



Department of Finance

Working Paper Series 1998

FIN-98-006

Optimal Retention in Principal/Agent Models

Jeffrey S Banks, Rangarajan K Sundaram

February 2, 1998

This working paper series has been generously supported by a grant from

CDC Investment Management Corporation

Optimal Retention in Agency Problems¹

Jeffrey S. Banks

Division of Humanities
and Social Sciences
California Institute of Technology
Pasadena, CA 91125
bnks@hss.caltech.edu

Rangarajan K. Sundaram

Department of Finance
Stern School of Business
New York University
New York, NY 10012
rsundara@stern.nyu.edu

February 3, 1998

¹This paper has benefitted from the comments of participants at a number of seminars, and from two exceptionally detailed referee reports. Both authors are grateful to the National Science Foundation for financial support; the first author would also like to thank the Sloan Foundation in this regard.

Abstract

This paper studies the interaction between a single long-lived principal and a series of short-lived agents in the presence of both moral hazard and adverse selection. We assume that the principal can influence the agents' behavior only through her choice of a retention rule; this rule is further required to be sequentially rational (i.e., no precommitment is allowed). We provide general conditions under which equilibria exist in which (a) the principal adopts a "cut-off" rule under which agents are retained only when the reward they generate exceeds a critical bound; and (b) agents separate according to type, with better agents taking superior actions. We show that in equilibrium, a retained agent's productivity is necessarily declining over time, but that retained agents are also more productive on average than untried agents due to selection effects. Finally, we show that for each given type, agents of that type are more productive in the presence of adverse selection than when there is pure moral hazard (i.e., when that type is the sole type of agent in the model); nonetheless, adding uncertainty about agent-types cannot benefit the principal except in uninteresting cases.

1 Introduction and Summary

Principal/Agent models in economics and allied fields have traditionally taken the form of the compensation contract to be the primary mechanism available to the principal to induce the agent to take desirable actions. The “standard” model adopts a Stackelberg approach to equilibrium: the agent’s reservation utility level (i.e., the utility level from alternative employment) is taken as given, and the principal maximizes her net utility subject to incentive and participation constraints for the agent.

Two considerations suggest that it may be worthwhile considering an alternative formulation of the principal/agent relationship. First, it is well known that equilibrium compensation contracts in principal/agent models under the traditional approach often take on complex and unintuitive forms. In contrast, even in cases where performance is relatively easy to measure, contracts actually observed in practice have very simple structures; in the mutual fund industry, for example, the overwhelming fee of choice is one that simply charges a fixed percentage of the assets under management.¹ Second, a wealth of anecdotal and empirical evidence suggests that a manager’s utility level from alternative employment depends non-trivially on the conditions under which current employment is terminated. In particular, in situations where the termination is initiated by the employer (i.e., the manager is “fired”), managers typically incur substantial personal costs in the form of lower wages and other lost benefits² (see, e.g., Gilson [15]). This suggests that the threat of dismissal, or equivalently, the choice of conditions under which an agent is retained, may itself be a powerful tool available to the principal.

This paper studies a principal/agent model that focusses attention on the issue of optimal agent-retention by the principal. Our model admits both moral hazard and adverse selection in its formulation; indeed, our primary aim is to understand the extent to which the threat of dismissal may be used to alleviate each of these problems, i.e., to induce agents to take more desirable actions, as well as to separate “better” agents from “worse” ones. Given this objective, we abstract from other considerations and assume the choice of the retention rule to be the sole policy instrument at the principal’s disposal. In particular, we take the agent’s compensation contract as given, although we place few restrictions on its possible form.

A second distinguishing characteristic of our paper is in our use of Nash, rather than Stackelberg, as the equilibrium concept. In a Stackelberg equilibrium, the agent (or, if there is adverse selection, the worst type of agent) receives only his reservation utility in equilibrium; the bulk of the gains from the principal/agent interaction accrue to the principal. Our use of Nash equilibrium is designed to capture the possibility that these gains could occur to both sides, and not just to the principal. We also require the principal’s strategy to be fully sequentially rational; the principal is not allowed to pre-commit to a retention rule.

A third and final dimension along which we differ from much of the literature in this area is

¹Lipper Analytical Services estimates that only 75 of the 5,400 stock and bond funds in the US charge fees that depend directly on the fund’s return performance; see Gould [16]. Indeed, it is also true in many cases that compensation contracts *within* an industry differ very little. For example, in a study of portfolio managers, Lakonishok, Shleifer, and Vishny [23] report that on accounts of \$50 million, half of their sample had fees lying between 43 and 56 basis points of the value of assets managed.

²For example, stock option grants that have not yet vested are usually lost when the employee is fired.

in admitting a divergence in the horizons of the principal and the agents. It is not an unusual occurrence in many “real-life” principal/agent relationships for principals (mutual fund investors, shareholders, voters) to be looking at a longer term than the concerned agents (fund managers, corporate executives, elected representatives). To capture such a situation in a tractable setting, we assume that the principal in our model is infinitely-lived, while all the agents have two-period lives. The assumption on the principal’s horizon is solely for expositional convenience; analogous results to those we derive here remain valid even if the principal has a finite life-span of arbitrary length. The assumption concerning agents’ life spans is a little trickier; some of the results we derive easily extend to the case where agents have arbitrary finite lives, but for some others, even the issue of existence of equilibrium appears hard to resolve.³

Our model, then, has the following structure. There is a single principal who must, in each period of an infinite horizon, decide on the choice of agent to be hired for that period. Agents are available for at most two periods. Thus, if the current incumbent agent (the one who was hired by the principal last period) has worked only one period, the principal has a choice between retaining him and replacing him with an untried agent (such an agent is always available); however, if the incumbent has already been employed for two periods, he must be replaced. Agents generate rewards for the principal when employed by her; the distribution of rewards depends on the action or effort level of the incumbent agent. Agents may also be one of a finite number of “types;” an agent’s type affects his utility from taking any particular action. Neither an agent’s type, nor his action when employed, are observable by the principal, whose only source of information concerning an agent lies in observation of the rewards generated by that agent when employed. Based on a conjecture of the strategies generating these rewards, the principal decides on the retention rule to be employed, i.e., the conditions, if any, on realized rewards under which an incumbent agent will be retained for a second period. As we have already mentioned, this choice of retention rule is required to be sequentially rational; no precommitment to a given rule is allowed. Agents choose their strategies as optimal reactions to the principal’s choice of retention rule. Equilibrium is then defined in the usual fashion.

A summary of our results, and some brief comments on the related literature, are provided in the two subsections immediately following. One final aspect of our model, however, merits mention here. A number of labor arrangements exist in the marketplace in which a single “promotion” or “retention” decision by an employer is of paramount importance (for example, the promotion of a lawyer in a firm to partner, or the granting of tenure to a university professor). Given that agents in our model have two-period lives, an obvious analogy arises between retaining an incumbent agent in our model, and a promotion decision of the sort described above. Thus, although the main motivation for our analysis lies in understanding the incentives inherent in the possibility of dismissal, our model and results speak to such arrangements also.⁴

³A companion paper, Banks and Sundaram [4], looks at the case where the principal and the agents are all infinitely-lived, and identifies a particularly appealing class of equilibria in that setting.

⁴The fact that the two periods in our model are equal in length may appear to make the analogy between our model and these situations imperfect. In fact, this is not a relevant issue. Using similar arguments to those employed in this paper, one can show the existence of similar equilibria even if the second “tenured” stage of agents’ lives is longer than the first.

1.1 Results

Our first result (Proposition 3.1) shows that, under suitable general conditions, there is an equilibrium in the model which is “symmetric” (i.e., the principal’s strategy treats all ex-ante identical agents identically, and all agents of a given type adopt identical strategies). This equilibrium possesses some strong desirable properties. Agent strategies are ordered in type in each period in the sense that “better” types of agents take (weakly) “better” actions. Further, the principal’s equilibrium retention strategy is a *cut-off rule* in which an incumbent who has completed one period is retained for a second if, and only if, the first-period reward exceeds a critical (“cut-off”) amount.⁵

Propositions 3.2 and 3.3 further characterize this equilibrium. Proposition 3.2 shows that in addition to the type-monotonicity identified above, agents’ strategies also have a *temporal monotonicity* property, viz., agents of each given type take (weakly) higher actions in the first period of employment than in the second. Under some strengthened assumptions, Proposition 3.3 shows that this inequality may be made strict. This temporal monotonicity in agents’ strategies is, of course, an equilibrium consequence of the principal’s retention strategy, in that the threat of dismissal for falling short of the cut-off provides an incentive for agents to work harder in the first period than they would otherwise.⁶

A feature of this equilibrium, which is implicit in the presentation above, but deserves special emphasis, is the following: even though all agents take lower actions in their second periods, the principal nonetheless finds it optimal, in equilibrium, to retain incumbents who generated a sufficiently large reward in their first period. Intuitively, the optimality of the cut-off strategy is simply a consequence of the principal’s beliefs shifting towards the better types of agents on account of the high first-period reward. Of course, one would expect that for retention to be sequentially rational for the principal, it should be the case that if an agent is retained, the principal’s expected one-period reward from the agent is greater than that received from an untried agent in his first period of employment. Proposition 3.4 shows that this is, indeed, the case. Thus, while all agents work harder in their first period employment than in their second, *retained* agents are, in equilibrium, more productive on average in their second periods than untried agents in their first.

