidence that, with increased competition, there was more rating inflation, and that ratings became poorer predictors of both bond prices and default probabilities.

Our paper differs in several respects. First, we focus on a different mechanism: differentiation of the rating scale. Following Lizzeri (1999), we assume that the rating agency perfectly observes seller's quality type and chooses how much information to disclose. Thus in the context of the model, more refined rating scale implies more informed ratings. When information acquisition is costly, competition may also lead to less informed rating agencies and/or less accurate ratings. It is, of course, quite possible that both mechanisms, differentiated rating scales and lower investments in the quality of risk modeling, could work simultaneously but in opposition. It would then be an empirical question on which effect dominates.

Second, our empirical evidence comes from insurance ratings and not bond ratings. Our results are supportive of the model where a new rating agency designs a rating

³¹E.g. see Benabou and Laroque (1992), Mailath and Samuelson (2001), Tadelis (2002), Bar-Isaac and Shapiro (2011). Also Hörner (2002) shows that competition can increase reputation incentives.

scale that allows higher than average firms in each rating category of the incumbent to signal their higher quality. The summary statistics in Table 8 suggest that insurers organized as publicly traded corporations, those that write more complex commercial lines of business, and insurers that are larger and that operate in more states and in more lines of business are more likely to request S&P's services. Given that there is no regulation that requires insurers to obtain the S&P rating, these characteristics are all consistent with large, successful but more opaque organizations seeking to differentiate themselves through their use of multiple ratings.

Third, the two sectors, insurance and bond ratings, differ in terms on the number of incumbent rating agencies. In case of bond ratings analyzed by Becker and Milbourn (2011), the industry transitioned from two to three CRAs, while in case of insurance the transition was from a monopoly CRA to a duopoly. The insurance industry data suggests that the transition is favorable and that competition from a new entrant does lead to improved prediction of insolvency risk. Also the ability of Moody's and S&P to sustain their position in corporate bond market in spite of Fitch entry can be due to a common requirement to hold at least two investment grade ratings. At the same time, there is no requirement of multiple ratings in insurance industry. Before S&P entry, majority of insurance companies were rated only by A.M. Best.

In general, a race to the bottom effect can also be limited by the fact that the value of a rating is heterogeneous across different types of issuers. While lower quality issuers may prefer to escape the rigor and purchase inaccurate ratings, higher quality issuers value accurate certification. Other examples with a similar trade-off include the regulation in the financial sector in London versus New York, and the differences in accounting rigor among countries.³² Santos and Scheinkman (2001) analyze competition among exchanges and show that in a market where traders are heterogeneous with respect to credit risk and intermediaries design contracts to attract trading volume, competing exchanges set higher collateral requirements than a monopoly.

³²We thank the referee for suggesting these examples.

The use of ratings in regulation may also have an adverse effect on the quality of ratings. Kisgen and Strahan (2009) investigated another form of entry when the SEC assigned NRSRO status to a fourth credit rating agency in the mid 1990s. The authors showed that the SEC's certification of the Dominion Bond Rating Service led to a statistically significant reduction in the cost of debt capital for the firms already rated by Dominion prior to its certification. The authors suggests that their results are driven not by an increase in the information available to market participants, but instead because there was an increase in demand for the debt issued by firms that received a bond rating from Dominion following the agency's certification by the SEC. This effect is not present in our study because, unlike bond issuers, insurance companies are not required by regulation to be rated to operate in the insurance market.

7 Conclusion

In this paper, we examined how the entry of a new credit rating agency to the market served by an incumbent monopolistic credit rating agency changes the information content of ratings. In doing so, we departed from the literature that has largely ignored the endogeneity of the CRA's rating standards and shown that an entrant's rating standards may be significantly different from those used by the incumbent agency. The theory presented here predicts new entrants have incentives to require higher standards relative to the incumbent CRA in order for a firm to achieve a similar rating. The empirical analysis of the market for insurance ratings is consistent with this hypothesis as we show, all else equal, insurers of the same quality received a lower rating by the new entrant relative to the incumbent.

From a policy perspective, transparency of rating standards and rating performance are crucial to achieve the benefits of competition. Otherwise, differences in rating standards across multiple rating agencies are likely to create confusion about the meaning of ratings and decrease the precision of information. Also it is important that the CRAs compete in providing the information content rather than maintaining regulatory licences.

Recent Dodd-Frank regulation has taken steps in these dimensions in requesting more disclosure about the ratings process and reducing the regulatory reliance on ratings.

References

- [1] Akerlof, George A, 1970. The Market for 'Lemons': Quality Uncertainty and the Market Mechanism. *Quarterly Journal of Economics* 84(3), 488-500.
- [2] Allen, Franklin, 1990. The Market for Information and the Origin of Financial Intermediation. *Journal of Financial Intermediation* 1(1), 3-30.
- [3] Altman, Edward I., 1968. Financial Ratios, Discriminant Analysis and the Prediction of Corporate Bankruptcy. *Journal of Finance* 23, 589-609.
- [4] A.M. Best, 1991. Best's Insurance Reports: Property/Casualty United States 1991 Edition (A.M. Best Company, Inc., Oldwick, NJ.)
- [5] A.M. Best, 1992. Best's Insurance Reports: Property/Casualty United States 1992 Edition (A.M. Best Company, Inc., Oldwick, NJ.)
- [6] A.M. Best. 2004. Best's Insolvency Study: Property/Casualty U.S. Insurers 1969-2002 (A.M. Best Company, Inc., Oldwick, NJ.)
- [7] A.M. Best, various years, Best's Key Rating Guide: Property/Casualty Edition (A.M. Best Company, Inc.; Oldwick NJ).
- [8] Bar-Isaac, Heski and Joel Shapiro, 2011. Credit Ratings Accuracy and Analyst Incentives. *American Economic Review* (Papers and Proceedings), forthcoming.
- [9] Becker, Bo, and Todd T. Milbourn, 2011. How did Increased Competition Affect Credit Ratings? *Journal of Financial Economics*, forthcoming.
- [10] Benabou, Roland and Guy Laroque, 1992. Using Privileged Information to Manipulate Markets: Insiders, Gurus, and Credibility. *Quarterly Journal of Economics* 107(3), 921-958.

