

Squintani, Francesco

Working Paper

Moral Hazard

Discussion Paper, No. 1269

Provided in Cooperation with:

Kellogg School of Management - Center for Mathematical Studies in Economics and Management Science, Northwestern University

Suggested Citation: Squintani, Francesco (1999) : Moral Hazard, Discussion Paper, No. 1269, Northwestern University, Kellogg School of Management, Center for Mathematical Studies in Economics and Management Science, Evanston, IL

This Version is available at:

<https://hdl.handle.net/10419/221625>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Discussion Paper No. 1269

Moral Hazard, Renegotiation and Forgetfulness

by
Francesco Squintani
Northwestern University^{*}

July 1999

<http://www.kellogg.nwu.edu/research/math>

Abstract

When a principal and an agent operate with simple contracts, at equilibrium, renegotiation will occur after the action is taken. Also, since renegotiation makes incentive contracts non-credible, the principal may prefer non-renegotiable monitoring options. Current literature does not fully reconcile these predictions with the observation of simple non-renegotiated incentive contracts. We model a principal-agent interaction in a social learning framework, and assume that when renegotiation is not observed, players may forget its feasibility, with *infinitesimal* probability. The *unique* stable state of our model predicts that the second-best simple incentive contracts occur with *non-negligible* positive frequency.

^{*} Department of Economics, 2003 Sheridan Rd. Evanston, IL, 60208-2009.

E-mail: squinzio@nwu.edu. I would like to thank Eddie Dekel, Jeff Ely, Steven Matthews, John Panzar, Lars Stole, and Juuso Valimaki. The usual caveat applies.

1 Introduction

This paper characterizes the stable outcome of a simple social learning model of a moral-hazard game. We introduce a minimal departure from standard assumptions: if renegotiation is not observed its feasibility may be forgotten with infinitesimal probability. We suggest that this exercise may be helpful in refining current literature’s predictions with respect to simple moral-hazard scenarios.

In the classic formulation (Mirrlees 1976) of the two-action principal-agent problem, for the relevant parameter values, the agent takes high effort when offered a second-best contract, which is an incentive scheme that makes her indifferent between taking high or low effort, and rejecting the offer. However, if the principal can renegotiate the contract after the agent has chosen her effort and before output is realized, Fudenberg and Tirole (1990) insightfully show that the second-best contract will be renegotiated, and thus will not elicit high effort. When the parties sign a “simple” contract (i.e, a contract that does not require the agent to report the action taken to the principal), one may show that the principal will offer renegotiation on the equilibrium path.¹ In fact, if contracts do not depend on messages, but just on the outcome, the only means to separate high-effort from low-effort workers is to offer a renegotiation that only the latter will accept.

Current literature does not fully reconcile this prediction with intuition about simple moral-hazard scenarios. Consider for instance, a contractor building a house for a young professional. The parties do not use a menu-contract, but rather a simple incentive scheme: if the project is delayed, the constructor is subject to a penalty. We believe that most often renegotiations are not proposed.

One may offer the conjecture that the principal forms a reputation that includes a

¹That result is folk-knowledge according to personal communication with Steve Matthews. We will formally prove it in the second section of this paper. Hermalin and Katz (1991) and Matthews (1995) prove the same point in different models.

commitment not to renegotiate incentive schemes, along the lines of the corporate-culture proposal by Kreps (1990). While the reputation framework (see Kreps and Wilson 1982, and Milgrom and Roberts 1982) plausibly explains the rarity of renegotiation in two-player repeated interactions, that need not be the case for complex societies modeled as random-matching games. In fact, when renegotiation is not observable by third parties, the existence of a reputation equilibrium requires the agent to report the principal when offered renegotiation.² But, as the principal makes the agent better-off when offering renegotiation, the agent should not respond by ruining the principal's name. It is instead intuitive that principal and agent will cooperate and secretly renegotiate the contract, to their mutual advantage.

This brings us back to the question of how to explain the rarity of renegotiation in simple moral-hazard scenarios. An implicit assumption in the original principal-agent model is that the principal will achieve a higher payoff by designing a (second-best) incentive scheme, than by suffering a dead-weight loss to monitor the agent's action. With the introduction of renegotiation, there may instead be scenarios in which the second-best outcome dominates the monitoring option, but implementing the latter is more profitable than the solution of Fudenberg and Tirole (1990). In fact monitoring is implicitly non-renegotiable when the dead-weight loss in profit is paid before, or while, the agent takes her action,³ thus, unlike renegotiation-proof contracts, it may elicit high effort. Therefore, monitoring can explain the rarity of renegotiation, but it does not account for the prevalence of second-best simple contracts.

In formulating our social learning model, we take a small departure from standard assumptions, based on the perception that the relevance of renegotiation is not im-

²For a folk theorem in random matching games, see Kandori (1992).

³Our homeowner, for instance, may go onto the site to monitor the progress of the construction. This form of monitoring is hardly renegotiable. Border and Sobel (1987), instead, consider situations in which the monitoring cost is paid after the player takes her action.

mediately apparent when faced with a principal-agent. Because of that, a model that includes renegotiation may be more difficult to formulate than a simpler moral-hazard model. Specifically, consider a large population of players that live two periods.⁴ At the second period, they are randomly paired to (myopically) play an interaction with a moral-hazard structure. At the first period they observe their parents play. The assignment in the role of principal or agent is anonymous, and independent over time.⁵

After being matched, each player in the population needs to construct a (subjective) model of the interaction, that model may or may not include the possibility to renegotiate contracts. If a player's model is incomplete, she will never propose a renegotiation or expect a renegotiation offer; but if she is offered renegotiation, we assume that she readily understands its value, and accepts the offer if that is convenient to her. Each player enters into play holding her parents' model, unless one of the two following possibilities occurs. If the model of a player's parents is incomplete, but her parents are proposed a renegotiation, the player will immediately *learn* that renegotiation is possible, and include it in her interpretational model. At the same time, if a player's parents are aware of renegotiation, but they are not offered that option, with small probability, a player may *forget* that renegotiation is feasible.

