

Accepted Manuscript

Moral Hazard in Strategic Decision Making

Martin C. Byford

PII: S0167-7187(17)30126-1
DOI: [10.1016/j.ijindorg.2017.10.001](https://doi.org/10.1016/j.ijindorg.2017.10.001)
Reference: INDOR 2395



To appear in: *International Journal of Industrial Organization*

Received date: 22 February 2017
Revised date: 18 September 2017
Accepted date: 1 October 2017

Please cite this article as: Martin C. Byford, Moral Hazard in Strategic Decision Making, *International Journal of Industrial Organization* (2017), doi: [10.1016/j.ijindorg.2017.10.001](https://doi.org/10.1016/j.ijindorg.2017.10.001)

This is a PDF file of an unedited manuscript that has been accepted for publication. As a service to our customers we are providing this early version of the manuscript. The manuscript will undergo copyediting, typesetting, and review of the resulting proof before it is published in its final form. Please note that during the production process errors may be discovered which could affect the content, and all legal disclaimers that apply to the journal pertain.

Highlights

- I develop a theory of moral hazard in which the agent is a strategic decision-maker.
- Career concerns cause the agent to manipulate the project's risk-return tradeoff.
- The moral hazard problem can be ameliorated by outcome contingent contracts.
- The optimal incentive contract involves 'in-the-money' options.
- The optimal tenure standard requires the agent to exceed expectations.

Moral Hazard in Strategic Decision Making*

Martin C. Byford[†]

Abstract

This paper develops a theory of moral hazard in which the agent takes the role of strategic decision-maker. Career concerns give rise to preferences over risk, which in turn create an incentive for the agent to manipulate the project's risk-return tradeoff to the disadvantage of the principal. The resultant moral hazard can be ameliorated by an incentive contract. The optimal non-decreasing wage involves granting 'in-the-money' options. In the context of academic tenure, the optimal tenure standard requires the agent to exceed expectations, and lies within one standard deviation of the expected outcome.

THIS VERSION: OCTOBER 5, 2017

J.E.L. Classification: D83, D86.

Keywords: moral hazard, strategic decision-making, career concerns, in-the-money stock-options, tenure.

* I am grateful to Michael Harris, Guillaume Roger and an anonymous referee for comments and advice on earlier versions of this paper.

[†] School of Economics, Finance and Marketing, RMIT University. Building 80 Level 11, 445 Swanston Street, Melbourne VIC 3000, Australia. Email: martin.byford@rmit.edu.au.

1 Introduction

In the principal-agent literature, moral hazard refers to the problem of directing an agent's actions when those actions are either unobservable, or cannot be contracted upon. The usual analysis focuses on the problem of motivating an agent to exert effort (early references include [Mirrlees, 1976, 1999](#); [Holmstrom, 1979](#); [Rogerson, 1985](#)). In contrast, moral hazard that emerges when the agent's role is to be the strategic decision-maker for a project, has received relatively little attention.

The problem of moral hazard in strategic decision-making is, perhaps, nowhere more evident than in the case of a firm's CEO. The CEO is hired by the owners of the firm to set and implement a strategic direction for their firm. Among the many strategic decisions facing the CEO are pricing the firm's products, selecting the size of production runs, developing an R&D strategy, and choosing which markets to contest. Any failure by the CEO to take these decisions in the best interests of shareholders can create substantial costs for the owners of the firm.

In this context, effort motivation is unlikely to be a primary concern for the firm's owners. As [Holmstrom and Ricart i Costa \(1986\)](#) argue, senior managers are not typically effort averse. Indeed, in many ways, the career path to senior management seems designed to identify individuals who experience a particularly low disutility of effort. Rather, the firm's owners are more likely to be concerned with the CEO's level of ability and attitude to risk.

A dedicated theory of moral hazard in strategic decision-making is required because strategic decision-making has three characteristics that distinguish it from the standard effort motivation problem:

1. *The strategy that maximises the expected return lies in the interior of the choice set:*

In strategic decision-making, more is not always better. For example, the optimal capacity for a factory is unlikely to be the maximum capacity that a firm can construct. Likewise, the research project with the highest expected return (academic or commer-

cial) is unlikely to be the most ambitious project available. In order to capture this characteristic, the assumption employed in this paper is that the expected outcome of a project is concave quadratic in the strategy chosen by the agent, with a maximum in the interior of the choice set.