Another way of viewing these results is as follows. We can think of the temporal monotonicity captured in Propositions 3.2 and 3.3 as saying something about the time-series data concerning an agent’s performance while employed, namely, that expected performance declines with time. On the other hand, suppose we had a number of these principal-agent relationships going concurrently. If we were to look at the data cross-sectionally, in a single period we would have a certain number of agents being employed for a second period, with the remainder being in their first period. In this case, the former would, on average, be more productive than the latter (by Proposition 3.4). These seemingly contradictory conclusions are rendered consistent by the selection bias inherent in “better” (from the principal’s perspective) types being retained sufficiently more often than “worse” types; in equilibrium, retained agents will be more productive in their second period than untried agents in their first period.

⁵Such cut-off strategies have been used by many authors (for instance, Barro [6] or Reed [31]), but have not—as far as we know—been derived as *equilibrium* strategies when there are no restrictions on choosing possible retention rules. The “review strategies” of Radner [30] are cut-off strategies if reviews are conducted every period, but are otherwise a weaker class of strategies.

⁶Empirical support for this phenomenon can be found in Besley and Case [7].

We then turn to a study of the special case of our model of pure moral hazard, i.e., where all agents are of the same known type. Our main result here (Proposition 3.5) shows that only “trivial” equilibria now survive: in every equilibrium of the pure moral hazard game, every agent plays a *myopic* action in every period. (A myopic action is one that maximizes the agent’s one-period utility; in taking such an action the agent effectively ignores its impact on future utility.) Therefore, in contrast to the earlier results, the absence of uncertainty about agent preferences implies there can be no equilibrium performance effect: the principal cannot credibly provide the incentive for agents to work hard in the first period.^{7,8}

The combination of Propositions 3.3 and 3.5 raises an interesting possibility. The former result showed that, under certain circumstances, there exist equilibria in which agents’ strategies separate over time, with agents of all types taking strictly higher actions in their first period than their last. Since all agents evidently take myopic actions in their last period (the future is now irrelevant to them), the question naturally arises whether there are conditions under which the principal can actually be made better off by introducing adverse selection into the model.⁹ Indeed, it appears plausible that under suitable circumstances, the principal’s payoffs could be improved even if the new types of agent-types introduced into a pure moral hazard problem are all “worse” than the single agent-type of the original problem. (For instance, the probability of the old type could be arbitrarily close to unity, while the new types may be only marginally worse.)

Unfortunately, and somewhat unexpectedly, this intuition turns out to be misguided. We show first in Proposition 3.6 that the set of equilibria varies upper-semicontinuously with the prior belief on types. This implies that adding a “little bit” of uncertainty about agent types will only induce a “small” movement away from myopic behavior. Proposition 3.7 then completes the argument by showing that this “small” movement away from myopic behavior can never be sizeable enough to compensate for the introduction of “worse” agent-types. More specifically, we show that the principal can *never* benefit from the introduction of “worse” types of agents, and, less surprisingly, can never lose from the introduction of “better” types. However, we do show that if adverse selection is introduced in more complex (but still intuitively appealing) ways than simply adding a single new type of agent, then it is certainly possible that the principal strictly benefits from the added uncertainty over the agent’s type.

⁷This is no longer true if the principal can credibly *pre-commit* to a given retention rule; see, e.g., Ferejohn [13] who examines precisely this situation.

⁸This result of adverse selection attenuating the effects of moral hazard should be contrasted with that found in Coate and Morris [9]. In their model, agents (politicians) can be “good” or “bad.” Agents have the opportunity to develop a project which may or may not be efficient to provide; agents know this information, but the principal (the voter) does not. If the agents’ types are known to the principal, the project is only undertaken when it is efficient to do so. However, when types are unknown, a bad type will, in equilibrium, always develop the project in order to hide a transfer of benefits to a special interest. (The bad types, unlike the good types, values the utility of special interests.) In this sense, moral hazard in their paper is greater in the presence of adverse selection.

⁹This question is reminiscent of the motivation in Kreps, et al [20], who show that admitting a “small” probability that one of the players in a finitely-repeated prisoners’ dilemma is “irrational” (say, is an automaton) could create equilibria under which both players are better off.

1.2 Comments on the Related Literature

Dynamic models of principal/agent interaction have been the focus of considerable attention in recent years (examples include Radner [28], [29], [30]; Rogerson [32]; Fudenberg, et al [14]; Baron and Besanko [5]; Dewatripont [10]; Laffont and Tirole [21], [22]; and Dutta and Radner [12]). As mentioned in the opening remarks our model differs from the standard approach taken by most (but not all; see the following paragraphs) of these papers in our focus on the retention rule, rather than the compensation contract, being the main motivating tool available to the principal. This change in perspective necessitates a change in modeling strategy. For example, in the traditional approach, it is typical to assume the existence of only one agent; even where multiple agents are implicitly involved in the analysis, they are rarely explicitly introduced. In contrast, the presence of multiple agents is central to our approach.

Some other papers in the literature have also focussed (in some cases, implicitly) on the retention decision of the principal. One example is Radner [30]. Radner examines a model in which the principal “reviews” the agent’s performance once every k periods; if the agent’s average performance over a review cycle falls short of a cut-off the agent is terminated. Radner’s model does not consider adverse selection problems; the model also does not concern itself with equilibrium, but rather on the incentive effects of different review strategies. Heinkel and Stoughton [18] provide a second example. They study portfolio management contracts in a two-period principal/agent model with both moral hazard and adverse selection problems. They show that the threat of dismissal at the end of one period can be a powerful motivation for the agent to take desirable actions.

A branch of the principal/agent literature that is closely related to our work is that on repeated elections, where the (representative) voter plays the role of principal and the potential candidates for election that of the set of available agents. In such a scenario, it is natural to abstract from compensation considerations, and to focus on the reelection decision of the voters as a function of the agent’s performance. Interpreted in this light, our model may be considered a model of repeated elections with the presence of term-limits (of two periods) on the agents.

Most previous research on repeated elections in the presence of informational asymmetries has studied the voter’s decision problem from either the moral hazard or the adverse selection perspective. Examples of the former include Barro [6], Ferejohn [13], and Austen-Smith and Banks [3], while those of the latter include Rogoff [33], and Reed [31]. The current model goes beyond these in that it incorporates both moral hazard and adverse selection considerations. It is closest in spirit and structure to that found in Banks and Sundaram [4], with the principal difference being that in Banks and Sundaram [4] no term limits were posited for the agents.

2 The Model

2.1 Timing and Information

We consider a discrete-time infinite horizon model with periods indexed by $t = 0, 1, 2, \dots$. In each period, an infinitely-lived *principal* must decide on the choice of a single *agent* for that period. Agents may be employed for at most two periods. Thus, if the current incumbent agent has one

period of employment remaining, then the choice the principal faces is between re-selecting the incumbent and replacing him with another agent; while if the incumbent has already completed two periods of employment, the principal must select a new agent. All untried agents appear ex ante identical to the principal; we assume, without loss, therefore, that at each date there is only a single untried agent available.¹⁰ This agent will be described by his date of appearance t .

Agents in the model are distinguished by their *type*. Each agent may be one of a finite number of possible types $\omega \in \Omega = \{\omega_1, \dots, \omega_n\}$. We will define an ordering on Ω by $\omega_1 < \dots < \omega_n$, and will refer to ω_k as a “better” type than ω_l if $k > l$. Thus, ω_n represents the best type of agent, and ω_1 the worst type. The justification for this ordering and terminology will become clear when we list our assumptions on the model below (see the discussion after Assumption A6).

Agents’ types are assumed to be realizations of i.i.d. draws from a common distribution $\pi = (\pi_1, \dots, \pi_n) \in P(\Omega)$. (For any set X , $P(X)$ will represent the set of all probability distributions on X .) The distribution π is taken to be common knowledge to the players in the model. Adverse selection is introduced into the model by the assumption that an agent’s “true” type is private information to the agent, and is unknown ex ante to the principal (provided, of course, that $\pi_k > 0$ for at least two distinct values of k).¹¹

Agents generate rewards $r \in \mathbb{R}$ for the principal while employed by her. The distribution of these rewards in any period depends on the action a the agent takes in that period. The set of actions available to any agent in any period is taken to be a non-degenerate compact interval $A = [a^{\min}, a^{\max}] \subset \mathbb{R}$. If the agent takes the action $a \in A$, then the reward r is realized according to the distribution $F(\cdot|a)$. The principal observes the reward r in each period, but not the action a which generated it. However, the principal does have complete knowledge of the set A and the family of conditional distributions $F(\cdot|a)$ for $a \in A$.

We make the following assumptions on the reward distributions:

- A1** For each $a \in A$, $F(\cdot|a)$ admits a density $f(\cdot|a)$ with respect to the Lebesgue measure.
- A2** The set $\{r : f(r|a) > 0\}$ is a bounded interval in $\mathbb{R}$ that is independent of $a \in A$; the closure of this interval is denoted S .
- A3** The density $f(\cdot|a)$ is jointly continuous on $S \times A$.
- A4** $f(\cdot|\cdot)$ satisfies the Monotone Likelihood Ratio Property (MLRP), i.e. for all $a, \hat{a} \in A$ with $a > \hat{a}$, the ratio $f(r|a)/f(r|\hat{a})$ is strictly increasing in r .

Assumptions A1 and A3 are self-explanatory. Assumption A2 ensures that the principal cannot, with certainty, rule out certain actions on the basis of observed rewards. Such an assumption is standard in the literature on imperfect monitoring (e.g. Abreu, et.al. [1]). The MLRP condition has

¹⁰There is an implicit assumption here that agents who are not hired in the period of their appearance are no longer available to the principal, and neither are incumbents who are replaced at any point. This assumption is without loss of generality for the class of equilibria we consider below.

