

The Dual Role of Ratings

Allen Vong*

April 1, 2019

Abstract

A conventional wisdom is that ratings exist to solve adverse selection and moral hazard problems. Raters often collect payments from their ratees. It is unclear whether rating schemes tailored to maximize ratees' payments solve adverse selection and moral hazard problems. I prove that ratings which fully extract the ratee's net surplus entirely solve moral hazard by leveraging the presence of adverse selection over time. I find a tension between rating transparency and economic efficiency — ratings that maximize the ratee's and the market's surplus are opaque. I illustrate the relationship between rating coarseness and moral hazard, as well as the implications of fully-extracting ratings for market beliefs and behaviors. I reconcile the conventional wisdom with critiques that ratings add little information to the markets.

Keywords:

ratings, repeated games, reputation, information intermediation.

JEL Classifications:

C72, C73, D82, D83, M52, G24.

*Yale University, 30 Hillhouse Ave., New Haven, CT 06520, USA, allen.vong@yale.edu.

I am deeply indebted to both Johannes Hörner and Larry Samuelson for their guidance and support. I am also grateful to Marina Halac for detailed comments and suggestions on the paper. Conversations with Ian Ball, Dirk Bergemann, Tilman Börgers, Jin-Wook Chang, Florian Ederer, Péter Eső, Jack Fanning, Simone Galperti, John Geanakoplos, Mira Frick, Ryota Iijima, Yuhta Ishii, Jan Knöpfle, Matthew Leister, George Mailath, Meg Meyer, Luís Carvalho Monteiro, Stephen Morris, Xiaosheng Mu, Dan Quigley, Anna Sanktjohanser, Juuso Välimäki, Kai Hao Yang and comments from many seminar and conference audiences are also very helpful.

Contents

1	Introduction	1
2	Model	3
2.1	Histories and Strategies	4
2.2	The Rater's Problem	5
3	Full Extraction	8
3.1	Necessity and Sufficiency	9
3.2	A Simple, Fully-Extracting Menu	11
3.3	Full Extraction Without Screening of Types	15
4	Equilibrium Implications	17
4.1	Implications of Moral Hazard	17
4.2	Implications of Adverse Selection	19
4.3	Implications of Fully-Extracting Ratings	20
4.3.1	Adverse Selection and Moral Hazard	20
4.3.2	Pure Adverse Selection	22
4.3.3	Pure Moral Hazard	23
5	Full Extraction with Higher Costs	24
5.1	Impossibility with Fixed Discounting	24
5.2	A Limit-Payoff Characterization	25
5.3	Towards a Complete Characterization	29
5.3.1	Prohibitively High Costs	29
5.3.2	Intermediate Costs	29
5.3.3	Summary	34
5.3.4	Worst-Case Scenario: Fully Solving Adverse Selection	34
6	Final Remarks	35
	Appendices	36
A	Omitted Details	36
A.1	The Rater Needs At Most Two Schemes	36
A.2	Equilibrium Existence	37
A.3	Fully-Revealing Rating System	38
A.4	The Menu Ξ^*	38
A.5	The Menu Ξ^{**}	38
A.6	The Menu Ξ_{MH}	38
A.7	Competent Firm's Maximal Profit	40
A.8	The Menu Ξ^∞	41
A.9	Convergence to the Average Payoff Criterion	42
A.10	Perfect Monitoring	42

B	Proofs	44
B.1	Proof of Proposition 1	44
B.2	Proof of Proposition 2	50
B.3	Proof of Proposition 3	52
B.4	Proof of Proposition 4	54
B.5	Proof of Proposition 5	54
B.6	Proof of Proposition 6	56
B.7	Proof of Proposition 7	56
B.8	Proof of Proposition 8	57
B.9	Proof of Proposition 9	59
B.10	Proof of Proposition 10	61
B.11	Proof of Proposition 11	61
B.12	Proof of Proposition 12	63
B.13	Proof of Proposition 13	63
	References	65

1 Introduction

Information intermediaries are central to markets with adverse selection and moral hazard. Examples include intermediaries who rate borrowers (*e.g.*, credit rating agencies), doctors (*e.g.*, RateMD), employers (*e.g.*, Glassdoor), hotels (*e.g.*, TripAdvisor), restaurants (*e.g.*, Yelp) or businesses in general (*e.g.*, Better Business Bureau). The signals they produce, broadly conceived as *ratings*, coordinate market beliefs of the ratees' abilities and their quality provisions over time.

A conventional wisdom is that ratings exist to solve adverse selection and moral hazard (*e.g.*, see Dellarocas (2005), Gonzalez et al. (2004), Portes (2008), Levich et al. (2012) and Belleflamme and Peitz (2018)). Hotel or restaurant ratings, or ratings of experience-good sellers in general, may reflect their ability to provide good services and may provide incentives to do so. Similarly, credit ratings may reflect a borrower's ability to repay loans (Kashyap and Kovrijnykh, 2015), and may provide her with incentives to manage her proceedings diligently for the lender's benefit (Boot et al., 2005).

Yet, intermediaries often collect upfront payments from their ratees. It is unclear whether rating schemes tailored to maximize ratees' payments solve adverse selection and moral hazard.¹ Transparent ratings tackle adverse selection, but may diminish inept ratees' revenues, limiting their willingness to pay the rater. They may also fail to tackle moral hazard by allowing competent ratees to successfully build their reputations and then rest on their laurels. In contrast, opaque ratings limit competent ratees' ability to build a reputation and hence their willingness to pay.

Do revenue-maximizing, ratee-pays ratings address adverse selection and moral hazard? What are their structure and implications for market beliefs and behaviors? How does the structure depend on the market conditions?

I develop a simple model in Section 2 to address these questions. It captures the essential features of ratee-pays rating relationships, without aspiring to describe closely any specific market. The model builds upon Mailath and Samuelson (2001). A firm repeatedly trades with a succession of short-lived consumers. Adverse selection and moral hazard are present: the firm has private information about its ability and choices of quality provision, and faces a myopic temptation to shirk against consumers. The firm may pay an upfront fee to participate in a rating scheme. A rating quantifies

¹See Sangiorgi and Spatt (2017) for a discussion on credit ratings, Crowe (2018) on RateMD, Henry (2014) on Glassdoor, Kugel (2016) on Yelp Ads and TripAdvisor placements, Fleming (2010) on the Better Business Bureau and Marriage and Thompson (2018) on corporate audits.

the firm’s reputation and links its past behavior to market expectations of its future behavior, affecting its future revenue. Ratings may therefore act as a screening device and a commitment device that (respectively) address adverse selection and moral hazard. The rater chooses a menu of rating schemes that maximizes her revenue via the upfront fees. I discuss and justify these features as the paper proceeds.

Sections 3 – 5 collect the results. I derive necessary and sufficient conditions for a menu to fully extract the ratee’s net surplus. Most importantly, fully-extracting ratings leverage adverse selection to fully solve moral hazard. These ratings maximize a firm’s value to consumers and in turn its willingness to pay the rater. To illustrate, I explicitly construct a fully-extracting menu. The menu attracts participation by firms capable *and* incapable of high quality provision, maintaining market uncertainty over a rated firm’s ability. The ratings coordinate consumer learning to ensure that reputation effects remains operative over time to solve moral hazard.

I also show that fully-extracting menus maximize market surplus. A self-interested rater acts as if she is benevolent, such as those who derive revenues through commission fees (*e.g.*, Expedia for accommodations and UpWork for freelancers) or advertising (*e.g.*, Avvo for lawyers, RateMyProfessors for professors or review blogs with Google AdSense), as maximizing social surplus maximizes their popularity to facilitate trades.

To understand the implications of each information friction, I characterize fully-extracting menus absent adverse selection or moral hazard. Ratings may contain no information absent moral hazard, but are maximally transparent absent adverse selection. The rater strictly benefits from adverse selection. To highlight the implications of ratings for beliefs and behaviors, as well as the role of the information frictions in these implications, I contrast the results to the canonical repeated-game counterparts without intermediation under pure moral hazard (Fudenberg and Levine, 1994), pure adverse selection (Tadelis, 1999) and both frictions (Mailath and Samuelson, 2001).

Broadly, this paper joins a growing literature on information intermediation in dynamic games (*e.g.*, Ekmekci (2011), Pei (2016), Che and Hörner (2017) and Hörner and Lambert (2018)). It identifies a new rationale behind the prevalence of opaque information, which is central to the information intermediation literature (*e.g.*, Lizzeri (1999), Dellacaros (2005)). This includes censoring past play to maintain market uncertainty, speaking to the reputation literature (Cripps et al., 2004). The resulting use of opaque information to coordinate efficient trades reconciles the conventional wisdom with critiques of credit ratings that suggest they add little information to markets (*e.g.*, Macey (2006), Fitzpatrick and Sagers (2009), Rhee (2015)).

2 Model

Consider a market populated by a long-lived firm, a sequence of short-lived consumers and a long-lived rater. Time is discrete, indexed by t , and the horizon is infinite.

In period $t = -1$, nature draws the firm's type θ , choosing "competent" with probability $\mu \in (0, 1)$ and "inept" with probability $1 - \mu$. The rater then offers the firm a menu of rating schemes $\Xi = \{\xi_C, \xi_I\}$.² The scheme $\xi_\theta = (f_\theta, R_\theta, S_\theta)$, intended for a firm of type θ , consists of a fee $f_\theta \in \mathbb{R}_+$, a countable set of possible ratings R_θ , and a rating system S_θ that maps each history of ratings and signals (described below) to a probability distribution on R_θ , from which a rating is drawn before each consumer enters. In period $t = 0$, the firm chooses once-and-for-all whether to participate in a scheme and which scheme.³ If the firm chooses a scheme ξ_θ , it pays the rater f_θ .

In each period $t = 0, 1, \dots$, a consumer enters the market, observes the recently drawn rating⁴ and pays the firm upfront her expected payoff. The rating is either drawn by the system in a selected scheme, or is a null rating $\emptyset$ if the firm does not participate. The firm then exerts effort. A competent firm chooses either high or low effort, and an inept firm only exerts low effort. High effort entails a cost $c > 0$, yielding good signal ($\bar{y}$) with probability $1 - \rho \in (\frac{1}{2}, 1)$ and bad signal ($\underline{y}$) with probability ρ . Low effort is costless, yielding good signal with probability ρ and bad signal with probability $1 - \rho$.⁵ A good signal gives the consumer a payoff of 1. A bad signal gives 0. High effort is efficient, so that the market surplus associated with high effort exceeds that with low effort:

$$1 - \rho - c > \rho. \quad (1)$$

The consumer then leaves the market and period $t + 1$ unfolds. The moves are summarized in Figure 1 below.

The calendar time and the menu are observed by all parties. The type and effort choices are the firm's private information. The firm observes its

²It is without loss of generality to assume that the rater needs at most two schemes in a menu if the rater can offer the firm a menu of lotteries over schemes. Appendix A.1 provides a formal argument.

³The results are unaffected if the firm can quit the scheme after each history of play, at a significant notational cost. See Remark 10 for a discussion.

⁴I discuss the case when consumers observe multiple previous ratings in Section 6.

⁵The symmetry assumption in the monitoring structure can be relaxed without affecting the results at the cost of additional notation. The rater's payoff when she fully extracts the net market surplus, the firm's participation constraint, and the relevant incentive constraints for effort can be easily adjusted to adopt an asymmetric monitoring structure.

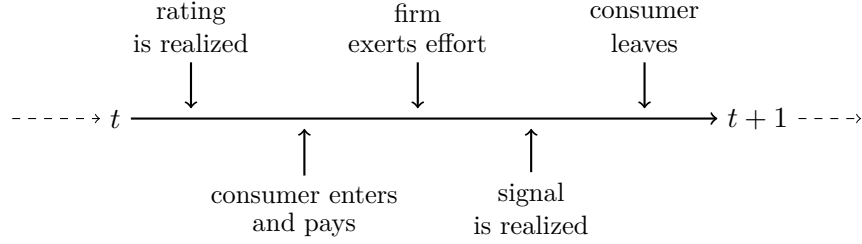


Figure 1: Timing of Moves in Each Period $t = 0, 1, \dots$

participation decision, the selected scheme, and the previous ratings and signals. Each consumer observes the rating in the period she enters, and also the signal before leaving the market.⁶ Each consumer uses the observed rating and Bayes' rule to update her belief of the firm being competent, interpreted as the firm's *reputation*. The rater observes the firm's participation decision, the selected scheme, and the histories of ratings and signals.

2.1 Histories and Strategies

Let $\Theta := \{C, I\}$ be the set of types, $D := \{N, \xi_C, \xi_I\}$ be the set of participation decisions where N denotes not participating (*i.e.*, the outside option), A be the set of efforts, $Y := \{\underline{y}, \bar{y}\}$ be the set of signals and $R := R_C \cup R_I \cup \{\emptyset\}$ be the set of all ratings. The rater's period- t history, denoted by h_r^t , belongs to the set $H_r^t := D \times (Y \times R)^t$. The firm's period- t history upon observing the period- t rating, denoted by h_f^t , belongs to the set $H_f^t := D \times (A \times Y \times R)^t \times R$. A period- t consumer's history upon observing the period- t rating r_t is simply r_t . Write $H_r := \cup_{t=0}^{\infty} H_r^t$ and $H_f := \cup_{t=0}^{\infty} H_f^t$.

A rating system in a scheme ξ_θ is a function $S_\theta : H_r \rightarrow \Delta(R_\theta)$, which draws a rating r_t with probability $S_\theta(r_t|h_r)$ after each rater's history h_r conditional on $d = \xi_\theta$.

A firm's strategies are a pair $\sigma = (\pi, \tau)$. The participation strategy $\pi : \Theta \rightarrow \Delta(D)$ specifies $\pi(d|\theta)$, the probability of choosing participation $d \in D$ by each type θ . The effort strategy $\tau : H_f \times \Theta \rightarrow [0, 1]$ specifies the probability of high effort after each history h_f for each type, with the restriction that the inept firm only chooses low effort, *i.e.*, $\tau(\cdot, I) = 0$.

Consumers' beliefs are defined over the set of outcomes of the game, $\Omega := \Theta \times D \times (A \times Y \times R)^\infty$. Given μ, σ induces a probability measure

⁶Whether the consumer observes the realized signal before leaving is formally irrelevant. Allowing so invites an interpretation that the consumer observes the signal and reports to the rater, which is a common phenomenon in many online rating platforms.

$P \in \Delta(\Omega)$. Let P_t^θ denote the marginal of P on H_1^t conditional on θ . The probability $P_t^\theta(r|d)$ of realizing a rating $r \in R$ in period t given that a type- θ firm chooses participation $d \in D$ is

$$P_t^\theta(r|\xi_{\theta'}) := \sum_{h_r^t \in H_1^t} S_{\theta'}(r|h_r^t) P_t^\theta(h_r^t|\xi_{\theta'}) \text{ and } P_t^\theta(\emptyset|N) = 1,$$

where $P_t^\theta(h_r^t|\xi_{\theta'})$ is the probability of realizing a history h_r^t in period t given that the firm has a type θ and chooses $\xi_{\theta'}$. A period- t consumer who observes a rating r forms a posterior of the firm being competent:

$$\varphi_t(r) = \frac{\mu \sum_{d \in D} \pi(d|C) P_t^C(r|d)}{\mathbf{E}^\mu[\sum_{d \in D} \pi(d|\theta) P_t^\theta(r|d)]}, \quad (2)$$

where $\mathbf{E}^\mu[\cdot]$ is an expectation over types θ with respect to μ . A period- t consumer's belief is a function $p_t^\sigma : R \rightarrow [0, 1]$ capturing the probability of receiving a good signal upon observing a rating, and equals her payment given the payoff normalizations.

The firm, with a discount factor $\delta \in (0, 1)$, chooses σ to maximize its payoff. The expected profit extraction from consumers by a type- θ firm upon participation d is

$$U(\sigma, \theta; d) := (1 - \delta) \mathbf{E}^P \left[\sum_{t=0}^{\infty} \delta^t (p_t^\sigma(r_t) - c(e_t)) \middle| \theta, d \right], \quad (3)$$

where $c(e_t)$ denotes the cost incurred by choosing effort e_t in period t , and the expectation $\mathbf{E}^P[\cdot|\theta, d]$ is taken with respect to the measure P induced by σ , conditional on θ and d . The firm's payoff is its profit minus any participation payment:

$$U^*(\sigma, \theta) := \sum_{d \in D} \pi(d|\theta) U(\sigma, \theta; d) - (1 - \delta) \sum_{\theta'} \pi(\xi_{\theta'}|\theta) f_{\theta'}, \quad (4)$$

Equilibrium refers to *perfect Bayesian equilibrium* (σ, φ) , so that σ is maximizing after each history of the firm of each type given beliefs $\varphi = (\varphi_t)_{t=0}^\infty$, and the beliefs are consistent with Bayes' rule given σ whenever possible.

2.2 The Rater's Problem

The rater seeks a revenue-maximizing menu. She has commitment, *i.e.*, the menu is chosen once and for all. For each menu Ξ , denote the set of all

equilibria in the induced continuation game by $B(\Xi)$.⁷ The rater's normalized payoff using a menu Ξ in equilibrium $(\sigma, \varphi) \in B(\Xi)$ is

$$W(\Xi, (\sigma, \varphi)) := (1 - \delta) \mathbf{E}^\mu \left[\sum_{\theta' \in \Theta} \pi(\xi_{\theta'} | \theta) f_{\theta'} \right]. \quad (5)$$

The focus is on rater-preferred equilibria, those that generate the highest rater's payoff. The rater's problem is therefore

$$\sup_{\Xi} \sup_{(\sigma, \varphi) \in B(\Xi)} W(\Xi, (\sigma, \varphi)). \quad (6)$$

Remark 1 (Full revelation). The rater can recover a standard repeated game setting using a *fully-revealing* rating system. A rating system is fully-revealing if each consumer puts probability one on the true history of signals upon seeing each rating from the system.⁸ The standard repeated game counterpart to this paper is Mailath and Samuelson (2001), the relation to which is discussed extensively in Section 4.3.

Remark 2 (Prices). The familiar assumption that consumers pay their expected payoffs (Holmström, 1999, Mailath and Samuelson, 2001, Board and Meyer-ter Vehn, 2013) permits a focus on the strategic interaction between the rater and the firm. Upon observing rating r in equilibrium (σ, φ) , a period- t consumer payment is

$$\begin{aligned} p_t^\sigma(r) &= \varphi_t(r) \mathbf{E}^P[\tau(h_f^t, C) | r_t = r, \theta = C](1 - \rho) + (1 - \varphi_t(r) \mathbf{E}^P[\tau(h_f^t, C) | r_t = r, \theta = C])\rho \\ &= \rho + (1 - 2\rho) \underbrace{\varphi_t(r)}_{\text{reputation}} \underbrace{\mathbf{E}^P[\tau(h_f^t, C) | r_t = r, \theta = C]}_{\text{expected effort}}. \end{aligned} \quad (7)$$

The conditional expectation of the competent firm's effort is taken over possible firm's histories h_f^t since the consumer is unaware of firm's history except the rating $r_t = r$. Expression (7) highlights the role ratings play in coordinating consumers' beliefs of the firm's type and effort choices. To maximize her payoff, the rater wishes to maximize both the firm's reputation and its expected effort in each period.

Remark 3 (Market structure). The market structure is familiar from the literature on seller reputation.⁹ Adverse selection arises because consumers (and the rater) are unsure of the firm's type. Moral hazard arises because

⁷Appendix A.2 shows that the set $B(\Xi)$ is non-empty for any menu Ξ .

⁸Appendix A.3 provides details on the construction of a fully-revealing rating system.

⁹See Bar-Isaac and Tadelis (2008) for a comprehensive survey.

the competent firm faces a myopic temptation to shirk against consumers. Absent intermediation, the competent firm would shirk upon receiving upfront payment from consumers, because consumers do not observe any past play. Each consumer thus pays the firm ρ in equilibrium. However, the firm would obtain a higher profit if it could convince consumers that it is competent and that it could commit to exert high effort. By adequately revealing information to consumers, ratings may act as a *screening* device to address adverse selection and as a *commitment* device to address moral hazard.

Remark 4 (Rating relationship). The upfront fee captures in a simple manner the source of a rater’s revenue under the ratee-pays paradigm, and is familiar from the literature (*e.g.*, Lizzeri (1999) and Farhi et al. (2013)). In practice, this monetary transfer from the ratee to the rater may involve a stream of constant payments fixed *ex ante*. The fee in the model can be viewed as a discounted sum of these payments. The once-and-for-all participation decision captures the prevalent feature that most rating relationships are long-lived due to their impact on both parties’ profits, and termination of the relationship is rare.¹⁰ Following the lead of the information intermediation literature, the rater has full commitment power.¹¹ This assumption is plausible when commitment arises from, for example, reputational concern of the intermediary.¹²

Remark 5 (Solution concept). The model admits two possible off-path events. The first corresponds to consumers identifying via the null rating that a firm deviates to its outside option in an equilibrium in which both types participate with probability one. A non-participating competent firm finds it optimal to always exert low effort, because subsequent consumers do

¹⁰Partnoy (2006) documents that the three largest credit rating agencies, Moody’s, Standard & Poor’s and Fitch, receive around 90% of their revenues from fees paid by issuers. A withdrawal of the rating relationship has severe adverse impact on the firm’s profit extraction from consumers. See Salvadè (2014, 2017) for related documents in the context of credit ratings, Luca (2016) in the context of restaurant ratings on Yelp, and Fickenscher (2017) in the context of the hotel industry.

¹¹For example, see Admati and Pfleiderer (1986, 1988), Lizzeri (1999), Albano and Lizzeri (2001), Farhi et al. (2013), Che and Hörner (2017) and Hörner and Lambert (2018).

¹²The rater is viewed as a “reputational intermediary,” who strives to publish credible signals to the market (Coffee, 1997). A famous example of a rater’s desire to signal its commitment is the expulsion of the Los Angeles affiliate of Council of Better Business Bureaus, after it was discovered in 2009 that several eateries in Southern California simply paid for high ratings (Better Business Bureau, 2013). In the context of credit ratings, the raters often disclose information regarding their rating methodologies *ex ante* and their rating processes are monitored by the U.S. Securities and Exchange Commission who publishes public annual reports regarding their performances.

not observe any information about its past play. Expecting this, consumers always pay the firm ρ upon identifying its decision to choose the outside option, regardless of how we specify consumers' off-path beliefs. To be clear, a firm who chooses the outside option may obtain a payoff above ρ . This is the case in an equilibrium in which a competent type participates with positive probability in a rating system that sends a null rating, and exerts high effort upon the null rating. The second off-path event happens when consumers identify that a firm deviates to be rated in an equilibrium that specifies both types to choose the outside option.¹³

3 Full Extraction

We are interested in revenue-maximizing menus that solve the rater's problem (6). In particular, we are interested in *fully-extracting* menus.