As forgetfulness occurs with infinitesimal probability, we may believe that it is irrelevant for the analysis: all players are aware of renegotiation, at least in the long-run, so that the second-best contract is still precluded. To the contrary, full awareness cannot be a steady state, so that second-best contracts (and hard-working agents) will be observed with positive frequency in the long-run. In fact, when it is common knowledge that the entire population is aware of the feasibility of renegotiation, players in the role of principals choose to monitor. Their opponents do not observe renegotiation, and thus

⁴As customary with evolutionary game theory, the model should not be taken too literally.

⁵That assumption is for analytical tractability only. Our results hold as long as offsprings may be in a role different than their parents with positive probability.

their offsprings will forget that it is possible.

Once established that full awareness cannot be a stationary description of the society, we proceed to characterize the actual stable state. Since forgetful agents believe that they are playing a simple game without renegotiation (and thus that their opponents are forgetful as well), they will work hard when offered the second-best contract. For the same reason, forgetful principals will offer a second-best contract, and will not offer to renegotiate it. Aware players on the other hand, will condition their strategy on their conjecture with respect to their opponents' awareness type: they offer a second-best contract with renegotiation when assessing their opponents to be forgetful, and choose to monitor if assessing them to be aware. While players cannot observe their opponents' state of mind, following a fairly customary approach in evolutionary game theory,⁶ we do assume that aware players correctly assess the *population* distribution of types.

This paper is concerned only with long-run predictions, and thus our assumption can be motivated by noticing that, in the long-run, aware players should not be systematically incorrect when assessing average population awareness. Suppose in fact that each aware player adopts her parents' assessment. When forgetfulness is infinitesimal, the number of consecutive periods in which a complete model is maintained across generations is very large. After her period of play, any aware agent will be able to perfectly infer her opponent's awareness and use that information to update her assessment: in fact, she knows that her opponent chose a non-renegotiated second-best contract if and only if she was a forgetful player. The inference problem for an aware principal is more complex, but since role assignment is independent across generations, almost all aware players will hold approximately the same conjecture in the long-run, and such a conjecture cannot be systematically and significantly wrong.

⁶See for example Stahl (1993), Saez-Marti' and Weibull (1999), Dekel, Ely and Yilankaya (1998), or Heller (1999).

Under the above assumption we can pin down aware players' strategies as a function of the average population awareness, and we can complete our analysis. Whenever too many players are aware of the possibility of renegotiation, aware principals choose to monitor. Their opponents do not observe renegotiation, and thus they may forget it. When the ratio of forgetful players is large enough, aware principals play the second-best contract and renegotiate it to offer the full-insurance contract. The system is at rest only when the amount of aware players forgetting renegotiation equals the number of forgetful players recalling it. Forgetful principals offer non-renegotiated second-best contracts, that thus appear with stable significantly positive frequency.

Somewhat unexpectedly, the stable frequency of second-best contracts is increasing in the payoff of the monitoring option. In fact, that frequency coincides with the value that makes aware principals indifferent between the second-best contract and the monitoring device. When the latter is more valuable, aware principals will choose to monitor for lower ratios of forgetful players in the population.

As in other works (e.g. Stahl 1993, or Saez Marti' and Weibull 1999), the different types in the population do not identify a strategy, but rather, a mental model. While this paper is concerned with different interpretational models of the game, the other works has analyzed different bounded levels of rationalizability. Unlike standard evolutionary analysis, our learning dynamics are not payoff-monotone,⁷ and the diffusion of a strategy in the population depends by how often it is used, rather than by its payoff.

The paper is presented as follows. The second section presents a traditional treatment of moral-hazard with simple contracts. The third section presents our social learning model. The fourth section derives the results. Some of the proofs are in the Appendix.

⁷Payoff-monotone dynamics were first analyzed by Nachbar (1990). Schlag (1998) formally derives imitative payoff-monotone dynamics. For a general review of evolutionary literature, see Fudenberg and Levine (1998), and Weibull (1995).

2 Moral-Hazard with Simple Contracts

A principal P may motivate a prospective agent A through an incentive contract C or through the use of a monitoring device M . The agent accepts, y , or refuses, n , the principal's proposal and then takes a privately observed action a , consisting of high effort H or low effort L . That action will influence the probability p_a that a high h or a low l output will be produced to the principal. High output will be more likely under high effort (i.e. $p_H > p_L$). The incentive contract C prescribes agent's compensation profile (c_h, c_l) dependent on the output realized. When using the monitoring device, the principal pays k to know the action taken by the agent, and compensates her with the profile $M = (m_H, m_L)$ dependent on her action.

The utility of the agent is $V = U(c) - e(a)$, where $e(H) = e > e(L) = 0$, c is the compensation received, and the reservation utility is $\bar{V} = 0$. The function $U(\cdot)$ is strictly increasing and strictly concave, while the agent is a risk-neutral profit-maximizer. When the parties sign an incentive contract, the agent's Von Neumann-Morgenstern expected utility, and the principal's profit are

$$\begin{aligned} V(C, a) &= (1 - p_a)U(c_l) + p_a U(c_h) - e(a) \\ \Pi(C, a) &= (1 - p_a)(l - c_l) + p_a(h - c_h). \end{aligned}$$

If the principal proposes and the agent accepts a monitoring device, the payoffs are

$$\begin{aligned} V(M, a) &= U(m_a) - e(a) \\ \Pi(M, a) &= (1 - p_a)l + p_a h - m_a - k. \end{aligned}$$

To make the problem non-trivial, we assume the principal to prefer to motivate the agent to work hard, or to monitor her, over letting her shirk and giving her no compensation (we denote such a contract by $C = 0$). That is, there exist a contract C s.t. $\Pi(C, H) >$

$\Pi(0, L)$ and $V(C, H) \geq V(C, L)$, and there exist a compensation M s.t. $V(M, H) \geq V(M, L)$ and $\Pi(M, H) > \Pi(0, L)$.