¹¹We do not require $\pi_k > 0$ for all $k = 1, \dots, n$. Indeed, since we wish to examine, among other things, the convergence of equilibria under adverse selection to those arising in the pure “moral hazard” case (i.e. when $\pi_k = 1$ for some k), it is essential to allow for the possibility that some components of π may be 0.

been extensively employed in a number of papers both in the principal-agent literature (e.g. Grossman and Hart [17]) and elsewhere (Milgrom [25]). An important implication of this assumption is that $F(\cdot|a)$ is ordered according to first-order stochastic dominance, i.e., for any given $r \in \text{int } S$, it is the case that $F(r|a) < F(r|\hat{a})$, whenever $a > \hat{a}$.

2.2 Preferences

The utility of an incumbent agent in either period of employment depends on the action taken by the agent that period and the agent's type, and is denoted $u(a, \omega)$; the utility level of agents while unemployed is normalized to 0. To make the choice problem faced by the agent non-trivial we shall assume throughout that for each type $\omega \in \Omega$ there exists an action $a(\omega)$ satisfying $u(a(\omega), \omega) > 0$. This captures the point that all agents find employment more desirable than unemployment. Agents discount future utilities by the common factor $\delta \in [0, 1]$.

Our next assumption on preferences requires a definition. Let $g(x, y)$ be a function from $\mathbb{R} \times \mathbb{R}$ into $\mathbb{R}$. Then, g is said to be *supermodular* in x and y if whenever $(x, y) \gg (\hat{x}, \hat{y})$, we have

$$g(x, y) - g(x, \hat{y}) > g(\hat{x}, y) - g(\hat{x}, \hat{y}).$$

If the strict inequality in the definition is replaced by a weak one, then g will be said to be *weakly supermodular*. In words, supermodularity requires that as y increases, the impact on g of increasing x increases as well. If g is twice-continuously differentiable, then the condition of supermodularity is equivalent to the requirement that $\partial^2 g(x, y)/\partial x \partial y > 0$; if g is differentiable only in x , then the condition is that $\partial g(x, y)/\partial x$ increases in y .

Supermodularity has been widely utilized in studying a variety of problems including adverse selection (e.g., Spence [35]) and games with strategic complementarities (e.g., Milgrom and Roberts [26]; see Milgrom and Shannon [27] for a detailed analysis of supermodularity and its variants, as well as further references of applications). Note that if g is additively separable in x and y then it is weakly supermodular, but not supermodular; and that if g is multiplicatively separable in x and y , and is also increasing in x and y , then g is supermodular. We will make use of both these observations shortly.

We make the following additional assumptions on the agents' preferences:

A5 $u(a, \omega)$ is supermodular in (a, ω) .

A6 $u(a, \cdot)$ is increasing on Ω .

A7 $u(\cdot, \omega)$ is continuous and strictly quasi-concave on A .

The content of Assumption A5 has been explained above. Assumption A6 states that “better” types of agents find it cheaper to take a given action than “worse” types, relative to the unemployment utility level; this is, in fact, the sense in which they are “better.” Assumption A7 guarantees the existence of unique best actions for the agent if employed for a second period.

The principal's preferences take on a simpler form: her utility in a period depends solely on the reward r generated by the incumbent agent that period, and is given by $v(r)$, where $v(\cdot)$ is a strictly increasing function. The principal's discount factor is given by $\alpha \in [0, 1]$.

A comment or two on the generality of this specification of preferences is, perhaps, in order here. One possibility is that we have $u(\cdot) = z - c(a, \omega)$, where z denotes the fixed direct benefits from being employed and $c(a, \omega)$ is interpreted as the opportunity cost of taking action $a \in A$ when the agent's type is $\omega \in \Omega$. As long as $c(\cdot, \cdot)$ is continuous in a , increasing in ω , and *submodular* in (a, ω) (i.e. $-c$ is supermodular in (a, ω)), the assumptions we have outlined above will be satisfied. Alternatively, it could be that the principal and the agents actually have *coincident* underlying preferences over rewards; for example u could be given by $u(a, \omega) = \int v(r) dF(r|a) - c(a, \omega)$ where $v(\cdot)$ is, as defined above, the principal's utility function. In this set-up, the principal and agent agree on what are good and bad outcomes; the heterogeneity comes about from the presumption that it is the *agent* who directly influences these outcomes.

We also could have a "canonical" version of the principal/agent model found in the economics literature: let the agent's per-period preferences over income and cost be given by the quasi-linear form $y - c(a, \omega)$, let $r \in \mathbb{R}$ denote monetary profits accruing to the principal, and let $s(r)$ denote the share of profits going to the agent.¹² Then, the agent's total utility in any one period of employment will be $s(r) - c(a; \omega)$; so the ex ante expected utility the agent obtains is

$$u(a, \omega) = \int s(r) dF(r | a) - c(a, \omega).$$

If c is submodular and $v(r) = r - s(r)$ is increasing in r , the assumptions on the principal's and agents' preferences will be satisfied.

2.3 Histories and Strategies

In general, a *t-history* for the game is a list of the agents who have been employed in periods $1, \dots, t-1$, the rewards they generated, and the current incumbent agent; a *strategy* ζ for the principal is a sequence of measurable maps $\{\zeta_t\}$, where for each t , ζ_t specifies, as a function of the *t-history* to date, whether the current incumbent agent is to be replaced; and a *strategy* λ for agent t is a sequence of measurable maps $\{\lambda_{t\tau}\}$, where $\lambda_{t\tau}$ specifies an action for agent t in period $(t + \tau)$ given the $(t + \tau)$ -history to date, $\tau = 0, 1$. In this paper, we focus on a class of strategies of special importance that we call *anonymous strategies*. Anonymous strategies, in words, are those which are symmetric in agent types (i.e., where the identity of an agent beyond his type is immaterial): all ex-ante identical agents are treated identically by the principal, and all agents of a given type employ the same strategy.

A *personal history* for a given agent is simply the reward $r \in S$ generated by the agent in his first period of employment. Since the principal must replace an agent if the latter is at the end of his two periods in office, an *anonymous strategy* σ for the principal is simply a retention rule employed by the principal between an agent's first and second periods that is based solely on the

¹²It is implicit here that the same sharing rule applies to both periods of employment. However, this is unnecessary. See the remark following Proposition 3.1 below.

agent's personal history. That is, an anonymous strategy is a measurable map $\sigma: S \rightarrow \{0, 1\}$, where 0 denotes the action of replacing, and 1 denotes the action of continuing with, the incumbent agent. Let Σ denote the set of all anonymous strategies for the principal.

An *anonymous strategy* for a generic agent is a first period action as a function of the agent's type, and a second period action as a function of type and the reward generated in the first period. Since (as we shall see) allowing the agent to mix in her choice of first period action is essential in guaranteeing existence, we will denote an anonymous strategy for an agent of type ω_k by (μ_k, γ_k) , where $\mu_k \in P(A)$ is a mixed action for an agent of type ω_k in the first period, and $\gamma_k: S \rightarrow A$ is a function specifying a type- ω_k 's second period action as a function of the realized first period reward. Let the vector $(\{\mu_k\}, \{\gamma_k\})_{k=1, \dots, n}$ be denoted by (μ, γ) , and let Θ denote the set of all possible (μ, γ) , i.e., the set of all possible anonymous strategies for the agents.

Two further definitions will be useful in the sequel. First, an anonymous strategy σ for the principal will be called a *cut-off strategy* if there exists an $\bar{r} \in S$ such that $\sigma(r) = 1$ if and only if $r \geq \bar{r}$. That is, a cut-off strategy is one in which an incumbent agent is retained for a second period if and only if the reward generated by the agent in the first period exceeds some critical level.

Second, an anonymous strategy (μ, γ) for the agents will be called *type-monotone* if it involves better types taking higher actions in each period, i.e., if:

1. The agents' first period actions $\mu_k \in P(A)$ are "ordered" in the sense that there exist intervals $[a_k, b_k] \subset A$ such that $b_k \leq a_{k+1}$ for $k = 1, \dots, n-1$, and $\mu_k([a_k, b_k]) = 1$ for all k .
2. The agents' second period actions are also ordered in the sense that for all $r \in \mathbb{R}$, we have $\gamma_k(r) \leq \gamma_{k+1}(r)$ for $k = 1, \dots, n-1$.

2.4 Best-Responses and Equilibrium

Since our focus in this paper is only on equilibrium in anonymous strategies, we avoid irrelevant generality and define equilibrium for that case alone. We begin with a description of the principal's best-response problem.

The principal starts with a prior π regarding an agent's true type, and updates her beliefs using the observed rewards and a conjectured (anonymous) strategy for the agents. Formally, let the conjectured agent strategy be $(\mu, \gamma) \in \Theta$, and suppose that the current incumbent has generated a reward of r in his first period of employment. Now for any distribution $m \in P(A)$, the conditional probability of observing r given m is

$$\varphi(r|m) = \int_A f(r|a)m(da). \quad (2.1)$$

Thus, upon observing r , the principal's updated belief that the incumbent agent is of type k , denoted $\beta_k(r)$, is

$$\beta_k(r) = \frac{\pi_k \varphi(r|\mu_k)}{\sum_{j=1}^n \pi_j \varphi(r|\mu_j)} \quad (2.2)$$

Given this updating rule, each strategy σ for the principal defines in the obvious (if notationally dense) manner, a t -th period expected utility against (μ, γ) denoted $w_t(\sigma; \mu, \gamma)$. The *worth* of the profile (σ, μ, γ) to the principal is then given by

$$W(\sigma, \mu, \gamma) = \sum_{t=0}^{\infty} \alpha^t w_t(\sigma; \mu, \gamma). \quad (2.3)$$

A strategy σ^* is said to be *optimal against* (μ, γ) for the principal if

$$W(\sigma^*, \mu, \gamma) \geq W(\sigma, \mu, \gamma) \quad (2.4)$$

for all $\sigma \in \Sigma$.