- [11] Biglaiser, Gary, 1993. Middlemen as Experts. *RAND Journal of Economics* 24(2), 212-223.
- [12] Bolton, Patrick, Xavier Freixas, and Joel Shapiro, 2008. The Credit Ratings Game. Unpublished working paper.
- [13] Bond, Eric W., and Keith J. Crocker, 1991. Smoking, Skydiving and Knitting: The Endogenous Categorization of Risks in Insurance Markets With Asymmetric Information. *Journal of Political Economy* 99, 177-200.
- [14] Boot, Arnoud W.A., Todd T. Milbourn, and Anjolein Schmeits, 2006. Credit Ratings as Coordination Mechanisms. *Review of Financial Studies* 19(1), 81-118.
- [15] Bradford, Michael, 2003. Big Changes at Kemper Prompt Buyer Concerns. *Business Insurance*, January 6.
- [16] Cantor, Richard, and Frank Packer, 1994. The Credit Rating Industry. *Quarterly Review*, Federal Reserve Bank of New York Summer, 1-26.
- [17] Cantor, Richard, and Frank Packer, 1997. Differences of Opinion and Selection Bias in the Credit Rating Agency. *Journal of Banking and Finance* 21, 1395-1417.
- [18] Cantor, Richard, Owain Ap Gwilym, and Stephen Thomas, 2007. The Use of Credit Ratings in Investment Management in the U.S. and Europe. *Journal of Fixed Income* 17(2), 13-26.
- [19] Crawford, Vincent P., and Joel Sobel, 1982. Strategic Information Transmission. *Econometrica* 50(6), 1431-1451.
- [20] Crocker, Keith J., and Arthur Snow, 2010. Multidimensional Screening in Insurance Markets with Adverse Selection. *Journal of Risk and Insurance*, forthcoming.

- [21] Cummins, J. David, Scott E. Harrington, and Robert W. Klein, 1995. Insolvency Experience, Risk-Based Capital, and Prompt Corrective Action in Property-Liability Insurance. *Journal of Banking and Finance* 19, 511-527.
- [22] Doherty, Neil A., and Richard D. Phillips, 2002. Keeping up with the Joneses: Changing Rating Standards and the Buildup of Capital by U.S. Property-Liability Insurers. *Journal of Financial Services Research* 21, 55–78.
- [23] Epermanis, Karen, and Scott E. Harrington, 2006. Market Discipline in Property/Casualty Insurance: Evidence from Premium Growth Surrounding Changes in Financial Strength Ratings. *Journal of Money, Credit and Banking* 38(6), 1515-1544.
- [24] Epstein, Larry G., and Martin Schneider, 2008. Ambiguity, Information Quality, and Asset Pricing. *Journal of Finance* 63(1), 197-228.
- [25] Farhi, Emmanuel, Josh Lerner, and Jean Tirole, 2008. Fear of Rejection? Tiered Certification and Transparency. Unpublished working paper. Harvard University.
- [26] Faure-Grimaud, Antoine, Eloic Peyrache, and Lucia Quesada, 2009. The Ownership of Ratings. *RAND Journal of Economics* 40(2): 234-257.
- [27] Government Accountability Office, 1994. Insurance Ratings: Comparison of Private Agency Ratings for Life/Health (U.S. Government Accountability Office, Washington D.C.)
- [28] Goel, Anand and Anjan Thakor, 2010. Credit Ratings and Litigation Risk. Unpublished working paper. Washington University in St. Louis.
- [29] Grace, Martin, Scott E. Harrington, and Robert W. Klein, 1995. An Analysis of the FAST Solvency Monitoring System, presented at the NAIC Financial Analysis Working Group meeting, Kansas City, MI.

- [30] Grossman, Sanford J., 1981. The Informational Role of Warranties and Private Disclosure About Product Quality. *Journal of Law and Economics*, 461–483.
- [31] Grossman, Sanford J., and Olivier Hart, 1980. Disclosure Laws and Takeover Bids. *Journal of Finance* 35, 323-334.
- [32] Heckman, James J., 1979. Sample Selection Bias as a Specification Error. *Econometrica* 47(1), 153-161.
- [33] Hörner, Johannes, 2002. Reputation and Competition. *The American Economic Review*, 92(3), 644-663.
- [34] Hoy, Michael, 1982. Categorizing Risks in the Insurance Industry. *The Quarterly Journal of Economics* 97(2), 321–336.
- [35] Hunt, Alister, Susan E. Moyer, and Terry Shevlin, 2000. Earnings Volatility, Earnings Management, and Equity Value. Unpublished working paper. University of Washington.
- [36] Kisgen, Darren J., and Philip E. Strahan, 2009. Do Regulations Based on Credit Ratings Affect a Firm’s Cost of Capital? Unpublished working paper. Boston College.
- [37] Lerner, Josh, and Jean Tirole, 2006. A Model of Forum Shopping. *American Economic Review* 96(4), 1091-1113.
- [38] Lewis, Christopher M., and Kevin C. Murdock, 1996. The Role of Government Contracts in Discretionary Reinsurance Markets for Natural Disasters. *Journal of Risk and Insurance* 63(4), 567-597.
- [39] Lizzeri, Alessandro, 1999. Information Revelation and Certification Intermediaries. *RAND Journal of Economics* 30(2), 214-231.
- [40] Mailath, George J., and Larry Samuelson, 2001. Who Wants a Good Reputation? *Review of Economic Studies* 68(2), 415-441.

- [41] Mariano, Beatriz, 2006, Conformity and Competition in Financial Certification. Unpublished working paper. Universidad Carlos III de Madrid.
- [42] Mathis, Jerome, James McAndrews, and Jean-Charles Rochet, 2009. Rating the Raters: Are Reputational Concerns Powerful Enough to Discipline Rating Agencies? *Journal of Monetary Economics* 56(5), 657-674.
- [43] Mayers, David, and Clifford W. Smith, Jr., 1987. Corporate Insurance and the Underinvestment Problem. *Journal of Risk and Insurance* 54(1), 45-54.
- [44] McDaniel, Raymond. W., President, Moody's Investors Service, Before the United States Securities and Exchange Commission, November 21, 2002, <http://www.sec.gov/news/extra/credrate/moodys.htm>.
- [45] Milgrom, Paul R., 1981. Good News and Bad News: Representation Theorems and Applications. *Bell Journal of Economics* 12(2), 380-391.
- [46] Morrison, Andrew, and Lucy White, 2005. Crises and Capital Requirements in Banking. *American Economic Review* 95(5), 1548-1572.
- [47] Ottaviani, Marco, and Peter Norman Sorensen, 2006a. Professional Advice. *Journal of Economic Theory* 126, 120-142.
- [48] Ottaviani, Marco, and Peter Norman Sorensen, 2006b. Reputational Cheap Talk. *RAND Journal of Economics* 37(1), 155-175.
- [49] Ottaviani, Marco, and Peter Norman Sorensen, 2006c. The Strategy of Professional Forecasting. *Journal of Financial Economics* 81, 441-466.
- [50] Pagano, Marco, and Paolo Volpin, 2008. Securitization, Transparency and Liquidity. Unpublished working paper. University of Naples Federico II.
- [51] Peyrache, Eloic, and Lucia Quesada, 2005. Intermediaries, Credibility and Incentives to Collude. Unpublished working paper. HEC Paris.