Definition 1 (Full extraction). *A menu Ξ is fully-extracting if there exists an equilibrium $(\sigma, \varphi) \in B(\Xi)$ such that $W(\Xi, (\sigma, \varphi)) = \mu(1 - 2\rho - c)$.*

In Definition 1, $\mu(1 - 2\rho - c)$ is an upper bound on the rater's expected payoff, representing the maximal surplus the rater can possibly capture from the market via the firm's payment. To see this, note first that in any equilibrium (σ, φ) , in period 0, the probability of a participating firm being competent is $\mu\pi(\Xi|C)$, while that of a participating firm being inept is $(1 - \mu)\pi(\Xi|I)$. For participation to be individually rational, the rater must compensate a participating firm with a reservation payoff of at least ρ (Remark 5). Because high effort is efficient, the maximal surplus that can be extracted from a competent firm is $1 - \rho - c$, while that from an inept firm is ρ . The net market surplus in period 0 that can be extracted by the rater is thus at most

$$\begin{aligned} & \mu\pi(\Xi|C)(1 - \rho - c) + (1 - \mu)\pi(\Xi|I)\rho - (\mu\pi(\Xi|C) + (1 - \mu)\pi(\Xi|I))\rho \\ &= \mu\pi(\Xi|C)(1 - 2\rho - c). \end{aligned}$$

¹³Equilibrium outcomes are therefore equivalent to those induced by Bayesian Nash equilibrium strategies, given two restrictions: (1) a non-participating firm obtains a payoff ρ , and (2) at least one type participates with positive probability. The first requirement is due to the fact that, in a Bayesian Nash equilibrium in which both types participate, a firm who deviates to its outside option need not shirk afterwards, and may earn a payoff different from ρ . Without the second restriction, there is a perfect Bayesian equilibria in which both types choose the outside option, and a firm who deviates to be rated by a menu is believed to be inept for sure, despite participating may be free-of-charge and the menu may adequately reveal information in favor of the firm. Choosing the outside option therefore need not constitute a Bayesian Nash equilibrium.

Similarly, given any posterior φ_t of a firm being competent in period t , the market surplus in period t that can be extracted by the rater is at most $\varphi_t \pi(\Xi|C)(1 - 2\rho - c)$. The expected discounted sum of market surplus that can be extracted by the rater is therefore at most

$$\begin{aligned} (1 - \delta) \mathbf{E}^P \left[\sum_{t=0}^{\infty} \delta^t \varphi_t \pi(\Xi|C)(1 - 2\rho - c) \right] &= \pi(\Xi|C)(1 - 2\rho - c)(1 - \delta) \sum_{t=0}^{\infty} \delta^t \mathbf{E}^P[\varphi_t] \\ &= \pi(\Xi|C)\mu(1 - 2\rho - c) \\ &\leq \mu(1 - 2\rho - c), \end{aligned}$$

where the second line follows from the martingale property of posteriors.

3.1 Necessity and Sufficiency

We begin with a general characterization of fully-extracting menus.

Proposition 1 (Characterization). *A menu*

$$\Xi = \{\xi_C, \xi_I\} = \{(f_C, S_C, R_C), (f_I, S_I, R_I)\}$$

achieves full extraction if and only if there is an equilibrium $(\sigma, \varphi) \in B(\Xi)$ in which

1. $\pi(\Xi|C) = 1$,
2. $\pi(\Xi|I) > 0$,
3. $\tau(h_f, C) = 1$ after every history h_f that occurs with positive probability conditional on the firm's type being competent,
4. for each type θ , if $\pi(\xi_{\theta'}|\theta) > 0$, then $(1 - \delta)f_{\theta'} = U(\sigma, \theta; \xi_{\theta'}) - \rho$.

The rater's payoff under full extraction is $\mu(1 - 2\rho - c)$.

The first condition says that the competent type participates with probability one. Second, the inept type participates with positive probability. Third, the competent firm consistently exerts high effort upon participation. Together, they illustrate that the rater leverages the presence of adverse selection to fully solve moral hazard. Finally, if a firm participates in a scheme with positive probability, then the fee extracts all its profits above ρ . Observe also that the proposition is not an existence result. There are parameters (μ, δ, ρ, c) such that fully-extracting menus do not exist (Proposition 10).

The rater achieves full extraction in an equilibrium using some menu if and only if several conditions hold. First, to generate the maximal surplus

$1 - \rho - c$ in the event that the firm is competent, equilibrium calls for *consistent high effort* by the competent type.¹⁴ Second, the rater extracts a profit from the competent firm with probability equal to μ if and only if the competent firm participates with probability one. Third, each type participates with positive probability. Suppose on the contrary that there is an equilibrium in which only the competent firm participates with positive probability, and it consistently exerts high effort. The competent firm's reputation equals one upon participation and consumer payments always equal $1 - \rho$ by (7). This destroys the putative equilibrium, as the competent firm faces an irresistible temptation to shirk after each history upon participation, saving on effort cost yet receiving the same future revenue. On the other hand, a competent firm who chooses the outside option finds it optimal to shirk in each period (Remark 5).

Remark 6 (The value of inept participation). The idea that attracting both types is necessary for optimality may appear surprising at first glance. When μ is relatively small, for example, consumer payments may remain quite small despite consistent high effort from the competent firm, limiting the surplus that the rater can extract. A reasonable conjecture is that the rater wishes to set a sufficiently high rating fee that screens away the inept type (*i.e.*, rules out equilibria in which the inept type participates with positive probability), raising consumers' beliefs of a participating firm's type and hence their willingness to pay. Lemma 3 in Appendix B.1 shows precisely that this seemingly compelling intuition is wrong. Specifically, in any equilibrium in which only the competent firm participates with positive probability, the rater obtains at most

$$\mu \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right), \quad (8)$$

which is strictly smaller than $\mu(1 - 2\rho - c)$.¹⁵ Screening away the inept firm presents two drawbacks on the rater's expected payoff. First, the rater cannot extract any surplus from the inept firm. Second, consistent high effort is no longer possible. At best, ratings coordinate *frequent* high effort. Lemma 3 constructs a menu that gives the rater the payoff (8) by coordinating consumers' expectations on the competent firm's effort upon different ratings. Specifically, they expect high effort upon rating 1 and thus pay the firm $1 - \rho$

¹⁴That is, the competent type exerts high effort after every history that occurs with positive probability in equilibrium.

¹⁵The lemma explicitly constructs a menu and an equilibrium in the induced game in which only the competent firm participates that give the rater a payoff equal to (8).

but expect low effort upon rating 0 and pay ρ . The bound (8) is attained when rating 1 is visited frequently enough, but not too frequently to diminish the threat of transiting to rating 0 to maintain incentives for effort.

Remark 7 (Inept participation probability). In view of Remark 6, it is curious that it suffices to have the inept firm participating with some positive probability. After all, if an inept type participates with a higher probability, the rater extracts the expected surplus from the inept firm with a higher probability. However, the higher is the inept firm's participation probability, the lower are consumer payments and the lower is the surplus that can be extracted by the rater. Proposition 1 shows precisely that these two counter-veiling forces on the rater's payoff balance out in expectation.

Remark 8 (The social value of intermediation). Given fully-extracting menus, the firm obtains the same payoff as if no intermediation is available, *i.e.*, ρ . The important difference in the two settings is the firm's behavior and hence the firm's value to consumers. Full-extracting ratings make the competent firm as valuable to consumers as it can be. Importantly, because by construction consumers always get an expected payoff of 0 and the payment to the rater by the firm is simply a monetary transfer, fully-extracting menus are *socially efficient*, despite the rater's self-interest.

3.2 A Simple, Fully-Extracting Menu

To illustrate Proposition 1, I now construct a menu $\Xi^* = \{\xi_C^*, \xi_I^*\}$ that achieves full extraction for the rater in an equilibrium $(\sigma^*, \varphi^*) \in B(\Xi^*)$, whenever

$$c \leq \bar{c}(\mu, \delta, \rho) := \delta(1 - 2\rho)^2 \left(\frac{1 - \mu}{1 - \mu + \mu\rho} \right). \quad (9)$$

Condition (9) permits a very simple construction to illustrate the underlying economic forces. We consider settings in which this condition does not hold in Section 5. Not surprisingly, this condition refines the static, benchmark efficiency condition (1). To motivate high effort, its benefit must outweigh the cost. The benefit, however, is only partially captured through future consumer payments due to discounting and incomplete information over the firm's type while the cost is incurred immediately.

Condition (9) relaxes when the firm is relatively patient (high δ), or if monitoring is relatively precise (small ρ), or if consumers are skeptical of the firm's ability (small μ). One would expect ratings to provide effective inter-temporal incentives when the firm cares about its future profits influenced

by the ratings, and when monitoring is less noisy. Finally, ratings add little to affecting consumers' beliefs when consumers are born confident about the firm being competent.

In the scheme ξ_C^* , the rating set $R_C^* = \{0, 1\}$ is binary. The rating system S_C^* announces rating 1 upon a good signal, and announces rating 0 otherwise:

$$S_C^*(h_r^t) = \begin{cases} 1 \circ \{1\}, & \text{if } y_{t-1} = \bar{y}, \\ 1 \circ \{0\}, & \text{otherwise.} \end{cases} \quad (10)$$

In the scheme ξ_I^* , the rating set $R_I^* = \{0\}$, and the system S_I^* always gives a rating 0:

$$S_I^*(h_r) = 1 \circ \{0\}, \text{ for every } h_r. \quad (11)$$

The fees f_C^* and f_I^* are given in Appendix A.4. They extract all profits above ρ of the participating firm in the equilibrium (σ^*, φ^*) .

The equilibrium strategy $\sigma^* = (\pi^*, \tau^*)$ calls for the firm to participate with probability one in the scheme intended for its type in the menu, and the competent type always exerts high effort upon participation:

$$\pi^*(\xi_C^*|C) = 1, \quad \pi^*(\xi_I^*|I) = 1, \quad (12)$$

$$\tau^*(h_f, C) = \begin{cases} 1, & \text{if } d = \xi_C^*, \\ 0, & \text{otherwise.} \end{cases} \quad (13)$$

Consumers believe that a firm is inept for sure upon seeing a null rating, $\varphi_t^*(\emptyset) = 0$ for every t . Specifying this off-path belief is without loss in view of Remark 5, and a firm's outside option payoff is ρ .

Proposition 2 (Simple characterization). *Given (9), $(\sigma^*, \varphi^*) \in B(\Xi^*)$, and the menu Ξ^* is fully-extracting in (σ^*, φ^*) .*

By (7), prices satisfy $p_t^{\sigma^*}(r) = \rho + (1 - 2\rho)\varphi_t^*(r)$ in equilibrium. Each consumer puts probability one on the firm being competent upon observing the rating 1, because it must be drawn from the system S_C^* . This leads to the maximal payment $p_t^{\sigma^*}(1) = 1 - \rho$ for any period $t \geq 1$. In contrast, the rating 0 is *opaque* in the sense that consumers cannot infer the firm's type via the rating. In any period $t \geq 1$, an entering consumer who observes a rating 0 infers that either the firm is competent but delivers a bad signal in the last period or the firm is inept. Consumer payment then equals

$$p_t^{\sigma^*}(0) = \rho + (1 - 2\rho)\varphi_t^*(0) = \rho + (1 - 2\rho)\left(\frac{\mu\rho}{\mu\rho + 1 - \mu}\right).$$

The rating 0 in period 0 does not reflect any past signal, thus $\varphi_0^*(0) = \mu$, and consumer payment equals $p_0^{\sigma^*}(0) = \rho + (1 - 2\rho)\mu$.

The competent firm consistently exerts high effort to chase the high consumer payment associated with rating 1 and to avoid the low payment associated with rating 0. Because the system S_C^* has one-period memory, the discounted benefits of high effort in any period t realized in period $t + 1$ must outweigh the immediate cost incurred:

$$\delta \left[\underbrace{((1 - \rho)p_{t+1}^{\sigma^*}(1) + \rho p_{t+1}^{\sigma^*}(0))}_{\text{expected payment after high effort}} - \underbrace{(\rho p_{t+1}^{\sigma^*}(1) + (1 - \rho)p_{t+1}^{\sigma^*}(0))}_{\text{expected payment after low effort}} \right] \geq (1 - \delta)c.$$

This incentive constraint for high effort is satisfied whenever (9) holds.

In any equilibrium, say that a rating system is *coarse* if the number of signal histories exceeds the number of ratings that can be announced with positive probability in some period. The rating system S_C^* is coarse: only two ratings are sent in each period $t \geq 1$, despite a firm's exponentially growing number of possible track records. The ratings limit consumer learning of a rated firm's track record, but is crucial for fully solving moral hazard. They ensure that consumer posteriors are sufficiently responsive to every signal delivered by the firm, sustaining reputation effects on the incentives for high effort over the long run. To see this, consider a fully-revealing system (Remark 1). In an equilibrium in which both types are rated by this system, consistent high effort is impossible. Otherwise, consumer posteriors would eventually become very close to 1 (conditional on the firm being competent), and further signals have virtually no impact on the posteriors. The competent type then finds it profitable to shirk, destroying the putative equilibrium. Similarly, if the rating 0 in the inept scheme is instead labelled as 2, then it ceases to be opaque and both ratings 1 and 0 perfectly reveal the firm's type being competent. This effectively fully solves the adverse selection problem. In view of Remark 6, the moral hazard problem cannot be fully solved. The coarseness and opacity exhibited by the ratings in Ξ^* thus illustrate a clear tension between rating transparency and economic efficiency.

Together with the ratings, the fees ensure that the firm finds it incentive-compatible to choose the scheme intended for its own type. Because the competent type consistently exerts high effort, an inept type who deviates to participate in the competent scheme obtains the higher payment $p_t^{\sigma^*}(1)$ less frequently and thus obtains lower expected profits than the competent firm does. This deviation must give the inept firm a payoff strictly less than ρ , because the competent firm obtains a payoff of exactly ρ after paying the fee f_C^* . On the other hand, the competent firm is indifferent between

both schemes. By deviating to the “inept” scheme, the competent firm extracts the same profits from consumers as the inept type does because the inept scheme ξ_I^* is fully coarse: there is only one rating and thus one possible payment. Since the inept firm obtains a payoff ρ after paying the fee f_I^* , so does the competent firm upon the deviation. Participation is also individually rational, since deviating to the outside option gives each type a payoff ρ .

Remark 9 (Contrast to the intermediation literature). The information intermediation literature often asks why intermediaries often provide coarse ratings, withholding some collected information. The present result that the revenue-maximizing ratings are coarse and opaque is familiar from Lizzeri (1999) in the context of static certification absent moral hazard.¹⁶ In contrast, in the current setting the rater does not observe the participating firm’s private type. The ratings thus also serve a *screening* role, ensuring that the firm picks the scheme intended for its own type.

The ratings also serve a *sanctioning* role, because they need to tackle moral hazard and incentivize consistent high effort by the competent firm. To carry out this role, the ratings censor past play from consumers to ensure reputation effects over the competent firm’s effort choices remain effective over the long run. This property of limited memory to maintain long-run reputation effects is familiar from Ekmekci (2011).¹⁷ He characterizes a rating system with one-period memory that achieves frequent high effort by a long-lived firm in a repeated product-choice game, where the firm faces a succession of short-lived consumers and the firm might be a Stackelberg type. Contrary to the equilibrium he constructs, which features low effort upon some ratings,¹⁸ the assumption here that the firm is possibly an inept type but never a Stackelberg type provides the competent firm with a perpetual incentive to exert high effort to “separate” from the inept type.

Remark 10 (No opt-out). Even if the model allows the firm to choose to quit the rating relationship in each period, the firm never has a strict incentive to do so in any equilibrium, as continuing allows the firm to be at least as well off (Remark 5). On the other hand, any equilibrium in which the

¹⁶See also Harbaugh and Rasmusen (2018) and the references therein.

¹⁷Relatedly, see Liu (2011) and Liu and Skrzypacz (2014).

¹⁸In the equilibrium constructed, the firm exerts *low* effort yet consumers pay the firm a high price upon the best rating, giving the firm a high payoff. Upon the worst rating, however, the firm exerts low effort and consumers pay a low price, giving the firm a low payoff. The normal type of firm exerts high effort whenever other ratings emerge in an attempt to achieve the best rating.

firm quits the relationship must not be rater-preferred, because the rater's payoff falls short of the upper bound.

3.3 Full Extraction Without Screening of Types

The menus Ξ^* relies on the rater offering distinct schemes to screen a firm's type. It is instructive to examine the rater's optimal strategy if she is restricted to use only singleton menus. Here, I sketch a singleton menu $\Xi^{**} = \{\xi^{**}\}$ that achieves full extraction for the rater whenever

$$c \leq \bar{c}_1(\mu, \delta, \rho) = \frac{\delta(1-\mu)\mu(1-2\rho)^3}{(1-\rho-\mu(1-2\rho))(\rho+\mu(1-2\rho))}. \quad (14)$$

Note that (14) is stronger than (9): $\bar{c}_1 < \bar{c}$.¹⁹ We will see that this difference comes from dramatically different belief dynamics induced by the two menus Ξ^* and Ξ^{**} . In the scheme $\xi^{**} = (f^{**}, R^{**}, S^{**})$, the set of ratings $R^{**} = \{0, 1\}$ is binary, and the rating system delivers a rating 1 with probability α if the last signal is good, and delivers a rating 0 otherwise:

$$S^{**}(h_r^t) = \begin{cases} \alpha \circ \{1\} + (1-\alpha) \circ \{0\}, & \text{if } y_{t-1} = \bar{y}, \\ 1 \circ \{0\}, & \text{otherwise,} \end{cases} \quad (15)$$

for each history h_r^t , where $\alpha \equiv \alpha(\mu, \delta, \rho, c)$ is a well-defined probability whenever (14). The probability α and the fee f^{**} are presented in Appendix A.5. The fee is set to fully extract the surplus above ρ from *both* types who participate in an equilibrium $(\sigma^{**}, \varphi^{**})$ specified below.

The strategy $\sigma^{**} = (\pi^{**}, \tau^{**})$ specifies that both types participate in the scheme ξ^{**} with probability one and the competent firm exerts high effort consistently upon participation:

$$\pi^{**}(\xi^{**}|C) = \pi^{**}(\xi^{**}|I) = 1, \quad (16)$$

$$\tau^{**}(h_f, C) = \begin{cases} 1, & \text{if } d = \xi^{**}, \\ 0, & \text{if } d = N, \end{cases} \quad (17)$$

for all histories $h_f \in H_f$. Consumers believe that a firm is inept for sure upon seeing a null rating, $\varphi_t^{**}(\emptyset) = 0$ for every t .

Proposition 3 (Singleton menu). *Given (14), $(\sigma^{**}, \varphi^{**}) \in B(\Xi^{**})$, and the menu Ξ^{**} achieves full extraction in $(\sigma^{**}, \varphi^{**})$.*

¹⁹ To see this, note that the inequality can be simplified to $(1-\rho)/(1-2\rho) > (2-\mu)\mu$. The left side is strictly larger than 1 for $\rho \in (0, \frac{1}{2})$. The right side is bounded above by 1.

The rating 0 in period 0 does not reflect any past signal, thus $\varphi_0^{**}(0) = \mu$. In each period $t \geq 1$, the two ratings statistically reveal the signal in the last period, inducing two possible levels of reputations:

$$\varphi_t^{**}(0) = \frac{\mu(1 - \alpha(1 - \rho))}{\mu(1 - \alpha(1 - \rho)) + (1 - \mu)(1 - \alpha\rho)} < \varphi_t^{**}(1) = \frac{\mu(1 - \rho)}{\mu(1 - \rho) + (1 - \mu)\rho}.$$

Consistent high effort again implies that $p_t^{\sigma^{**}}(r) = \rho + (1 - 2\rho)\varphi_t^{**}(r)$. Thus, $p_t^{\sigma^{**}}(1) > p_t^{\sigma^{**}}(0)$ for each $t \geq 1$. When (14) holds, the competent firm consistently exerts high effort for the higher payment $p_t^{\sigma^{**}}(1)$.

Both ratings 0 and 1 are opaque. The firm's reputation cycle is largely dampened, in contrast to that in the equilibrium (σ^*, φ^*) in the game induced by Ξ^* . The worst reputation a firm can acquire by participating in the menu is $\varphi_1^{**}(0)$, which reflects either a bad signal or with some probability a good signal in the past period. The comparison is depicted in Figure 2 below.

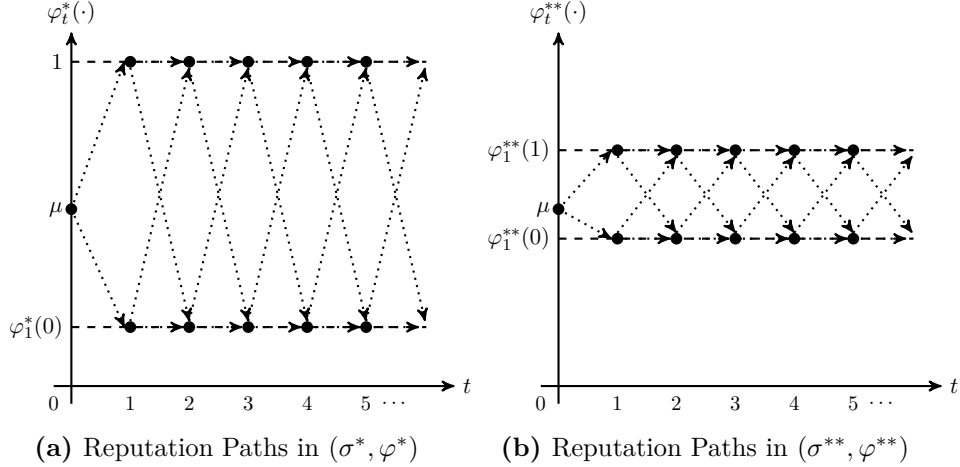


Figure 2: Amplified and Dampened Reputation Dynamics

The figure depicts all possible reputation paths of the competent firm in each of the rater-preferred equilibria in the induced games of interest. Specifically, it illustrates how the menu Ξ^* amplifies both the reputation building and dissipation processes, and how the menu Ξ^{**} dampens the processes on the contrary. Note that the reputation dynamics on the right panel applies also to the inept type. On the left panel, however, an inept firm's reputation stays at the lower bound at all times except period 0.