The agent's best response problem is easier to define. If an agent of type ω_k who generated a reward of r in the first period is retained for a second period, the second-period action $\gamma_k(r)$ will be optimal for the agent if and only if

$$\gamma_k(r) \in \arg \max \{u(a, \omega_k) \mid a \in A\} \quad (2.5)$$

for all $r \in S$. To identify the optimal first-period action for an agent of type ω_k when the principal is using the anonymous strategy σ , let

$$U(\omega_k) = \max \{u(a, \omega_k) \mid a \in A\}, \quad k = 1, \dots, n, \quad (2.6)$$

and let

$$R_\sigma = \{r \in S \mid \sigma(r) = 1\} \quad (2.7)$$

denote the set of first-period rewards for which the agent is retained. Then the first-period strategy μ_k is *optimal against* σ (for an agent of type ω_k) if and only if $\mu(A(\sigma, \omega_k)) = 1$, where

$$A(\sigma, \omega_k) = \arg \max_{a \in A} \left\{ u(a, \omega_k) + \delta \int_{R_\sigma} U(\omega_k) dF(r|a) \right\}. \quad (2.8)$$

Finally, we say that the profile (σ, μ, γ) constitutes an *equilibrium* for this model if σ and (μ, γ) are optimal against each other.

It is important to note, in closing, that while we require best-responses to be chosen from the class of anonymous strategies, this is not really a restriction. If the principal bases the criteria on which a particular agent t is retained in office only on t 's performance, t 's optimization problem is unaffected by the history of previous play; thus, if at all it admits any solution, it must admit a solution in which the optimal strategy is based solely on t 's own history. Conversely, if all agents use anonymous strategies, the principal cannot benefit by conditioning her response against any particular agent on the performance of any other agent. It follows that equilibria in anonymous strategies remain equilibria even when some or all players may condition their actions on the entire history of play.

3 Results

Our first result establishes the existence of anonymous equilibria in the model under the assumptions we have outlined above:

Proposition 3.1 *There exists an equilibrium in anonymous strategies $(\sigma^*, \mu^*, \gamma^*)$ where σ^* is a cut-off strategy, and (μ^*, γ^*) is a type-monotone strategy.*

Proof See Appendix A. □

Remark A perusal of Appendix A reveals that the proof of the proposition nowhere uses either the convexity of A or the (weak or strict) quasi-concavity of $u(\cdot, \omega)$ on A . In addition, trivial modifications to the proof show that the same result can be obtained even if the reward distributions and the agents' per-period utility functions depend on whether the agent was in his first or second period of employment, as long as both sets of functions continued to satisfy Assumptions A1–A7. In particular, the value of office for an agent, the productivity of the agent, and the compensation structure for the agent could all be functions of an agent's time of employment without affecting Proposition 3.1. However, generalizing the model in these directions creates significant notational complexity, without adding anything of corresponding value to our understanding of the problem. □

The proof of Proposition 3.1 consists of three steps. In the first step, it is established that when the principal uses a cut-off strategy, then the agents' best-responses satisfy type-monotonicity. Then, in the second step, it is shown that when the agents employ type-monotone strategies, the principal's optimal response will have a cut-off form. In the third and final step, continuity properties of these best-responses are used in conjunction with a fixed-point theorem to provide the desired result.

The combination of supermodularity and the monotone likelihood ratio property of the reward densities plays a big role in these steps. The intuitive arguments are quite simple; we will sketch these arguments here using the notation of subsection 2.4 (see, especially, expressions (2.1)–(2.8)). Consider the first step. By the supermodularity of u , better types of agents will always take higher actions than worse types in their second period of employment; i.e., regardless of the first-period reward r , we will have

$$\gamma_1(r) \leq \dots \leq \gamma_n(r). \tag{3.1}$$

Moreover, by Assumption A6, better types must have higher second-period utility also:

$$U(\omega_1) \leq \dots \leq U(\omega_n). \tag{3.2}$$

Now, suppose the principal employs a cut-off strategy with cut-off point $\bar{r}$. Given the action $a \in A$, the probability of failing to make the cut-off is then $F(\bar{r}|a)$; thus, a type- ω_k agent's first-period utility from the action a is simply

$$H(\bar{r}, a, \omega) = u(a, \omega_k) + \delta U(\omega_k)[1 - F(\bar{r}|a)]. \tag{3.3}$$

By the MLRP property, higher actions result in stochastically dominant reward distributions. Combining this with monotonicity of $U(\cdot)$ in ω_k (equation (3.2)), and the supermodularity of u , it can be shown that H is itself supermodular in (a, ω) for each $\bar{r}$. It is immediate from this that *first-period* actions for the agents will also be type-monotone, completing the first step of the proof.

The second step of the proof obtains by combining several observations. Under MLRP, high rewards are more likely to have come from high actions. Thus, if the agents employ type-monotone strategies, then a cut-off strategy for the principal makes it more likely that only better types of agents will be retained. Retaining only the better types of agents is desirable for the principal, because (by hypothesis) the second-period actions of the agents are ordered according to type. Intuitively, it is apparent from all this that a cut-off strategy is optimal for the principal against type-monotone strategies; the formal arguments only require some additional technical details.

From existence of equilibrium, we now turn to the issue of characterization. Throughout this discussion, we confine attention to anonymous-strategy equilibria of the type identified in Proposition 3.1 (i.e., cut-off strategies for the principal, type-monotone strategies for the agents). For the purposes of these results, we will also define the support of a distribution $m \in P(A)$ to be the smallest closed subset A' of A such that $m(A') = 1$.

Proposition 3.2 *Let $(\sigma^*, \mu^*, \gamma^*)$ be an equilibrium as in Proposition 3.1. Then, for $k = 1, \dots, n$, and any $r \in S$, it is the case that $a_k \geq \gamma_k^*(r)$, where a_k is the lower bound on the support of μ_k^* . That is, agents' actions are (weakly) decreasing over time.*

Proof Fix any $k \in \{1, \dots, n\}$. As earlier, let $U(\omega_k) = \max\{u(a, \omega_k) \mid a \in A\}$ denote second-period payoff to a type- ω_k agent. By hypothesis, $U(\omega_k) > 0$. Let r^* denote the cut-off employed by the principal. For $q \in \{0, 1\}$, define

$$U^*(q, a, \omega_k) = u(a, \omega_k) + q\delta[1 - F(r^*|a)]U(\omega_k)$$

Then, $U^*(0, a, \omega_k)$ gives the maximand for the agent's second period maximization problem, and $U^*(1, a, \omega_k)$ that for the first period. Now if r^* is equal to either $\inf S$ or $\sup S$, or if $\delta = 0$, the agent will choose the same action in either period (the uniqueness of this action follows from strict quasi-concavity of u), and the claim in the proposition obtains trivially. So suppose, therefore, that $r^* \notin \{\inf S, \sup S\}$ and that $\delta > 0$.

As a first step, we will show that $U^*(q, a, \omega_k)$ is supermodular in (q, a) . Let $a > \hat{a}$, so $F(r^*|\hat{a}) > F(r^*|a)$. Since $\delta U(\omega_k) > 0$, we have that

$$u(a, \omega_k) - u(\hat{a}, \omega_k) + \delta U(\omega_k)[F(r^*|\hat{a}) - F(r^*|a)] > u(a, \omega_k) - u(\hat{a}, \omega_k), \quad (3.4)$$

which is equivalent to

$$U^*(1, a, \omega_k) - U^*(1, \hat{a}, \omega_k) > U^*(0, a, \omega_k) - U^*(0, \hat{a}, \omega_k).$$

This says precisely that U^* is supermodular in (q, a) . Of course, it follows immediately that the set of maximizers will be ordered across time; that is, letting

$$M^*(q, \omega_k) = \arg \max_{a \in A} U^*(q, a, \omega_k),$$

it is the case that if $a \in M^*(1, \omega_k)$ and $\hat{a} \in M^*(0, \omega_k)$,¹³ then $a \geq \hat{a}$. In equilibrium, we must have $\mu_k^* \in P(M^*(1, \omega_k))$; Proposition 3.2 follows. $\square$

By adding some smoothness and interiority assumptions, we can generate “strict” rather than “weak” predictions. Define

$$a_k^m = \arg \max_{a \in A} u(a, \omega_k) \quad (3.5)$$

Then, a_k^m is type- ω_k ’s optimal myopic action, and is unique by the strict quasi-concavity of u . Say that the environment is *nice* if the following conditions are satisfied:

1. u and F are continuously differentiable on A .
2. $\delta > 0$.
3. a_k^m is in the interior of A for each k .
4. $u(a^{\max}, \omega_k) < 0$ for all k .

Condition 1 allows us to use differentiable techniques in appealing to supermodularity. Condition 2 implies that an agent’s two-period decision problem looks different from the last-period problem as long as the cut-off $r \notin \{\inf S, \sup S\}$. Conditions 3 and 4 keep an agent’s optimal actions away from the boundaries of A .

Proposition 3.3 *Let $(\sigma^*, \mu^*, \gamma^*)$ be an equilibrium as in Proposition 3.1. Suppose the environment is nice. Suppose also that $\pi_k > 0$ for at least two values of k . Let a_k and b_k denote the minimum and maximum points in the support of μ_k , $k = 1, \dots, n$. Then, for $k = 1, \dots, n-1$, we have:*

1. $\gamma_{k+1}^* > \gamma_k^*$.
2. $a_{k+1} > b_k$.
3. $a_k > \gamma_k^*$.