- [52] Pottier, Stephen W., and David W. Sommer, 1999. Property-Liability Insurer Financial Strength Ratings: Differences Across Rating Agencies. *Journal of Risk and Insurance* 66(4), 621-642.
- [53] Santos, Tano, and José A. Scheinkman, 2001. Competition among Exchanges. *Quarterly Journal of Economics* 116(3), 1027-1061.
- [54] Scharfstein, David, and Jeremy Stein, 1990. Herd Behavior and Investment. *American Economic Review* 80, 465-479.
- [55] Securities and Exchange Commission, 2011. Annual Report on Nationally Recognized Statistical Rating Organizations, <http://www.sec.gov/divisions/marketreg/ratingagency/nrsroannrep0111.pdf>.
- [56] Shumway, Tyler, 2001. Forecasting Bankruptcy More Accurately: A Simple Hazard Model. *Journal of Business* 74(1), 101-124.
- [57] Skreta, Vasiliki, and Laura Veldkamp, 2009. Rating Shopping and Asset Complexity: A Theory of Ratings Inflation. *Journal of Monetary Economics* 56(5), 678-695.
- [58] Standard & Poor's, 1991. S&P Insurer Solvency Review: Property/Casualty Edition, Spring 1991 (Standard & Poor's Corporation, New York, NY)
- [59] Standard & Poor's, 1995. Standard and Poor's Insurance Company Rating Guide, 1995 Edition (McGraw-Hill, New York, NY)
- [60] Strausz, Roland, 2005. Honest Certification and the Threat of Capture. *International Journal of Industrial Organization* 23, 45-62.
- [61] Tadelis, Steven, 2002. The Market for Reputations as an Incentive Mechanism. *Journal of Political Economy* 110(4), 854-882.

- [62] Thiery, Yves, and Caroline Van Schoubroeck, 2006. Fairness and Equality in Insurance Classification. *Geneva Papers on Risk & Insurance - Issues and Practice* 31(2), 190-211.
- [63] Tucker, Jennifer W. and Paul A. Zarowin, 2006. Does Income Smoothing Improve Earnings Informativeness? *The Accounting Review* 81(1), 251–270.
- [64] Veronesi, Pietro, 2000. How does Information Quality Affect Stock Returns? *Journal of Finance* 55, 807-837.
- [65] White, Lawrence, J., 2002. The Credit Rating Industry: An Industrial Organization Analysis in R.M. Levich, G. Majnoni, and C.M. Reinhart, ed.: *Ratings, Rating Agencies and the Global Financial System* (Kluwer, Boston)
- [66] Winter, Ralph A. 1991. The Liability Insurance Market. *Journal of Economic Perspectives* 5(3), 115-136.

Table 1
Coverage of the U.S. Property-Liability Insurance Industry
By A.M. Best and Standard & Poor's 1989 - 2000

Table 1 displays the number of property-liability insurers operating in the United States, the number of firms rated by A.M. Best Company, and the number of firms that requested a rating from Standard & Poor's over the years 1989 - 2000. The table also displays the total assets of the industry and the total assets of the firms rated by A.M. Best and Standard & Poor's. The numbers in parantheses show the percentage of firms (or assets) of the industry receiving a rating from either firm.

Year	Number of Companies			Total Assets (\$ billions)		
	NAIC	A.M. Best	S&P	NAIC	A.M. Best	S&P
1989	1904	1110 (58.3%)	3 (0.2%)	534.6	491.6 (91.9%)	0.6 (0.1%)
1990	1898	1176 (62.0%)	12 (0.6%)	575.2	520.2 (90.4%)	42.1 (7.3%)
1991	1973	1266 (64.2%)	25 (1.3%)	618.0	567.9 (91.9%)	44.9 (7.3%)
1992	2031	1370 (67.5%)	26 (1.3%)	688.6	626.7 (91.0%)	69.6 (10.1%)
1993	2068	1444 (69.8%)	30 (1.5%)	703.3	640.9 (91.1%)	51.3 (7.3%)
1994	2068	1517 (73.4%)	29 (1.4%)	731.5	670.9 (91.7%)	53.9 (7.4%)
1995	2097	1563 (74.5%)	36 (1.7%)	786.8	726.9 (92.4%)	30.2 (3.8%)
1996	2127	1603 (75.4%)	281 (13.2%)	842.6	786.4 (93.3%)	411.4 (48.8%)
1997	2115	1617 (76.5%)	314 (14.8%)	916.1	865.9 (94.5%)	447.7 (48.9%)
1998	2136	1660 (77.7%)	354 (16.6%)	980.7	928.8 (94.7%)	495.3 (50.5%)
1999	2072	1647 (79.5%)	297 (14.3%)	970.6	920.5 (94.8%)	283.4 (29.2%)
2000	1964	1582 (80.5%)	258 (13.1%)	942.2	892.0 (94.7%)	258.7 (27.5%)

Figure T1
U.S. Property-Liability Industry Assets Covered by
A.M. Best vs. Standard & Poor's 1989 - 2000