A renegotiation R is a new contract, proposed after the action is taken and before the output is realized, that assigns new wages (r_l, r_h) . If the agent accepts the renegotiation, the payoffs will be $V(R, a)$ and $\Pi(R, a)$. The option not to renegotiate is denoted by N . Unlike the incentive contract, the monitoring option is not renegotiable, because the dead-weight cost is paid before the agents chooses effort.

First we consider the case when the principal may not propose renegotiation. In equilibrium, it is well known that the principal offers the agent a “second-best contract” $C^* = (c_h^*, c_l^*)$: an incentive scheme that guarantees her reservation utility if she takes high effort, and that makes her indifferent between low and high effort:

$$V(C^*, L) := (1 - p_L)U(c_l^*) + p_L U(c_h^*) = V(C^*, H) := (1 - p_H)U(c_l^*) + p_H U(c_h^*) - e = 0.$$

The agent accepts the contract C^* and chooses high effort (H). By construction,

$$\Pi(C^*, H) > \Pi(0, L) > \Pi(C^*, L).$$

When renegotiation is allowed, the second-best outcome cannot be achieved at equilibrium anymore. Consider the principal’s decision after contract C^* has been proposed and accepted, and the agent has taken her action. The agent has played H , thus the principal knows that the output h will occur with probability p_H . Moreover, the agent’s certainty equivalent $U^{-1}(e)$ is less than her expected compensation $(1 - p_H)c_l^* + p_H c_h^*$ because she is risk averse. Thus the principal can increase her expected profit by renegotiating the second-best contract C^* and offering the agent full-insurance $R^* := (U^{-1}(e), U^{-1}(e))$ before the output is revealed. The profit of the full-insurance contract R^* after the agent played H is

$$\Pi(R^*, H) = (1 - p_H)l + p_H h - U^{-1}(e) > \Pi(C^*, H)$$

because $U'' < 0$. The agent's utility is again $V(R^*, H) = 0$. But then, when offered the second-best contract C^* , the agent knows that the principal will eventually renegotiate it and propose R^* . If she plays L instead of H , she obtains

$$V(R^*, L) = e > V(R^*, H) = 0:$$

the principal gets “cheated” and her expected profit is

$$\Pi(R^*, L) = (1 - p_L)l + p_L h - U^{-1}(e).$$

At equilibrium, the principal first proposes a contract making the agent indifferent between working or shirking, the agent shirks with positive probability, and finally the principal offers in renegotiation a full-insurance contract that only the shirking agent accepts. The following proposition is proven in the Appendix.

Proposition 1 *At any Perfect Bayesian Equilibrium, the principal initially proposes the second-best contract $C^* = (c_l^*, c_h^*)$. For any initial contract C s.t. $V(C, H) = V(C, L)$, the agent chooses the high effort action with probability $\sigma_A(H|C) \in [0, \frac{U^{-1}(e)}{(1-p_H)c_l + p_H c_h}]$. For any other contract C , $\sigma_A(H|C) = 0$. The principal renegotiates offering $M = (U^{-1}(V(C, L)), U^{-1}(V(C, L)))$ the agent accepts the offer if and only if she has taken the low action.*

The equilibrium principal's payoff is:⁸

$$\Pi^* = \frac{\left(\begin{aligned} &U^{-1}(e)[(1 - p_H)(l - c_l^*) + p_H(h - c_h^*)] + \\ &[(1 - p_H)c_l^* + p_H c_h^* - U^{-1}(e)][(1 - p_L)l + p_L h - U^{-1}((1 - p_L)c_l^* + p_L c_h^*)] \end{aligned} \right)}{(1 - p_H)c_l^* + p_H c_h^*}.$$

It is easy to obtain that $\Pi(C^*, H) > \Pi^*$.

⁸As customary in the renegotiation literature, we select the equilibrium where the agent takes the most favourable action to the principal.

We finally consider the monitoring option. If, in equilibrium, the principal opts for monitoring, she will choose the contract

$$M^* = (m_H^*, m_L^*) \text{ such that } 0 = V(m_H^*, H) \geq V(m_L^*, L).$$

In fact by assumption, there exist a compensation M such that $V(M, H) \geq V(M, L)$ and $\Pi(M, H) > \Pi(0, L)$. By definition of M^* , $U(m_H^*) = e$ so M^* maximizes $\Pi(M, H)$. Suppose agent played L with positive probability after accepting M^* , or did not accept M^* . Then the principal's best-response would be empty, as $\forall m_H > m_H^*$, the agent's sequentially rational response is to accept M and work hard. We restrict attention to those games where $\Pi(M^*, H) > \Pi^*$: the principal's payoff for (optimally) monitoring the agent is larger than the payoff for the equilibrium of any subgame following an incentive contract with renegotiation. Thus the unique Perfect Bayesian Equilibrium is such that the principal plays M^* , and the agent replies by playing H (off-path beliefs are as specified in Proposition 1).

3 Social Learning Model

At any time t , a continuous population is randomly matched to play the principal-agent game with renegotiation and monitoring; we denote by ρ_t the fraction of players aware of renegotiation. Forgetful players believe that they are playing against forgetful players only. Aware players believe that they are matched with a aware opponent with probability ρ_t . We assume that, when presented with an unforeseen renegotiation offer, forgetful agents do not revise their probability assessment on the occurrence of h and l . In each period, we assume players play a Perfect Bayesian Equilibrium, restricted by their possibly partial model of the interaction. As is customary in evolutionary game theory, we assume myopic play⁹ Before we define the equilibrium requirement, it is

⁹It can be shown, however, that such an assumption is irrelevant for our results.