Finally, it must be the case that r^* is in the interior of S , where r^* is the cut-off associated with the strategy σ^* .

Proof By definition, for each r , $\gamma_k^*(r)$ must solve $\max_{a \in A} u(a, \omega_k)$. Since a_k^m is the unique solution to this problem, we must have $\gamma_k^*(r) = a_k^m$ for all r . By the supermodularity of u in (a, ω) , it is immediate that we must have $a_k^m < a_{k+1}^m$ for each k , so the first result obtains.

Next, we will show that $r^* \in \text{int } S$. Suppose to the contrary that we had $r^* \in \{\inf S, \sup S\}$. Then, of course, agents will take their myopic action a_k^m in both periods. However, since these actions are strictly monotone in type, the principal’s one-period reward from retaining the incumbent

¹³These maximizers exist by the continuity of U^* and the compactness of A .

agent is strictly increasing in the first period reward (this follows from the monotone-likelihood property of the reward densities; see lemmata A.5 and A.6 in the proof of Proposition 3.1 in Appendix A). For a sufficiently high first period reward, this second period reward from retaining the agent will dominate that obtained from replacing him. Consequently, $r^* \in \{\inf S, \sup S\}$ cannot be optimal for the principal, and the claimed equilibrium breaks down. This establishes $r^* \in \text{int } S$ in any equilibrium.

Next, as in the proof of Proposition 3.2, define

$$U^*(q, a, \omega_k) = u(a, \omega_k) + q\delta[1 - F(r^*|a)]U(\omega_k)$$

and let $M^*(q, \omega_k) = \arg \max\{U^*(q, a, \omega_k) \mid a \in A\}$. Then, as we showed in Proposition 3.2, $\delta > 0$ and $r^* \in \text{int } S$ together imply that U^* is supermodular in (q, a) for each fixed ω_k . Of course, a first-period action a_k is optimal for type ω_k if and only if $a_k \in M^*(1, \omega_k)$. By Proposition 3.2, any such a_k must also satisfy $a_k \geq \gamma_k^*(r) = a_k^m$. By Condition 4 in the definition of a nice environment, we cannot have $a_k = \sup S$, so any second period optimal action must lie in the interior of S . By the differentiability assumptions on u and F , a_k must, therefore, satisfy

$$\frac{\partial u}{\partial a}(\hat{a}, \omega) - qU(\omega) \frac{\partial F}{\partial a}(r^*|\hat{a}) = 0.$$

We cannot, therefore, have $a_k = a_k^m$, so we must have $a_k > a_k^m$ for each k . Since $\gamma_k^*(r) = a_k^m$ for all r , the third claim is established. Similarly, the supermodularity of U^* in (a, ω) and its differentiability imply $a_{k+1} > b_k$ for each k , establishing the second claim and completing the proof. $\square$

Propositions 3.2 and 3.3 describe what may be called a “performance effect” of the principal’s retention decision, in that each type of agent works harder in the first period than in the second period. Of course, if an agent will necessarily work less hard in his second period of employment, it raises the question of why the principal would find it in her interest to at times retain such an agent. The answer can be found in the “selection effect” taking place in equilibrium:¹⁴ not all agent types are equally likely to be retained. Indeed, higher types (who, by type-monotonicity, are more productive) are more likely to be retained than lower types. To see this selection effect in action, let $v_1(\mu, \gamma)$ denote the principal’s expected one-period reward from an agent’s first period of employment given the anonymous strategies (μ, γ) for the agents. Similarly let $v_2(r; \mu, \gamma)$ be the principal’s expected second period payoff; of course, this payoff will, in general, depend on the agent’s first-period reward r through the principal’s Bayesian-updated beliefs.

Proposition 3.4 *Let $(\sigma^*, \mu^*, \gamma^*)$ be an equilibrium as in Proposition 3.1, and let r^* be the cut-off associated with σ^* . Then, for all $r \geq r^*$, it is the case that $v_2(r, \mu^*, \gamma^*) \geq v_1(\mu^*, \gamma^*)$.*

Proof Let V be the principal’s expected payoff from the game (or, equivalently, from the continuation game at any point, conditional on beginning afresh with an untried agent at that point). By definition of r^* , for any $r \geq r^*$, the strategy of keeping the incumbent agent for another period and

¹⁴The terms “performance effect” and “selection effect” are due to Reed [31].

then replacing him generates an expected payoff of $v_2(r, \cdot) + \alpha V$. Since r^* is the equilibrium cut-off this payoff must be higher than simply replacing the incumbent today. Consider the alternative strategy in which the principal discards the incumbent agent, and discards the next agent regardless of the reward latter generates, and then returns to the optimal strategy. This will generate an expected payoff of $v_1(\cdot) + \alpha V$. Since this alternative strategy cannot generate a higher payoff for the principal than the optimal strategy, it must be the case that

$$v_2(r, \cdot) + \alpha V \geq v_1(\cdot) + \alpha V,$$

and so we must have

$$v_2(r, \cdot) \geq v_1(\cdot). \square$$

Proposition 3.4 establishes that an incumbent agent is retained precisely when the shift in beliefs concerning the agent's type is sufficient to overcome the fact that a retained agent will necessarily work less hard. Equivalently put, while all agents worker harder in their first period employment than in their second, the equilibrium selection bias implies that *retained* agents are on average more productive in their second periods than untried agents in their first.

The temporal monotonicity result in Proposition 3.3 required the existence of at least two types of agents, since otherwise no “learning” on the part of the principal takes place. It turns out that it is a simple matter to identify the anonymous equilibria when this condition fails to hold i.e. in the “pure moral hazard” version of our model in which $\pi = e_k$ for some k , where e_k is the k -th unit vector in $\mathbb{R}^n$.

Proposition 3.5 *Suppose $\pi = e_k$ for some k . Then, in any anonymous equilibrium:*

1. *Agents take the myopic action a_k^m in both periods.*
2. *If the environment is nice, then $r^* \in \{\inf S, \sup S\}$, where r^* is the equilibrium cut-off.*

Proof By Proposition 3.2, each type's (mixed) actions are weakly decreasing over time, so first-period actions are (weakly) larger than second-period actions. Thus, it suffices to show they cannot be strictly larger with positive probability. Now, if μ_k^* places positive probability on any action greater than a_k^m , the principal's unique best response is to always replace an incumbent agent, i.e., to set r^* equal to $\sup S$; but then the agent's unique best first period action is a_k^m , contradicting the hypothesis. This establishes Part 1.

To see Part 2, note that if $r^* \notin \{\inf S, \sup S\}$, then, in a nice environment, an agent's first period action will be strictly higher than her second period action. (This is a simple consequence of the agents' best-response problem.) However, this would in turn imply that we should have $r^* = \sup S$, a contradiction. $\square$

Proposition 3.5 shows that in the absence of adverse selection (i.e., of uncertainty concerning agents' types), all equilibria are “trivial.” Without adverse selection, there is no possibility of selection effects, so there will not exist performance effects either.

Our next result examines the continuity properties of the equilibrium correspondence to see if “small” movements away from the pure moral hazard case can have “large” effects. Unsurprisingly, we find that the answer is in the negative. Some notational simplification will aid in the presentation of this result. Since the optimal second-period actions taken by agents are always their myopic actions, we drop dependence on this part of their strategies, and denote a typical equilibrium by simply (σ^*, μ^*) . Thus, given any vector π of prior beliefs, let $\mathcal{E}(\pi)$ be the set of all anonymous equilibria of Proposition 3.1.

Proposition 3.6 *$\mathcal{E}$ is an upper-semicontinuous correspondence in π .*

Proof See Appendix B. □

Proposition 3.6 implies, in particular, that as the prior π converges to e_k for some k , the equilibria converge to equilibria of the “pure moral hazard” case wherein the principal knows the agent to be of type ω_k with probability one. Therefore by Proposition 3.5, the equilibrium cut-off values r^* must be converging to either $\inf S$ or $\sup S$; in either case, the first-period actions of a type ω_k agent are converging to the myopic action a_k^m . Hence, adding a “little bit” of uncertainty concerning agent types will only move the agents a little bit away from their myopic actions.

The combination of Propositions 3.3 and 3.5 nevertheless raises an intriguing possibility. Proposition 3.3 demonstrated that agents take higher actions in the first period of employment than in the second in the presence of adverse selection. Proposition 3.5 shows that this fails to hold when there is no adverse selection. The question naturally arises: might it be the case that the principal is actually better off *with* adverse selection than without? In particular, suppose we compare the principal’s equilibrium payoff under two scenarios, the first where only the “best” type of agents exist (i.e. $\pi_n = 1$) and the second where there is some uncertainty about agents’ types ($\pi_n < 1$). From Proposition 3.5, we know that under the former the principal receives the reward generated by a_n^m , the best type’s myopic action, in every period. On the other hand, Proposition 3.3 informs us that performance effects are present in the latter setting, implying all types take higher actions in their first period of employment. Certainly, it appears plausible that under some appropriate set of circumstances—say, the probability of type ω_n is arbitrarily close to unity, and the “worse” types are only marginally worse than ω_n —the principal could be made better off.

Somewhat surprisingly, this intuition is misguided. The following result shows that the principal can *never* become better off if “worse” types are added, and can never become worse off if “better” types are added. Recall that e_k denotes the k -th unit vector in $\mathbb{R}^n$; let Δ^{n-1} denote the unit simplex in $\mathbb{R}_+^n$.