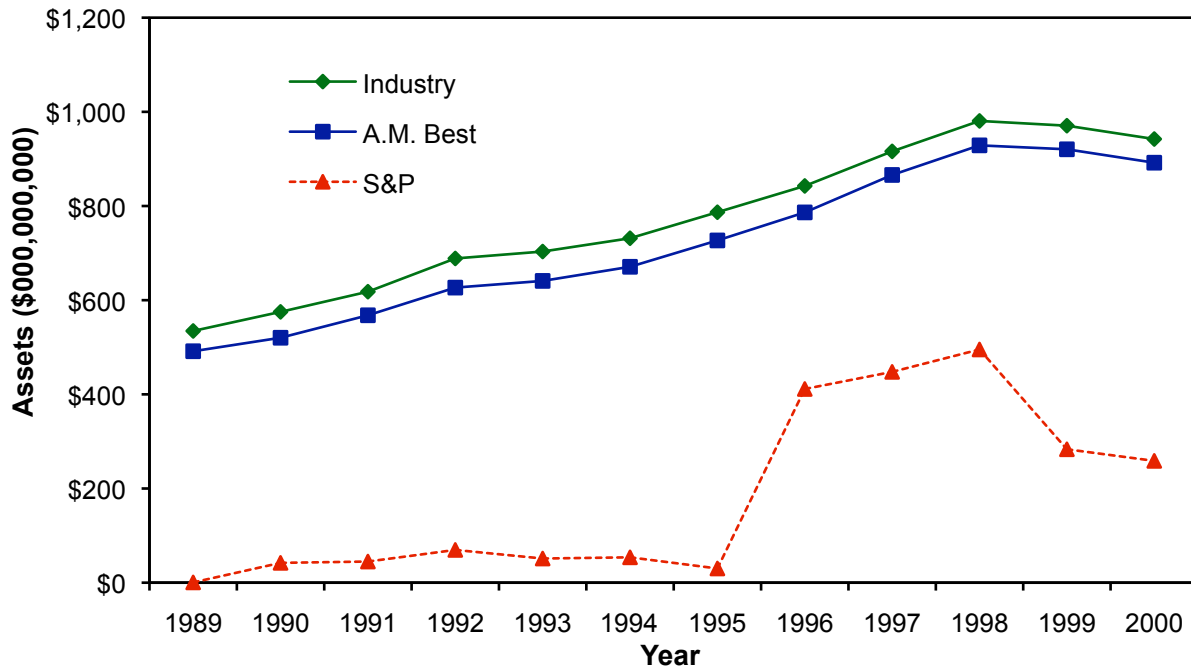


Table 2
Insurer Rating Categories: A.M Best and Standard & Poor's

Table 2 displays the mapping used in this study to compare financial strength rating categories across A.M. Best and Standard & Poor's. The terms shown in the column labeled "Verbal Description" are the descriptive words that each agency uses in their rating manuals to characterize a specific rating. The terms shown in the column labeled "Category" are the two broad terms each agency uses to group the individual ratings. Sources: A.M. Best (1991, 1992) and Standard & Poor's (1995).

Number	A.M. Best			Standard & Poor's		
	Verbal Description	Rating	Category	Verbal Description	Rating	Category
3	Superior	A++,A+	Secure	Extremely Strong	AAA	Investment Grade
2	Excellent	A, A-	Secure	Very Strong	AA, A	Investment Grade
1	Good	B++,B+	Secure	Good	BBB	Investment Grade
0	Marginal	B and below	Vulnerable	Marginal	BB and below	Non-Investment Grade

Table 3: FAST Ratio and Control Variable Summary Statistics: Solvent versus Insolvent Insurers 1989 - 2000

The table displays summary statistics of the variables used to estimate the one year default probabilities using the discrete-time hazard model. The statistics are shown separately for the solvent insurers and the insolvent insurer samples. All insurers are included in the analysis except insurers that have insufficient data or those that fail for which data is not available either one year or two years prior to the first regulatory action being taken against the firm. There are 217 firm-year observations in the insolvent sample and 24,236 in the solvent sample.

FAST Ratios and Other Control Variables	Solvent Insurers		Insolvent Insurers		Test Statistic
	μ_{sol}	σ_{sol}	μ_{ins}	σ_{ins}	$H_0: \mu_{sol} = \mu_{ins}$
Kenney Ratio: Net Premiums Written to Equity Capital	1.13	0.85	1.87	1.12	9.654
Insurance Reserves to Equity Capital	1.03	0.94	1.64	1.24	7.229
1 Yr. Growth in Net Premiums Written (%)	11.85	41.58	11.18	60.92	0.160
1 Yr. Growth in Gross Premiums Written (%)	11.90	37.48	10.86	52.45	0.292
Aid to Equity Capital due to Reinsurance	2.05	4.34	6.04	7.50	7.838
Investment Yield (%)	5.71	1.38	5.42	1.54	2.778
1 Yr. Growth in Equity Capital (%)	8.81	16.33	-8.74	19.84	12.990
Adverse Reserve Development to Equity Capital (%)	-2.72	10.84	4.18	11.73	8.635
Gross Expenses to Gross Premiums Written	0.58	0.76	0.55	0.64	0.839
1 yr. Change in Gross Expenses (%)	0.05	0.46	0.08	0.57	0.877
1 yr. Change in Liquid Assets (%)	1.17	2.67	0.35	1.78	6.694
Investments in Affiliates to Equity Capital	0.58	1.32	0.95	1.75	3.105
Receivables from Affiliates to Equity Capital	0.02	0.04	0.04	0.05	5.237
Misc. Recoverables to Equity Capital	0.03	0.05	0.07	0.08	6.735
Non-investment Grade Bonds to Equity Capital	0.69	2.51	0.76	2.70	0.353
Other Invested Assets to Equity Capital	0.01	0.03	0.02	0.04	3.567
Dummy = 1 if insurer has a large single agent	0.12	0.33	0.22	0.42	3.414
Dummy = 1 if insurer has a large single agent they control	0.08	0.28	0.12	0.32	1.437
Losses, Exp's, Div's and Taxes Paid to Premiums Collected	1.29	0.73	1.60	0.84	5.364
Total Assets (000000's in 2000 \$)	435.88	2218.95	134.63	695.43	6.109
Ind. = 1 if Insurer is Part of a Mutual Group	0.26	0.44	0.08	0.28	9.070

Table 4: Discrete-Time Hazard Bankruptcy Model Regression Results**Panel A:**

Table displays the results of the discrete-time hazard regression model. The dependent variable $y_{it} = 1$ for each insurer that has a formal regulatory action taken against the insurer in year $t+1$. Otherwise $y_{it} = 0$ for all other observations. All U.S. property-liability insurers from 1989 - 2000 are included assuming they have adequate data. There are 24,236 healthy firm-year observations and 217 insolvent company observations.