Proposition 3 is surprising in view of Proposition 1, which says that the competent firm must consistently exert high effort and the fee must fully extract the participating firm in the rater-preferred equilibrium. For a

singleton menu to be fully-extracting, the single fee must bind *both* types' participation constraints. This requires both types to extract from consumers identical expected profits, but consistent high effort arises from the competent firm's incentive to separate itself from the inept type and to secure higher expected profits. Indeed, the probability α is chosen so that the competent firm's incentive constraint for high effort binds after every history upon participating in the scheme ξ^{**} . In particular, by creating the possibility of reaching a bad rating 0 even after a good signal, the rating system S^{**} pools the two reputation levels and their induced payments closer together. The competent firm is indifferent between each effort level after every such history, because the potential benefit of acquiring a good rating 1 balances the cost incurred after each history upon participation. The competent type thus extracts the same amount of profits from consumers as an inept type does, so that the fee f^{**} binds the participation constraints of both types. Specifically, $U(\sigma^{**}, C; \xi^{**}) = U(\sigma^{**}, I; \xi^{**}) = \rho + \mu(1 - 2\rho - c)$. Effectively, the rater "transfers" all the rents from the competent type to the inept type, without disrupting the competent type's incentive for high effort. After paying the fee f^{**} , each type obtains a payoff ρ . Participation is individually rational, since deviating to the outside option again gives each type ρ .

It is now not surprising that $\bar{c}_1 < \bar{c}$. The menu Ξ^{**} needs to ensure reputation effects remain strong enough to motivate consistent high effort when reputation building is largely dampened.

4 Equilibrium Implications

This section explores how adverse selection and moral hazard shape the structure of fully-extracting menus. It also examines the implications of fully-extracting menus for market beliefs and behaviors.

4.1 Implications of Moral Hazard

I now show that the presence of moral hazard forces fully-extracting ratings to contain some information content. We first characterize fully-extracting menus absent moral hazard (*i.e.*, in a pure adverse selection setting), in which $c = 0$ and the competent firm only exerts high effort, $\tau(\cdot, C) = 1$.

Proposition 4 (Pure adverse selection). *In a pure adverse selection setting, a menu Ξ is fully-extracting if and only if either condition is true:*

- A. *there exists an equilibrium $(\sigma, \varphi) \in B(\Xi)$ such that*

1. $\pi(\Xi|C) = 1$,
2. $\pi(\Xi|I) > 0$,
3. for each type θ , if $\pi(\xi_{\theta'}|\theta) > 0$, then $(1 - \delta)f_{\theta'} = U(\sigma, \theta; \xi_{\theta'}) - \rho$.

B. there exists an equilibrium $(\sigma, \varphi) \in B(\Xi)$ such that

1. $\pi(\Xi|C) = 1$,
2. $\pi(\Xi|I) = 0$,
3. if $\pi(\xi_{\theta}|C) > 0$, then $(1 - \delta)f_{\theta} = U(\sigma, C; \xi_{\theta}) - \rho$.

The rater's payoff under full extraction is $\mu(1 - 2\rho)$.

Not surprisingly, the set of fully-extracting menus expands in comparison to Proposition 1, as the rater no longer needs to motivate high effort. The conditions in Part A are familiar from Proposition 1. The additional fully-extracting menus, characterized by Part B, induces an equilibrium in which only the competent firm participates, and does so with probability one. In the schemes to which it assigns positive probability for participation, the fees fully extract its profits above ρ .

Analogous to Proposition 1, this is not an existence result. But there are indeed fully-extracting menus that satisfy both sets of conditions. I present a menu $\Xi_{AS} = \{\xi_{AS}\}$ with an equilibrium satisfying the conditions in Part A, in which ratings contain no information content. In the scheme, the rating set contains only a rating 0, and the system always announces the rating 0. The fee f_{AS} satisfies $(1 - \delta)f_{AS} = \mu(1 - 2\rho)$. The equilibrium calls for participation by both types with probability one, and consumers pays a firm ρ upon observing a null rating. This is an equilibrium because the fee f_{AS} binds the participation constraints of *both* types: analogous to the discussion following Proposition 1, the expected market surplus that can be extracted by the rater is bounded above by $\rho + \mu(1 - 2\rho)$. In this equilibrium, consumer posteriors always equal the prior, leading to identical payments over time, as if intermediation is absent. Intermediation here adds no value to the market. In contrast, when moral hazard is present, intermediation adds value (see Remark 8) because information must be revealed to create differential payments to sustain incentives for high effort.

There is a menu $\Xi'_{AS} = \{\xi'_{AS}\}$ with an equilibrium satisfying the conditions in Part B. In the scheme ξ'_{AS} , the rating set contains only one non-null rating, and the system always announces that rating. The fee f'_{AS} satisfies $(1 - \delta)f'_{AS} = 1 - 2\rho$. In the equilibrium, only the competent firm participates, and does so with probability one. Neither type wishes to deviate from their participation choices because it would earn a payoff of ρ either way. An inept

firm is immediately exposed to consumers once they observe a null rating and obtains ρ . The competent firm obtains $1 - \rho$ from consumers, and the fee f'_{AS} extracts all its profit above ρ . The rater obtains $\mu(1 - 2\rho)$.

4.2 Implications of Adverse Selection

This section illustrates the implications of adverse selection to the rater and the competent firm. To set the stage, I characterize a fully-extracting menu in a pure moral hazard setting, in which the firm is commonly known to be competent, that is, $\mu = 1$.

Proposition 5 (Pure moral hazard). *Consider a pure moral hazard setting, and fix (δ, ρ) . The competent firm's equilibrium profit extraction from consumers is bounded above by*

$$1 - \rho - c - \frac{\rho c}{1 - 2\rho}. \quad (18)$$

The rater's expected payoff is bounded above by

$$1 - 2\rho - c - \frac{\rho c}{1 - 2\rho}. \quad (19)$$

There exists $\bar{c}_{MH}(\delta, \rho)$ such that for every $c \leq \bar{c}_{MH}(\delta, \rho)$, there exists an optimal menu $\Xi_{MH} = \{\xi_{MH}\}$ and an equilibrium in the induced game that give the rater exactly the payoff (19) and the competent firm exactly the profit (18) from consumers.

The construction of the menu Ξ_{MH} and the equilibrium is in Appendix A.6. The threshold of interest is

$$c \leq \frac{\delta(1 - 2\rho)^2}{1 - \delta\rho} =: \bar{c}_{MH}(\delta, \rho). \quad (20)$$

Not surprisingly, it is weaker than (9), because adverse selection which depresses the benefits of exerting high effort is no longer present.

Absent adverse selection, consistent high effort is not possible in equilibrium. Optimal ratings here reward the competent firm by revealing good signals to motivate high effort as frequently as possible.²⁰ When there is adverse selection and hence the possibility to induce consistent high effort in equilibrium, the rater can achieve the maximal payoff $\mu(1 - 2\rho - c)$ by Proposition 1, exceeding (19) when μ is sufficiently large.

In contrast, the competent firm is worse off given adverse selection:

²⁰This result is familiar from Dellarocas (2005), who considers a product-choice game with pure moral hazard and finds that effort-maximizing ratings coordinate different effort choices in reward and punishment phases.

Proposition 6. *In the setting with both adverse selection and moral hazard, the competent firm's equilibrium profit extraction from consumers is strictly less than (18).*

Proposition 6 also illustrates a contrast with canonical reputation games under imperfect monitoring in which the Stackelberg type can arise (Fudenberg and Levine, 1992). In those cases, the competent player subject to binding moral hazard may be assured a payoff in the game with incomplete information over the firm's type in excess of any equilibrium payoff in the complete-information game. Incomplete information opens the possibility of a horizon during which the consumers are uncertain of the firm's type, and in which the play does not resemble any equilibrium play of the complete-information game and is *in favor of* the competent firm. Here, the possible presence of an inept firm poses a negative impact on consumers' payments and the competent firm's profits.

4.3 Implications of Fully-Extracting Ratings

I now relate the results to the literature on repeated games. In the canonical settings, there are no ratings and no rater, and consumers observe the history of signals. To make a reasonable comparison, I focus on the competent firm's expected profit extraction from consumers for fixed (c, μ, δ, ρ) .

4.3.1 Adverse Selection and Moral Hazard

Consider first the main setting of interest in which both adverse selection and moral hazard are present. The corresponding canonical setting is Mailath and Samuelson (2001). In the proof of Proposition 1, I show that the competent firm's expected profit given the menu Ξ^* in equilibrium (σ^*, φ^*) is

$$\bar{U}_{AS+MH}^{\text{rating}}(\mu) := \rho + \mu(1 - 2\rho) + \frac{\delta(1 - \mu)^2(1 - \rho)(1 - 2\rho)}{1 - \mu(1 - \rho)} - c. \quad (21)$$

In the canonical setting, it is standard to show that, when c is sufficiently small, the upper bound on a firm's equilibrium profit is²¹

$$\bar{U}_{AS+MH}^{\text{public}}(\mu) := \rho + (1 - \delta)(1 - 2\rho) \sum_{t=0}^{\infty} \frac{\delta^t(1 - \rho)^t \mu}{(1 - \rho)^t \mu + \rho^t(1 - \mu)} - c - \frac{\rho c}{1 - 2\rho}. \quad (22)$$

²¹See Appendix A.7 for the omitted calculations.

When c is sufficiently small, there exists an equilibrium in the canonical setting in which the competent firm frequently but never consistently exerts high effort (Mailath and Samuelson, 2015):²²

Proposition. (Mailath and Samuelson, 2015) *There exists an equilibrium in the canonical setting in which the competent firm chooses high effort in the beginning and continues to do so as long as signal $\bar{y}$ is realized. Upon producing a signal $\underline{y}$ after history h_f , the competent firm switches to low effort for a finite horizon with length $L(h_f) \geq 0$, before play resumes with the firm choosing high effort.*

When the length $L(h_f)$ is chosen small enough for each history h_f without disrupting incentives, the firm obtains an equilibrium profit close to (22). The comparison between the two bounds is depicted in Figure 3 for

$$\mu \leq \bar{\mu}(c, \delta, \rho) := 1 - \frac{c\rho}{\delta - (c + 4\delta\rho)(1 - \rho)},$$

which is obtained by rearranging (9), with $c = 0.1$, $\delta = 0.9$ and $\rho = 0.25$.²³ The effect of consistent high effort is particularly dramatic when μ is sufficiently close to $\bar{\mu}$. In this event, there is a reversal in the competent firm’s profit extraction from consumers in the two settings. Specifically, it extracts a higher profit when the rater obscures information rather than publicly disclosing the complete history of signals. Consumer payments remain relatively large during punishment phases (*i.e.*, upon “bad” rating) because consumers’ posteriors are only slightly less than μ , pushing the firm’s profit beyond the canonical bound. When μ is low, the competent firm manages to extract more profits from consumers than in the canonical setting because the scheme ξ_C^* reveals the competent firm’s true type *immediately* and *perfectly* every time after a good signal. In the canonical setting, it takes time for the competent firm to build up its reputation to enjoy high consumer payments under frequent high effort. For intermediate values of μ , the competent firm’s payoff upper bound in the canonical setting may exceed that in the rating case when the firm is relatively patient. In the canonical setting, the initial periods of relatively low payments then contribute less to the firm’s profits, and the later periods of higher consumer payments when it gradually builds up its reputation contribute more under frequent high effort. When posteriors are large, a bad signal has little impact on the posterior. In the rating game, however, each bad signal leads immediately to the low payment associated with rating 0, regardless of other past performances.

²²The failure to sustain consistent high effort here is an instance of a more general result due to Cripps et al. (2004).

²³The infinite sum in (22) is numerically computed by `NSum` in Mathematica[©].

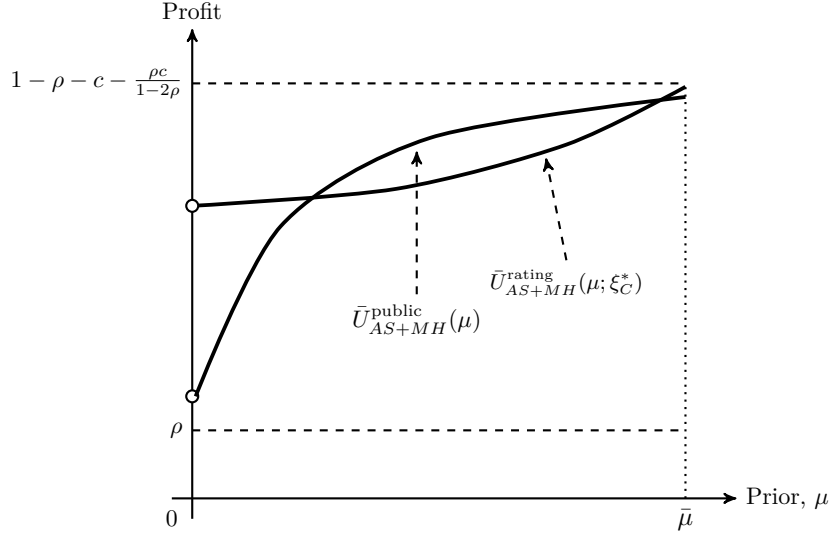


Figure 3: Implications under both Adverse Selection and Moral Hazard

This figure plots the maximal equilibrium profit extraction from consumers by the competent firm as a function of the prior in the rating game induced by the menu Ξ^* ($\bar{U}_{AS+MH}^{\text{rating}}(\mu; \xi_C^*)$) and in the canonical setting ($\bar{U}_{AS+MH}^{\text{public}}(\mu; \xi_C^*)$) with both adverse selection and moral hazard respectively. When the prior is relatively low, the firm does better in the rating game because the rating 1 allows an immediate revelation of its type, whereas it takes time to build a reputation in the canonical setting. In the intermediate region, it does better in the canonical setting because a bad signal has little impact on consumer posteriors once it builds a reputation, but the rating 0 can lead to a relatively large fall in profits in any period in the rating game. Finally, there is a reversal when the prior is sufficiently large: the fact that consistent high effort is feasible in the rating game but not in the canonical setting makes the firm better off in the rating game induced by the menu Ξ^* .

4.3.2 Pure Adverse Selection

Consider next the comparison under pure adverse selection. The canonical setting with pure adverse selection corresponds to the market interaction in Tadelis (1999). The equilibrium profit upper bound can be obtained by taking $c = 0$ in (22), giving

$$\bar{U}_{AS}^{\text{public}}(\mu) := (1 - \delta)(1 - 2\rho) \sum_{t=0}^{\infty} \delta^t \frac{(1 - \rho)^t \mu}{(1 - \rho)^t \mu + \rho^t (1 - \mu)} + \rho. \quad (23)$$

In the rating game with pure adverse selection, the optimal menu Ξ'_{AS} characterized in Section 4.1 induces an equilibrium that reveals the firm's

type to consumers upon participation. The competent firm thus extracts the maximal profit $1 - \rho$ in every period. For every $\mu \in (0, 1)$, this once-and-for-all reputation establishment allows the firm to extract strictly higher profit, as illustrated in Figure 4 below.

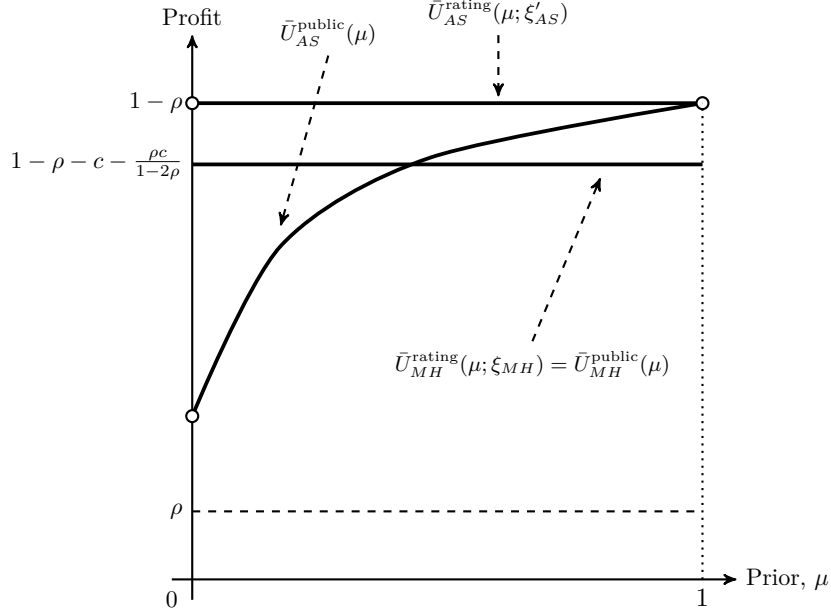


Figure 4: Implications under Pure Adverse Selection or Pure Moral Hazard

This figure plots the maximal equilibrium profit extraction by the competent firm against the prior in the settings with pure adverse selection and with pure moral hazard. With pure adverse selection, the menu Ξ'_{AS} helps the competent firm to establish a reputation beginning from period 0, allowing it to do strictly better in the rating game ($\bar{U}_{AS}^{\text{rating}}(\mu; \xi'_{AS})$) than in the canonical game ($\bar{U}_{AS}^{\text{public}}(\mu; \xi'_{AS})$). With pure moral hazard, the competent firm is equally well off in the rating game and in the canonical game, *i.e.* $\bar{U}_{MH}^{\text{rating}}(\mu; \xi_{MH}) = \bar{U}_{MH}^{\text{public}}(\mu)$. This follows because the rater essentially maximizes the firm's profit extraction from consumers in order to extract the largest surplus from the firm.

4.3.3 Pure Moral Hazard

Consider finally the comparison under pure moral hazard. Here, the canonical setting corresponds to Fudenberg and Levine (1994). The maximal equilibrium profits in the rating game and in the canonical setting must coincide, because maximizing the rating fee that can be extracted amounts

to maximizing the competent firm's profits. Proposition 5 implies that the profit upper bound in both settings is

$$\bar{U}_{MH}^{\text{rating}}(\mu; \xi_{MH}) = \bar{U}_{MH}^{\text{public}}(\mu) := 1 - \rho - c - \frac{\rho c}{1 - 2\rho}, \quad (24)$$

again illustrated in Figure 4. Further, the proposition shows that the bound is tight in the rating game using a scheme with two ratings, and transitions between the ratings across each period depends only on the most recent signal realized. Indeed, the bound is also tight in the canonical setting. One can interpret the rating system as a two-state automaton depicting an equilibrium in the canonical setting with frequent high effort, such that one rating corresponds to the “good” state in which the firm exerts high effort and obtains the maximal stage profit $1 - \rho - c$, and the other rating corresponds to the “bad” state in which the firm exerts low effort and obtains the minimal stage profit ρ . The ratings here have no impact on the maximal equilibrium profit and firm's behaviors.

5 Full Extraction with Higher Costs

For full extraction, the menu Ξ^* relies on (9) that $c \leq \bar{c}$, leaving open whether fully-extracting menus exist for larger costs. Proposition 7 makes precise an intuition that full extraction fails for large costs, because consistent high effort fails in equilibrium (Proposition 1).

Proposition 7 (Necessary condition for existence of full extraction). *If the rater's equilibrium payoff equals $\mu(1 - 2\rho - c)$, then $c < \delta(1 - 2\rho)^2$.*

Because $\delta(1 - 2\rho)^2 > \bar{c}$, there is a region of costs in which full extraction may be achievable by some menu, but not by Ξ^* . Section 5.1 first shows that there is no such menu when c is too close to $\delta(1 - 2\rho)^2$ for fixed $\delta \in (0, 1)$, motivating Section 5.2 to consider a patient-limit setting and characterize a fully-extracting menu for $c \leq \lim_{\delta \uparrow 1} \delta(1 - 2\rho)^2 = (1 - 2\rho)^2$. We also discuss the relationship between rating coarseness and the severity of the moral hazard problem. Section 5.3 turns to a broader picture and discusses the rater's optimal behavior given any c that satisfies (1).

5.1 Impossibility with Fixed Discounting

Proposition 8 (Impossibility with fixed discounting). *There exists $\bar{\eta} \equiv \bar{\eta}(\mu, \delta, \rho) > 0$ such that for every $\eta \in (0, \bar{\eta}]$, full extraction is impossible in equilibrium whenever $c \geq \delta(1 - 2\rho)^2 - \eta$.*

When c is close to $\delta(1 - 2\rho)^2$, the firm has a stronger incentive to shirk. To incentivize consistent high effort by a firm when c is arbitrarily close to $\delta(1 - 2\rho)^2$ in an equilibrium in which both types participate, punishment upon a bad signal *after any history* must be made as harsh as possible. Because of discounting, such harsh punishments include the use of ratings that induce a sufficiently low reputation of the firm soon enough after a bad signal. To induce a low reputation, consumers must put a small enough probability on the event that the participating firm is competent conditional on the rating. But if these ratings are realized soon and frequently enough upon a bad signal after any history by the competent type, consumer posteriors about the firm being competent cannot be too small, yielding a contradiction.

5.2 A Limit-Payoff Characterization

In view of Proposition 8, we consider a setting in which the firm chooses σ to maximize the *limit* (unnormalized) payoff

$$\liminf_{\delta \uparrow 1} \frac{U(\cdot, \theta)}{1 - \delta},$$

i.e., the firm is infinitely patient. I characterize a menu that achieves full extraction in the limit-payoff setting whenever $c \leq \lim_{\delta \uparrow 1} \delta(1 - 2\rho)^2 = (1 - 2\rho)^2$. We assume that monitoring is sufficiently precise whenever the market is relatively confident:

$$\mu > \frac{1}{2} \implies \rho < \frac{1}{\mu} - 1. \quad (25)$$

Imposing the patient limit circumvents the contradiction observed in Proposition 8 by obviating the need to visit the “bad” ratings soon and frequently enough upon a bad signal after any history.

Consider the menu $\Xi^\infty = \{\xi_C^\infty, \xi_I^\infty\}$, where $\xi_\theta^\infty = (f_\theta^\infty, R_\theta^\infty, S_\theta^\infty)$, defined as follows. The competent rating set $R_C^\infty = \{0, 1, 2, \dots\}$ is the set of non-negative integers, while the inept rating set $R_I^\infty = \{1, 2, \dots\}$ is the set of positive integers. Let

$$\phi \equiv \phi(\mu, \rho) := \begin{cases} 1, & \text{if } \mu \in (0, \frac{1}{2}], \\ \frac{1 - \mu(1 + \rho)}{(1 - \mu)(1 - \rho)}, & \text{if } \mu \in (\frac{1}{2}, 1), \end{cases} \quad (26)$$

be a transition probability. By (25), $\phi \in (0, 1]$. Rating transitions given by the rating system S_θ^∞ in each scheme are depicted below in Figure 5. The

systems S_C^∞ and S_I^∞ , as well as the fees f_C^∞ and f_I^∞ , are formally specified in Appendix A.8. The fees fully extract all profits above ρ of the participating firm in the equilibrium $(\sigma^\infty, \varphi^\infty) \in B(\Xi^\infty)$, where $\sigma^\infty = (\pi^\infty, \tau^\infty)$, defined below.