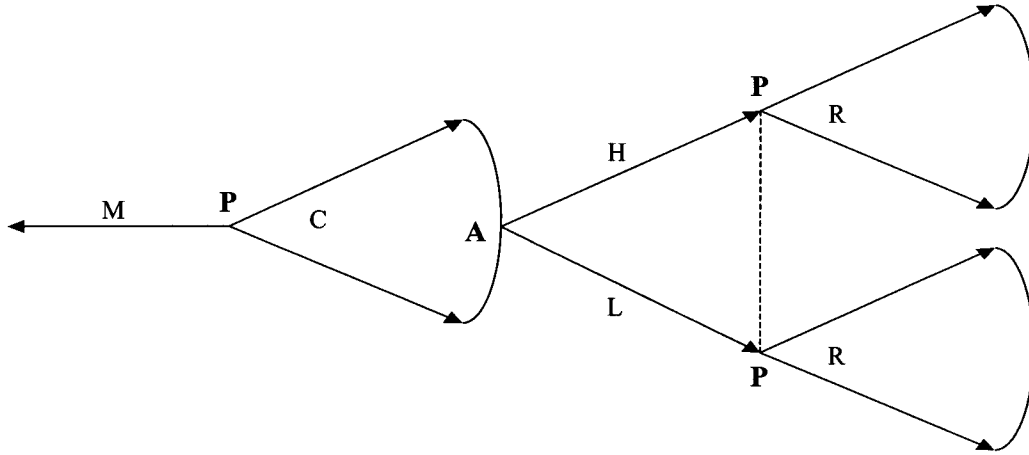


Figure 1: Reduced Game G

expositionally useful to slightly reduce the model, without loss of generality.

First, we note that by a standard argument, there are no equilibria where any agent refuses any contract that makes her indifferent. Thus we can assume that any principal offers only contract that make agents at least indifferent, and that her opponent accepts such contracts. That allows us to eliminate that agent's choice from the game tree. In particular, we are saying that if a forgetful agent is presented with an (unforeseen) renegotiation offer by a principal, she will accept it if at least as good as the initial contract. Secondly, we formulate the payoffs as expected payoff before nature decides whether output is h or l . Thus the principals' decision at the renegotiation stage map directly into the end-nodes. Let the reduced model be called game G , see Figure 1.

In order to model equilibrium play of a player with an incomplete description of the game, we shall make use of the Harsanyi model of games of incomplete information with subjective priors (see Harsanyi 1967), and specifically of a game-structural concept called *elaboration*, (see also Fudenberg, Kreps and Levine 1988). Informally, we shall assign to each unaware player the belief that if a renegotiation is ever proposed, both players

will be very harshly punished. With respect to equilibrium strategies, that is equivalent to impose that they will never offer renegotiation, or expect it to be offered. Aware players know that renegotiation is not punished at all, and believe to play against an aware opponent with probability ρ . In order to accommodate forgetful players incorrect beliefs, and aware players' second-order beliefs, we need to formulate a state space that independently specify each player's belief, and the nature's choice, on whether renegotiation is punished or not. That introduces the need for an elaboration of the game.

The formal definition of elaboration is as follows. Given a game Γ , a state space S , a type space T , and a system of (possibly subjective) priors p , an elaboration is constructed attaching a copy $\Gamma(s)$ of Γ to each $s \in S$, (possibly) assigning different payoffs to the end-nodes of different $\Gamma(s)$, and completing the information structure so as to maintain the type space T . In our model, the state space $S = \{000, 001, 010, 111\}$ will be appropriate. The states are expressed in binary notation. A 0-digit means that renegotiation is legal, a 1-digit that it is very harshly punished. The first digit represent nature's choice, the second one the principal's belief, the third one the agent's belief.

Definition 1 *Let the Augmented Game AG be the subjective-priors elaboration of G with state space S and the following specification.*

- *The principal's types are $A_P := \{000, 001\}$, and $F_P := \{010, 111\}$, and the principal's prior is $p_P(000) = \rho^2$, $p_P(001) = (1 - \rho)\rho$, $p_P(111) = 1 - \rho$.*
- *The agent's types and prior are symmetrically defined.*
- *Nature's choice¹⁰ is $p(000) = \rho^2$, $p(010) = p(001) = (1 - \rho)\rho$, $p(111) = (1 - \rho)^2$.*

¹⁰Even though renegotiation is never really punished, if both players believe it is, at equilibrium it is the same as if they were correct, so we may assign positive probability on the state 111, instead of introducing the state 011.

- *The payoffs of each $G(s)$ assign $V = -\infty, \Pi = -\infty$ to all paths where $R \neq C$, in all states s where the first digit is 1, otherwise they are the same as in G .*

In each period t , we assume the players to coordinate on a Perfect Bayes Equilibrium $(\sigma_{AP}, \sigma_{FP}; \sigma_{AA}, \sigma_{FA})$ of the augmented game AG under the type distribution ρ_t .

The evolution of the population is as follows. Each offspring observes her actions and payoffs. At the next round, the offspring will be randomly matched to play the game. Matching and role assignment are anonymous, and independent across generations. Each player enters into play holding her parents' model of the game, except in one of the two following cases. The offspring of a forgetful agent who receives a renegotiation offer will become aware. The offspring of an aware agent who does not receive a renegotiation offer, will become forgetful with probability ε . The offsprings of players in the role of principal do not observe directly the opponents' actions. We thus assume that they maintain their parents' model of the game. Alternatively we could assume that they forget the feasibility of renegotiation with probability ε if their parents do not offer it, our results would be unchanged under that specification.

The equilibrium given a type distribution, together with the evolution of the type distribution define a stochastic process that we can approximate with a deterministic system by Theorem 6.4, Alos-Ferrer (1999).¹¹ Our results concern the long-run aggregate distributions of play f , calculated by compounding the strategies played by the different types, with the distribution of types. Formally, let $\zeta_{(T,Q)} = \sigma_T \sigma_Q$ be the distribution induced on the terminal nodes of G by a principal of type T and an agent of type Q .

¹¹Alos-Ferrer 1999 constructs matching schemes under which a continuous population stochastic evolution may be approximated with a dynamic system. Boylan 1992, proposes a similar result for finite large population, with an argument often referred to as a "Law of Large Numbers" in evolutionary games.

The aggregate distribution of play is

$$f = \rho^2 \zeta_{(A,A)} + (1 - \rho) \rho [\zeta_{(F,A)} + \zeta_{(A,F)}] + (1 - \rho)^2 \zeta_{(F,F)}.$$

4 Long-Run Results

We now show that in the long run, the non-renegotiated second-best contract C^*C^* appears with significant positive frequency. Moreover, that frequency turns out to be increasing in the value of the monitoring option.