Independent Variable	Coefficient Estimate	Standard Error	χ^2 Statistic	
Intercept	-0.2458	0.9069	0.0734	
Kenney Ratio: Net Premiums Written to Equity Capital	0.6619	0.0995	44.2155	***
Insurance Reserves to Equity Capital	0.2637	0.0910	8.4058	***
1 Yr. Growth in Net Premiums Written (%)	0.0019	0.0018	1.1025	
1 Yr. Growth in Gross Premiums Written (%)	0.0007	0.0022	0.0969	
Aid to Equity Capital due to Reinsurance	0.0453	0.0114	15.9287	***
Investment Yield (%)	-0.0538	0.0534	1.0180	
1 Yr. Growth in Equity Capital (%)	-0.0469	0.0056	69.7371	***
Adverse Reserve Development to Equity Capital (%)	0.0282	0.0069	16.9760	***
Gross Expenses to Gross Premiums Written	-0.1024	0.1112	0.8489	
1 yr. Change in Gross Expenses (%)	0.1509	0.1580	0.9120	
1 yr. Change in Liquid Assets (%)	-0.0821	0.0421	3.7955	*
Investments in Affiliates to Equity Capital	0.1527	0.0481	10.0745	***
Receivables from Affiliates to Equity Capital	3.3804	1.4911	5.1395	**
Misc. Recoverables to Equity Capital	1.3329	1.0189	1.7112	
Non-investment Grade Bonds to Equity Capital	0.0144	0.0285	0.2559	
Other Invested Assets to Equity Capital	4.6585	1.9498	5.7083	**
Dummy = 1 if insurer has a large single agent	0.4375	0.2001	4.7827	**
Dummy = 1 if insurer has a large single agent they control	-0.2045	0.2580	0.6285	
Losses, Exp's, Div's and Taxes Paid to Premiums Collected	0.5949	0.1028	33.4640	***
Ln(Total Assets in \$2000)	-0.3856	0.0519	55.2763	***
Ind. = 1 if Insurer is Part of a Mutual Group	-1.1267	0.2653	18.0353	***
Log Likelihood Function Value	-876.46			
Pseudo R ²	26.21%			

*** - significant at the 1 percent level; ** - significant at the 5 percent level; * - significant at the 10 percent level

The pseudo R² equals 1 minus the ratio of the log likelihood function value divided by the log likelihood function value where all coefficients are constrained to be zero (see Greene 1997 p. 891).

Panel B:

Table displays summary statistics of the predicted one-year probability of default for solvent firm-year observations and for bankrupt firm-year observations.

Firm Type	Num	Ave.	Median	Standard Deviation	1 st Percentile	99 th Percentile
Solvent	24,236	0.81%	0.20%	2.50%	0.01%	11.48%
Insolvent	217	9.22%	4.36%	12.55%	0.08%	68.69%

Panel C:

The classification table displays the number of actual insolvent and solvent firm-year observations versus the number of predicted insolvent and solvent firm-year observations using the estimated hazard model. Observations were predicted to be insolvent (solvent) that had one-year probabilities of default estimated to be greater than the population average over this time period (0.89 percent).

Firm Type	Predicted Firm Type		Totals
	Insolvent	Solvent	
Insolvent	177	40	217
Solvent	4,240	19,996	24,236

Table 5
Summary Statistics One-Year Probability of Default for
U.S. Property-Liability Insurers Rated By
A.M. Best and Standard & Poor's: 1989 - 2000

Table 5 displays the average and median probability of default of the firms that receive ratings by A.M. Best and from Standard & Poor's. The t-test column reports the results of the null hypothesis of equal means for the probability of default for A.M. Best rated insurers versus S&P rated insurers assuming unequal variances against the alternative hypothesis the average probability of default for S&P rated insurers is lower than the average for A.M. Best rated insurers. The chart below displays the average probability of default time series for each agency over time period of this study.

Year	A.M. Best			Standard & Poor's			Test Statistic	
	Num	Mean	Median	Num	Mean	Median	$H_0: \mu_{\text{Best}} = \mu_{\text{S\&P}}$	
1989	1110	0.39%	0.14%	3	0.18%	0.24%		
1990	1175	0.59%	0.26%	12	0.32%	0.11%	-1.991	**
1991	1261	0.50%	0.14%	25	0.13%	0.11%	-7.735	***
1992	1352	0.63%	0.15%	26	0.22%	0.12%	-5.070	***
1993	1437	0.47%	0.12%	34	0.10%	0.05%	-6.770	***
1994	1515	0.46%	0.14%	39	0.20%	0.11%	-5.536	***
1995	1551	0.42%	0.10%	46	0.17%	0.04%	-3.117	***
1996	1577	0.87%	0.18%	290	0.34%	0.18%	-6.346	***
1997	1598	0.51%	0.14%	320	0.32%	0.14%	-3.503	***
1998	1620	0.68%	0.16%	372	0.43%	0.15%	-3.013	***
1999	1620	0.79%	0.19%	401	0.99%	0.18%	1.291	
2000	1570	0.68%	0.20%	365	0.74%	0.23%	0.444	

***, **, or * - significant at the 1, 5, or 10 percent level, respectively

Figure T5
Average One-Year Probability of Default for U.S. Property-Liability
Insurers Rated By A.M. Best and Standard & Poor's 1989-2000