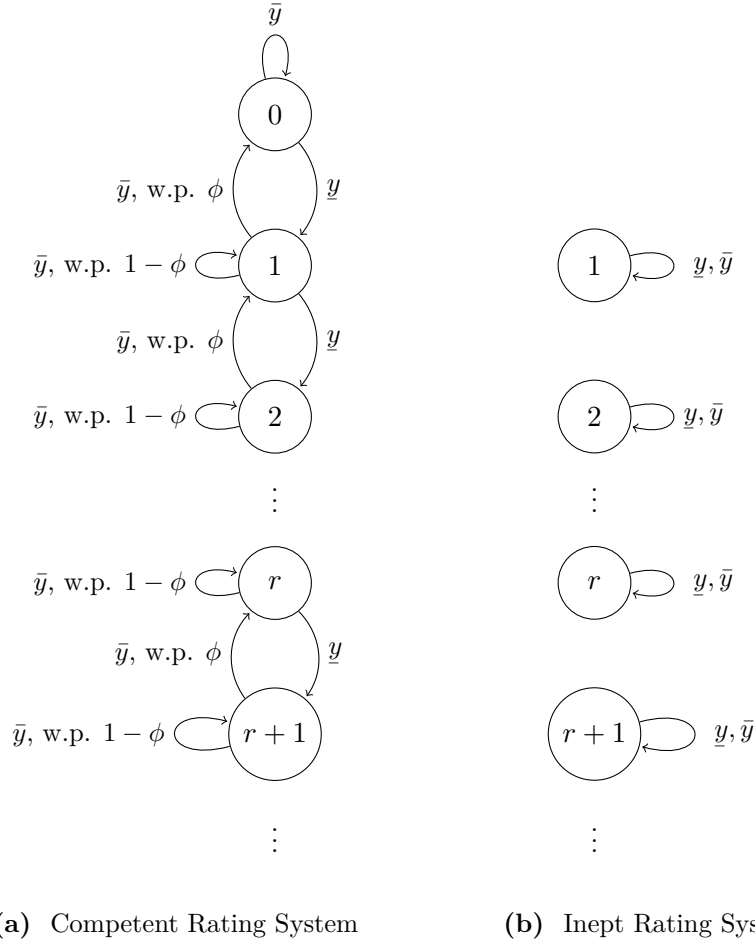


Figure 5: Rating Systems in the Menu Ξ^∞

In the competent system, conditional on the present rating being 0, a good signal $\bar{y}$ leads to rating 0 in the next period, while a bad signal $\underline{y}$ leads to rating 1. Conditional on the present rating being $r > 0$, a good signal $\bar{y}$ leads to rating $r - 1$ with probability ϕ , to r with probability $1 - \phi$, while a bad signal $\underline{y}$ leads to $r + 1$ for sure. The initial distributions λ^C and λ^I over ratings in each scheme are specified formally in Appendix A.8.

The firm participates in the scheme for its type with probability one:

$$\pi^\infty(\xi_\theta^\infty|\theta) = 1, \quad \text{for each } \theta = C, I, \quad (27)$$

The competent firm always exerts high effort upon participating in ξ_C^∞ :

$$\tau^\infty(h_f, C) = \begin{cases} 1, & \text{if } d = \xi_C^\infty, \\ 0, & \text{otherwise.} \end{cases}$$

By construction, consumer posterior of the firm being competent is one upon observing a rating 0 in the equilibrium, because the rating comes from a competent system S_C^∞ . The initial distribution $\lambda^C = (\lambda_r^C)_{r=0}^\infty$ over ratings in the competent system is the stationary distribution induced by the transitions of the system in the equilibrium. The initial distribution $\lambda^I = (\lambda_r^I)_{r=1}^\infty$ over ratings in the inept system S_I^∞ is chosen such that in the equilibrium, consumer posterior of the firm being competent upon observing a positive rating $r > 0$ in any period t equals

$$\underline{\mu} := \varphi_t(r) = \frac{\mu\lambda_r^C}{\mu\lambda_r^C + (1-\mu)\lambda_r^I} = \begin{cases} \frac{\rho}{1-\mu(1-\rho)}, & \text{if } \mu \in (0, \frac{1}{2}], \\ \frac{\mu\rho}{1-\mu}, & \text{if } \mu \in (\frac{1}{2}, 1). \end{cases} \quad (28)$$

Again, by (25), $\underline{\mu} \in (0, 1)$. The time subscript is omitted as the beliefs are stationary. All but rating 0 are opaque, maintaining market uncertainty over time. The distributions λ^C and λ^I are formally specified in Appendix A.8.

Proposition 9 below shows that full extraction obtains.

Proposition 9 (Limit-payoff characterization). *Suppose that (25) holds, and that $c \leq (1-2\rho)^2$. Then $(\sigma^\infty, \varphi^\infty) \in B(\Xi^\infty)$, and the menu Ξ^∞ achieves full extraction in $(\sigma^\infty, \varphi^\infty)$.*

The fact that the menu Ξ^∞ is fully-extracting in the equilibrium $(\sigma^\infty, \varphi^\infty)$ follows directly from Proposition 1, because the relevant conditions are satisfied.²⁴ We discuss why σ^∞ is an equilibrium strategy.

Because ratings are stationary (in the sense that each rating induces identical consumer posterior over time) and there are two induced beliefs, 1 and $\underline{\mu}$, there are only two possible payments in any period t :

$$p_t^{\sigma^\infty}(r) = \begin{cases} \bar{p} := 1 - \rho, & \text{if } r = 0, \\ \underline{p} := \rho + (1-2\rho)\underline{\mu}, & \text{if } r > 0. \end{cases} \quad (29)$$

²⁴Strictly speaking, Proposition 1 applies to settings for fixed $\delta \in (0, 1)$. It is straightforward to see from its proof that it remains true in the present setting with limit payoff.

Note that $\bar{p} > \underline{p}$. The belief μ obtains when the consumer infers that the participating firm is either inept, or is competent and potentially generated countably many bad signals. Regardless of a competent firm's track record, once it generates a bad signal, it earns the low payment $\underline{p}$ in the next period. Consistent high effort follows from the competent firm's incentive to "climb up the ladder" for the payment $\bar{p}$ and avoid $\underline{p}$.

When ρ is relatively small, signals are less noisy, making it relatively easy for the competent firm to climb up and to remain at the "good" rating 0 with payment $\bar{p}$ via consistent high effort. Moreover, shirking likely leads to the "bad" ratings $r > 0$ with the low payment $\underline{p}$. Moreover, for small ρ , $\underline{p}$ differs significantly from $\bar{p}$.

As ρ increases, signals are noisier, and $\underline{p}$ approaches $\bar{p}$. It also becomes harder to remain at rating 0, weakening effort incentives. Nonetheless, it also becomes harder to climb up, lengthening the punishment phase associated with the payment $\underline{p}$, compensating for the weakening of incentives.

When $\mu \leq \frac{1}{2}$, consumers are skeptical that the firm is competent. Observing a positive rating reinforces this prior belief, yielding a small $\underline{p}$. There is a sufficiently large wedge between $\bar{p}$ and $\underline{p}$, making pure transitions between ratings sufficient for motivating consistent high effort for $c \leq (1 - 2\rho)^2$.

When $\mu > \frac{1}{2}$, consumers are relatively confident that the firm is competent. Positive ratings have less influence on consumer posteriors, and $\underline{p}$ remains relatively large. The resulting small wedge between $\bar{p}$ and $\underline{p}$ is insufficient to motivate consistent high effort by pure rating transitions. The competent rating system therefore features mixed transitions, in which the probability ϕ of climbing up one rating (26) strictly decreases in μ , effectively lengthening the punishment phase. When monitoring becomes so noisy that (25) is violated, there is no probability ϕ that can compensate for this weakening of incentives to support consistent high effort.

The limit-payoff assumption is critical. Upon receiving a sufficiently large rating, it takes a large number of periods for the firm to climb up and obtain the maximal payment $\bar{p}$. With a fixed $\delta \in (0, 1)$, this future benefit may be largely discounted, disrupting incentives for high effort.

Finally, when verifying that the participation strategy π^∞ is an equilibrium phenomenon, we need to compute each type- θ firm's limit profit extraction from consumers upon choosing a participation decision d , *i.e.*, $\lim_{\delta \uparrow 1} U(\sigma^\infty; \theta, d)/(1 - \delta)$. Indeed, subject to a boundedness condition on the continuation profits verified in Appendix A.9, the limit profit extraction

equals the long-run expected average profit per period (Ross, 2014), satisfying

$$\liminf_{\delta \uparrow 1} \frac{U(\sigma^\infty, \theta; d)}{1 - \delta} = \liminf_{T \uparrow \infty} \frac{\mathbf{E}^P[\sum_{t=0}^T p_t^{\sigma^\infty}(r_t) - c|\theta, d]}{T + 1}, \quad (30)$$

It is then analogous to Proposition 2 to show that π^∞ constitutes an equilibrium.

Remark 11 (Contrast to standard intermediation results). The menu Ξ^∞ features *countably many* stationary ratings, making it clear that fully-extracting ratings need not be coarse but must be opaque, contrary to standard results in the intermediation literature (see Remark 9). This contrast is due to the non-trivial interaction between adverse selection and moral hazard in consumer payments in the dynamic environment. An increasing number of ratings allows the construction of a more effective punishment phase to sustain consistent high effort.

5.3 Towards a Complete Characterization

For fixed δ , full extraction obtains for $c \leq \bar{c}$ (Proposition 2). In this section, I discuss the rater's optimal behavior for any c satisfying (1) and $c > \bar{c}$. Because the analysis inevitably requires keeping track of several cost cutoffs, Table 1 below collects all the critical cutoffs in ascending order and summarizes their significances for the reader's convenience.

5.3.1 Prohibitively High Costs

First, when c is too high, the rater extracts no positive surplus because she fails to motivate high effort in any period.

Proposition 10 (Zero extraction for prohibitively high costs). *Suppose that $c \in (\bar{c}_{MH}, 1 - 2\rho)$, where $\bar{c}_{MH}$ is defined in (20). The competent firm never exerts high effort in equilibrium. The rater's optimal payoff is 0.*

5.3.2 Intermediate Costs

Consider the intermediate range of costs $(\bar{c}, \bar{c}_{MH}]$. One may conjecture that if the rater can incentivize consistent high effort using some menu, then she can do so using two stationary ratings, thereby creating bang-bang incentives via two payments. This is, however, not true by Proposition 8. I now present two examples in which the rater achieves full extraction for a larger range of costs by using more than two ratings. The examples illustrate the benefit

Cost cutoffs given (μ, δ, ρ)	Significance
$\bar{c}_1$, defined by (14)	Ξ^{**} is fully-extracting (Proposition 3)
$\bar{c}$, defined by (9)	Ξ^* is fully-extracting (Proposition 2); necessary for full extraction with binary stationary ratings (Proposition 11)
$\delta(1 - 2\rho)^2$	necessary condition for full extraction (Proposition 7); more ratings weakly improve rater's ability for full extraction (Examples 1–3); impossibility result (Proposition 8)
$\bar{c}_{MH}$, defined by (20)	necessary condition for positive extraction (Proposition 10); Ξ_{MH} is fully-extracting given pure moral hazard (Proposition 5)
$1 - 2\rho$	high effort is efficient, see (1)
Cost cutoff in limit setting given (μ, ρ) and (25)	Significance
$(1 - 2\rho)^2$	Ξ^∞ is fully-extracting (Proposition 9)

Table 1: Cost cutoffs

The top table collects the cost cutoffs considered given (μ, δ, ρ) in ascending order, together with their significances. The bottom table presents the the cost cutoff considered in the limit-payoff setting given (μ, ρ) and (25), together with its significance.

of increasing the number of ratings when the market faces a more severe moral hazard problem, in line with Remark 9. The examples, together with Proposition 8, suggest that a tight threshold on c above which full extraction is not achievable is in general a non-trivial function of (μ, δ, ρ) . Deriving this threshold poses tractability concerns, as it would require, for each (μ, δ, ρ) , characterizing a menu that maximally relaxes all incentive constraints for high effort in the induced game among all menus.

Proposition 11 first shows that if the rater is restricted to use at most two ratings in a menu, and the ratings are stationary, full extraction is not achievable for any $c > \bar{c}$:

Proposition 11 (Full-extraction cost upper bound for two ratings). *Suppose that each menu contains at most two ratings, $|R_C \cup R_I| \leq 2$, and ratings are stationary in equilibrium. Then for every $c > \bar{c}$, the rater's payoff is strictly below $\mu(1 - 2\rho - c)$.*

The rater, however, can possibly achieve full extraction for some $c > \bar{c}$ using three stationary ratings. In the examples below, I focus on the competent firm's incentive constraints for high effort faced upon participating in the competent scheme. Analogous to the construction of Ξ^* and (σ^*, φ^*) , it is straightforward to set the fees in the respective schemes and construct an equilibrium (σ, φ) in the induced game in which the firm chooses the scheme intended for its type and the rater achieves full extraction, provided consistent high effort by the competent firm.

Example 1 (Three ratings with a ladder structure). Fix $\delta = \frac{9}{10}$, $\rho = \frac{1}{10}$ and $\mu = \frac{9}{10}$, so that $\bar{c} \approx 0.303$. Consider a menu $\Xi = \{\xi_C, \xi_I\}$, in which the competent scheme contains 3 ratings, labeled 0, 1 and 2. The inept scheme contains the ratings 0 and 1. The competent rating system, with transitions depicted below in Figure 6, has an initial distribution over the ratings $\{0, 1, 2\}$ equal to the stationary distribution $(\frac{1}{91}, \frac{9}{91}, \frac{81}{91})$.

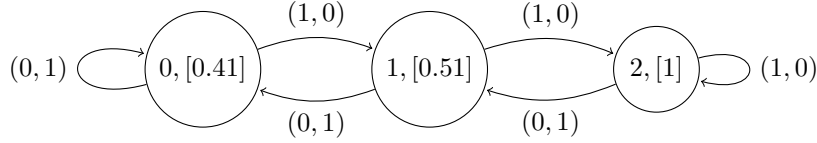


Figure 6: Transitions in the Competent System

Each node is labeled by the rating, followed by the corresponding equilibrium consumer belief in squared brackets. The transition probabilities are written in the form (x, y) and are rounded up to two decimal places, where x denotes the transition upon a good signal and y denotes the transition upon a bad signal.

In the inept system depicted by Figure 7 below, the initial distribution over the ratings $\{0, 1\}$ is $(0.14, 0.86)$. The ratings are absorbing so that, once drawn, the rating remains the same forever.

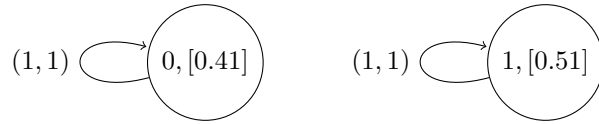


Figure 7: Transitions in the Inept System

The continuation profit V_r^σ upon each rating $r = 0, 1, 2$ by the competent firm in the equilibrium (that features consistent high effort) satisfies the

following system of Bellman equations:

$$V_r^\sigma = (1 - \delta)(p^\sigma(r) - c) + \delta((1 - \rho)V_{\min(r+1,2)}^\sigma + \rho V_{\max(r-1,0)}^\sigma), \quad r = 0, 1, 2.$$

I omit the time subscripts in the payments given stationarity in consumer beliefs. The incentive constraint for high effort upon each rating r is

$$\frac{\delta}{1 - \delta}(1 - 2\rho)(V_{\min(r+1,2)}^\sigma - V_{\max(r-1,0)}^\sigma) \geq c.$$

The maximal cost below which consistent high effort holds is therefore

$$\frac{\delta(1 - 2\rho)}{1 - \delta} \min(V_1^\sigma - V_0^\sigma, V_2^\sigma - V_0^\sigma, V_2^\sigma - V_1^\sigma) \approx 0.31 > \bar{c}.$$

◆

In fact, the rater can achieve full extraction using three stationary ratings for some costs higher than 0.31 in the above example, if she dispenses with the ladder structure in the competent system. Example 2 below illustrates.

Example 2 (Three ratings with mixed transitions). Keep $\delta = \frac{9}{10}$, $\rho = \frac{1}{10}$ and $\mu = \frac{9}{10}$. Consider a menu $\Xi = \{\xi_C, \xi_I\}$, in which the competent scheme again contains ratings 0, 1 and 2, and the inept scheme contains ratings 0 and 1. Transitions in the competent system are depicted in Figure 8 below, with the initial distribution over $\{0, 1, 2\}$ being the stationary distribution (0.01, 0.8, 0.19). Transitions in the inept system are depicted in Figure 9 below, with the initial distribution over $\{0, 1\}$ being (0.21, 0.79).

Let α_{rs} denote the transition probability from rating r to rating s upon a good signal in the competent system, and let β_{rs} denote the counterpart upon a bad signal. The continuation profit V_r^σ upon each rating $r = 0, 1, 2$ by the competent firm satisfies

$$V_r^\sigma = (1 - \delta)(p^\sigma(r) - c) + \delta \sum_{s=0}^2 ((1 - \rho)\alpha_{rs} + \rho\beta_{rs})V_s^\sigma.$$

Using the fact that $\alpha_{r2} = 1 - \alpha_{r0} - \alpha_{r1}$ and $\beta_{r2} = 1 - \beta_{r0} - \beta_{r1}$, the incentive constraints for high effort by the competent firm upon rating r is

$$\frac{\delta(1 - 2\rho)}{1 - \delta} \left[\sum_{s=0,1} (\alpha_{rs} - \beta_{rs})(V_s^\sigma - V_2^\sigma) \right] \geq c.$$

The maximal cost below which consistent high effort holds is therefore

$$\min_{r \in \{0,1,2\}} \frac{\delta(1 - 2\rho)}{1 - \delta} \left[\sum_{s=0,1} (\alpha_{rs} - \beta_{rs})(V_s^\sigma - V_2^\sigma) \right] \approx 0.35 > \bar{c}.$$

◆

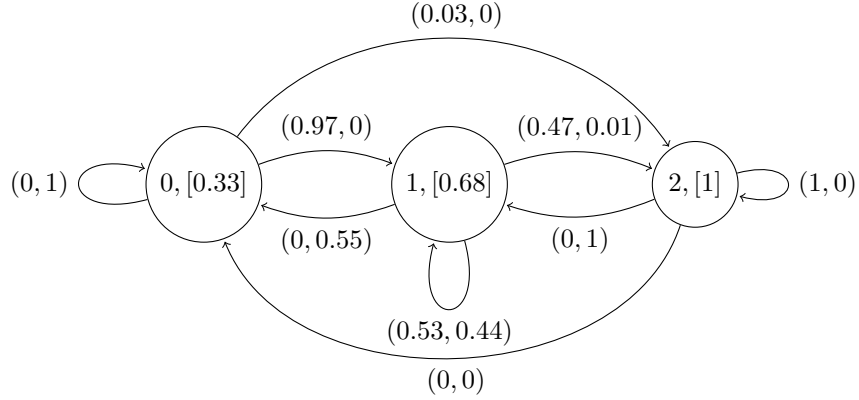


Figure 8: Transitions in the Competent System

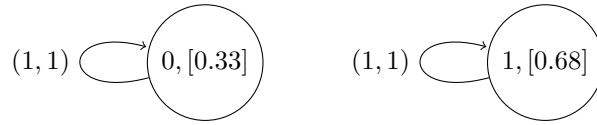


Figure 9: Transitions in the Inept System

To see why more ratings may relax the cost threshold, it is instructive to compare the transitions in Figure 8 to the transitions induced by the competent scheme ξ_C^* in the menu Ξ^* , depicted in Figure 10 below.

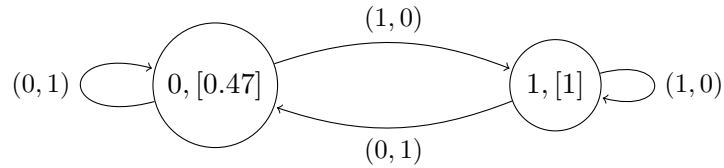


Figure 10: Transitions in the System S_C^*

In line with Remark 9, the use of more ratings allows the rater to “push” consumer posteriors towards 0 by statistically revealing more bad signals in a firm’s track record, lowering consumer payments associated with “bad” ratings to create stronger effort incentives. The downside is the inevitable introduction of additional incentive constraints for high effort upon observing

the additional ratings. Ratings that induce intermediate beliefs may adversely create a “cushion” that weakens effort incentives. Of course, the rater can always use a smaller amount of ratings than the available ones by not allowing some ratings to be realized with positive probability. The preceding two examples show that full extraction can be achieved for some $c > \bar{c}$ when μ is relatively large so that the spread of posterior beliefs is valuable. The following example shows that this is not the case when μ is relatively low so that the extent to which ratings can depress consumer posteriors is limited.

Example 3 (Three ratings with mixed transitions and low prior).

Keep $\delta = \frac{9}{10}$ and $\rho = \frac{1}{10}$, but suppose instead that $\mu = \frac{1}{100}$, so that $\bar{c} \approx 0.575$. Keeping the structure of the above competent scheme and inept schemes unchanged but choosing new transition probabilities $(\alpha_{rs}, \beta_{rs})_{r,s=0,1,2}$ and new initial distribution over inept ratings $(\lambda_0^I, \lambda_1^I)$ (and fixing the initial distribution over competent ratings to be the stationary distribution given the new transitions), the maximal cost below which consistent high effort holds is

$$\max_{(\alpha_{rs}, \beta_{rs})_{r,s=0,1,2}, (\lambda_r^I)_{r=0,1}} \min_{r \in \{0,1,2\}} \frac{\delta(1-2\rho)}{1-\delta} \left[\sum_{s=0,1} (\alpha_{rs} - \beta_{rs})(V_s^\sigma - V_2^\sigma) \right] \approx 0.575,$$

which does not improve over $\bar{c}$. ◆

Similar tractability difficulties arise when deriving the rater’s optimal payoff given parameters such that full extraction is unachievable, which requires keeping track of histories after which the rater would like the firm to exert high effort and those after which the rater would like the firm to shirk in the game induced by each menu, and optimizing over all menus.

5.3.3 Summary

To summarize the discussion, the rater’s optimal payoff $W(c)$ as a function of c is also illustrated below in Figure 11.