2. *There exists a trade-off between risk and expected return:* Another common feature of strategic decision-making is that the chosen strategy jointly determines a project's risk and expected return. To capture this feature, it is assumed that the standard deviation of the project's outcome is linear and increasing the strategy. When combined with the interior optimum, this assumption lends itself to the characterisation that low strategies are 'too safe,' while high strategies are 'too risky.'
3. *The agent is indifferent between the available strategies:* In principle, the strategy chosen by the agent need have no direct impact on her utility. In terms of the agent's time and effort, one strategy need be no more costly to implement than another. While this is certainly a strong assumption, it is useful insofar as it isolates the indirect effects of the agent's choice of strategy. Specifically, the way in which the choice of strategy impacts on the agent's career concerns and compensation.¹

One might reason that moral hazard would not be a significant factor in strategic decision-making. Given that a CEO's future career prospects tend to be linked to the performance of their firm, both the CEO and shareholders share a common incentive to maximise the firm's profits. However, as [Holmstrom and Ricart i Costa \(1986\)](#) (see also [Holmstrom, 1999](#)) demonstrate, a CEO's career concerns can also give rise to preferences over the firm's level of risk. This in turn creates an incentive for the CEO to manipulate the firm's risk-return tradeoff to the detriment of the firm's owners. Empirical evidence of such an effect has been

¹Contrast these features with those of the standard effort motivation model. In the standard treatment, more effort is assumed to always improve the outcome of a project in the sense of first-order stochastic dominance (see [Milgrom, 1981](#), for a discussion). The effect of effort on project risk is thus irrelevant as any second-order effects are dominated by the first-order effects. Moreover, the agent is assumed to be effort-averse, creating a tension between the preferences of the principal and agent out the outset.

found by [Chevalier and Ellison \(1999\)](#), amongst others.

The two papers in the prior literature that come closest to addressing the problem of moral hazard in strategic decision-making, are [Holmstrom and Ricart i Costa \(1986\)](#) and [Manso \(2011\)](#). In [Holmstrom and Ricart i Costa \(1986\)](#), a manager must choose whether or not to commission a risky project. They show that, absent appropriate incentives, the manager may reject a project with positive expected profits if failure would provide a negative signal of the manager's talent, harming the manager's future earnings. In a similar vein, [Manso \(2011\)](#) investigates the problem of motivating a manager to explore risky new technologies, when it is possible to exploit existing technologies with well known characteristics. In both papers the authors show that the optimal contract must tolerate failure (at least early on in the agent's career) in order to encourage the agent to take risks as desired by the principal. These results are predicated on the assumption that the principal can and will guarantee continued employment, pay and conditions for the agent, regardless of her performance; an assumption that is at odds with this paper's motivating examples of firm CEOs and tenure track faculty.

This paper builds on these earlier works by considering a richer class of strategic decision-making problems, more akin to the types of strategic decisions that are regularly studied in the Industrial Organization and Strategy literatures. Whereas [Holmstrom and Ricart i Costa \(1986\)](#) and [Manso \(2011\)](#) each consider binary decisions by the agent, in this paper the agent chooses a strategy from a continuum. This allows for the marginal effect of the agent's actions on the project's risk-return tradeoff, to be analysed.

A further difference from the prior literature is that the only tools available to the principal are assumed to be single-period outcome contingent contracts. Neither the principal nor the agent can make ex-ante commitments to future employment on which they may subsequently wish to renege. Thus the principal cannot 'insure' the agent's career against failure in the manner of [Holmstrom and Ricart i Costa \(1986\)](#) and [Manso \(2011\)](#). Instead, the remedies for the moral hazard problem must be housed entirely within the current period's contract.

The non-monotonic likelihood ratio that arises from the model, at once predicts ‘in-the-money’ option contracts for executives, and tenure standards that require faculty to exceed expectations.

The paper proceeds as follows: The career concerns model, adapted from Dewatripont, Jewitt and Tirole (1999a,b), is presented in section 2. A risk-averse agent lives and works for two periods. Each period the agent seeks employment in the labour market. The employers in the labour market are the owners of projects who require the specialised skills of an agent to select and implement a strategy for their projects.

In section 3 it is shown that the agent’s indifference between the strategies in the choice set implies that the first-best outcome is achievable in the second and final period of the game. If the agent is offered a fixed wage contract then, insulated from the project’s risk, she will weakly prefer to implement the strategy that maximises the expected outcome.