Proposition 3.7 *For any $\pi \in \Delta^{n-1}$, let $V(\pi)$ denote the payoff the principal obtains in some anonymous equilibrium. Then, $V(e_n) \geq V(\pi) \geq V(e_1)$.*

Proof Pick any $\pi \in \Delta^{n-1}$, $\pi \neq e_n$. If $\pi = e_k$ for some $k \neq n$, the result is trivial, so suppose for $\pi \neq e_k$ for any k . Let K be the set of k for which $\pi_k > 0$, let $(\sigma^* \mu^* \gamma^*)$ be any equilibrium of the game given π , and for any first period reward r let $\beta(r) = (\beta_k(r))_{k \in K}$ represent the principal’s beliefs regarding the incumbent agent’s type in this equilibrium after observing r . In their last

period, all agents necessarily take their myopic action a_k^m ; thus, given that the current incumbent has generated a reward of r , the principal solves

$$\max\{V(\pi), \sum_{k \in K} \beta_k(r) E(a_k^m) + \alpha V(\pi)\},$$

where $E(a_k^m) = \int v(\hat{r}) dF(\hat{r}|a_k^m)$. We consider two cases. First, suppose the set

$$\{r \mid \sum_{k \in K} \beta_k(r) E(a_k^m) \geq (1 - \alpha)V(\pi)\}$$

is non-empty. Then since $E(a_n^m) \geq E(a_k^m)$ for all k , we also have $E(a_n^m) \geq (1 - \alpha)V(\pi)$, and since $E(a_n^m)/(1 - \alpha) = V(e_n)$ by Proposition 3.5, we are done.

So suppose the second case holds, i.e., for all r , we have

$$(1 - \alpha)V(\pi) > \sum_{k \in K} \beta_k(r) E(a_k^m).$$

Then, necessarily, all agents are dismissed after one period. But this in turn means that in their first period, all agents will take their myopic actions, implying

$$(1 - \alpha)V(\pi) = \sum_{k \in K} \pi_k E(a_k^m),$$

and hence that again $V(\pi) \leq V(e_n)$.

The second part of the Proposition, that $V(\pi) \geq V(e_1)$ is proved analogously. $\square$

Therefore the performance effects present in any adverse selection scenario are never enough to overcome the possibility of “adverse” selection, namely, the possibility of retaining an incumbent who is less than the best type.

An intuitively rather appealing case we could consider—and one regarding which Proposition 3.7 has nothing to offer—is one where we introduce adverse selection by making the agent-type of the moral hazard model, an “average” agent-type of the general problem. Unfortunately, “average” has no single obvious meaning in our context. One option, for instance, is to suppose that the myopic actions of the agents in the adverse selection model imply the same expected payoffs for the principal as that from the single agent-type in the moral hazard model; that is, we have $\sum_{k=1}^n \pi_k E(a_k^m) = E(a_\ell^m)$, where ω_ℓ is the single type of the moral hazard model.

In this case, it is not too difficult to see that there are conditions under which the principal can improve from the presence of adverse selection. Pick any equilibrium of the type in Proposition 3.1, and suppose $r^* \in S$ is the cut-off in this equilibrium. Then, since all agents choose (weakly) higher actions than their myopic actions in the first period, the principal’s first period reward, denoted v_1 , is at least $E(a_\ell^m)$. Let $v_2(r)$ represent the principal’s second period reward from retaining an

incumbent who produced a reward of r in the first period. By Proposition 3.4, we have $v_2(r^*) \geq v_1$, and of course, by hypothesis, we have $v_1 \geq E(a_\ell^m)$, so we must have $v_2(r^*) \geq E(a_\ell^m)$. Since $v_2(\cdot)$ is a non-decreasing function, we are done. Note in particular, that if the environment is nice, the principal does *strictly* better in any cut-off equilibrium under adverse selection (provided, of course, that adverse selection is introduced in this particular manner).

A Proof of Proposition 3.1

Proposition 3.1 is proved via several lemmata. For ease of exposition, we present the proof in three steps, each contained in a separate subsection.

A.1 The Agents' Best-Response Problem

Let σ be any cut-off strategy for the principal, with cut-off point $\bar{r} \in S$. Henceforth, we will refer to this strategy simply by the point $\bar{r}$.

The Agents' Second Period Problem

In the second period, an agent of type ω_k who has been retained after generating a reward of r in the first period solves the following problem:

$$\max\{u(a, \omega_k) \mid a \in A\}.$$

Since u is continuous and A is compact, a maximum exists. Let

$$U(\omega_k) = \max\{u(a, \omega_k) \mid a \in A\}$$

$$M(\omega_k) = \arg \max\{u(a, \omega_k) \mid a \in A\}$$

Lemma A.1 *Let $a_k \in M(\omega_k)$, $a_l \in M(\omega_l)$, where $k > l$. Then $a_k \geq a_l$.*

Proof By definition of a_k and a_l , we have

$$u(a_k, \omega_k) \geq u(a_l, \omega_k)$$

$$u(a_l, \omega_l) \geq u(a_k, \omega_l).$$

Thus,

$$u(a_k, \omega_k) - u(a_k, \omega_l) \geq u(a_l, \omega_k) - u(a_l, \omega_l).$$

By the supermodularity of u , this last inequality implies that we must have $a_k \geq a_l$. □

Lemma A.2 *$U(\omega_k)$ is increasing in k .*

Proof This is immediate from Assumption A6. □

The Agents' First-Period Problem

In his first period, an agent of type ω_k solves

$$\max\{u(a, \omega_k) + \delta \int_{r \geq \bar{r}} U(\omega_k) dF(r|a) \mid a \in A\}.$$

For notational simplicity let

$$G(\bar{r}, a, \omega_k) = \delta \int_{r \geq \bar{r}} U(\omega_k) dF(r|a) = \delta U(\omega_k)[1 - F(\bar{r}|a)].$$

Then the problem is to solve

$$\max\{u(a, \omega_k) + G(\bar{r}, a, \omega_k) \mid a \in A\}.$$

Lemma A.3 *G is weakly supermodular in (a, ω_k) for each fixed $\bar{r}$.*

Proof By Lemma A.2, $U(\cdot)$ is increasing in the index k . Moreover, for any $\bar{r} \in S$, $[1 - F(\bar{r}|a)]$ is weakly increasing in a by Assumption A4 (the monotone-likelihood ratio property of the reward densities). As the product of weakly increasing functions of ω_k and a , respectively, the result follows. □

Lemma A.4 *G is continuous in its arguments.*

Proof This is an immediate consequence of the atomlessness of F and the joint continuity of f on $S \times A$. □

Now define

$$M^*(r, \omega_k) = \arg \max\{u(a, \omega_k) + G(r, a, \omega_k) \mid a \in A\}.$$

By Lemma A.4 and the compactness of A , M^* is non-empty for each (r, ω_k) . By Lemma A.3 and the supermodularity of u on $A \times \Omega$, the elements of $M^*(r, \cdot)$ are ordered on Ω for each fixed r ; that is, $k < l$, $a_k \in M^*(r, \omega_k)$, $a_l \in M^*(r, \omega_l)$ implies $a_k < a_l$. Finally, by the Maximum Theorem, M^* is an upper-semicontinuous correspondence on $S \times \Omega$. Thus, letting $\mathcal{M}^*(r, \omega_k) = P(M^*(r, \omega_k))$ be the set of mixed best actions for an agent of type ω_k (given the cut-off r), we obtain the following:

1. $\mathcal{M}^*$ is a non-empty valued, convex valued, upper-semicontinuous correspondence on $S \times \Omega$.

2. For each fixed r , the elements of $\mathcal{M}^*(r, \cdot)$ are ordered in ω in the sense that if $k < l$, $\mu_k \in \mathcal{M}^*(r, \omega_k)$, and $\mu_l \in \mathcal{M}^*(r, \omega_l)$, then there exist intervals $[a_k, b_k]$ and $[a_l, b_l]$ in A such that $b_k \leq a_l$ and $\mu_k([a_k, b_k]) = \mu_l([a_l, b_l]) = 1$.

The Principal's Best-Response Problem

Let anonymous strategies $(\mu, \gamma) = (\mu_k, \gamma_k)_{k=1}^n$ for the agents be given, where (i) $\mu_k \in P(A)$ is a mixed action that is the first period action of a type ω_k agent; and (ii) $\gamma_k: S \rightarrow A$ is a function that specifies for each realized first period reward $r \in S$, a second period action $\gamma_k(r) \in A$ for an agent of type ω_k . We suppose throughout that the strategies are monotone in type, i.e., that

1. For $k = 1, \dots, n-1$, $\gamma_{k+1}(r) \geq \gamma_k(r)$ for all $r \in S$.
2. There exist (possibly degenerate) intervals $[a_k, b_k]$ in A such that $a_{k+1} \geq b_k$ for $k = 1, \dots, n-1$, and

$$\mu_k([a_k, b_k]) = 1, \quad k = 1, \dots, n.$$

We will show that the principal's best response to such behavior by the agents is to use a cut-off strategy; and that the set of optimal cut-offs is an upper-semicontinuous convex valued correspondence in the agents' strategies.

So suppose that the first period reward witnessed was r , and let the principal's updated beliefs regarding the incumbent's type be given by $(\beta_k(r))$. Then, given the agents' strategies, the expected second period utility to the principal from retaining the incumbent, denoted $v_2(r, \mu, \gamma)$, is given by

$$v_2(r, \mu, \gamma) = \sum_{k=1}^n \left[\int_S \beta_k(r) v(\hat{r}) f(\hat{r} | \gamma_k(r)) d\hat{r} \right].$$

Now, recall from Section 2 that for any $m \in P(A)$, $\varphi(r|m)$ is the probability of observing the reward r under the mixed-action m :

$$\varphi(r|m) = \int_A f(r|a) m(da).$$

The following lemma implies almost immediately that $v(r, \mu, \gamma)$ is increasing in r .