5.3.4 Worst-Case Scenario: Fully Solving Adverse Selection

At any rate, the bottom line is that the rater can always fully solve the adverse selection problem in period 0 by implementing the menu Ξ_{MH} , inducing a continuation game with pure moral hazard and guaranteeing herself a positive payoff. Specifically:

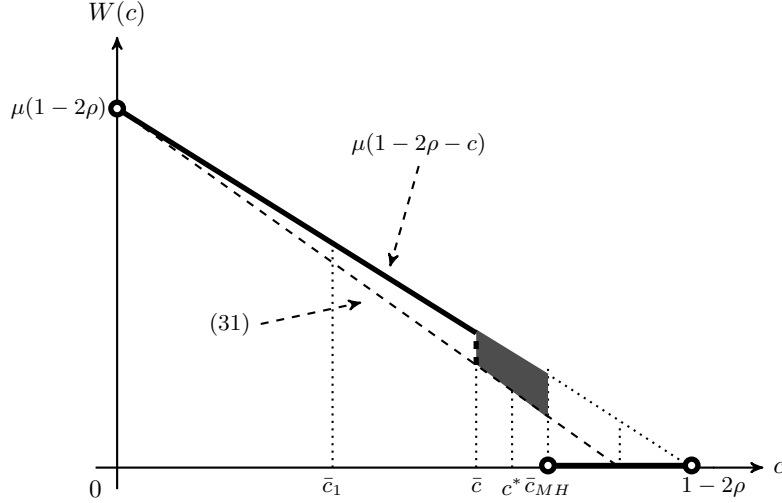


Figure 11: Rater's Optimal Expected Payoff

When $c \leq \bar{c}$, the rater achieves full extraction using the menu Ξ^* in the equilibrium σ^* , as shown by Proposition 2. The dashed line represents a lower bound on the rater's payoff that she can guarantee for $c \leq \bar{c}_{MH}$ by fully solving the adverse selection problem in the beginning of the game (see Proposition 12). Let $c^* = \delta(1 - 2\rho)^2$. When $c \in (\bar{c}, c^*)$, full extraction is in general impossible: the ability of doing so relies on the parameters (μ, δ, ρ) (Proposition 8 and Examples 1–3 in Section 5.3.2). When $c \in [c^*, \bar{c}_{MH}]$, full extraction is impossible (Proposition 7), and the rater's payoff is strictly below $\mu(1 - 2\rho - c)$. When $c > \bar{c}_{MH}$, the rater extracts nothing from the market. Observe that the fully-extracting payoff and the lower bound (31) converges as monitoring becomes increasingly accurate, *i.e.* as ρ falls. The intermediate range vanishes as $\rho \downarrow 0$. At $\rho = 0$ so that monitoring is perfect, a complete characterization of the rater's payoff obtains, see Appendix A.10.

Proposition 12 (Lower bound). *The rater's optimal payoff is at least*

$$\mu \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right). \quad (31)$$

for $c \in (0, \bar{c}_{MH}]$.

6 Final Remarks

The preceding analysis has some features that invite further discussion. The rater faces no restriction on the amount of past data to utilize. This assumption is appealing for the question of interest, given the rapidly growing

data storage technologies in modern information systems. Yet, it also makes the stark simplicity of the menu Ξ^* rather surprising. The system S_C^* , defined by (10), requires one-period memory while the system S_I^* , defined by (11), is memoryless. Moreover, both menus Ξ^* and Ξ^∞ induce a stationary environment, facilitating a tractable analysis.

Such simplicity relies on the assumption that consumers have limited information. If consumers are infinitely-lived or if the entering short-lived consumers observe all previous ratings, they may learn purely from the signals or the ratings and allow the competent firm to eventually rest on its laurels. Once an equilibrium calls for the competent type to shirk after some history, full extraction is no longer possible. Nonetheless, it does not overturn the fundamental insight that revenue maximization by the rater calls for maximizing the firm's expected value to consumers, though the rater's optimal payoff might now become much harder to keep track of.

Appendices

A Omitted Details

A.1 The Rater Needs At Most Two Schemes

The assumption of a binary menu is without loss of generality if we allow the rater to offer lotteries. Consider a setting in which the rater can offer a menu of countably many schemes. In the game induced by this menu, consider an equilibrium σ in which type- θ firm chooses $\pi(\xi|\theta) > 0$ for $\xi \in \Xi_\theta = \{\xi_1, \xi_2, \dots\} \subseteq \Xi$, where $\xi_i = (f_i, R_i, S_i)$ for each i .

Consider now a menu of lotteries $\Xi = \{\xi_C, \xi_I\}$ that induces another continuation game. Here, each ξ_θ is a *lottery* over Ξ_θ , choosing each scheme ξ with probability $\pi(\xi|\theta)$. Both the rater and the firm observes the realization. Each lottery ξ_θ can therefore be represented as a scheme $(f_\theta, R_\theta, S_\theta)$, where

$$f_\theta = \sum_{i: \xi_i \in \Xi_\theta} \pi(\xi_i|\theta) f_i, \quad R_\theta = \bigcup_{i: \xi_i \in \Xi_\theta} R_i.$$

and the system S_θ , conditional on the realization ξ_i , equals S_i .

Consider now a firm's effort strategy σ' which, conditional on the realization ξ_i is identical to σ conditional on choosing ξ_i from Ξ_θ in the original game. They are also identical conditional on choosing $d = N$. The resulting play and information flow are therefore identical given σ and σ' in both

games, so are consumer beliefs over time. Conditional on the participation decision, the effort strategy in σ' constitutes an equilibrium. The firm's expected profits given σ and σ' are identical in both settings.

Because the firm mixes its participation choices in the original game, it must be indifferent between each of these schemes and the participation constraints for each of these schemes are satisfied:

$$\sum_{i:\xi_i \in \Xi_\theta} \pi(\xi_i|\theta)(U(\sigma, \theta; \xi_i) - (1 - \delta)f_i) = U(\sigma, \theta; \xi_i) - (1 - \delta)f_i \geq U(\sigma, \theta; N).$$

By construction, the participation constraint is satisfied in the new game:

$$\begin{aligned} U(\sigma', \theta; \xi_\theta) - (1 - \delta)f_\theta &= \sum_{i:\xi_i \in \Xi_\theta} \pi(\xi_i|\theta)[U(\sigma, \theta; \xi_i) - (1 - \delta)f_i] \\ &\geq U(\sigma, \theta; N) = U(\sigma', \theta; N), \end{aligned}$$

As a result, choosing ξ_θ with probability one in the new game constitutes an equilibrium. Because the rater's payoff is linear in the rating fee she collects, her expected payoffs are the same in both σ and σ' .

Remark 12 (Observable realization). If the firm does not observe the scheme ξ_i drawn, then the information structure in the new game is not identical to the one in the original game. The firm cannot condition its future play on the scheme ξ_i , which is, however, feasible in the original game. Nonetheless, all the results in the characterizations continue to hold if the rater is assumed to offer a menu of countably many rating schemes.

A.2 Equilibrium Existence

Lemma 1. *For any menu Ξ , the set $B(\Xi)$ is non-empty.*

Proof. Fix an arbitrary menu Ξ . In the game induced by Ξ , the profile (σ, φ) where $\pi(C) = \pi(I) = 0$ and $\tau(h_2, C) = \tau(h_2, I) = 0$ for each history h_2 , and in which consumers form beliefs $\varphi_t(r) = 0$ upon seeing any non-null rating r in any period t , is an equilibrium. In the equilibrium, consumers always expect low effort from the firm because consumers do not observe past play. Each type of firm receives a payoff ρ . By deviating to participate in some scheme, two possibilities arise. Upon a null rating, consumers do not detect the deviation and still pay the firm ρ . Upon observing any non-null rating, consumers believe the firm is inept for sure and pay the firm ρ . Moreover, by participating, the firm pays the rater at least 0 upfront. Its payoff following a deviation is therefore at most ρ . ■

A.3 Fully-Revealing Rating System

The fully-revealing rating system $S^{FR} : H_1 \rightarrow \Delta(R^{FR})$ is constructed as follows. Because the set of signal histories in each period t is countable, and the union of countably many countable sets are countable, one can construct the rating set R^{FR} that is countable by first partitioning the set H_1^t for each period t into equivalence classes, each of which contains histories h_r^t with identical signal histories, and second labeling each equivalence class with a distinct element in $\mathbb{R}$, letting R_t^{FR} be a set that collects these labels, and finally letting $R^{FR} := \bigcup_{t=0}^{\infty} R_t^{FR}$. The system S^{FR} is a function that maps each history h_r , which contains a particular signal history (in addition to the history of ratings), to the label of the equivalence class identified by that signal history in R^{FR} with probability one.

A.4 The Menu Ξ^*

It remains to specify the rating fees, given by

$$f_C^* = \frac{1}{1-\delta} \left[\mu(1-2\rho) + \frac{\delta(1-\mu)^2(1-\rho)(1-2\rho)}{1-\mu(1-\rho)} - c \right], \quad (32)$$

$$f_I^* = \frac{1}{1-\delta} \left[(1-\delta)(\rho + \mu(1-2\rho)) + \frac{\delta\rho(1-\mu\rho)}{1-\mu(1-\rho)} - \rho \right]. \quad (33)$$

A.5 The Menu Ξ^{**}

It remains to specify the transition probability α and the fee f^{**} . The transition probability is given by

$$\alpha := \frac{c(\rho + \mu(1-2\rho))}{c(\rho + \mu(1-2\rho))^2 + \delta(1-\mu)\mu(1-2\rho)^3}, \quad (34)$$

which belongs to $(0, 1]$ when Condition (14) holds. The fee is

$$f^{**} = \frac{1-2\rho}{1-\delta} \left[\mu - \frac{\delta\alpha(1-\mu)\varphi^2(1-2\rho)^2}{(\rho + \mu(1-2\rho))(1-\alpha(\rho + \mu(1-2\rho)))} \right]. \quad (35)$$

A.6 The Menu Ξ_{MH}

The menu $\Xi_{MH} = \{\xi_{MH}\}$, where $\xi_{MH} = (f_{MH}, R_{MH}, S_{MH})$, is defined as follows. The rating set $R_{MH} = \{0, 1\}$ is binary, and the rating system S_{MH} is defined piecewise depending on the value of c . If

$$c \leq \hat{c} := \frac{\delta(1-2\rho)^2}{1 + \delta(1-2\rho)}, \quad (36)$$

then

$$S_{MH}(h_r^t) = \begin{cases} (1 - \beta) \circ \{1\} + \beta \circ \{0\}, & \text{if } y_{t-1} = \underline{y} \text{ and } r_{t-1} = 1, \\ 1 \circ \{1\}, & \text{otherwise,} \end{cases}$$

where

$$\beta := \frac{c}{\delta((1 - 2\rho)^2 - c(1 - \rho))}. \quad (37)$$

In words, the rating system begins with rating 1, announces rating 0 with probability β if the previous rating is 1 and a bad signal takes place, and announces rating 1 otherwise. If $c > \hat{c}$, then

$$S_{MH}(h_r^t) = \begin{cases} 1 \circ \{1\}, & \text{if } y_{t-1} = \bar{y} \text{ and } r_{t-1} = 1, \\ \kappa \circ \{1\} + (1 - \kappa) \circ \{0\}, & \text{if } y_{t-1} = \underline{y} \text{ and } r_{t-1} = 0, \\ 1 \circ \{0\}, & \text{otherwise,} \end{cases}$$

where

$$\kappa := \frac{\delta(1 - 2\rho)^2 - c(1 - \delta\rho)}{c\delta(1 - \rho)}. \quad (38)$$

In words, the system announces rating 1 with probability κ if only good signals occur in the past, and announces rating 0 otherwise. The fee f_{MH} is

$$f_{MH} = \frac{1}{1 - \delta} \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right).$$

The probabilities β and κ are well-defined given (20). The corresponding rater-preferred equilibrium is $(\sigma_{MH}, \varphi_{MH})$, where $\sigma_{MH} = (\pi_{MH}, \tau_{MH})$ satisfies

$$\begin{aligned} \pi_{MH}(\xi_{MH}|C) &= 1, \\ \tau_{MH}(h_f, C) &= \begin{cases} 1, & \text{if } r_t = 1 \text{ and } d = \xi_{MH}, \\ 0, & \text{otherwise,} \end{cases} \end{aligned}$$

so that the competent type participates for sure. It chooses high effort with probability one following a rating 1 upon participating in the menu. Otherwise, it exerts low effort with probability one. The belief system φ_{MH} is trivial in a pure moral hazard setting.

A.7 Competent Firm's Maximal Profit

Let $V^\sigma(h_f)$ denote the competent firm's continuation profit in some equilibrium σ after a history h_f , and let $V^\sigma(y; h_f)$ denote the expected continuation profit upon a signal realization y in the game after history h_f . Exerting high effort, the continuation profit is

$$V^\sigma(h_f^t) = (1 - \delta)(p_t^\sigma - c) + \delta((1 - \rho)V^\sigma(\bar{y}; h_f^t) + \rho V^\sigma(\underline{y}; h_f^t)).$$

Shirking gives

$$V^\sigma(h_f^t; \underline{e}) = (1 - \delta)p_t^\sigma + \delta(\rho V^\sigma(\bar{y}; h_f^t) + (1 - \rho)V^\sigma(\underline{y}; h_f^t)).$$

The incentive constraint for high effort after each history h_f^t is $V^\sigma(h_f^t) - V^\sigma(h_f^t; \underline{e}) \geq 0$, which can be written as

$$V^\sigma(\bar{y}; h_f^t) - V^\sigma(\underline{y}; h_f^t) \geq \frac{(1 - \delta)c}{\delta(1 - 2\rho)}. \quad (39)$$

For sufficiently small c , the competent firm's equilibrium expected profit is

$$\begin{aligned} V^\sigma(h_f^0) &\leq (1 - \delta)(\rho + \mu(1 - 2\rho) - c) + \delta((1 - \rho)V^\sigma(\bar{y}; h_f^0) + \rho V^\sigma(\underline{y}; h_f^0)) \\ &\leq (1 - \delta)(\rho + \mu(1 - 2\rho) - c) + \delta\left(V^\sigma(\bar{y}; h_f^0) - \frac{(1 - \delta)\rho c}{\delta(1 - 2\rho)}\right) \\ &\leq (1 - \delta)(\rho + \mu(1 - 2\rho) - c) + \delta\left((1 - \delta)\left(\rho + \frac{\mu(1 - \rho)(1 - 2\rho)}{\mu(1 - \rho) + (1 - \mu)\rho} - c\right.\right. \\ &\quad \left.\left.+ \delta((1 - \rho)V^\sigma(\bar{y}; h_f^1) + \rho V^\sigma(\underline{y}; h_f^1)) - \frac{(1 - \delta)\rho c}{\delta(1 - 2\rho)}\right)\right) \\ &\leq \dots \\ &\leq (1 - \delta)\left(\sum_{t=0}^{\infty} \delta^t \rho + \sum_{t=0}^{\infty} \delta^t \frac{(1 - 2\rho)(1 - \rho)^t \mu}{\mu(1 - \rho)^t + (1 - \mu)\rho^t} - \sum_{t=0}^{\infty} \delta^t c - \sum_{t=0}^{\infty} \delta^t \frac{\rho c}{1 - 2\rho}\right) \\ &= \rho + (1 - \delta)(1 - 2\rho) \sum_{t=0}^{\infty} \delta^t \frac{(1 - \rho)^t \mu}{\mu(1 - \rho)^t + (1 - \mu)\rho^t} - c - \frac{\rho c}{1 - 2\rho}, \end{aligned}$$

as desired.

A.8 The Menu Ξ^∞

It remains to formally specify the rating fees and the rating systems. The fees f_C^∞ and f_I^∞ are

$$f_C^\infty = \begin{cases} \frac{1}{1-\delta} \left[1 - \frac{\rho(2(1-\mu(1-2\rho)) - 3\rho)}{1-\rho-\mu(1-2\rho)} - \rho - c \right], & \text{if } \mu \in (0, \frac{1}{2}), \\ \frac{1-2(1-\rho)\rho - \rho - c}{1-\delta}, & \text{if } \mu \in (\frac{1}{2}, 1). \end{cases}$$

$$f_I^\infty = \begin{cases} \frac{\mu(1-2\rho)\rho}{(1-\delta)(1-\rho-\mu(1-2\rho))}, & \text{if } \mu \in (0, \frac{1}{2}), \\ \frac{\mu\rho(1-2\rho)}{(1-\delta)(1-\mu)}, & \text{if } \mu \in (\frac{1}{2}, 1). \end{cases}$$

The rating systems $S_C^\infty : H_r \rightarrow \Delta(R_C^\infty)$ and $S_I^\infty : H_r \rightarrow \Delta(R_I^\infty)$ follow

$$S_C^\infty(h_r^t) = \begin{cases} \sum_{r=0}^\infty \lambda_r^C \circ \{r\}, & \text{if } t = 0, \\ 1 \circ \{0\}, & \text{if } r_{t-1} = 0 \text{ and } y_{t-1} = \bar{y}, \\ \phi \circ \{r_{t-1} - 1\} + (1 - \phi) \circ \{r_{t-1}\}, & \text{if } r_{t-1} > 0 \text{ and } y_{t-1} = \bar{y}, \\ 1 \circ \{r_{t-1} + 1\}, & \text{if } y_{t-1} = \underline{y}. \end{cases} \quad (40)$$

$$S_I^\infty(h_r^t) = \begin{cases} \sum_{r=0}^\infty \lambda_r^I \circ \{r\}, & \text{if } t = 0, \\ 1 \circ \{r_{t-1}\}, & \text{if } t > 0, \end{cases} \quad (41)$$

where the distributions $\lambda^C = (\lambda_r^C)_{r=0}^\infty$ and $\lambda^I = (\lambda_r^I)_{r=0}^\infty$ are defined as follows. The initial distribution λ^C in the competent system is the stationary distribution induced by the transitions in the equilibrium σ^∞ :

$$\lambda_r^C = \begin{cases} \left(\frac{1-2\rho}{1-\rho} \right) \left(\frac{\rho}{1-\rho} \right)^r, & \text{if } \mu \in (0, \frac{1}{2}] \\ \left(\frac{1-\mu-\rho}{1-\mu-\mu\rho} \right) \left(\frac{(1-\mu)\rho}{1-\mu-\mu\rho} \right)^r, & \text{if } \mu \in (\frac{1}{2}, 1) \end{cases} \quad (42)$$

for each rating $r \in R_C^\infty$. The initial distribution over ratings λ^I in the inept system satisfies

$$\lambda_r^I = \begin{cases} \left(\frac{1-2\rho}{1-\rho} \right) \left(\frac{\rho}{1-\rho} \right)^{r-1}, & \text{if } \mu \in (0, \frac{1}{2}] \\ \frac{1-\mu-\rho}{(1-\mu)\rho} \left(\frac{(1-\mu)\rho}{1-\mu-\mu\rho} \right)^r, & \text{if } \mu \in (\frac{1}{2}, 1) \end{cases} \quad (43)$$

for each rating $r \in R_I^\infty$.

A.9 Convergence to the Average Payoff Criterion

Let $V_r^{\sigma^\infty}(\xi; \theta)$ denote a type- θ firm's continuation profit upon participating in ξ_C^∞ and receiving a rating r , given the strategy profile σ^∞ . The regularity condition for (30) to hold is that, there exists $K < \infty$ such that

$$\left| \frac{V_r^{\sigma^\infty}(\xi; \theta) - V_{r'}^{\sigma^\infty}(\xi; \theta)}{1 - \delta} \right| < K, \quad (44)$$

for some fixed r' that can be realized upon a participation decision $d = \xi$ and all r that can be realized upon a participation decision $d = \xi$, and for all $\delta \in (0, 1)$ (Theorem 2.2, Ross (2014), Chapter 5). This condition clearly holds for any type θ who chooses $d = N$, because the left side equals 0. Because consumer beliefs and payments are identical upon observing all ratings in ξ_I^∞ , the left side also equals 0 for each type θ who chooses $d = \xi_I^\infty$. We now show that it also holds when a type- θ firm chooses $d = \xi_C^\infty$.

Let $Q_{rs}^{\sigma^\infty}(\xi; \theta)$ be the transition probability from rating r to s in a scheme ξ . Note first that $V_0^{\sigma^\infty} \geq V_r^{\sigma^\infty}$ for each r . This is a straightforward implication of Proposition 3.2 in Ross (2014), Chapter II, given that $p^{\sigma^\infty}(r)$ decreases in r , and $\sum_{s=0}^{k-1} Q_{rs}^{\sigma^\infty}(\xi_C^\infty; \theta)$ for each k decreases in r . It thus follows from the Bellman equation for $V_r^{\sigma^\infty}(\xi_C^\infty; \theta)$ that,

$$\frac{V_r^{\sigma^\infty}(\xi_C^\infty; \theta)}{1 - \delta} \leq p^{\sigma^\infty}(r) - c + \frac{\delta V_0^{\sigma^\infty}(\xi_C^\infty; \theta)}{1 - \delta} < 1 - \rho - c + \frac{V_0^{\sigma^\infty}(\xi_C^\infty; \theta)}{1 - \delta},$$

ensuring (with $r' = 0$ in (44))

$$\left| \frac{V_r^{\sigma^\infty}(\xi_C^\infty; \theta) - V_0^{\sigma^\infty}(\xi_C^\infty; \theta)}{1 - \delta} \right| < 1 - \rho - c =: K < \infty.$$

A.10 Perfect Monitoring

This section characterizes optimal rating menus when $\rho = 0$, *i.e.* when there is perfect monitoring of the firm's effort choice by the rater. A good signal perfectly reveals that the firm has exerted high effort, and a bad signal perfectly reveals that the firm has exerted low effort. The condition for high effort being efficient reduces to $c < 1$.

Proposition 13 (Necessary and Sufficient Conditions). *If $c \leq \delta$, a menu Ξ is optimal if and only if one of the following conditions is true:*

- A. *there exists an equilibrium $(\sigma, \varphi) \in B(\Xi)$ such that*

1. $\pi(\Xi|C) = 1$,
2. $\pi(\Xi|I) > 0$,
3. $\tau(h_f^t, C) = 1$ after every history h_f^t that occurs with positive probability for each $t \geq 0$ conditional on the firm's type being competent,
4. for each type θ , if $\pi(\xi_{\theta'}|\theta) > 0$, then $(1 - \delta)f_{\theta'} = U(\sigma, \theta; \xi_{\theta'}) - \rho$.

B. there exists an equilibrium $(\sigma, \varphi) \in B(\Xi)$ such that

1. $\pi(\Xi|C) = 1$,
2. $\pi(\Xi|I) = 0$,
3. $\tau(h_f^t, C) = 1$ after every history h_f^t that occurs with positive probability for each $t \geq 0$ conditional on the firm's type being competent,
4. if $\pi(\xi_{\theta'}|C) > 0$, then $(1 - \delta)f_{\theta'} = U(\sigma, C; \xi_{\theta'}) - \rho$.

The rater's fully-extracting payoff is $\mu(1 - c)$.

Observe that the cost threshold δ can be obtained by taking the limit $\rho \downarrow 0$ in (9). The conditions in Part A are familiar from Proposition 1. The set of optimal menus expand relative to Proposition 1 because it is now possible to induce consistent high effort in an equilibrium in which only the competent type participates, contrary to the imperfect monitoring case. Specifically, there is a singleton menu with an equilibrium in the induced game satisfying the conditions in Part B. The corresponding rating system assigns “bad” ratings inducing low effort *forever* whenever *one* bad signal is observed, and a “good” rating otherwise. When $c \leq \delta$, the participating competent firm consistently exerts high effort to avoid the perpetual punishment of a low payment associated with low effort.