In contrast, the agent’s first-period behaviour will be influenced by the risk preferences derived from her career concerns. The labour market uses the outcome of the agent’s first-period project to update its beliefs about the agent’s talent. In turn, these beliefs determine the agent’s second-period wage. In this way, the risk-averse agent is exposed to the risk associated with her first-period project, creating an incentive for her to select a risk-averse strategy at the expense of the expected outcome.²

The problem of utilising an incentive contract to ameliorate the first-period moral hazard problem is considered in section 4. An incentive contract influences the agent’s behaviour by rewarding outcomes that are indicative of the agent choosing the desired action. In the strategic decision-making problem, extreme outcomes, both high and low, become more likely if the agent implements a risky strategy. Conversely, moderate outcomes are indicative of the agent adopting a risk-averse strategy.

There are a number of practical problems associated with a contract that pays the agent a high wage when her project produces a low outcome, and a low wage when the project

²In contrast, Dewatripont, Jewitt and Tirole (1999a,b) explore the extent to which career concerns motivate the agent to exert effort in the first period, absent an incentive contract.

performs more or less as expected. Notably, a wage that rewards failure in this way, creates an incentive for the agent to sabotage the project or implement a perverse strategy. A straightforward solution is to impose an additional constraint on the contracting problem, requiring that the agent's wage be non-decreasing in the project outcome.

The optimal non-decreasing wage involves granting the agent 'in-the-money' options on the project's outcome. The discount at which the options are offered is increasing in both the risk associated with the project, and the magnitude of the moral hazard problem.³ This result also highlights the costs of tax policies that penalise in-the-money options. Such policies prevent the principal from exploiting the information present in a range of project outcomes, and result in incentive contracts with higher cost to the principal or lower incentive power (or both).

Another principal-agent relationship in which there may exist moral hazard in strategic decision-making, is the relationship between university and academic. High quality research output requires risk-taking. However, tenure-track faculty may prefer to 'play it safe' with their research agenda, minimising the risk of a failed project that would damage their future employment prospects. In this relationship, the prize of academic tenure creates incentive that can be used ameliorate the moral hazard problem. The optimal tenure-standard is derived in section 5. The optimal standard requires the agent to exceed expectations, and lies within one standard deviation of the expected outcome.

The paper concludes with a discussion of the results.

2 The career concerns model

The career concerns model employed in this paper is adapted from [Dewatripont, Jewitt and Tirole \(1999a,b\)](#). An *agent* lives and works for two periods. Each period the agent seeks employment in the labour market. The employers in the labour market are the owners of

³Within the moral hazard literature, option contracts have previously been shown to be optimal in the presence of loss aversion ([de Meza and Webb, 2007](#)) and bounds on payments ([Jewitt, Kadan and Swinkels, 2008](#)). In each case the problem under consideration was motivating an agent to exert effort.

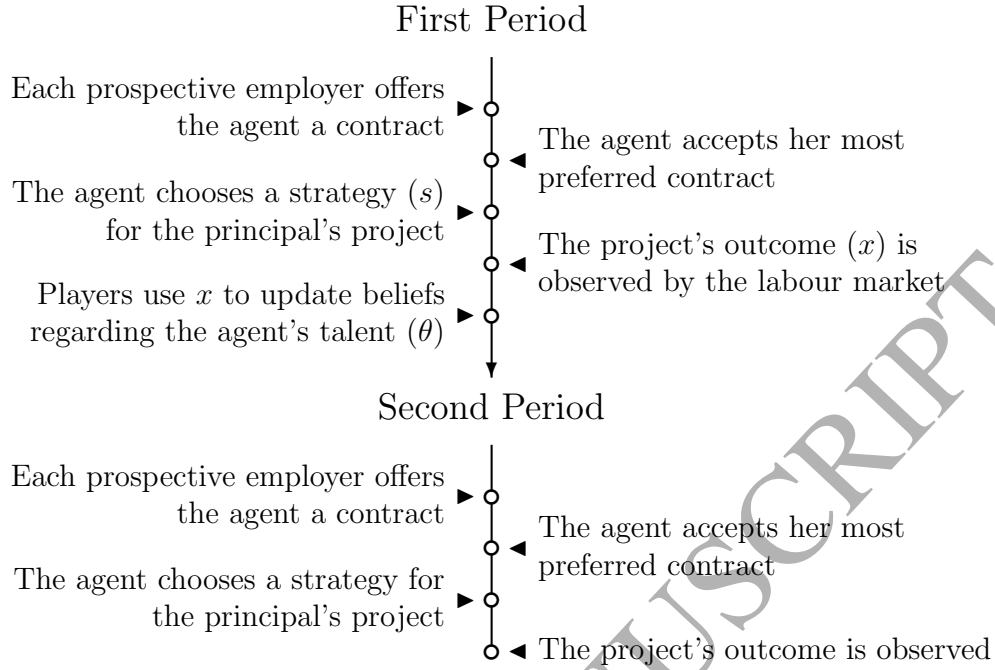


Figure 1: The Manager's Career

projects, the principals in the principal-agent problem. Each *principal* requires the specialised skills of an agent to oversee the completion of a project.