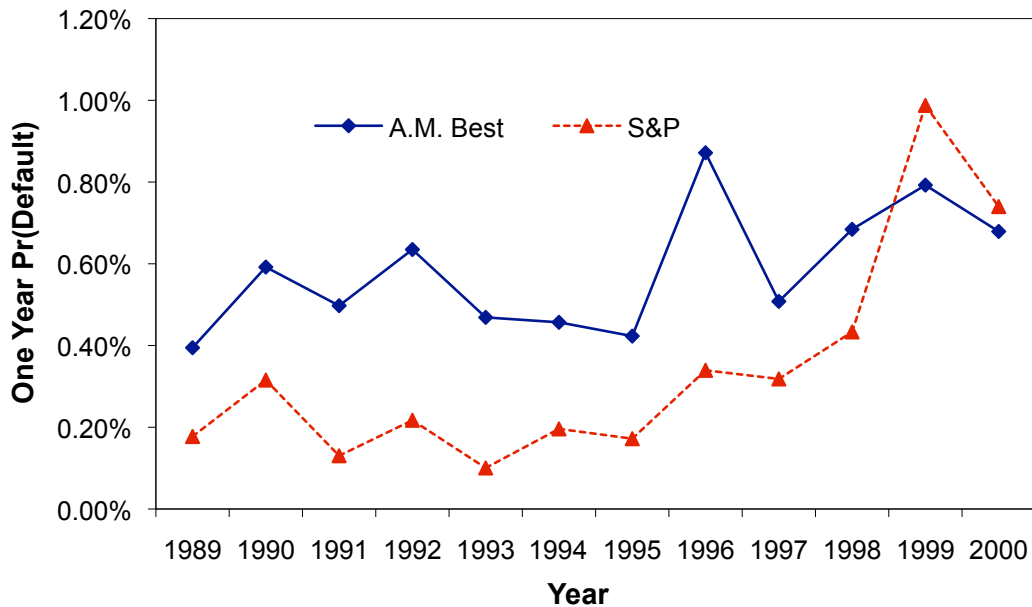


Table 6
Average Rating Assigned by A.M. Best and Standard & Poor's by Year: 1989 - 2000

Table displays summary statistics of the ratings assigned by A.M. Best and S&P to U.S. property-liability insurers over the years 1989 - 2000. The final two columns display t-statistics reporting the difference in means test assuming unequal variances

Year	A.M. Best					Standard & Poor's					Difference in Means Test Statistics S&P - A.M. Best	
	Num	μ_{Best}	σ_{Best}	Min	Max	Num	$\mu_{\text{S\&P}}$	$\sigma_{\text{S\&P}}$	Min	Max	$H_0: \mu_{\text{Best}} = \mu_{\text{S\&P}}$	
1989	1110	2.16	0.82	1	3	3	1.33	1.15	0	2	-	
1990	1176	2.13	0.86	1	3	12	2.17	0.72	1	3	0.18	
1991	1266	2.04	0.87	1	3	25	2.20	0.58	1	3	1.37	*
1992	1370	2.02	0.86	1	3	26	2.12	0.65	0	3	0.70	
1993	1444	2.00	0.85	1	3	34	1.85	0.93	0	3	0.89	
1994	1517	1.92	0.86	1	3	35	1.71	0.86	0	3	1.41	*
1995	1563	1.89	0.84	1	3	42	2.00	0.58	0	3	1.19	
1996	1603	1.90	0.84	1	3	283	2.11	0.44	1	3	6.17	***
1997	1617	1.92	0.81	1	3	319	2.07	0.45	1	3	4.65	***
1998	1660	1.97	0.81	1	3	367	2.10	0.43	1	3	4.38	***
1999	1647	1.95	0.84	1	3	390	1.97	0.57	0	3	0.49	
2000	1582	1.93	0.85	1	3	279	1.85	0.54	0	3	2.03	**

***, **, or * - significant at the 1, 5, or 10 percent level, respectively

Figure T6
Average Rating for U.S. Property-Liability Insurers Rated By
A.M Best and Standard & Poor's 1989-2000

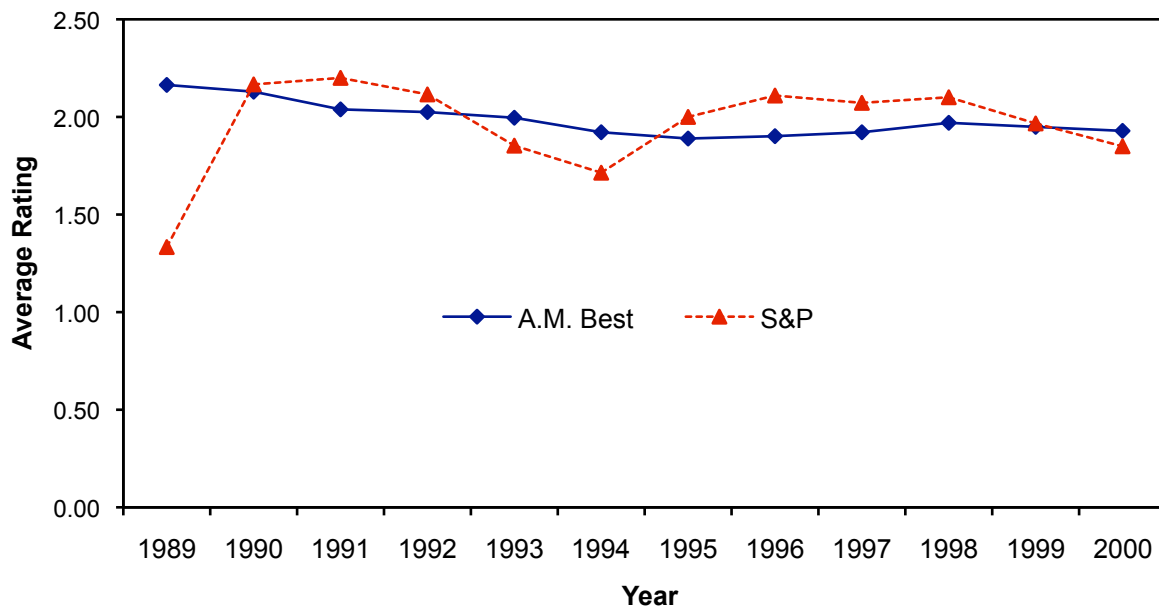


Table 7
Comparing A.M. Best and Standard & Poor's
Ratings for Common Insurers: 1989 - 2000

Table compares the ratings assigned by A.M. Best and S&P for insurers that receive a rating from both firms. The data includes all firm-year observations over the years 1989 - 2000 that received a rating from A.M. Best and a rating from S&P. Each cell of the matrix, c_{ij} , equals the number of firm-year observations that receive rating i from A.M. Best and rating j from S&P where $i, j \in \{\text{Superior, Excellent, Good, Marginal}\}$.