Proposition 2 *For any ε sufficiently small, the unique stable distribution of play f satisfies:*

$$\begin{aligned} f(C^*C^*) &\approx \frac{\Pi(M^*, H) - \Pi(R^*, L)}{\Pi(R^*, H) - \Pi(R^*, L)}, \quad \frac{\partial[f(C^*C^*)]}{\partial[\Pi(M^*, H)]} > 0, \\ f(C^*R^*) &\approx 0, \quad f(M^*) \approx \frac{\Pi(R^*, H) - \Pi(M^*, H)}{\Pi(R^*, H) - \Pi(R^*, L)}. \end{aligned}$$

The stable distribution is reached in finite time, regardless of the initial distribution.

A technical issue arises because the game AG allows for multiple Equilibria when $\rho = [\Pi(M^*, H) - \Pi(R^*, L)] / [\Pi(R^*, H) - \Pi(R^*, L)]$. Depending on the equilibria selected, a stable distribution of play may or may not exist. In the following proof we will nevertheless show that, even if a stable distribution does not exist, we can calculate the average frequency of non-renegotiated second-best contracts over time, and show that:

$$\lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^T f_t(C^*C^*) = \frac{\Pi(M^*, H) - \Pi(R^*, L)}{\Pi(R^*, H) - \Pi(R^*, L)}.$$

Proof of Proposition 2. From here onwards, for brevity, we shall use the following short-hand notation.

$$b = \Pi(R^*, H), \quad m = \Pi(M^*, H), \quad s = \Pi(R^*, L).$$

Lemma 3, proven in the Appendix, finds all the possible PBE of the game AG for any distribution of types ρ . As in Proposition 1, we select the equilibrium where the agent takes the action most favorable to the principal. For brevity we do not report agents' behavior on (off-path) subgames following contracts different from C^* and M^* , the only contracts ever offered. We let σ_H be the probability that a aware agent plays H under the second-best contract, and σ_M, σ_{CR} , the probabilities that a aware principal offers respectively monitoring M^* ; or the second-best contract, with renegotiation to full-insurance (C^*R^*). Denote by (C^*C^*) the choice of offering a second-best contract without renegotiating it. We shall henceforth drop star subscripts, to avoid burdening the notation.

Lemma 3 *In all Perfect Bayes Equilibria, all agents play H if monitored. Forgetful principals choose CC . Forgetful agents choose H when offered C . Aware players' choice depends on ρ as follows:*

$$\left\{ \begin{array}{ll} \sigma_{CR} = 1, \sigma_H = 1 & \text{if } \rho = 0 \\ \sigma_{CR} = 1, \sigma_H = 0 & \text{if } \rho \leq \frac{b-m}{b-s} \\ \sigma_{CR} \in [0, 1], \sigma_H = 0 & \text{if } \rho = \frac{b-m}{b-s} \\ \sigma_M = 1, \sigma_H = \frac{m-b}{\rho(b-s)} + 1 & \text{if } \rho \geq \frac{b-m}{b-s} \end{array} \right. \quad (1)$$

The aggregate principal's distribution of play at any period t is therefore:

$$f_t(M) = \rho_t \sigma_M(\rho_t), f_t(CR) = \rho_t \sigma_{CR}(\rho_t), f_t(CC) = (1 - \rho_t) + \rho_t \sigma_{CC}(\rho_t).$$

The intuitive demonstration is as follows: forgetful principals offer the second-best contract, and forgetful agents respond by working hard. Both these types play as if renegotiation were not possible. Most importantly, aware principals do not initially offer any contract different from the second best or from the implementation of the monitoring device. The monitoring option dominates the equilibrium of the subgame starting after an incentive contract when the players recognize the possibility to renegotiate. If, with

large enough probability, the agent does not foresee renegotiation, the second-best contract is optimal. Once the forgetful agent is motivated to work hard, it is a dominant action to renegotiate and offer full insurance. Because of that, aware agents shirk with positive probability if offered the second-best contract. If there are too many aware agents, aware principals prefer to monitor.

After imposing our assumptions on individual transition, by Theorem 6.4, Alos-Ferrer (1999), we can approximate the stochastic evolution of the large population with a simple difference equation in ρ_t . For expositional purposes, we analyze the problem in a continuous-time, rather than discrete-time, dynamics. As is customary, see Hale (1969), that is done by subdividing the length of each time interval into k sub-intervals, and assuming that in each subinterval, a ratio $1/k$ of the players is called to play (so that all the population plays once in each time interval). Letting $k \rightarrow \infty$ one obtains the following differential equation.

$$\dot{\rho}_t = \frac{f_t(CR) - \varepsilon \rho_t [f_t(CC) + f_t(M)]}{2} \quad (2)$$

The dynamic analysis for any small ε is reported in Lemma 4 below, which we prove in the Appendix.

A technical detail arises because the evolution of the equation 2 at time t depends on f_t which in turns depends on the equilibrium $\sigma(\rho_t)$ characterized in Lemma 3. By Lemma 3, for $\rho = \frac{b-m}{b-s}$, any value $\sigma_{CR} \in [0, 1]$ may be obtained. In order to have a well-defined differential equation, one needs to select a single $\sigma_{CR}(\frac{b-m}{b-s})$.

Lemma 4 *Let $0 < \varepsilon < 1/2$. If $\sigma_{CR}(\frac{b-m}{b-s}) = \frac{\varepsilon[b-s]}{m(1-\varepsilon)+\varepsilon b-s}$, then $\rho = \frac{b-m}{b-s}$ is the unique stable state, and it is global attractor reached in finite time.*

Otherwise the system reaches a small neighborhood of $\rho = \frac{b-m}{b-s}$, from any initial point, in finite time, but does not have any stationary state.

Now we can derive the relative aggregate distribution of play. While in any case, the unique stable frequency of second-best non-renegotiated contracts is:

$$f(CC) = 1 - \rho = \frac{m - s}{b - s},$$

the values for $f(CR)$ and $f(M)$ depend on the particular σ_{CR} selected at $\rho = \frac{b-m}{b-s}$.