Moreover, after every history, the equilibrium continuation profit of the competent firm is at most $1 - c$, and it is at least 0. An almost identical proof to Proposition 10 yields:

Proposition 14. *If $c > \delta$, the competent firm never exerts high effort in any equilibrium, and the rater's optimal payoff is 0.*

Thus, putting together Propositions 13 and 14, a complete characterization of the rater's payoff $W^p(c)$ obtains:

$$W^p(c) = \begin{cases} \mu(1 - c), & \text{if } c \in (0, \delta], \\ 0, & \text{if } c \in (\delta, 1). \end{cases}$$

B Proofs

When there is no risk of ambiguity, I write the competent firm's effort strategy $\tau(\cdot, C)$ simply as $\tau(\cdot)$, because an inept type only exerts low effort. Second, for each $i = f, r$ and two histories $h_i, h'_i \in H_i$, the notation $h_i h'_i$ denotes the concatenation of h_i followed by h'_i , and also belongs to the set H_i .

B.1 Proof of Proposition 1

The proof proceeds via a succession of lemmas. Lemma 2 considers the rater's problem in equilibria in which each type participates in the menu with positive probability. It shows that $\mu(1 - 2\rho - c)$ is an upper bound on the rater's payoff in these equilibria. Lemma 3 derives an upper bound on the rater's payoff given any menu and equilibrium in which at most one type participates in a scheme with positive probability. This upper bound is strictly smaller than $\mu(1 - 2\rho - c)$. Lemma 4 argues that the rater achieves $\mu(1 - 2\rho - c)$ if and only if the four conditions in Proposition 1 hold.

To set the stage, for each menu Ξ define

$$B_2(\Xi) := \{(\sigma, \varphi) \in B(\Xi) : \pi(\Xi|C) > 0, \pi(\Xi|I) > 0\} \quad (45)$$

as the set of equilibria in which each type participates in at least one scheme with positive probability. Next, define

$$W_1 := \sup_{\Xi} \sup_{(\sigma, \varphi) \in B(\Xi) \setminus B_2(\Xi)} W(\Xi, \sigma), \quad (46)$$

$$W_2 := \sup_{\Xi: B_2(\Xi) \neq \emptyset} \sup_{(\sigma, \varphi) \in B_2(\Xi)} W(\Xi, \sigma). \quad (47)$$

The payoff W_1 specifies the value of the rater's problem (6), restricting attention to equilibria in which at least one type chooses the outside option with probability one. The payoff W_2 is defined analogously, restricting to equilibria in which each type participates in a scheme with positive probability. Note that $B_2(\Xi)$ may be empty, for example when the fees are set sufficiently high. We first bound W_2 from above. The relevant rater's problem is

$$\sup_{\Xi} \sup_{(\sigma, \varphi) \in B_2(\Xi)} W(\Xi, (\sigma, \varphi))$$

subject to individual rationality (IR) and incentive compatibility (IC):

$$U(\sigma, C; \xi_C) - (1 - \delta)f_C \geq U(\sigma, C; N), \quad (\text{IR}_C)$$

$$U(\sigma, I; \xi_I) - (1 - \delta)f_I \geq U(\sigma, I; N), \quad (\text{IR}_I)$$

$$U(\sigma, C; \xi_C) - (1 - \delta)f_C \geq U(\sigma, C; \xi_I) - (1 - \delta)f_I, \quad (\text{IC}_C)$$

$$U(\sigma, I; \xi_I) - (1 - \delta)f_I \geq U(\sigma, I; \xi_C) - (1 - \delta)f_C. \quad (\text{IC}_I)$$

Lemma 2. *Given any feasible candidate solution (Ξ, σ) of the above problem,*

$$W(\Xi, (\sigma, \varphi)) \leq \mu(1 - 2\rho - c).$$

Proof. The rater's payoff given any feasible candidate solution (Ξ, σ) of the above problem must satisfy

$$\begin{aligned} W(\Xi, (\sigma, \varphi)) &= (1 - \delta)\mathbf{E}^\mu \left[\sum_{\theta'} \pi(\xi_{\theta'}|\theta) f_{\theta'} \right] \\ &\leq \mathbf{E}^\mu \left[\sum_{\theta'} \pi(\xi_{\theta'}|\theta) (U(\sigma, \theta; \xi_{\theta'}) - U(\sigma, \theta; N)) \right] \\ &\leq \mathbf{E}^\mu \left[\sum_{\theta'} \pi(\xi_{\theta'}|\theta) (U(\sigma, \theta; \xi_{\theta'}) - \rho) \right], \end{aligned} \quad (48)$$

where (48) follows from the IR and IC constraints: for $\theta = \theta'$, (IR_θ) gives $(1 - \delta)f_\theta \leq U(\sigma, \theta; \xi_\theta) - U(\sigma, \theta; N) \leq U(\sigma, \theta; \xi_\theta) - \rho$. Here, the second inequality follows because $U(\sigma, \theta; N) \geq \rho$ (Remark 5). For $\theta \neq \theta'$, if $\pi(\xi_{\theta'}|\theta) > 0$, then it must be true that (IC_θ) binds. By (IR_θ) , $(1 - \delta)f_{\theta'} \leq U(\sigma, \theta; \xi_{\theta'}) - U(\sigma, \theta; N) \leq U(\sigma, \theta; \xi_{\theta'}) - \rho$. Continuing from (48),

$$\begin{aligned} &W(\Xi, (\sigma, \varphi)) \\ &\leq (1 - \delta) \sum_{t=0}^{\infty} \delta^t \left\{ \mathbf{E}^\mu \left[\sum_{\theta'} \pi(\xi_{\theta'}|\theta) \sum_{h_f^t} P_t^\theta(h_r^t|\xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t|h_r^t) p_t^\sigma(r_t) \right] - \rho \right\} \\ &= (1 - \delta) \sum_{t=0}^{\infty} \delta^t \left\{ \mathbf{E}^\mu \left[\sum_{\theta'} \pi(\xi_{\theta'}|\theta) \sum_{h_f^t} P_t^\theta(h_r^t|\xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t|h_r^t) \right. \right. \\ &\quad \times \underbrace{(\rho + (1 - 2\rho)\varphi_t(r_t)\mathbf{E}^P[\tau(h_f^t)|r_t, \theta = C] - \rho)}_{p_t^\sigma(r_t)} \left. \right] \\ &\quad \left. - \mu \sum_{\theta'} \pi(\xi_{\theta'}|C) \sum_{h_f^t} P_t^C(h_r^t|\xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t|h_r^t) \tau(h_f^t)c \right\} \end{aligned} \quad (49)$$

$$\begin{aligned}
&= (1 - \delta) \sum_{t=0}^{\infty} \delta^t \left\{ \mathbf{E}^{\mu} \left[\sum_{\theta'} \pi(\xi_{\theta'} | \theta) \sum_{h_f^t} P_t^{\theta}(h_r^t | \xi_{\theta'}) \right. \right. \\
&\quad \times (1 - 2\rho) \sum_{r_t} S_{\theta'}(r_t | h_r^t) \underbrace{\frac{\mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t)}{\mathbf{E}^{\mu}[\sum_{\theta'} \pi(\xi_{\theta'} | \theta) \sum_{h_f^t} P_t^{\theta}(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t)]}}_{=\varphi_t(r_t)} \\
&\quad \times \mathbf{E}^P[\tau(h_f^t) | r_t, \theta = C] \left. \right] \\
&\quad - \mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t | h_r^t) \tau(h_f^t) c \Big\} \\
&= (1 - \delta) \sum_{t=0}^{\infty} \delta^t \left\{ (1 - 2\rho) \sum_{r_t} \mathbf{E}^{\mu} \left[\sum_{\theta'} \pi(\xi_{\theta'} | \theta) \sum_{h_f^t} P_t^{\theta}(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \right] \times \right. \\
&\quad \frac{\mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t)}{\mathbf{E}^{\mu}[\sum_{\theta'} \pi(\xi_{\theta'} | \theta) \sum_{h_f^t} P_t^{\theta}(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t)]} \mathbf{E}^P[\tau(h_f^t) | r_t, \theta = C] \\
&\quad \left. - \mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t | h_r^t) \tau(h_f^t) c \right\} \tag{51}
\end{aligned}$$

$$\begin{aligned}
&= (1 - \delta) \sum_{t=0}^{\infty} \delta^t \sum_{r_t} \left\{ (1 - 2\rho) \mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \right. \\
&\quad \times \underbrace{\frac{\sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \tau(h_f^t; \xi_{\theta'})}{\sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t)}}_{=\mathbf{E}^P[\tau(h_f^t) | r_t, \theta = C]} \\
&\quad \left. - \mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \tau(h_f^t) c \right\} \tag{52} \\
&= (1 - \delta) \sum_{t=0}^{\infty} \delta^t \sum_{r_t} \left\{ (1 - 2\rho) \mu \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \tau(h_f^t) \right. \\
&\quad \left. - \mu c \sum_{\theta'} \pi(\xi_{\theta'} | C) \sum_{h_f^t} P_t^C(h_r^t | \xi_{\theta'}) S_{\theta'}(r_t | h_r^t) \tau(h_f^t) \right\} \tag{53}
\end{aligned}$$

$$= \mu(1 - 2\rho - c)(1 - \delta) \sum_{t=0}^{\infty} \delta^t \sum_{\theta'} \pi(\xi_{\theta'}|C) \sum_{h_f^t} P_t^C(h_r^t|\xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t|h_r^t) \underbrace{\tau(h_f^t)}_{\leq 1} \quad (54)$$

$$\leq \mu(1 - 2\rho - c)(1 - \delta) \sum_{t=0}^{\infty} \delta^t \sum_{\theta'} \pi(\xi_{\theta'}|C) \underbrace{\sum_{h_f^t} P_t^C(h_r^t|\xi_{\theta'}) \sum_{r_t} S_{\theta'}(r_t|h_r^t)}_{=1} \quad (55)$$

$$= \mu(1 - 2\rho - c) \underbrace{\sum_{\theta'} \pi(\xi_{\theta'}|C)}_{\leq 1} \leq \mu(1 - 2\rho - c), \quad (56)$$

as desired. In the above derivation,

- (49) follows by substituting $p_t^\sigma(r)$ using (7),
- (50) follows by substituting $\varphi_t(r_t)$ using (2),
- (51) rearranges the order of summations,
- (52) expresses the expected effort explicitly according to the corresponding conditional distribution,
- (53) follows from simplifying (52),
- (54) follows from factorizing out the term $\mu(1 - 2\rho - c)$,
- (55) follows because the effort probability is bounded above by one,
- and (56) follows as the participation probability is at most one.

■

Lemma 3. *It holds that*

$$W_1 \leq \mu \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right). \quad (57)$$

Proof. First, fix a menu Ξ and a candidate equilibrium $(\sigma, \varphi) \in B(\Xi)$ with participation $\pi(\Xi|I) > 0$ and $\pi(\Xi|C) = 0$. In this candidate equilibrium, consumers believe that a participating firm is inept for sure, so that $\varphi_t(r) = 0$ and $p_t^\sigma(r) = \rho$ for any rating $r \in \bigcup_{\theta} \text{supp} S_{\theta}(H_r^t)$ and any period t . Respecting the inept type's participation constraint $U(\sigma, I; \xi_{\theta}) - (1 - \delta)f_{\theta} \geq U(\sigma, I; N)$, and the fact that $U(\sigma, I; N) \geq \rho$ (Remark 5), the rater's payoff satisfies

$$(1 - \delta)(1 - \mu) \sum_{\theta} \pi(\xi_{\theta}|I) f_{\theta} \leq (1 - \mu) \sum_{\theta} [U(\sigma, I; \xi_{\theta}) - \rho] = 0.$$

Next, fix a menu Ξ and a candidate equilibrium $(\sigma, \varphi) \in B(\Xi)$ with participation $\pi(\Xi|C) > 0$ and $\pi(\Xi|I) = 0$. In this equilibrium, consumers' beliefs satisfy $\varphi_t(r) = 1$ for each non-null rating r and each period t . Upon participating in a scheme ξ_θ ,

$$U(\sigma, C; \xi_\theta) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho}. \quad (58)$$

We now show (58). To obtain an upper bound on the competent firm's profit upon participating in ξ_θ , it is without loss to assume that the rating system S_θ does not send a null rating, which would lower consumer beliefs and payments. The continuation profit of the competent firm after a history h_f^t , conditional on participating in a scheme ξ_θ , is given by

$$\begin{aligned} & V_C^\sigma(h_f^t; \xi_\theta) \\ &:= (1 - \delta) \mathbf{E}^P \left[\sum_{s=t}^{\infty} \delta^s (p_s^\sigma(r_s) - c(e_s)) \middle| h_f^t, \theta = C, \xi_\theta \right] \\ &= (1 - \delta) (p_t^\sigma(r_t) - \tau(h_f^t)c) + \delta \mathbf{E}^P [V_C^\sigma(h_f^t, e_t, y_t) | h_f^t, e_t, \theta = C, \xi_\theta]. \end{aligned} \quad (59)$$

Next, observe that the firm's flow profit after each h_f^t conditional on participating in ξ_θ depends only on h_r^t (which is embedded in h_f^t by definition), because the histories of ratings determine consumer payments and the evolution of ratings depends on h_r^t . The firm essentially faces a Markov decision problem, with the set of states given by the set of the rater's histories H_r , and so has a Markov best reply. One can then write for each period t and each history h_f^t that

$$V_C^\sigma(h_f^t; \xi_\theta) = V_C^\sigma(h_r^t, r_t; \xi_\theta). \quad (60)$$

It is essential to derive the incentive constraint for high effort. High effort after history (h_r^t, r_t) gives the firm a continuation profit

$$\begin{aligned} V_C^\sigma(h_r^t, r_t; \xi_\theta) &= (1 - \delta) (p_t^\sigma(r_t) - c) \\ &\quad + \delta [(1 - \rho) V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta) + \rho V_C^\sigma(h_r^t, r_t, \underline{y}; \xi_\theta)]. \end{aligned} \quad (61)$$

Shirking, on the other hand, gives a continuation profit

$$\begin{aligned} V_C^\sigma(h_r^t, r_t; \underline{e}; \xi_\theta) &= (1 - \delta) (p_t^\sigma(r_t) - c) \\ &\quad + \delta [\rho V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta) + (1 - \rho) V_C^\sigma(h_r^t, r_t, \underline{y}; \xi_\theta)]. \end{aligned} \quad (62)$$

The incentive constraint for high effort after (h_r^t, r_t) , that (61) weakly exceeds (62), can be simplified to

$$V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta) - V_C^\sigma(h_r^t, r_t, \underline{y}; \xi_\theta) \geq \frac{(1 - \delta)c}{\delta(1 - 2\rho)}. \quad (63)$$

Now,

$$\begin{aligned}
U(\sigma, C; \xi_\theta) &= \mathbf{E}^P[V_C^\sigma(h_f^0; \xi_\theta) | \theta = C, \xi_\theta] \\
&= \mathbf{E}^P \left[(1 - \delta)(\rho + \mathbf{E}^P[\tau(h_f^0) | r_0, \theta = C, \xi_\theta](1 - 2\rho - c)) \right. \\
&\quad \left. + \delta((1 - \rho)V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta) + \rho V_C^\sigma(h_r^t, r_t, \underline{y}; \xi_\theta)) \middle| \theta = C, \xi_\theta \right] \\
&\leq (1 - \delta)(1 - \rho - c) + \mathbf{E}^P \left[\delta \left(V_C^\sigma(h_r^0, r_0, \bar{y}; \xi_\theta) - \frac{(1 - \delta)\rho c}{\delta(1 - 2\rho)} \right) \middle| \theta = C, \xi_\theta \right] \\
&= (1 - \delta) \left(1 - \rho - c - \frac{\rho c}{1 - 2\rho} \right) + \delta \mathbf{E}^P[V_C^\sigma(h_r^0, r_0, \bar{y}; \xi_\theta) | \theta = C, \xi_\theta]. \quad (64)
\end{aligned}$$

The inequality follows from efficiency (1) and the incentive constraint (63) at $t = 0$. Now, $V_C^\sigma(h_r^0, r_0, \bar{y}; \xi_\theta)$, and more generally $V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta)$ for each t , can be analogously expressed as

$$\begin{aligned}
V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta) &\leq (1 - \delta) \left(1 - \rho - c - \frac{\rho c}{1 - 2\rho} \right) \\
&\quad + \delta \mathbf{E}^P[V_C^\sigma(h_r^{t+1}, r_{t+1}, \bar{y}; \xi_\theta) | h_r^t, r_t, \theta = C, \xi_\theta], \quad (65)
\end{aligned}$$

Continuing from (64), and recursively substituting $V_C^\sigma(h_r^t, r_t, \bar{y}; \xi_\theta)$ using (65) for each period t ,

$$U(\sigma, C; \xi_\theta) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho}, \quad (66)$$

as desired. The competent firm's participation constraint in ξ_θ is $U(\sigma, C; \xi_\theta) - (1 - \delta)f_\theta \geq U(\sigma, C; N) \geq \rho$. The rater's expected payoff therefore satisfies

$$\begin{aligned}
(1 - \delta)\mu \sum_{\theta \in \Theta} \pi(\xi_\theta | C) f_\theta &\leq \mu \sum_{\theta \in \Theta} \pi(\xi_\theta | C) (U(\sigma, C; \xi_\theta) - \rho) \\
&\leq \mu \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right), \quad (67)
\end{aligned}$$

as desired. Finally, and trivially, in any equilibrium in which neither type participates, the rater's payoff is zero. $\blacksquare$

Lemma 4 below completes the proof of Proposition 1.

Lemma 4. *For every menu Ξ and equilibrium $(\sigma, \varphi) \in B(\Xi)$, $W(\Xi, (\sigma, \varphi)) = \mu(1 - 2\rho - c)$ if and only if the conditions stated in Proposition 1 are satisfied.*

Proof. Fix a menu Ξ . Suppose there is no equilibrium that satisfies all of the stated conditions in the proposition. Fix one such equilibrium $(\sigma, \varphi) \in B(\Xi)$. First, if $(\sigma, \varphi) \notin B_2(\Xi)$, then the menu Ξ cannot be fully-extracting by Lemma 3, because $W_1 < \mu(1 - 2\rho - c)$. Second, if $(\sigma, \varphi) \in B_2(\Xi)$ but at least one of Conditions 2–4 does not hold, then at least one of the inequalities in the derivation in Lemma 2 becomes strict, giving $W(\Xi, (\sigma, \mu)) < \mu(1 - 2\rho - c)$. Conversely, fix a menu Ξ and an equilibrium $(\sigma, \varphi) \in B(\Xi)$ that satisfy the conditions. Then all the inequalities in the derivation in Lemma 2 become equalities, giving $W(\Xi, (\sigma, \mu)) = \mu(1 - 2\rho - c)$. ■

B.2 Proof of Proposition 2

It is straightforward to compute that

$$W(\Xi^*, (\sigma^*, \varphi^*)) = (1 - \delta)[\mu f_C^* + (1 - \mu)f_I^*] = \mu(1 - 2\rho - c).$$

I show that $(\sigma^*, \varphi^*) \in B(\Xi^*)$. Fix the candidate profile (σ^*, φ^*) . Consider first the continuation game after a competent firm chooses participation $d = \xi_C^*$. Because the system S_C^* relies only on the most recent signal, and there are only two possible ratings, 0 and 1, and because the firm's effort is identical conditional on each rating, the firm has only two possible stage payoffs for each period $t \geq 1$:

$$p_t^{\sigma^*}(0) - c = \rho + (1 - 2\rho)\varphi_t^*(0) - c = \rho + \frac{(1 - 2\rho)\mu\rho}{\mu\rho + 1 - \mu} - c, \quad (68)$$

$$p_t^{\sigma^*}(1) - c = \rho + (1 - 2\rho)\varphi_t^*(1) - c = 1 - \rho - c. \quad (69)$$

The play by the competent firm from $t \geq 1$ upon participation $d = \xi_C^*$ can be represented by a two-state automaton, with the set of states given by $R = \{0, 1\}$. Here, a state r collects all firm's histories with the most recent realized rating being r . Let $V_C^{\sigma^*}(r; \xi_C^*)$ be the competent firm's continuation profit in state r given (σ^*, φ^*) conditional on $d = \xi_C^*$. The competent firm's continuation profit in state r is

$$V_C^{\sigma^*}(r; \xi_C^*) = (1 - \delta)(p_t^{\sigma^*}(r) - c) + \delta[(1 - \rho)V_C^{\sigma^*}(1; \xi_C^*) + \rho V_C^{\sigma^*}(0; \xi_C^*)].$$

A deviation to shirk yields

$$V_C^{\sigma^*}(r; \varepsilon; \xi_C^*) := (1 - \delta)p_t^{\sigma^*}(r) + \delta[\rho V_C^{\sigma^*}(1; \xi_C^*) + (1 - \rho)V_C^{\sigma^*}(0; \xi_C^*)].$$

The incentive constraint $V_C^{\sigma^*}(r; \xi_C^*) - V_C^{\sigma^*}(r; \varepsilon; \xi_C^*) \geq 0$ in state r needs to hold, for each $r \in \{0, 1\}$. It can be simplified to $\delta(1 - 2\rho)(p_{t+1}^{\sigma^*}(1) - p_{t+1}^{\sigma^*}(0)) \geq c$,

or equivalently,

$$\delta(1-2\rho)^2 \left(1 - \frac{\mu\rho}{1-\mu+\mu\rho}\right) \geq c. \quad (70)$$

The left hand side equals $\bar{c}(\mu, \delta, \rho)$, thus the constraint holds by (9). In period 0, each type of firm receives a payment of

$$p_0^{\sigma^*}(0) = \rho + (1-2\rho)\mu, \quad (71)$$

upon participation, and the competent firm faces the same incentive constraint (70).

Suppose instead the competent firm chooses $d = \xi_I^*$. Because only rating 0 is ever announced, consumers pay the firm $p_t^{\sigma^*}(0)$ in each period t . The competent firm thus finds it optimal to exert low effort in each period in the continuation game. Similarly, suppose the competent firm chooses $d = N$. Upon seeing the null signal, consumers pay the firm ρ every period. The competent firm finds it optimal to exert low effort in every period in the continuation game.

It remains to verify that participation π^* constitutes an equilibrium. Upon choosing ξ_C^* and paying the fee f_C^* , the competent firm's payoff equals

$$\begin{aligned} & U(\sigma^*, C; \xi_C^*) - (1-\delta)f_C^* \\ &= (1-\delta) \left(p_0^{\sigma^*}(0) - c + \sum_{t=1}^{\infty} \delta^t [(1-\rho)(p_t^{\sigma^*}(1) - c) + \rho(p_t^{\sigma^*}(0) - c)] \right) - (1-\delta)f_C^* \\ &= \rho, \end{aligned}$$

where the derivation from the second to the last line uses (68), (69), (71) and (32). By deviating to choose ξ_I^* , paying the fee f_I^* and shirking afterwards, the competent firm's payoff equals

$$U(\sigma^*, C; \xi_I^*) - (1-\delta)f_I^* = (1-\delta) \left(p_0^{\sigma^*}(0) + \sum_{t=1}^{\infty} \delta^t p_t^{\sigma^*}(0) \right) - (1-\delta)f_I^* = \rho,$$

where the derivation again uses (68), (71) and (33). By deviating to choose $d = N$, the competent firm does not pay an upfront fee but receive a consumer payment of ρ every period. Its payoff therefore equals ρ . Consequently, the competent type finds no profitable deviation from ξ_C^* to either ξ_I^* or N .