Lemma A.5 *If $k > l$, then the ratio $\varphi(r|\mu_k)/\varphi(r|\mu_l)$ is non-decreasing in r .*

Proof We can express the given ratio as

$$\frac{\varphi(r|\mu_k)}{\varphi(r|\mu_l)} = \int_A \left[\frac{f(r|\hat{a})}{\int_A f(r|a) \mu_l(da)} \right] \mu_k(d\hat{a}).$$

Now

$$\left[\frac{f(r|\hat{a})}{\int_A f(r|a)\mu_l(da)} \right]^{-1} = \int_{a_l}^{b_l} \left[\frac{f(r|a)}{f(r|\hat{a})} \right] \mu_l(da).$$

Since the support of μ_k is contained in $[a_k, b_k]$, we may take $\hat{a} \in [a_k, b_k]$ without loss. Since $k > l$, we have $a_k \geq b_l$ by hypothesis, so $\hat{a} \geq b_l$. Thus, the bracketed term on the *RHS* of the last expression is non-increasing in r by Assumption A4 (the monotone likelihood ratio property). It follows easily that $\varphi(r | \mu_k)/\varphi(r | \mu_l)$ is non-decreasing in r . $\square$

Lemma A.6 $v_2(r, \mu, \gamma)$ is increasing in r .

Proof Let $r, \hat{r} \in S$ with $r > \hat{r}$. Let p denote the updated posterior belief $\beta(r)$ regarding the agent's type after witnessing r , and let $\hat{p}$ similarly denote $\beta(\hat{r})$. We claim that p “stochastically dominates” $\hat{p}$ in the sense that for any $\ell \in \{1, \dots, n\}$, we have

$$\sum_{k=\ell}^n p_k \geq \sum_{k=\ell}^n \hat{p}_k.$$

In words, this expression says that upon witnessing higher rewards, beliefs shift towards higher types. To see this, fix any ℓ . From the definition of p and $\hat{p}$, we are required to show that

$$\frac{\sum_{m=\ell}^n \pi_m \varphi(r|\mu_m)}{\sum_{k=1}^n \pi_k \varphi(r|\mu_k)} \geq \frac{\sum_{m=\ell}^n \pi_m \varphi(\hat{r}|\mu_m)}{\sum_{k=1}^n \pi_k \varphi(\hat{r}|\mu_k)}.$$

Cross-multiplying and cancelling common terms, this is the same as showing

$$\left(\sum_{k=1}^{\ell-1} \pi_k \varphi(\hat{r}|\mu_k) \right) \left(\sum_{m=\ell}^n \pi_m \varphi(r|\mu_m) \right) \geq \left(\sum_{k=1}^{\ell-1} \pi_k \varphi(r|\mu_k) \right) \left(\sum_{m=\ell}^n \pi_m \varphi(\hat{r}|\mu_m) \right),$$

or, what is the same thing,

$$\sum_{k=1}^{\ell-1} \sum_{m=\ell}^n [\pi_k \pi_m \varphi(\hat{r}|\mu_k) \varphi(r|\mu_m)] \geq \sum_{k=1}^{\ell-1} \sum_{m=\ell}^n [\pi_k \pi_m \varphi(r|\mu_k) \varphi(\hat{r}|\mu_m)].$$

But for each $k \in \{1, \dots, \ell-1\}$ and $m \in \{\ell, \dots, n\}$, it is the case by Lemma A.5 that

$$\varphi(\hat{r}|\mu_k) \varphi(r|\mu_m) \geq \varphi(r|\mu_k) \varphi(\hat{r}|\mu_m).$$

This establishes the required inequality, and therefore, that p stochastically dominates $\hat{p}$ as claimed.

Now, in the second period, all agents are going to take their myopically optimal actions, and by Lemma A.1, these actions are higher for higher types. Since beliefs shift towards higher types upon observing higher rewards, this means immediately that second-period expected rewards are increasing in the reward r generated in the first period. This proves Lemma A.6. $\square$

It is intuitively obvious from Lemma A.6 that a cut-off strategy is optimal for the principal. For, suppose $V^*(\mu, \gamma)$ represents the value function for the principal each time she hires a new agent. (By the stationarity of the problem, $V^*(\mu, \gamma)$ is independent of when this occurs.) Then, since any agent must be replaced after two periods, the choice facing the principal after a first period reward of r from the incumbent is of the form

$$\max\{v_2(r, \mu, \gamma) + \alpha V^*(\mu, \gamma), V^*(\mu, \gamma)\}.$$

Since $v_2(r, \mu, \gamma)$ is non-decreasing in r , it is immediate from this expression that a cut-off strategy must be optimal.

Unfortunately, formalizing these ideas (in particular, showing that $V^*(\mu, \gamma)$ is well-defined) is notationally tedious. On the other hand, this formalizing is necessary to show that this optimal cut-off varies continuously with (μ, γ) .

Given (μ, γ) , the best response problem facing the principal is a stationary dynamic programming problem (SDP) whose components are given by the following:

1. The *state space* is the Cartesian product of $\{0, 1\}$, S , and the unit simplex in $\mathbb{R}^n$, Δ^{n-1} . A typical state is denoted (n, r, p) , where n is the number of periods the current incumbent has left, r is the reward produced by the incumbent in the period immediately preceding, and p is the current belief regarding the incumbent's type.
2. The *action space* A is the two point set $\{\text{"keep"} (\kappa), \text{"dismiss"} (d)\}$, whose components are self-explanatory. The set of feasible actions $\Psi(n, r, p)$ at the state (n, r, p) is given by

$$\Psi(n, r, p) = \begin{cases} A, & \text{if } n = 1 \\ \{d\}, & \text{if } n = 0 \end{cases}$$

3. The *reward function* λ is defined as follows: when action d is taken at any state (n, r, p) , the immediate reward received is

$$\lambda((n, r, p), d) = \sum_{k=1}^n \left\{ \pi_k \left(\int_S \hat{r} \varphi(\hat{r} | \mu_k) d\hat{r} \right) \right\}$$

while the action κ at the state (n, r, p) leads to a reward of

$$\lambda((n, r, p), \kappa) = \sum_{k=1}^n \left\{ p_k \left(\int_S \hat{r} f(\hat{r} | \gamma_k(r)) d\hat{r} \right) \right\}$$

4. Finally, the *transition probabilities* $q(\cdot|(n, r, p), d)$ and $q(\cdot|(n, r, p), \kappa)$ are defined as follows. When the action d is taken at any state (n, r, p) , the new state is

$$(1, \hat{r}, \beta_k(\hat{r})) \quad \text{with probability} \quad \sum_k (\pi_k \varphi(\hat{r}|\mu_k)) d\hat{r}$$

If the action κ is taken at state $(1, r, p)$, then the new state is

$$(0, r, \pi) \quad \text{with probability} \quad 1.$$

This completes the description of the principal's best response problem. Since the state and action spaces are compact, the feasible action Ψ correspondence is continuous, the reward function is continuous, the reward function λ is continuous, and the transition probabilities q are weakly continuous, standard results in stationary dynamic programming show that this problem is well-defined; an optimal strategy exists; and the value function $W(s, r, v)$ is continuous in its arguments.

It is trivial that the Bellman equation for the problem at the state $(1, r, p)$ can be written as

$$W(1, r, p) = \max\{v_2(r, \mu, \gamma) + \alpha W(0, r, \pi), W(0, r, \pi)\},$$

where $v_2(r, \mu, \gamma)$ was defined earlier. Now let

$$C(\mu, \gamma, \pi) = \begin{cases} \sup S & \text{if } v_2(r, \mu, \gamma) < (1 - \alpha)W(0, r, \pi) \text{ for all } r \in S \\ \inf S & \text{if } v_2(r, \mu, \gamma) > (1 - \alpha)W(0, r, \pi) \text{ for all } r \in S \end{cases}$$

and let $C(\mu, \gamma, \pi) = \{r \mid v_2(r, \mu, \gamma) = (1 - \alpha)W(0, r, \pi)\}$, otherwise. Since $v_2(r, \mu, \gamma)$ is non-decreasing in r by Lemma A.6, it is clear that the principal's set of optimal actions given (μ, γ) is $C(\mu, \gamma; \pi)$, i.e., given any $\bar{r} \in C(\mu, \gamma; \pi)$, an optimal strategy for the principal when facing (μ, γ) is to retain the agent if and only if the first period reward exceeds $\bar{r}$.

Finally, note that in the SDP defining the principal's best response problem, the reward function λ is jointly continuous in states, actions, and the triple $(\mu, \gamma; \pi)$; while the transition probabilities are jointly weakly continuous in states, actions, and $(\mu, \gamma; \pi)$. It follows from Theorem 1 of Dutta, Majumdar, and Sundaram [11] that treating $(\mu, \gamma; \pi)$ as the parameters of this SDP, the solution (i.e. the optimal action correspondence) is jointly upper-semicontinuous in states and parameters. In particular, $C(\mu, \gamma; \pi)$ is upper-semicontinuous in $(\mu, \gamma; \pi)$.

Summing up, the principal's best response to the anonymous agent strategy (μ, γ) is to use a cut-off rule. The set of optimal cut-off rules $C(\mu, \gamma; \pi)$ define an upper-semicontinuous, convex valued correspondence in (μ, γ) and in π .