In the case that

$$\sigma_{CR} = \varepsilon \frac{b - s}{m(1 - \varepsilon) + \varepsilon b - s},$$

we obtain:

$$f(CR) = \rho \sigma_{CR} = \frac{\varepsilon(b - m)}{m(1 - \varepsilon) + \varepsilon b - s}, \quad f(M) = (1 - \rho) \sigma_{CR} = \frac{(1 - \varepsilon)(b - m)(m - s)}{(b - s)(m(1 - \varepsilon) + \varepsilon b - s)}.$$

In any other case, the actual values of $f(M)$ and $f(CR)$ are undetermined (Nevertheless, for $\sigma_{CR} > [\varepsilon b - \varepsilon c]/[m(1 - \varepsilon) + \varepsilon b - s]$ we show in the that the time-average of $f(M)$ is larger than 0).

Taking $\varepsilon \rightarrow 0$, we finally prove proposition 2.

5 Proofs

Proof of Proposition 1. First it is easy to convince one's self that there is no equilibrium in which $\sigma_A(H) = 1$.

Suppose otherwise: under the arbitrary, incentive compatible contract C ,

$$V(C, H) = (1 - p_H)U(c_l) + p_H U(c_h) - e$$

$$\Pi(C, H) = (1 - p_H)(l - c_l) + p_H(h - c_h).$$

The principal can offer the renegotiation $R = (U^{-1}(V(C, H)), U^{-1}(V(C, H)))$, the agent accepts it, and, as $U'' < 0$, the principal is better-off. However the only agent's best-reply vs. R is $\sigma_A(H|R) = 0$: contradiction.

For any initial contract C , let the agent's best-reply be $\sigma_A(H|C)$: when taking her decision the agent can guarantee herself utility $\bar{U} = U^{-1}(V(C, \sigma_A(H|C)))$. Thus for any initial contract C , there is an equilibrium in which $\sigma_A(H) = 0$ and $R^* = (U^{-1}(V(C, \sigma_A(H|C))), U^{-1}(V(C, \sigma_A(H|C))))$.

We finally consider the case for $\sigma_H \in (0, 1)$, requires the agent to best-respond to a contract X such that $V(X, H) = V(X, L)$.

Suppose first that X was the initial contract: $X = C$.

Given equilibrium strategy $\sigma_A(H)$, the principal is left the option to renegotiate, offering contract R . At the moment at which she offers the renegotiation, the agent knows what action she has taken, i.e. she knows the realization of her mixed strategy.

The principal's utility is thus

$$\begin{aligned} \Pi(\sigma_A(H)) &= \sigma_A(H)[\Pi(C, H)\chi_{(V(C, H) > V(R, H))} + \Pi(R, H)\chi_{(V(C, H) \leq V(R, H))}] + \\ &\quad (1 - \sigma_A(H))[\Pi(C, L)\chi_{(V(C, L) > V(R, L))} + \Pi(R, L)\chi_{(V(C, L) \leq V(R, L))}]. \end{aligned}$$

Due to the linearity of $\Pi(\cdot, \sigma_A(H))$ and the strict concavity of $V(\cdot, \sigma_A(H))$, it follows that whenever the principal prefers to renegotiate, she will choose a perfect insurance contract $R^* = (r^*, r^*)$. As $V(R^*, H) = U(r^*) - e$.

Therefore $V(R^*, H) \geq V(C, H)$ iff $r^* \geq U^{-1}(V(C, H)) + e$.

Analogously, as $V(R^*, L) = U(r^*)$, $V(R^*, L) \geq V(C, L)$ iff $r^* \geq U^{-1}(V(C, L))$.

Thus the solution of the principal's problem is such that she will always insure the low type, and insure the high type only when $\sigma_A(H)$ is high enough. That is, she offers $R_L^* : r_L^* = U^{-1}(V(C, L))$ if

$$\sigma_A(H)\Pi(C, H) + (1 - \sigma_A(H))\Pi(R_L^*, L) \geq \sigma_A(H)\Pi(R_H^*, H) + (1 - \sigma_A(H))\Pi(R_H^*, L)$$

and she offers $R_H^* : r_H^* = U^{-1}(V(C, H)) + e$ otherwise.

After some algebra, one obtains that the principal offers R_L^* iff

$$\sigma_A(H) \in [0, U^{-1}(e)/((1 - p_H)r_l + p_H r_h)].$$

Whenever R_H^* is offered, $V(R_H^*, H) - V(R_H^*, L) = e > 0$ thus it is not possible that the agent randomizes, after all. Instead, when R_L^* is offered, $V(R_L^*, L) = V(C, L) = V(C, H)$ and the agent's randomization is verified.

Now consider the case in which the original contract C does not have the agent randomize: $V(C, H) \neq V(C, L)$. It could still be possible to get the agent to randomize if the renegotiated contract allowed for randomization: $R = X$.

However we have just proven above that, when the agent randomizes, optimal renegotiation results in an insurance contract. As there does not exist any $X = (x_h, x_l)$ s.t. $V(X, H) = V(X, L)$ and $x_h = x_l$, it follows that the only equilibrium for these subgames after C is s.t. $\sigma_A(H|C) = 0$.

Now we need to consider the principal's choice for the initial contract C .

We assume that in any subgame, the players play the equilibrium most favorable to the principal: $\sigma_A(H|C) = \bar{\sigma}_A(H|C) = U^{-1}(e)/((1 - p_H)c_l + p_H c_h)$.

After some substitutions, we obtain that the principal initially offers a contract $C = (c_l, c_h)$, that solves

$$\max_{\{c_l, c_h: V(C, H) = V(C, L) \geq 0\}} \bar{\sigma}_A(H|C)[(1 - p_H)(l - c_l) + p_H(h - c_h)] + (1 - \bar{\sigma}_A(H|C))[(1 - p_L)l + p_L h - U^{-1}(V(C, L))].$$

As $V(C, H) = V(C, L)$ translates as $c_h = U^{-1}(e/[p_H - p_L] + U(c_l))$, it follows that the problem can be rewritten as

$$\max_{\{c_l: (1 - p_L)U(c_l) + p_L(e/[p_H - p_L] + U(c_l)) \geq 0\}} U^{-1}(e)/[(1 - p_H)c_l + p_H U^{-1}(e/[p_H - p_L] + U(c_l))]$$

which is maximized for

$$c_l : (1 - p_L)U(c_l) + p_L(e/[p_H - p_L] + U(c_l)) = 0$$

That solves for the second-best contract C^* .