A.M. Best Rating	S&P Full Rating			
	Marginal	Good	Excellent	Superior
Marginal	23	5	6	0
Good	12	35	42	0
Excellent	1	63	786	23
Superior	0	0	649	247

Total number of firm-year observations: 1892

S&P gives the same rating as A.M. Best: 57.66%

S&P rates higher than A.M. Best: 4.02%

S&P rates lower than A.M. Best: 38.32%

Table 8
Summary Statistics of A.M. Best Rated Insurers that also Request a Rating from Standard & Poor's: 1990 - 2000

The sample includes insurer-year observations that receive an A.M. Best rating over the years 1990 - 2000. The columns labeled "A.M. Best and Standard & Poor's" displays summary statistics for just those insurer-year observations that requested a rating by both the incumbent and the new entrant agency. The columns labeled "A.M. Best Only" are insurer-year observations for firms that requested a rating only from A.M. Best. The variable labeled "[Median Pr(Def. | A.M. Best Rating) - Insurer Pr(Def.)]" equals the median probability of default for all insurers within the A.M. Best rating category as firm *i* minus insurer *i*'s own estimated default probability.

	A.M. Best Only		A.M. Best and Standard & Poor's		Test Statistics
	μ_{BestOnly}	σ_{BestOnly}	μ_{both}	σ_{both}	$H_0: \mu_{\text{BestOnly}} = \mu_{\text{both}}$
S&P Rating	-	-	1.989	0.506	-
A.M. Best Rating	1.955	0.835	2.339	0.684	20.2 ***
S&P Rating - A.M. Best Rating	-	-	-0.350	0.564	-
[Median Pr(Def. A.M. Best Rating) - Insurer Pr(Def.)]	-0.003	0.010	-0.004	0.011	3.0 ***
Ind. = 1 if Insurer is Part of a Mutual Group	0.319	0.466	0.171	0.377	14.0 ***
State of Business of Herfindahl	0.536	0.373	0.307	0.326	25.5 ***
Line of Business Herfindahl	0.454	2.133	0.325	0.593	5.4 ***
Total Assets (000000's in 2000 \$)	494.8	2339.9	1,860.4	5,963.9	8.8 ***
% Net Premiums Written in Retail Lines of Insurance	0.373	0.362	0.288	0.311	9.9 ***
Ind. = 1 if Insurer had Corporate Debt Rating from S&P	0.196	0.397	0.717	0.451	42.9 ***
Ind. = 1 if Insurer Requested S&P rating year t-1	0.017	0.128	0.684	0.465	55.4 ***
Ind. = 1 if Insurer was Assigned Qualified Rating by S&P Year t-1	0.388	0.487	0.144	0.351	24.5 ***
Ind. = 1 for Marginal A.M. Best Rating	0.071	0.257	0.024	0.153	10.4 ***
Ind. = 1 for Good A.M. Best Rating	0.159	0.365	0.050	0.218	16.9 ***
Ind. = 1 for Excellent A.M. Best Rating	0.515	0.500	0.489	0.500	1.9 **
Ind. = 1 for Superior A.M. Best Rating	0.256	0.436	0.437	0.496	13.6 ***
Ind. = 1 if year = 1990	0.079	0.270	0.007	0.081	23.1 ***
Ind. = 1 if year = 1991	0.082	0.274	0.016	0.125	16.5 ***
Ind. = 1 if year = 1992	0.091	0.287	0.015	0.123	18.7 ***
Ind. = 1 if year = 1993	0.096	0.295	0.015	0.123	19.9 ***
Ind. = 1 if year = 1994	0.100	0.300	0.017	0.128	19.8 ***
Ind. = 1 if year = 1995	0.105	0.307	0.023	0.151	17.4 ***
Ind. = 1 if year = 1996	0.090	0.287	0.174	0.380	8.3 ***
Ind. = 1 if year = 1997	0.086	0.280	0.185	0.388	9.6 ***
Ind. = 1 if year = 1998	0.086	0.281	0.212	0.409	11.6 ***
Ind. = 1 if year = 1999	0.092	0.290	0.174	0.379	8.1 ***
Ind. = 1 if year = 2000	0.092	0.289	0.162	0.369	7.2 ***

*** - significant at the 1 percent level; ** - significant at the 5 percent level; * - significant at the 10 percent level.

The number of firm-year observations that only received a rating from A.M. Best was 13,353. The number of firm-year observations that received both an A.M. Best and Standard & Poor's rating was 1503.

Table 9
Probit Regression Results Predicting Whether Insurer Requested a
Rating from Standard & Poor's: 1990 - 2000

Table displays Probit regression results where the dependent variable equals 1 when insurer i was assigned a rating by Standard & Poor's in year t and 0 otherwise. Panel A displays the estimated coefficients on the independent variables. Panel B displays the marginal effects. The model also contains year indicator variables but the results are suppressed to save space. The full results with the estimated coefficients for the year indicator variables are available upon request.

Panel A: Regression Results Independent Variable	Did Insurer Request a Rating from S&P?			
	Model 1	Model 2	Model 3	Model 4
Intercept	-1.019 *** (0.040)	-4.640 *** (0.222)	-4.090 *** (0.277)	-4.265 *** (0.306)
[Median Pr(Def. A.M. Best Rating) - Insurer Pr(Def.)]	47.233 *** (9.552)	53.604 *** (11.650)	41.192 *** (12.727)	43.149 *** (13.113)
[Median Pr(Def. A.M. Best Rating) - Insurer Pr(Def.)] x (Year - 1988)	-4.843 *** (0.915)	-5.652 *** (1.108)	-4.531 *** (1.254)	-4.932 *** (1.295)
Ln(Assets in \$2000)		0.179 *** (0.011)	0.114 *** (0.014)	0.108 *** (0.014)
Ind. = 1 if Insurer is Part of a Mutual Group		-0.165 *** (0.047)	-0.159 *** (0.058)	-0.168 *** (0.058)
Ind. = 1 if Insurer had Corporate Debt Rating from S&P		1.049 *** (0.037)	0.687 *** (0.047)	0.656 *** (0.048)
% Net Premiums Written in Retail Lines of Insurance		-0.189 *** (0.054)	-0.149 ** (0.068)	-0.143 ** (0.069)
State of Business of Herfindahl		-0.304 *** (0.056)	-0.294 *** (0.068)	-0.269 *** (0.069)
Ind. = 1 if Insurer Requested S&P rating year t-1			2.465 *** (0.063)	2.461 *** (0.063)
Ind. = 1 if Insurer was Assigned Qualified Rating by S&P Year t-1 for Years 1989-1994			0.088 (0.124)	0.075 (0.124)
Ind. = 1 if Insurer was Assigned Qualified Rating by S&P Year t-1 for Years 1995-1999			0.104 * (0.056)	0.098 * (0.057)
Ind. = 1 for Good A.M. Best Rating				0.158 (0.151)
Ind. = 1 for Excellent A.M. Best Rating				0.304 ** (0.139)
Ind. = 1 for Superior A.M. Best Rating				0.344 ** (0.143)
Log-likelihood function value	-4186.7	-3200.3	-2042.1	-2037.5
Pseudo R ²	13.99%	34.25%	58.05%	58.14%