Next, upon choosing ξ_I^* and paying f_I^* , the inept firm's payoff equals

$$U(\sigma^*, I; \xi_I^*) - (1-\delta)f_I^* = (1-\delta) \left(p_0^{\sigma^*}(0) + \sum_{t=1}^{\infty} \delta^t p_t^{\sigma^*}(0) \right) - (1-\delta)f_I^* = \rho,$$

where the derivation uses (68), (71) and (33). By deviating to choose ξ_C^* and paying f_C^* , its payoff equals

$$\begin{aligned} & U(\sigma^*, I; \xi_C^*) - (1 - \delta)f_C^* \\ &= (1 - \delta) \left(p_0^{\sigma^*}(0) + \sum_{t=1}^{\infty} \delta^t [\rho p_t^{\sigma^*}(1) + (1 - \rho)p_t^{\sigma^*}(0)] \right) - (1 - \delta)f_C^* \\ &= \frac{c - \delta(1 - \mu)(1 - 2\rho)^2 - c\mu(1 - \rho)}{1 - \mu(1 - \rho)} \leq \rho, \end{aligned}$$

whenever $c \leq \bar{c}(\mu, \delta, \rho)$, where the derivation from the second to the last line uses (69), (68), (71) and (32). Finally, by deviating to choose $d = N$, the inept firm does not pay an upfront fee and gets a consumer payment of ρ every period. Its payoff thus equals ρ . Hence, the inept type also finds no profitable deviation from ξ_I^* to either ξ_C^* or N . The proof is complete.

B.3 Proof of Proposition 3

By Proposition 1, the rater's fully-extracting payoff is $\mu(1 - 2\rho - c)$ given (9). Because $\bar{c}_1(\mu, \delta, \rho) < \bar{c}(\mu, \delta, \rho)$, to prove the proposition, it suffices to show that given (14), $(\sigma^{**}, \varphi^{**}) \in B(\Xi^{**})$ and $W(\Xi^{**}, (\sigma^{**}, \varphi^{**})) = \mu(1 - 2\rho - c)$.

It is again straightforward to compute

$$W(\Xi^{**}, (\sigma^{**}, \varphi^{**})) = (1 - \delta)[\mu f^{**} + (1 - \mu)f^{**}] = \mu(1 - 2\rho - c).$$

I show $(\sigma^{**}, \varphi^{**}) \in B(\Xi^{**})$. Fix the candidate profile $(\sigma^{**}, \varphi^{**})$. Consider first the continuation game after the competent firm chooses ξ^{**} . Analogous to the proof of Proposition 2, the firm has only two possible stage payoffs in each period $t \geq 1$:

$$p_t^{\sigma^{**}}(0) - c = \rho + (1 - 2\rho)\varphi_t^{**}(0) - c = \rho + \frac{(1 - 2\rho)\mu(1 - (1 - \rho)\alpha)}{\mu(1 - (1 - \rho)\alpha) + (1 - \mu)(1 - \rho\alpha)} - c, \quad (72)$$

$$p_t^{\sigma^{**}}(1) - c = \rho + (1 - 2\rho)\varphi_t^{**}(1) - c = \rho + \frac{(1 - 2\rho)\mu(1 - \rho)\alpha}{\mu(1 - \rho)\alpha + (1 - \mu)\rho\alpha} - c. \quad (73)$$

Again analogous to the proof of Proposition 2, the firm's continuation play from $t \geq 1$ can be represented by a two-state automaton, with the set of states given by $\{0, 1\}$. State r collects the firm's histories with the most recent realized rating being r . Let $V_C^{\sigma^{**}}(r; \xi^{**})$ be the continuation value of the competent firm in a state r of the automaton given the profile σ^{**}

conditional on participating in ξ^{**} . The competent firm's continuation profit in state r is

$$V_C^{\sigma^{**}}(r; \xi^{**}) = (1 - \delta)(p_t^{\sigma^{**}}(r) - c) + \delta[(1 - \rho)\alpha V_C^{\sigma^{**}}(1; \xi^{**}) + (1 - (1 - \rho)\alpha)V_C^{\sigma^{**}}(0; \xi^{**})].$$

A deviation to shirk yields

$$V_C^{\sigma^{**}}(r; \varepsilon; \xi^{**}) := (1 - \delta)p_t^{\sigma^{**}}(r) + \delta[\rho\alpha V_C^{\sigma}(1; \xi^{**}) + (1 - \rho\alpha)V_C^{\sigma}(0; \xi^{**})].$$

The incentive constraint $V_C^{\sigma^{**}}(r; \xi^{**}) - V_C^{\sigma^{**}}(r; \varepsilon; \xi^{**}) \geq 0$ in state r needs to hold, for each $r \in \{0, 1\}$. The constraint can be simplified as

$$\delta\alpha(1 - 2\rho)^2 \left(\frac{\mu(1 - \rho)}{\mu(1 - \rho) + (1 - \mu)\rho} - \frac{\mu(1 - \alpha(1 - \rho))}{\mu(1 - \alpha(1 - \rho)) + (1 - \mu)(1 - \alpha\rho)} \right) \geq c. \quad (74)$$

Some algebraic manipulation reveals that, with α defined by (34), the left hand side equals c .²⁵ Therefore, the incentive constraint holds. In period 0, each type of firm receives a payment

$$p_0^{\sigma^{**}}(0) = \rho + \mu(1 - 2\rho) \quad (75)$$

upon participation, and the competent firm faces the same incentive constraint (74).

Suppose instead the competent firm chooses $d = N$. Upon observing a null rating, consumers pays the firm ρ every period. The competent firm finds it optimal to exert low effort in every period in the continuation game. This also implies that each type's outside option payoff is ρ .

It remains to verify π^{**} constitutes an equilibrium. By construction, the fee f^{**} given by (35) satisfies the inept firm's participation constraint:

$$\begin{aligned} & U(\sigma^{**}; I, \xi^{**}) - (1 - \delta)f^{**} \\ &= (1 - \delta) \left(p_0^{\sigma^{**}}(0) + \sum_{t=1}^{\infty} \delta^t [\alpha\rho p_t^{\sigma^{**}}(1) + (1 - \alpha\rho)p_t^{\sigma^{**}}(0)] \right) - (1 - \delta)f^{**} = \rho, \end{aligned}$$

²⁵Substituting α defined by (34) on the left hand side of the inequality gives

$$\frac{c\delta(1 - 2\rho)^2(\rho + \mu(1 - 2\rho))}{c(\rho + \mu(1 - 2\rho))^2 + \delta\mu(1 - \mu)(1 - 2\rho)^3} \left(\frac{\mu(1 - \rho)}{\mu(1 - \rho) + (1 - \mu)\rho} - \frac{\delta\mu(1 - 2\rho)^2 - c\rho - c\mu(1 - 2\rho)}{\delta(1 - 2\rho)^2} \right) = c.$$

where the derivation uses (72), (73), (75) and (35). By construction of α , the incentive constraint (74) for high effort bind for each state r . The competent firm is therefore indifferent between high or low effort after each history upon participation, and $U(\sigma^{**}, C; \xi^{**}) = U(\sigma^{**}, I; \xi^{**})$. Hence, the competent firm's participation constraint also binds. This completes the proof.

B.4 Proof of Proposition 4

For any menu Ξ and equilibrium $(\sigma, \varphi) \in B_2(\Xi)$, it follows from a derivation analogous to that in Lemma 2 that $W(\Xi, \sigma) \leq \mu(1 - 2\rho)$, with the restrictions $c = 0$ and $\tau(\cdot, C) = 1$. The first set of conditions in the Proposition is a direct counterpart of Proposition 1. It is analogous to Lemma 4 to show that any menu Ξ and equilibrium $\sigma \in B_2(\Xi)$ that satisfy the first set of conditions give the rater a payoff $\mu(1 - 2\rho)$. Suppose that the second set of conditions holds for a menu Ξ and an equilibrium $(\sigma, \varphi) \in B(\Xi) \setminus B_2(\Xi)$. The firm consistently exerts high effort, generating an expected market surplus $1 - \rho$. Because participates with probability one, by fully extracting the firm's net surplus (that is, compensating the competent firm with the outside option payoff ρ), the rater obtains her fully-extracting payoff $\mu(1 - 2\rho)$. Conversely, if the rater's payoff is strictly less than $\mu(1 - 2\rho)$, it is straightforward to verify that at least one set of the conditions stated in the Proposition does not hold.

B.5 Proof of Proposition 5

For any scheme ξ and equilibrium strategy σ , it follows analogously from Lemma 3 that

$$U(\sigma, C; \xi) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho}.$$

Because participation must be individually rational and the firm's outside option payoff is at least ρ , the rater's equilibrium payoff is at most (19). It remains to show that $(\sigma_{MH}, \varphi_{MH}) \in B(\Xi_{MH})$, and that $W(\Xi_{MH}, (\sigma_{MH}, \varphi_{MH}))$ equals (19), whenever $c \leq \bar{c}_{MH}$.

Consider first the case $c \leq \hat{c} < \bar{c}_{MH}$, where $\hat{c}$ is defined by (36). I show that $(\sigma_{MH}, \varphi_{MH}) \in B(\Xi_{MH})$. Given σ_{MH} , after any history of play from period $t \geq 1$, upon each rating $r \in \{0, 1\}$, the firm's stage profit and its continuation strategy by the firm are identical. Consequently the play upon participation by the competent firm can be represented by a two-state automaton, with states given by the ratings $r \in \{0, 1\}$ and transitions depending on the current signals, depicted by Figure 12 below.

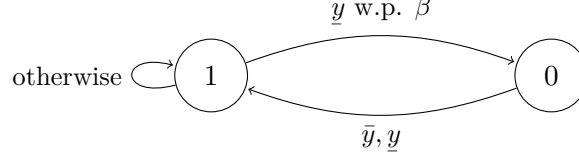


Figure 12: The two-state automation

The initial state is 1, in which the firm chooses high effort with probability one; in state 0, the firm chooses low effort with probability one.

Denote the competent firm's continuation profit extraction from consumers in state $r \in \{0, 1\}$ by $V_r^{\sigma_{MH}}$. These continuation profits satisfy:

$$\begin{aligned} V_1^{\sigma_{MH}} &= (1 - \delta)(1 - \rho - c) + \delta((1 - \rho\beta)V_1^{\sigma_{MH}} + \rho\beta V_0^{\sigma_{MH}}), \\ V_0^{\sigma_{MH}} &= (1 - \delta)\rho + \delta V_1^{\sigma_{MH}}. \end{aligned}$$

Solving for $V_1^{\sigma_{MH}}$ and $V_0^{\sigma_{MH}}$, and direct computation reveals that $\delta(1 - 2\rho)\beta(V_1^{\sigma_{MH}} - V_0^{\sigma_{MH}}) \geq (1 - \delta)c$ given (20). That is, the incentive constraint for high effort in state 1 is satisfied. Clearly, the incentive constraint for high effort in state 0 is violated, so that the firm finds it optimal to exert low effort. Straightforward algebraic manipulation shows

$$V_1^{\sigma_{MH}} = 1 - \rho - c - \frac{\rho c}{1 - 2\rho}.$$

Next, by deviating to not participate, the firm finds it optimal to always exert low effort, because future consumers do not observe past signals. Each entering consumer thus pays the firm exactly ρ upon observing a null rating, giving $U(\sigma_{MH}, C; N) = \rho$.

It remains to verify that the competent firm does not deviate from participating. By participating, the competent firm obtains a payoff

$$\begin{aligned} &U(\sigma_{MH}, C; \xi_{MH}) - (1 - \delta)f_{MH} \\ &= 1 - \rho - c - \frac{\rho c}{1 - 2\rho} - \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho}\right) = \rho = U(\sigma_{MH}, C; N), \end{aligned}$$

as desired. Consider next the case $c \in (\hat{c}, \bar{c}_{MH}]$. A similar two-state automaton can be constructed in Figure 13 below. The continuation profits $V_1^{\sigma_{MH}}, V_0^{\sigma_{MH}}$ satisfy:

$$V_1^{\sigma_{MH}} = (1 - \delta)(1 - \rho - c) + \delta((1 - \rho)V_1^{\sigma_{MH}} + \rho V_0^{\sigma_{MH}}),$$

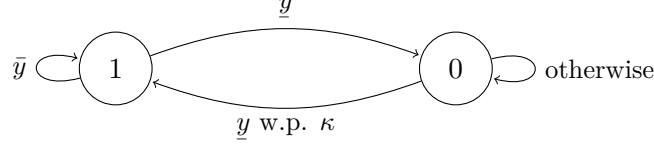


Figure 13: The two-state automation

The initial state is 1, in which the firm chooses high effort with probability one; in state 0, the firm chooses low effort with probability one.

$$V_0^{\sigma_{MH}} = (1 - \delta)\rho + \delta((1 - \rho)\kappa V_1^{\sigma_{MH}} + (1 - (1 - \rho)\kappa)V_0^{\sigma_{MH}}).$$

Again, direct computation reveals that the incentive constraint for high effort in state 1 is satisfied, because $\delta(1 - 2\rho)(V_1^{\sigma_{MH}} - V_0^{\sigma_{MH}}) \geq (1 - \delta)c$ given (20). Also, the incentive constraint for high effort in state 0 is violated, so that the firm finds it optimal to exert low effort. Straightforward algebraic manipulation again reveals that

$$V_1^{\sigma_{MH}} = 1 - \rho - c - \frac{\rho c}{1 - 2\rho}.$$

It is now analogous to the first case to show that the competent firm does not profit from deviating to not participate.

B.6 Proof of Proposition 6

The proof is omitted as it is almost identical to showing (58) in Lemma 3.

B.7 Proof of Proposition 7

By Proposition 1, it suffices to show that consistent high effort fails in equilibrium whenever $c \geq \delta(1 - 2\rho)^2$. A necessary condition for equilibrium consistent high effort is that the incentive constraint (39) after each history h_f holds. One can recursively substitute the incentive constraints for high effort after each history, as in the proof of Lemma 3, to obtain

$$V^\sigma(\bar{y}; h_f) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho},$$

$$V^\sigma(\underline{y}; h_f) > \rho - c + \frac{(1 - \rho)c}{1 - 2\rho}.$$

The second inequality is strict because the competent type always obtain a price strictly bigger than ρ in each period. These two inequalities, together with the incentive constraint (39), imply that $c < \delta(1 - 2\rho)^2$ must hold for equilibrium consistent high effort.

B.8 Proof of Proposition 8

If the rater achieves full extraction in an equilibrium (σ, φ) , then (σ, φ) satisfies the four conditions in Proposition 1. Suppose that the competent firm participates in a scheme ξ_θ with probability $\pi(\xi_\theta|C) > 0$. Moreover, the incentive constraint (39) necessarily holds after each firm's history h_f in the equilibrium (σ, φ) , conditional on participating in the scheme ξ_θ . One can recursively substitute the incentive constraints for high effort after each firm's history h_f , as in the proof of Lemma 3, to obtain that after the period-0 firm's history h_f^0 (conditional on $d = \xi_\theta$),

$$V^\sigma(\bar{y}; h_f^0) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho}.$$

Similarly,

$$\begin{aligned} & V^\sigma(\underline{y}; h_f^0) \\ &= \mathbf{E}^P \left[(1 - \delta)(p_1^\sigma(r_1) - c) + \delta((1 - \rho)V^\sigma(\bar{y}; h_f^0, (\underline{y}, r)) + \rho V^\sigma(\underline{y}; h_f^0, (\underline{y}, r))) \middle| h_f^0, \underline{y}, C \right] \\ &\geq \mathbf{E}^P \left[(1 - \delta)(p_1^\sigma(r_1) - c) + \delta \left(V^\sigma(\underline{y}; h_f^0, (\bar{y}, r_0)) + \frac{(1 - \delta)(1 - \rho)c}{\delta(1 - 2\rho)} \right) \middle| h_f^0, \underline{y}, C \right] \\ &\geq \mathbf{E}^P \left[(1 - \delta) \left(p_1^\sigma(r_1) - c + \frac{(1 - \rho)c}{1 - 2\rho} \right) + \delta V^\sigma(\underline{y}; h_f^0, (\bar{y}, r_0)) \middle| h_f^0, \underline{y}, C \right] \\ &\dots \\ &\geq (1 - \delta) \sum_{t=0}^{\infty} \delta^t \mathbf{E}^P[p_{t+1}^\sigma(r_{t+1}) | h_f^0, \underline{y}^{t+1}, C] - c + \frac{(1 - \rho)c}{1 - 2\rho} \\ &= \rho + (1 - 2\rho)(1 - \delta) \sum_{t=0}^{\infty} \delta^t \mathbf{E}^P[\varphi_{t+1}(r_{t+1}) | h_f^0, \underline{y}^{t+1}, C] - c + \frac{(1 - \rho)c}{1 - 2\rho}, \end{aligned}$$

where $\varphi_t(r_t)$ denotes consumer posterior of the firm being competent upon observing a rating r_t in period t . The incentive constraint after history h_f^0 and the preceding inequalities together imply that a necessary condition for full extraction is

$$c \leq \delta(1 - 2\rho)^2 \left(1 - (1 - \delta) \sum_{t=0}^{\infty} \delta^t \mathbf{E}^P[\varphi_{t+1}(r_{t+1}) | h_f^0, \underline{y}^{t+1}, C] \right). \quad (76)$$

Suppose, towards a contradiction, for any $\eta > 0$, the rater achieves full extraction for $c \in [\delta(1 - 2\rho)^2 - \eta, \delta(1 - 2\rho)^2]$. The necessary condition (76) reduces to

$$\eta \geq \delta(1 - \delta)(1 - 2\rho)^2 \left(\sum_{t=0}^{\infty} \delta^t \mathbf{E}^P[\varphi_{t+1}(r_{t+1}) | h_f^0, \underline{y}^{t+1}, C] \right),$$

implying that

$$\eta > \delta(1 - \delta)(1 - 2\rho)^2 \mathbf{E}^P[\varphi_1(r_1) | h_f^0, \underline{y}, C]. \quad (77)$$

The expectation in (77) satisfies

$$\begin{aligned} & \mathbf{E}^P[\varphi_1(r_1) | h_f^0, \underline{y}, C] \\ & \geq \sum_{r_1 \in R} S_\theta(r_1 | h_f^0, \underline{y}; \eta) \frac{\mu(\rho \sum_{\theta'} \pi(\xi_{\theta'} | C) S_{\theta'}(r_1 | h_f^0, \underline{y}; \eta))}{\mu(\rho \sum_{\theta'} \pi(\xi_{\theta'} | C) S_{\theta'}(r_1 | h_f^0, \underline{y}; \eta)) + (1 - \mu) \times 1} \\ & > \sum_{r_1 \in R} S_\theta(r_1 | h_f^0, \underline{y}; \eta) \mu \rho \sum_{\theta'} \pi(\xi_{\theta'} | C) S_{\theta'}(r_1 | h_f^0, \underline{y}; \eta) \\ & \geq \sum_{r_1 \in R} S_\theta(r_1 | h_f^0, \underline{y}; \eta) \mu \rho \pi(\xi_\theta | C) S_\theta(r_1 | h_f^0, \underline{y}; \eta) \\ & > \pi(\xi_\theta | C) \mu \rho \sum_{r \in R} S_\theta(r | h_f^0, \underline{y}; \eta)^2, \end{aligned} \quad (78)$$

where the distribution $S_\theta(\cdot; \eta)$ is parameterized by η to make explicit the possibility that the rating systems may depend on η . By definition of a distribution function,

$$\sum_{r \in R} S_\theta(r | h_f^0, \underline{y}; \eta) = 1. \quad (79)$$

This implies that

$$\sum_{r \in R} S_\theta(r | h_f^0, \underline{y}; \eta)^2 > 0.$$

Thus (78) and consequently the right side of (77) are strictly positive. Now, if η is arbitrarily close to 0, so is the left side of (77). But the right side of (77) cannot be arbitrarily close to 0. Otherwise, $S_\theta(r | h_f^0, \underline{y}; \eta)$ must be arbitrarily close to 0 for each r , violating (79). It therefore contradicts with the claim that (77) holds for every $\eta > 0$. There must exist $\bar{\eta} \equiv \bar{\eta}(\mu, \delta, \rho) > 0$ such that for every $c \in [\delta(1 - 2\rho)^2 - \bar{\eta}, \delta(1 - 2\rho)^2]$, (77) is violated and therefore full extraction fails in equilibrium.