Proof of Lemma 3. The discussion in the second section yields the behavior of forgetful types, and implies that we can rule out any monitoring contract other than M^* and say that the agent plays H after M^* .

By construction, M^* strictly dominates all initial contracts $C : V(C, L) > V(C, H)$ unless they are renegotiated into an incentive compatible contract.

Thus we can restrict to initial contracts $C : V(C, L) \leq V(C, H)$.

As in the proof of Lemma 1, the principal renegotiates such contracts: she either sell insurance to the agents who took L or to both types of agents.

Let the population distribution of play of action H be $\mu_H := 1 - \rho + \rho\sigma_H$.

After an incentive contract C , the principal prefers to renegotiate and to offer full insurance to the low type if and only if $\mu_H \leq U^{-1}(e)/((1 - p_H)c_l + p_H c_h)$. However, when that condition holds, the principal prefers to play M^* than to offer an incentive contract C and then renegotiate to the low type.

So, in equilibrium, when an aware principal chooses an incentive contract, she then renegotiates it to result in full insurance for both types (we denote that action by R^* in this proof). Given that, her optimal incentive contract is the second-best contract C^* .

Given that she observed C^* , each aware agent believes her opponent to play R^* with probability

$$\mu_{R^*} = [\rho\sigma_P(C^*)\sigma_P(R^*)]/[1 - \rho + \rho\sigma_P(C^*)].$$

She plays H if $\mu_{R^*} < 0$, she is indifferent for $\mu_{R^*} = 0$ and takes L for $\mu_{R^*} > 0$.

Set

$$y = \Pi(C^*, H), \quad z = \Pi(C^*, L),$$

$$\Pi(C^*) = \mu_H(\sigma_P(R^*)\Pi(R^*, H) + (1 - \sigma_P(R^*))y) + (1 - \mu_H)(\sigma_P(R^*)s + (1 - \sigma_P(R^*))l),$$

the principal offers C^* if $\Pi(C^*) > m$, plays M^* when $\Pi(C^*) < m$ and is indifferent otherwise.

Therefore, we have shown that the only contracts ever offered will be C^*C^* , C^*R^* , and M^* . Specifically, we obtain the following characterization for the set of equilibria (we indicate the ρ restrictions in square parenthesis). For the ease of the reader we omit stars.

case 0	$\sigma_H \in [0, 1]$ $[\rho = 0]$	$\sigma_{CR} = 1$ $[(1 - \rho + \rho\sigma_H)b + \rho(1 - \sigma_H)s \geq m]$
case 1	$\sigma_H = 0$ $[\rho \geq 0]$	$\sigma_{CR} = 1$ $[(1 - \rho)b + \rho s \geq m]$
case 2	$\sigma_H = 0$ $[\rho \geq 0]$	$\sigma_{CR} \in [0, 1]$ $[(1 - \rho)b + \rho s = m]$
case 3	$\sigma_H = 0$ $[\frac{0}{1-\rho} \geq 0]$	$\sigma_{CR} = 0$ $[(1 - \rho)b + \rho s \leq m]$
case 4	$\sigma_H \in [0, 1]$ $[\frac{0}{1-\rho} = 0]$	$\sigma_{CR} = 0$ $[(1 - \rho + \rho\sigma_H)b + \rho(1 - \sigma_H)s \leq m]$

Solving out the restriction on ρ that allows the different equilibria to exist, we obtain for case 0, $\rho = 0$, for case 1, $\rho \leq \min\{\frac{b-y}{b-y+z-s}, \frac{b-m}{b-s}\}$, for case 2, $\rho = \frac{b-m}{b-s}$ and $\frac{m-b}{\rho(b-s)} \geq \frac{y-b}{\rho(b-y+z-s)} + 1$, for case 3, $\rho \geq \frac{b-m}{b-s}$, and for case 4, $\rho \geq \frac{b-m}{b-s}$ and $\frac{m-b}{\rho(b-s)} \geq \frac{y-b}{\rho(b-y+z-s)} + 1$.

Substituting for ρ in the equilibria, we derive the following characterization.

case 0	$\sigma_H \in [\frac{m-b}{\rho(b-s)} + 1, 1]$	$\sigma_{CR} = 1$	$\rho = 0$
case 1	$\sigma_H = 0$	$\sigma_{CR} = 1$	$\rho \leq \frac{b-m}{b-s}$
case 2	$\sigma_H = 0$	$\sigma_{CR} \in [0, 1]$	$\rho = \frac{b-m}{b-s}$
case 3-4a	$\sigma_H \in [0, \frac{m-b}{\rho(b-s)} + 1]$	$\sigma_M = 1$	$\frac{b-m}{b-s} \leq \rho \leq \frac{b-y}{b-y+z-s}$
case 4b	$\sigma_H = \frac{m-b}{\rho(b-s)} + 1$	$\sigma_M = 1$	$\rho \geq \frac{b-y}{b-y+z-s}$

The Proposition is then proven by selecting the equilibria where the agent takes the action most favorable to the principal.

Proof of Lemma 4. The evolution of equation 2 at time t depends on f_t which in turns depends on the equilibrium σ_t characterized in Lemma 3. Since Lemma 3 yields different unique equilibria, for $\rho < \frac{b-m}{b-s}$, and for $\rho > \frac{b-m}{b-s}$, the equation is discontinuous. However, as the system is piece-wise continuous, one can apply standard techniques to the areas $\rho > \frac{b-m}{b-s}$ and $\rho < \frac{b-m}{b-s}$, and then complete the analysis considering the discontinuity set $\rho = \frac{b-m}{b-s}$. For $\rho \in (0, \frac{b-m}{b-s})$,

$$2\dot{\rho} = \rho - \varepsilon\rho[(1 - \rho)] > 0.$$

So that any state $\rho < \frac{b-m}{b-s}$ is unstable.