Panel B: Estimated Marginal Effects				
[Median Pr(Def. A.M. Best Rating) - Insurer Pr(Def.)]	5.946 ***	3.629 ***	2.128 ***	2.190 ***
[Median Pr(Def. A.M. Best Rating) - Insurer Pr(Def.)] x (Year - 1988)	-0.610 ***	-0.383 ***	-0.234 ***	-0.250 ***
Ln(Assets in \$2000)		0.012 ***	0.006 ***	0.005 ***
Ind. = 1 if Insurer is Part of a Mutual Group		-0.011 ***	-0.008 ***	-0.009 ***
Ind. = 1 if Insurer had Corporate Debt Rating from S&P		0.071 ***	0.035 ***	0.033 ***
% Net Premiums Written in Retail Lines of Insurance		-0.013 ***	-0.008 **	-0.007 **
State of Business of Herfindahl		-0.021 ***	-0.015 ***	-0.014 ***
Ind. = 1 if Insurer Requested S&P rating year t-1			0.127 ***	0.125 ***
Ind. = 1 if Insurer was Assigned Qualified Rating by S&P Year t-1 for Years 1989-1994			0.005	0.004
Ind. = 1 if Insurer was Assigned Qualified Rating by S&P Year t-1 for Years 1995-1999			0.005 *	0.005 *
Ind. = 1 for Good A.M. Best Rating				0.008
Ind. = 1 for Excellent A.M. Best Rating				0.015 **
Ind. = 1 for Superior A.M. Best Rating				0.017 **

*** - significant at the 1 percent level; ** - significant at the 5 percent level; * - significant at the 10 percent level

The pseudo R² equals 1 minus the ratio of the log likelihood function value divided by the log likelihood function value where all coefficients are constrained to be zero (see Greene 1997 p. 891). Standard errors are shown in parentheses.

Table 10
Regression Results Explaining Rating Difference
Between Standard & Poor's and A.M. Best Ratings: 1990 - 2000

Table displays OLS regression results where the dependent variable is the difference between the rating assigned by Standard & Poor's minus the rating assigned by A.M. Best in year t . The inverse Mill's terms in each model were calculated using the results of the corresponding Probit regression shown in Table 9.

Independent Variable	Did Insurer Request a Rating from S&P?			
	Model 1	Model 2	Model 3	Model 4
Intercept	-0.5832 *** (0.075)	-0.5066 *** (0.033)	-0.4213 *** (0.021)	-0.4337 *** (0.021)
Inverse Mills Ratio	0.1531 *** (0.048)	0.1344 *** (0.026)	0.0983 *** (0.020)	0.1156 *** (0.020)
R ²	0.67%	1.82%	1.53%	2.09%
Expected change in rating due to insurer strategic choice	0.233	0.157	0.071	0.084
Expected change in ratings due difference in S&P vs. A.M. Best standards	-0.583	-0.507	-0.421	-0.434
Average Rating Difference S&P Rating - A.M. Best Rating	-0.350	-0.350	-0.350	-0.350

Table 11
Probability of Default for Categories that were Split by A.M. Best in 1991
Relative to Standard & Poors: 1989-2000

Table compares the average probability of default for insurers rated AAA (BBB) by Standard & Poor's with A.M. Best's categories A++/A+ (B++/B+) before and after A.M. Best split the A+(B+) ratings into finer classifications in 1991. Then number of insurers assigned a particular rating is shown in parantheses.

Year	S&P AAA	A.M. Best A++	A.M. Best A+	Ratio		
				AAA/A++	AAA/A+	A++/A+
1989			0.19% (425)			
1990	0.07% (4)		0.29% (438)		0.232	
1991	0.13% (7)	0.12% (71)	0.19% (338)	1.120	0.686	0.612
1992	0.10% (6)	0.23% (94)	0.24% (326)	0.430	0.420	0.976
1993	0.11% (7)	0.17% (97)	0.17% (313)	0.630	0.615	0.976
1994	0.22% (3)	0.15% (107)	0.17% (280)	1.433	1.274	0.889
1995	0.11% (6)	0.09% (111)	0.12% (235)	1.276	0.970	0.760
1996	0.14% (45)	0.24% (123)	0.28% (240)	0.602	0.520	0.864
1997	0.21% (45)	0.20% (113)	0.24% (253)	1.044	0.896	0.858
1998	0.21% (54)	0.23% (114)	0.38% (312)	0.901	0.546	0.607
1999	0.20% (41)	0.18% (111)	0.46% (332)	1.138	0.436	0.383
2000	0.29% (8)	0.21% (131)	0.49% (269)	1.353	0.589	0.435
Average	0.16%	0.18%	0.27%	0.993	0.653	0.736

Year	S&P BBB	A.M. Best B++	A.M. Best B+	Ratio		
				BBB/B++	BBB/B+	B++/B+
1989			0.62% (127)			
1990	0.71% (2)		1.23% (116)		0.579	
1991	0.28% (2)	0.45% (53)	0.59% (107)	0.626	0.475	0.759
1992	0.34% (1)	0.58% (90)	0.80% (92)	0.583	0.425	0.729
1993	0.03% (2)	0.33% (99)	0.64% (101)	0.093	0.048	0.517
1994	0.07% (1)	0.48% (120)	0.72% (140)	0.140	0.093	0.664
1995	0.27% (4)	0.40% (131)	0.63% (120)	0.676	0.432	0.639
1996	0.42% (14)	0.50% (129)	1.40% (128)	0.832	0.297	0.358
1997	0.27% (22)	0.72% (137)	0.86% (144)	0.375	0.314	0.837
1998	0.50% (17)	0.62% (151)	1.93% (139)	0.806	0.261	0.324
1999	1.75% (20)	0.84% (185)	1.25% (134)	2.090	1.408	0.673
2000	1.18% (20)	0.82% (141)	0.78% (122)	1.441	1.510	1.047
Average	0.53%	0.57%	0.95%	0.766	0.531	0.655