B.9 Proof of Proposition 9

By Proposition 1, it suffices to check that $(\sigma^\infty, \varphi^\infty) \in B(\Xi^\infty)$ and the four conditions in Proposition 1 hold. We first check the competent firm's incentives for consistent high effort. Upon choosing ξ_C^∞ and receiving rating r , the competent firm's (unnormalized) continuation profit $V_r^{\sigma^\infty}/(1-\delta)$ follows

$$\begin{aligned} \lim_{\delta \uparrow 1} \frac{V_0^{\sigma^\infty}}{1-\delta} &= p^{\sigma^\infty}(0) - c + \delta \left((1-\rho) \lim_{\delta \uparrow 1} \frac{V_0^{\sigma^\infty}}{1-\delta} + \rho \lim_{\delta \uparrow 1} \frac{V_1^{\sigma^\infty}}{1-\delta} \right) \\ \lim_{\delta \uparrow 1} \frac{V_r^{\sigma^\infty}}{1-\delta} &= p^{\sigma^\infty}(r) - c + \delta \left((1-\rho) \phi \lim_{\delta \uparrow 1} \frac{V_{r-1}^{\sigma^\infty}}{1-\delta} \right. \\ &\quad \left. + (1-\rho)(1-\phi) \lim_{\delta \uparrow 1} \frac{V_r^{\sigma^\infty}}{1-\delta} + \rho \lim_{\delta \uparrow 1} \frac{V_{r+1}^{\sigma^\infty}}{1-\delta} \right), \quad r > 0, \end{aligned}$$

where the time subscript is dropped given stationarity in consumer beliefs (and hence payments). The incentive constraints for high effort upon each rating $r \in R_C^\infty$, derived analogously as in the proof of Proposition 2, are

$$\begin{aligned} (1-2\rho) \left(\lim_{\delta \uparrow 1} \frac{V_0^{\sigma^\infty} - V_1^{\sigma^\infty}}{1-\delta} \right) &\geq c, \\ (1-2\rho) \left[\phi \left(\lim_{\delta \uparrow 1} \frac{V_{r-1}^{\sigma^\infty} - V_{r+1}^{\sigma^\infty}}{1-\delta} \right) + (1-\phi) \left(\lim_{\delta \uparrow 1} \frac{V_r^{\sigma^\infty} - V_{r+1}^{\sigma^\infty}}{1-\delta} \right) \right] &\geq c, \quad r > 0. \end{aligned}$$

Given consistent high effort, the left side of the above incentive constraints are independent of cost c . Set $c = (1-2\rho)^2$, so that if the constraints hold, they hold for any $c \leq (1-2\rho)^2$. Define

$$\begin{aligned} k_0 &= 0, \\ k_r &= \lim_{\delta \uparrow 1} \frac{V_r^{\sigma^\infty} - V_0^{\sigma^\infty}}{1-\delta}, \quad \text{for each } r > 0. \end{aligned}$$

The incentive constraints are satisfied if and only if

$$\min \left\{ -k_1, \min_{s=1,2,\dots} \left[\phi(k_{s-1} - k_{s+1}) + (1-\phi)(k_s - k_{s+1}) \right] \right\} - (1-2\rho) \geq 0. \quad (80)$$

Given a boundedness condition on the continuation profits verified in Appendix A.9, the variables $(k_r)_{r=0,1,\dots}$ satisfy the following system of equations (81)–(82) (see Ross (2014), Chapter V.1 and V.2),

$$g = p^{\sigma^\infty}(0) + \rho k_1, \quad (81)$$

$$g + k_r = p^{\sigma^\infty}(r) - c + (1 - \rho)\phi k_{r-1} + (1 - \rho)(1 - \phi)k_r + \rho k_{r+1}, \quad r > 0, \quad (82)$$

where

$$\begin{aligned} g &= \mathbf{E}^P[p^{\sigma^\infty}(r) - c], \\ p^{\sigma^\infty}(0) &= 1 - \rho, \\ p^{\sigma^\infty}(r) &= \rho + (1 - 2\rho)\varphi_r = \rho + (1 - 2\rho)\frac{\mu\lambda_r^C}{\mu\lambda_r^C + (1 - \mu)\lambda_r^I} \\ &= \begin{cases} \frac{(1 - \rho)\rho}{1 - \rho - \mu(1 - 2\rho)}, & \text{if } \mu \in (0, \frac{1}{2}) \\ \frac{\rho(1 - 2\mu\rho)}{1 - \mu}, & \text{if } \mu \in (\frac{1}{2}, 1) \end{cases}, \quad r > 0. \end{aligned}$$

Thus,

$$\begin{aligned} g &= \mathbf{E}^P[p^{\sigma^\infty}(r) - c] = \lambda_0^C(1 - \rho - c) + (1 - \lambda_0^C)\left(\frac{\rho(1 - 2\mu\rho)}{1 - \mu} - c\right) \\ &= \begin{cases} 1 - \frac{\rho(2(1 - \mu(1 - 2\rho)) - 3\rho)}{1 - \rho - \mu(1 - 2\rho)} - c, & \text{if } \mu \in (0, \frac{1}{2}), \\ 1 - 2(1 - \rho)\rho - c, & \text{if } \mu \in (\frac{1}{2}, 1). \end{cases} \end{aligned}$$

Solving the system of equations (81)–(82) recursively for $(k_r)_{r=1}^\infty$ gives

$$k_r = \begin{cases} \frac{-r(1 - 2\rho)(1 - \mu)}{1 - \rho - \mu(1 - 2\rho)}, & \text{if } \mu \in (0, \frac{1}{2}), \\ -r(1 - 2\rho), & \text{if } \mu \in (\frac{1}{2}, 1). \end{cases}$$

Therefore,

$$\begin{aligned} -k_1 - (1 - 2\rho) &= \begin{cases} \frac{\rho(1 - 2\mu)(1 - 2\rho)}{1 - \rho - \mu(1 - 2\rho)}, & \text{if } \mu \in (0, \frac{1}{2}), \\ 0, & \text{if } \mu \in (\frac{1}{2}, 1), \end{cases} \\ &\geq 0. \end{aligned}$$

and for each $r > 0$,

$$\phi(k_{r-1} - k_{r+1}) + (1 - \phi)(k_r - k_{r+1}) - (1 - 2\rho)$$

$$\begin{aligned}
&= \begin{cases} \frac{2(1-\mu)(1-2\rho)}{1-\rho-\mu(1-2\rho)}, & \text{if } \mu \in (0, \frac{1}{2}), \\ \frac{(1-\mu(1+\rho))(1-2\rho)}{(1-\mu)(1-\rho)}, & \text{if } \mu \in (\frac{1}{2}, 1), \end{cases} \\
&> 0,
\end{aligned}$$

given (25), ensuring therefore that (80) holds. To complete the proof, it remains to verify that each type of firm obtains an expected payoff ρ in the equilibrium, and each type does not find it profitable to deviate from the specified participation decision. These steps are identical to those in the proof of Proposition 2 and are therefore omitted.

B.10 Proof of Proposition 10

In the setting with both adverse selection and moral hazard, the firm's incentive constraint for high effort after each history h_f is (39). Analogously from the proof of Lemma 3,

$$V(\bar{y}; h_f) \leq 1 - \rho - c - \frac{\rho c}{1 - 2\rho}. \quad (83)$$

Further, $V(y; h_f) \geq \rho$, because the firm can guarantee a continuation profit ρ by consistently shirking. The incentive constraint (39) therefore implies

$$\delta(1 - 2\rho) \left(1 - 2\rho - c - \frac{\rho c}{1 - 2\rho} \right) \geq (1 - \delta)c,$$

must hold. Because the left side strictly decreases in c and the right side strictly increases in c and they are equal at $c = \bar{c}_{MH}$, the incentive constraint for high effort must be violated whenever $c > \bar{c}_{MH}$.

Thus, in any equilibrium when $c > \bar{c}_{MH}$, the firm's profit equals ρ . Respecting individual rationality for participation, the rater must compensate a participating firm with at least ρ and thus obtain a payoff at most 0. Since the rater can guarantee a payoff of at least 0 (*e.g.*, by compensating a participating firm with exactly ρ), the rater obtains a payoff 0.

B.11 Proof of Proposition 11

To prove the claim, without loss of generality, label the two possible ratings in the competent scheme as 0 and 1. Let α_r be the probability of reaching rating 1 conditional on a good signal and the present rating being r in the competent scheme. Similarly, let β_r be the counterpart conditional on a bad

quality. Fix a profile (σ, φ) in which the competent firm participates in the competent scheme ξ_C with probability one, and it consistently exerts high effort, and that the inept firm participates in the inept scheme ξ_I with some positive probability. The profile and the transition probabilities induce a stationary distribution λ_0^C, λ_1^C over the ratings 0 and 1. Fix the transition probabilities so that the continuation value by the competent firm in the competent scheme at state 1 is strictly higher than the continuation value at state 0 (1 is the good rating): $V_1 > V_0$. Let $(\lambda_0^I, \lambda_1^I)$ denote the stationary distribution over the ratings in the inept scheme. The price $p^\sigma(r)$ upon each rating r is

$$p^\sigma(1) = \mu + (1 - 2\rho) \frac{\mu \lambda_1^C}{\mu \lambda_1^C + (1 - \mu) \lambda_1^I},$$

$$p^\sigma(0) = \mu + (1 - 2\rho) \frac{\mu(1 - \lambda_1^C)}{\mu(1 - \lambda_1^C) + (1 - \mu)(1 - \lambda_1^I)}.$$

The equilibrium continuation profit V_r^σ of the competent firm upon participation and rating r satisfies

$$V_1^\sigma = (1 - \delta)(p_1^\sigma - c) + \delta(((1 - \rho)\alpha_1 + \rho\beta_1)V_1^\sigma + (1 - (1 - \rho)\alpha_1 - \rho\beta_1)V_0^\sigma) \quad (84)$$

$$V_0^\sigma = (1 - \delta)(p_0^\sigma - c) + \delta(((1 - \rho)\alpha_0 + \rho\beta_0)V_1^\sigma + (1 - (1 - \rho)\alpha_0 - \rho\beta_0)V_0^\sigma) \quad (85)$$

Equilibrium consistent high effort requires that the incentive constraint for high effort upon each rating r holds:

$$\delta(1 - 2\rho)(\alpha_r - \beta_r)(V_1^\sigma - V_0^\sigma) \geq (1 - \delta)c.$$

Solving for V_1^σ and V_0^σ from (84) and (85), the incentive constraints can be written as

$$\min_{r=0,1} \left[\frac{\delta(1 - 2\rho)(\alpha_r - \beta_r)(p^\sigma(1) - p^\sigma(0))}{1 - \delta((\alpha_1 - \alpha_0)(1 - \rho) + (\beta_1 - \beta_0)\rho)} \right] \geq c. \quad (86)$$

We choose parameters $(\alpha_0, \alpha_1, \beta_0, \beta_1, \lambda_0^I, \lambda_1^I)$ to maximize the left side of (86) to obtain the maximal value of cost c such that (86) holds. It is easily verified that the left side of (86) is maximized at

$$(\alpha_0, \alpha_1, \beta_0, \beta_1, \lambda_0^I, \lambda_1^I) = (1, 1, 0, 0, 1, 0),$$

giving the value

$$\delta(1 - 2\rho)^2 \left(\frac{1 - \mu}{1 - \mu + \mu\rho} \right) = \bar{c}(\mu, \delta, \rho).$$

B.12 Proof of Proposition 12

Fix the menu Ξ_{MH} and an equilibrium in which only the competent participates in the menu with positive probability. In addition, it participates with probability one. The competent firm's effort strategy follows τ_{MH} . As argued before (for example in the proof of Proposition 5), if it chooses the outside option, then it always exerts low effort. By Lemma 3, the rater's expected payoff in an equilibrium in which only the competent firm participates with positive probability is bounded above by (31). It is analogous to the proof of Proposition 5 to show that the prescribed profile constitutes an equilibrium. In addition, the rater's expected payoff in this equilibrium given the menu Ξ_{MH} is precisely (31), completing the proof.

B.13 Proof of Proposition 13

Notice first that an identical proof of Lemma 2 holds with $\rho = 0$, so that $W_2 \leq \mu(1 - c)$, where W_2 is defined by (47). In addition, given the menu Ξ^* , the profile (σ^*, φ^*) remains an equilibrium whenever (9) holds at $\rho = 0$, *i.e.* whenever $c \leq \delta$, and $W(\Xi^*, (\sigma^*, \varphi^*)) = \mu(1 - c)$. Because the profile features participation by both types with positive probability, it follows that:

Lemma 5. *If $c \leq \delta$, then $W_2 = \mu(1 - c)$.*

Next, consider equilibria in which only one type participates with positive probability:

Lemma 6. *It holds that $W_1 \leq \mu(1 - c)$, where W_1 is defined by (46).*

Proof. Fix a menu Ξ and a candidate equilibrium $(\sigma, \varphi) \in B(\Xi)$ with participation $\pi(\Xi|I) > 0$ and $\pi(\Xi|C) = 0$. Consumers believe that a participating firm is inept for sure, so that $\varphi_t(r) = 0$ and $p_t^r(r) = 0$ for any non-null rating r and any period t . The outside option payoff of the inept firm is at least $\rho = 0$ (Remark 5). Respecting the inept type's participation constraint $U(\sigma, I; \xi_\theta) - (1 - \delta)f_\theta \geq U(\sigma, I; N) \geq 0$, the rater's expected payoff satisfies

$$(1 - \delta)(1 - \mu) \sum_{\theta} \pi(\xi_\theta|I) f_\theta \leq (1 - \mu) \sum_{\theta} U(\sigma, I; \xi_\theta) = 0.$$

Next, fix a menu Ξ and a candidate equilibrium $(\sigma, \varphi) \in B(\Xi)$ with participation $\pi(\Xi|C) > 0$ and $\pi(\Xi|I) = 0$. Consumer beliefs satisfy $\varphi_t(r) = 1$ for each non-null rating r and each period t . We show that upon participating in a scheme ξ_θ ,

$$U(\sigma, C; \xi_\theta) \leq 1 - c. \tag{87}$$

The competent firm extracts the maximal amount from consumers when consumers believe the firm to be competent for sure and expect the firm to consistently exert high effort, paying the firm $1 - c$ each period. This establishes (87). Finally, and trivially, in any equilibrium in which neither type participates, the rater's payoff is zero. Because the firm is competent with probability μ , $W_1 \leq \mu(1 - c)$. ■

I now show that if $c \leq \delta$, $W_1 = \mu(1 - c)$. I construct a menu $\Xi^p = \{\xi^p\}$, where the scheme $\xi^p = (f^p, R^p, S^p)$ specifies that

$$\begin{aligned} f^p &= \frac{1 - c}{1 - \delta}, \\ R^p &= \{0, 1\}, \\ S^p(h_r^t) &= \begin{cases} 1 \circ \{1\}, & \text{if there does not exist a } s < t \text{ such that } y_s = \underline{y}, \\ 1 \circ \{0\}, & \text{otherwise.} \end{cases} \end{aligned}$$

In words, so long as the participating firm never delivered a bad signal, the rating system announces rating 1. It sends rating 0 otherwise. Consider an equilibrium $(\sigma^p, \varphi^p) \in B(\Xi^p)$, where $\sigma^p = (\pi^p, \tau^p)$, in which the competent firm participates in ξ^p with probability one, the inept firm chooses the outside option with probability one, and the competent firm exerts high effort upon participation and receiving rating 1, and shirks otherwise:

$$\begin{aligned} \pi^p(\xi^p|C) &= 1, \quad \pi^p(\xi^p|I) = 0, \\ \tau^p(h_f^t, C) &= \begin{cases} 1, & \text{if } d = \xi^p \text{ and } r_t = 1, \\ 0, & \text{otherwise.} \end{cases} \end{aligned}$$

Lemma 7. $(\sigma^p, \varphi^p) \in B(\Xi^p)$ and $W(\Xi^p, (\sigma^p, \varphi^p)) = \mu(1 - c)$.

Proof. Direct computation reveals that

$$W(\Xi^p, (\sigma^p, \varphi^p)) = \mu(1 - \delta)f^p = \mu(1 - c).$$

I show that $(\sigma^p, \varphi^p) \in B(\Xi^p)$. Analogous to the proof of Proposition 2, the play upon participation by the competent firm can be represented by a two-state automaton, with states given by the ratings $r \in \{0, 1\}$ and transitions depending on the current signals, depicted by Figure 14 below. State r collects the competent firm's histories with the most recent realized rating being r . It is easy to verify that (σ^p, φ^p) is an equilibrium using the automaton. Denoting the competent firm's continuation profit extraction from consumers in state $r \in \{0, 1\}$ by $V_r^{\sigma^p}$, it holds that

$$V_1^{\sigma^p} = (1 - \delta)(1 - c) + \delta V_1^{\sigma^p}, \text{ and } V_0^{\sigma^p} = 0.$$

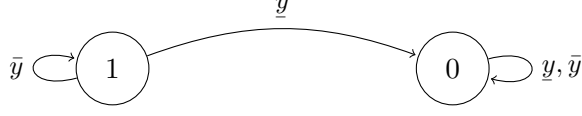


Figure 14: The two-state automation

The initial state is 1, in which the competent firm chooses high effort with probability one; in state 0, the competent firm chooses low effort with probability one.

Solving the system yields $V_1^{\sigma^p} = 1 - c$. The incentive constraint in state 1, namely $\delta(V_1^{\sigma^p} - V_0^{\sigma^p}) \geq (1 - \delta)c$, is satisfied whenever $c \leq \delta$. Also, clearly, the competent firm finds it optimal to exert low effort in state 0. It remains to verify that both types do not find it profitable to deviate from the specified participation choice. By deviating to the outside option, the competent firm finds it optimal to always exert low effort, because future consumers do not acquire any information about its past play. Each consumer thus pays the firm $\rho = 0$ upon a null rating, giving $U(\sigma^p, C; N) = 0$. Given the fee f^p in the scheme, by choosing ξ^p the competent firm obtains a payoff

$$U(\sigma^p, C; \xi^p) - (1 - \delta)f^p = 1 - c - (1 - c) = 0 = U(\sigma^p, C; N).$$

Similarly, by not participating, the inept firm obtains $U(\sigma^p, I; N) = 0$. If it deviates to participate in ξ^p , then it obtains

$$U(\sigma^p, I; \xi^p) - (1 - \delta)f^p = (1 - \delta)1 - (1 - c) \leq 0 = U(\sigma^p, I; N).$$

where the inequality follows given $c \leq \delta$. Therefore $(\sigma^p, \varphi^p) \in B(\Xi^p)$. ■

It remains to argue that the rater achieves the payoff upper bound if and only if the stated conditions in the proposition hold. Consider first equilibria in which each type participates in some scheme with positive probability. It follows analogously from Proposition 1 that the rater achieves the upper bound if and only if the first set of the conditions hold. Next, observe that in any equilibrium in which only one type participates with positive probability, the rater obtains a payoff $\mu(1 - c)$ if and only if first, the competent type participates with probability one; second, it consistently exerts high effort upon participation, and the rating fees in the schemes in which the competent type participates with positive probability fully extract all its profits above the minimal compensation ρ . This completes the proof.

References

- ADMATI, A. R. AND PFLEIDERER, P. 1986. A monopolistic market for information. *Journal of Economic Theory* 39:400–438.
- ADMATI, A. R. AND PFLEIDERER, P. 1988. Selling and trading on information in financial markets. *The American Economic Review* 78:96–103.
- ALBANO, G. L. AND LIZZERI, A. 2001. Strategic certification and provision of quality. *International economic review* 42:267–283.
- BAR-ISAAC, H. AND TADELIS, S. 2008. Seller reputation. *Foundations and Trends® in Microeconomics* 4:273–351.
- BELLEFLAMME, P. AND PEITZ, M. 2018. Inside the engine room of digital platforms: Reviews, ratings, and recommendations.
- BETTER BUSINESS BUREAU 2013. Better business bureau expels los angeles organization for failure to meet standards.
- BOARD, S. AND MEYER-TER VEHN, M. 2013. Reputation for quality. *Econometrica* 81:2381–2462.
- BOOT, A. W., MILBOURN, T. T., AND SCHMEITS, A. 2005. Credit ratings as coordination mechanisms. *The Review of Financial Studies* 19:81–118.
- CHE, Y.-K. AND HÖRNER, J. 2017. Recommender systems as mechanisms for social learning. *The Quarterly Journal of Economics* 133:871–925.
- COFFEE, J. C. 1997. Brave new world?: The impact (s) of the internet on modern securities regulation. *The Business Lawyer* pp. 1195–1233.
- CRIPPS, M. W., MAILATH, G. J., AND SAMUELSON, L. 2004. Imperfect monitoring and impermanent reputations. *Econometrica* 72:407–432.
- CROWE, K. 2018. Who’s rating doctors on ratemds? the invisible hand of ‘reputation management’.
- DELLACAROS, C. 2005. Reputation mechanisms.
- DELLAROCAS, C. 2005. Reputation mechanism design in online trading environments with pure moral hazard. *Information Systems Research* 16:209–230.

- EKMEKCI, M. 2011. Sustainable reputations with rating systems. *Journal of Economic Theory* 146:479–503.
- FARHI, E., LERNER, J., AND TIROLE, J. 2013. Fear of rejection? tiered certification and transparency. *The RAND Journal of Economics* 44:610–631.
- FICKENSCHER, L. 2017. Hyatt threatens to check out of Expedia.
- FITZPATRICK, T. J. AND SAGERS, C. 2009. Faith-based financial regulation: A primer on oversight of credit rating organizations. *Admin. L. Rev.* 61:557.
- FLEMING, T. 2010. Pay for play scandal at the better business bureau leads to consumer mistrust of the business rating organization. *Loy. Consumer L. Rev.* 23:445.
- FUDENBERG, D. AND LEVINE, D. 1992. Maintaining a reputation when strategies are imperfectly observed. *The Review of Economic Studies* 59:561–579.
- FUDENBERG, D. AND LEVINE, D. K. 1994. Efficiency and observability with long-run and short-run players. *Journal of Economic Theory* 62:103–135.
- GONZALEZ, F., HAAS, F., PERSSON, M., TOLEDO, L., VIOLI, R., WIELAND, M., AND ZINS, C. 2004. Market dynamics associated with credit ratings: a literature review.
- HARBAUGH, R. AND RASMUSEN, E. 2018. Coarse grades: Informing the public by withholding information. *American Economic Journal: Microeconomics* 10:210–35.
- HENRY, Z. 2014. Secrets of a very opaque glassdoor.
- HOLMSTRÖM, B. 1999. Managerial incentive problems: A dynamic perspective. *The review of Economic studies* 66:169–182.
- HÖRNER, J. AND LAMBERT, N. S. 2018. Motivational ratings.
- KASHYAP, A. K. AND KOVRIJNYKH, N. 2015. Who should pay for credit ratings and how? *The Review of Financial Studies* 29:420–456.
- KUGEL, S. 2016. Do those travel search results look fishy? here’s why. *The New York Times* .

- LEVICH, R. M., MAJNONI, G., AND REINHART, C. 2012. Ratings, rating agencies and the global financial system, volume 9. Springer Science & Business Media.
- LIU, Q. 2011. Information acquisition and reputation dynamics. *The Review of Economic Studies* 78:1400–1425.
- LIU, Q. AND SKRZYPACZ, A. 2014. Limited records and reputation bubbles. *Journal of Economic Theory* 151:2–29.
- LIZZERI, A. 1999. Information revelation and certification intermediaries. *The RAND Journal of Economics* pp. 214–231.
- LUCA, M. 2016. Reviews, reputation, and revenue: The case of yelp. com.
- MACEY, J. R. 2006. The politicization of american corporate governance. *Va. L. & Bus. Rev.* 1:10.
- MAILATH, G. AND SAMUELSON, L. 2001. Who wants a good reputation? *The Review of Economic Studies* 68:415–441.
- MAILATH, G. J. AND SAMUELSON, L. 2015. Reputations in repeated games, pp. 165–238. *In Handbook of Game Theory with Economic Applications*, volume 4. Elsevier.
- MARRIAGE, M. AND THOMPSON, B. 2018. Uk watchdog launches investigation into audit market. *Financial Times* .
- PARTNOY, F. 2006. How and why credit rating agencies are not like other gatekeepers.
- PEI, H. 2016. Reputation with strategic information disclosure. *Available at SSRN* .
- PORTES, R. 2008. Ratings agency reform. *The First Global Financial Crisis of the 21st Century* p. 145.
- RHEE, R. J. 2015. Why credit rating agencies exist. *Economic Notes: Review of Banking, Finance and Monetary Economics* 44:161–176.
- ROSS, S. M. 2014. Introduction to stochastic dynamic programming. Academic press.
- SALVADÈ, F. 2014. Why does a credit rating withdrawal matter. Technical report, Working Paper.

- SALVADÈ, F. 2017. The event of credit rating withdrawal: What happened? what follow? Technical report, Working Paper.
- SANGIORGI, F. AND SPATT, C. 2017. The economics of credit rating agencies. *Foundations and Trends® in Finance* 12:1–116.
- TADELIS, S. 1999. What's in a name? reputation as a tradeable asset. *The American Economic Review* 89:548–563.