For $\rho \in (\frac{b-m}{b-s}, 1)$,

$$2\dot{\rho} = -\varepsilon\rho[(1 - \rho) + \rho] < 0$$

Note that $\lim_{\rho \uparrow \frac{b-m}{b-s}} \dot{\rho} > 0$, and $\lim_{\rho \downarrow \frac{b-m}{b-s}} \dot{\rho} < 0$. That is, the length of the gradient does not vanish around $\frac{b-m}{b-s}$, for any time-interval t of length 1. Also, for any k , the length of the gradient any at point ρ is smaller than $1/k$, for any time sub-interval of length $1/k$. That implies that for any $\epsilon > 0$, the open ball $B_\epsilon(\frac{b-m}{b-s})$ is reached in finite time, for k large enough. The same argument also implies that there exist a finite time (not necessarily the first time that $B_\epsilon(\frac{b-m}{b-s})$ is reached by the system) after which the system may never leave $B_\epsilon(\frac{b-m}{b-s})$.

So, the only candidate stable state left is $\rho = \frac{b-m}{b-s}$. By Lemma 3, for $\rho = \frac{b-m}{b-s}$, any value $\sigma_{CR} \in [0, 1]$ may obtain. In order to have a well-defined dynamic system, one needs to select one σ_{CR} . For the system to have a stable state, we need $\dot{\rho} = 0$ and $\rho = \frac{b-m}{b-s}$. We obtain $\sigma_{CR} = \frac{\varepsilon b - \varepsilon c}{m(1-\varepsilon) + \varepsilon b - s}$. When selecting that value, the above analysis shows that the state $\rho = \frac{b-m}{b-s}$, is stationary, stable and global attractor.

If $\sigma_{CR} \neq \frac{\varepsilon b - \varepsilon c}{m(1-\varepsilon) + \varepsilon b - s}$, the equation converges in finite time to $\rho = \frac{b-m}{b-s}$; but, each time it reaches that state, it is discontinuously “pushed” away: the state is not stationary. If $\sigma_{CR} > \frac{\varepsilon b - \varepsilon c}{m(1-\varepsilon) + \varepsilon b - s}$ in particular, the system is pushed in the region $\rho \in (\frac{b-m}{b-s}, 1)$, where $\sigma_{CR} = 0$.

References

- [1] Alos-Ferrer [1999]: “Dynamical Systems with a Continuum of Random of Randomly Matched Agents”, *Journal of Economic Theory* 86: 245-267
- [2] Boylan R. [1992]: “Laws of Large Numbers for Dynamical Systems with Randomly Matched Individuals” *Journal of Economic Theory* 57: 473-504
- [3] Border K. and J. Sobel [1987]: “Samurai Accountant: A Theory of Auditing and Plunder” *Review of Economic Studies* 54: 525-40
- [4] Dekel E, J. Ely, and O. Yilankaya [1998]: “The Evolution of Preferences”, *Northwestern University*, mimeo
- [5] Fudenberg D. and D.K. Levine [1998]: *The Theory of Learning in Games*, MIT Press
- [6] Fudenberg D, D. M. Kreps and D. K. Levine [1988]: “On the Robustness of Equilibrium Refinements”, *Journal of Economic Theory* 44: 354-380
- [7] Fudenberg D. and J. Tirole [1990]: “Moral Hazard and Renegotiation in Contracts” *Econometrica* 58: 1279-1319
- [8] Fudenberg D. and J. Tirole [1991]: *Game Theory* Cambridge MA, MIT press

- [9] Grossman S. and O. Hart [1983]: “An Analysis of the Principal-Agent Problem” *Econometrica* 51: 7-45
- [10] Hale J. [1969]: *Ordinary Differential Equations*, Wiley, NY
- [11] Hermalin B. E. and M.L. Katz [1991]: “Moral Hazard and Verifiability: the Effects of Renegotiation in Agency.” *Econometrica* 59: 1735-1753
- [12] Harsanyi J. [1967]: “Games with Incomplete Information Played by ‘Bayesian’ Players, part I, II, III” *Management Science* 14: 159-182, 320-334, 486-502
- [13] Heller D. [1999]: “An Evolutionary Analysis of the Returns to Learning in a Changing Environment” *University of Chicago* mimeo
- [14] Kandori M. [1992]: “Social Norms and Community Enforcement”, *Review of Economic Studies* 59: 63-80
- [15] Kreps D. [1990] “Corporate Culture and Economic Theory”, in *Perspectives on Positive Political Economy*, J. Alt and K. Shepsle eds. Cambridge University Press
- [16] Kreps D. and Wilson R. [1982] “Reputation and Imperfect Information” *Journal of Economic Theory* 27: 253-279
- [17] Matthews S. [1995] “Renegotiation of Sales Contracts” *Econometrica* 63: 567-589
- [18] Milgrom P. and J. Roberts [1982] “Limit Pricing and Entry under Incomplete Information”, *Econometrica* 53: 889-904
- [19] Mirrlees J. [1976]: “The Optimal Structure of Incentives and Authority within an Organization”, *Bell Journal of Economics* 7:

- [20] Nachbar J. [1990]: “Evolutionary Selection Dynamics in Games: Convergence and Limit Properties” *International Journal of Game-Theory* : 59-89
- [21] Nash J. [1950]: “Equilibrium Points in n-Persons Games” *Proceedings of the National Academy of Science* 36: 48-49
- [22] Saez Marti’ M. and J. Weibull [1999]: “Clever Agents in Young’s Evolutionary Bargaining Model” *Journal of Economic Theory* 86: 268-279
- [23] Schlag K. [1998]: “Why Imitate, and If So, How?” *Journal of Economic Theory* 78: 130-156
- [24] Stahl D. [1993]: “Evolution of Smart-n Players”, *Games and Economic Behavior* 10: 218-54
- [25] Weibull J. W. [1995]: *Evolutionary Game Theory* Cambridge MA, MIT press