The Fixed-Point Theorem

The principal's space of strategies may evidently taken to be S . The agent's second period action does not depend on the principal's cut-off strategy; thus we may take the agent's second period action vector $(\gamma_1, \dots, \gamma_n)$ to be any measurable selection satisfying

$$\gamma_k(r) \in M(\omega_k) \quad \text{for all } r \in S.$$

(Recall that $M(\cdot)$ is the set of optimal second-period actions for the agents.) Pick any selection $\gamma = (\gamma_1, \dots, \gamma_n)$ and hold it fixed through the remainder of the argument. Note that by the monotonicity of $M(\cdot)$, we must have $\gamma_k(r) \geq \gamma_l(r)$ for all r whenever $k > l$.

Now define a subset Z^n of the n -fold Cartesian product of $P(A)$ as follows: a typical point $(\mu_1, \dots, \mu_n)$ in Z^n is such that

1. $\mu_k \in P(A)$ for all k .
2. There are intervals $[a_k, b_k]$ in A satisfying
 - (a) $b_k \leq a_{k+1}$, $k = 1, \dots, n-1$.
 - (b) $\mu_k([a_k, b_k]) = 1$.

From the agent's best response problem, we know that when the principal adopts a cut-off strategy r , the set of optimal first period actions for an agent of type ω_k is given by $\mathcal{M}^*(r, \omega_k)$, where $\mathcal{M}^*$ is an upper-semicontinuous, convex-valued correspondence. Define

$$B^*(r) = (\mathcal{M}^*(r, \omega_1), \dots, \mathcal{M}^*(r, \omega_n)).$$

By Lemmata A.3 and A.4, $B^*(\cdot)$ maps into Z^n ; it is also evidently a convex-valued upper-semicontinuous correspondence.

From the principal's best response problem, we know that for any $\mu \in Z^n$, the principal has an optimal response which is a cut-off strategy. The set of optimal cut-offs is given by $C(\mu, \gamma)$, which is an upper-semicontinuous, convex-valued correspondence from Z^n into S .

Now, let $D: S \times Z^n \rightarrow S \times Z^n$ be defined as

$$D(r, \mu) = (C(\mu, \gamma), B^*(r)).$$

Then D is itself an upper-semicontinuous convex-valued correspondence from $S \times Z^n$ into itself. Further, S is compact and convex by hypothesis, and it is easy to check that (under the weak topology) Z^n is also compact, and evidently convex. Therefore, the existence of a fixed-point follows from Glicksberg's Theorem (see, e.g., Smart [34]).

Such a fixed-point is evidently an equilibrium of the game. By construction it satisfies the remaining properties outlined in Proposition 3.1. $\square$

B Proof of Proposition 3.6

Let $\pi_l \rightarrow \pi$. Suppose $(\sigma_l^*, \mu_l^*) \in \mathcal{E}(\pi_l)$ for each l . Define Z^n as in the proof of Proposition 3.1 (see Appendix A). Since Z^n and S are compact, we may assume, without loss of generality, that there are r^* in S , and $\mu^* \in Z^n$ such that $r_l^* \rightarrow r^*$ and $\mu_l^* \rightarrow \mu^*$. (Here, in obvious notation, r_l^* is the cut-off associated with σ_l^* .)

Now, define the correspondence C as in the proof of Proposition 3.1 (again see Appendix A). Recall that C was established in the course of that proof to be upper-semicontinuous on $Z^n \times A^n \times \Delta^{n-1}$. By definition, $r_l^* \in C(\mu_l^*, \pi_l)$ for each l , establishing $r^* \in C(\mu^*, \pi)$.

Finally, define the correspondence B^* as in the proof of Proposition 3.1. Then, we have $\mu_l^* \in B^*(r_l^*)$. From the proof of Proposition 3.1, we know that B^* is an upper-semicontinuous correspondence on Z^n . Therefore, $\mu^* \in B^*(r^*)$, and thus $(\sigma^*, \mu^*) \in \mathcal{E}(\pi)$ completing the proof. $\square$

References

- [1] Abreu, D.; D. Pearce, and E. Stacchetti (1990) Toward a Theory of Discounted Repeated Games with Imperfect Monitoring, *Econometrica* 58(5), 1041-1064.
- [2] Amir, R.; L. Mirman, and W. Perkins (1991) One-Sector Nonclassical Optimal Growth, *International Economic Review* 32, 625-644.
- [3] Austen-Smith, D. and J. Banks (1989) Electoral Accountability and Incumbency, in P. Ordeshook (ed.) *Models of Strategic Choice in Politics*. Ann Arbor: University of Michigan Press.
- [4] Banks, J.S. and R.K. Sundaram (1993) Moral Hazard and Adverse Selection in a Model of Repeated Elections, in W. Barnett et.al. (eds.) *Political Economy: Institutions, Information Competition, and Representation*. New York: Cambridge University Press.
- [5] Baron, D. and D. Besanko (1987) Commitment and Fairness in a Dynamic Regulatory Relationship, *Review of Economic Studies* 54, 413-436.
- [6] Barro, R. (1973) The Control of Politicians: An Economic Model, *Public Choice* 14, 19-42.
- [7] Besley, T. and A. Case (1995) Does Electoral Accountability affect Economic Policy? Evidence from Gubernatorial Term-Limits, *Quarterly Journal of Economics* 110, 769-798.
- [8] Cadot, O. and B. Sinclair-Despaigne (1992) Prudence and Success in Politics, *Economics and Politics* 4(2), 171-190.
- [9] Coate, S. and S. Morris (1995) On the Form of Transfers to Special Interests, *Journal of Political Economy* 103, 1210-1235.
- [10] Dewatripont, M. (1989) Renegotiation and Information Revelation Over Time in Optimal Labor Contracts, *Quarterly Journal of Economics* 104, 589-620.
- [11] Dutta, P.K.; M. Majumdar, and R.K. Sundaram (1994) Parametric Continuity in Dynamic Programming Problems, *Journal of Economic Dynamics and Control* 18, 1069-1092.
- [12] Dutta, P.K. and R. Radner (1987) Principal-Agent Games in Continuous Time, mimeo, AT&T Bell Laboratories; revised, 1991.
- [13] Ferejohn, J. (1986) Incumbent Performance and Electoral Control, *Public Choice* 50, 5-25.
- [14] Fudenberg, D.; B. Holmstrom, and P. Milgrom (1986) Short-Term Contracts and Long-Term Relationships, *Journal of Economic Theory* 51, 1-31.
- [15] Gilson, S. (1989) Management Turnover and Financial Distress, *Journal of Financial Economics* 25, 241-262.
- [16] Gould, C. (1997) Paying Fund Managers with Carrots and Sticks, *The New York Times*, 9 February 1997.
- [17] Grossman, S. and O. Hart (1983) An Analysis of the Principal-Agent Problem, *Econometrica* 51, 7-46.

- [18] Heinkel, R. and N. Stoughton (1994) The Dynamics of Portfolio Management Contracts, *Review of Financial Studies* 7, 351–387.
- [19] Kahn, C. and G. Huberman (1988) Two Sided Uncertainty and “Up-or-Out” Contracts, *Journal of Labor Economics* 6(4), 423-444.
- [20] Kreps, D.; P. Milgrom, J. Roberts, and R. Wilson (1982) Rational Cooperation in the Finitely Repeated Prisoners’ Dilemma, *Journal of Economic Theory* 27, 245-252.
- [21] Laffont, J.J. and J. Tirole (1986) The Dynamicsa of Incentive Contracts, *Econometrica* 53, 1173-1198.
- [22] Laffont, J.J. and J. Tirole (1990) Adverse Selection and Renegotiation in Procurement, *Review of Economic Studies* 75, 597-626.
- [23] Lakonishok, J; A. Shleifer, and R. Vishny (1992) The Structure and Performance of the Money Management Industry, *Brookings Papers on Economic Activity: Microeconomics* 24, 339–391.
- [24] O’Flaherty, B. and A. Siow (1991) On-the-Job Screening, Up or Out Rules, and Firm Growth, mimeo, Columbia University and the University of Toronto.
- [25] Milgrom, P. (1981) Good News and Bad News: Representation Theorems and Applications, *Bell Journal of Economics* 12, 380-391.
- [26] Milgrom, P. and J. Roberts (1990) Rationalizability, Learning, and Equilibrium in Games with Startegic Complementarities, *Econometrica* 49, 1127-1148.
- [27] Milgrom, P. and C. Shannon (1994) Monotone Comparative Statics. *Econometrica* 62, 157–180.
- [28] Radner, R. (1981) Monitoring Cooperative Agreements in a Repeated Principal-Agent Relationship, *Econometrica* 49, 1127-1148.
- [29] Radner, R. (1985) Repeated Principal-Agent Games with Discounting, *Econometrica* 53, 1173-1198.
- [30] Radner, R. (1986) Repeated Moral Hazard with Low Discount Rates, in *Essays in Honor of Kenneth Arrow* (Heller, W.; R. Starr, and D. Starrett, Eds.), Cambridge University Press, Cambridge.
- [31] Reed, W.R. (1991) Electoral Control of Heterogeneous Politicians, mimeo, Department of Economics, Texas A & M University.
- [32] Rogerson, W. (1985) Repeated Moral Hazard, *Econometrica* 53, 69-77.
- [33] Rogoff, K. (1990) Equilibirum Political Budget Cycles, *American Economic Review* 80, 21–37.
- [34] Smart, D. (1974) *Fixed Point Theorems*, Cambridge University Press, London/New York.
- [35] Spence, M. (1973) Job-Market Signalling, *Quaterly Journal of Economics* 87, 355-374.
- [36] Waldman, M. (1990) Up or Out Contracts: A Signalling Approach, *Journal of Labor Economics* 8, 230-250.