Contracts in the labour market apply for a single period. Neither the agent, nor the principal, can make commitments in the first period that dictate employment, wages or behaviour in the second period. However, because the agent must re-enter the labour market following the first period, her second-period career concerns create incentives that may influence her first-period behaviour.

2.1 Timing

The timing of the model, illustrated in figure 1, is as follows: The first-period begins with the labour market. Prospective employers approach the agent with take-it-or-leave-it offers of a contract. The agent reviews the contracts on offer and accepts her preferred contract.

Following the labour market, the agent commences work on the principal's project. The agent's task is to select a strategy for the project and oversee its implementation. The agent's

actions are hidden from both the principal and the wider labour market.

The project's outcome is realised at the end of the first period. The outcomes is stochastic and contingent on both the strategy implemented by the agent, and the agent's talent. All labour market participants, including the principal, observes the project's outcome and use the outcome to update their beliefs regarding the agent's talent.

The agent comes out of contract between periods, re-entering the labour market at the start of the second period. The second period proceeds with the same timing as the first. The only difference between periods is that there is no need for the market to update beliefs regarding the agent's talent at the conclusion of the second period as this is the final period of the agent's career.

Note that the assumption that the agent re-enters the labour market following the first period, does not preclude the agent from being employed by the same principal in both periods. Rather, it requires the principal to match the agent's best outside offer, following her first-period performance, if the principal is to retain the agent's services in the second period. Moreover, it allows for the agent to lose her position if first-period performance indicates that the principal would be better off replacing her.

2.2 The agent

The agent is risk-averse and motivated only by the wages she receives over her career. Specifically, the agent's lifetime utility takes the form,

$$U(w, v) = u(w) + \delta u(v), \quad (1)$$

where w and v are the agent's first- and second-period wages respectively, and $\delta \in (0, 1)$ is the agent's discount factor. The agent's utility is strictly increasing concave in the wages she receives ($u'(\cdot) > 0$ and $u''(\cdot) < 0$).

A key assumption of this paper, outlined in the introduction and implied by (1), is that the agent's utility is independent of the actions she takes on a project. In other words, it is assumed that the agent does not have innate preferences over the set of available strategies.

Also implicit in (1) is the assumption that the agent can neither borrow against future earnings, nor save for future consumption. This assumptions is not necessary, but does simplify the analysis considerably.

The agent's productivity on a project is, in part, determined by her innate talent $\theta \in \mathbb{R}$. However, knowledge of the agent's talent is imperfect. At the start of the first period, all labour market participants share the agent's prior belief that θ is drawn from the normal distribution with mean $\bar{\theta}$ and variance σ^2 . The assumption of a common prior is standard in the career concerns literature as it removes adverse selection from the model (Holmstrom, 1999; Harris and Holmstrom, 1982; Holmstrom and Ricart i Costa, 1986; Dewatripont, Jewitt and Tirole, 1999a,b).

2.3 The agent's role in the project

The agent's task on a project is to select and implement a strategy s from the set of all viable strategies S , where S is an interval on $\mathbb{R}_{++}$. The outcome of the project $x \in \mathbb{R}$, is then a function of the agent's choice of strategy s , the agent's talent θ , and a random shock $\varepsilon \sim \mathcal{N}(0, 1)$.⁴ Specifically,

$$x = s\omega - \frac{s^2}{2} + \theta + s\varepsilon, \quad (2)$$

where the parameter ω is assumed to lie in the interior of S . It follows that x is, itself, a normally distributed random variable with probability density function (PDF),

$$f(x|s, \theta) = \frac{1}{\sqrt{2\pi}s} e^{-\frac{(x-s\omega+\frac{1}{2}s^2-\theta)^2}{2s^2}}.$$

Given the ex-ante uncertainty over the agent's talent, the PDF of the marginal distribution can be written,

$$\hat{f}(x|s) = \frac{1}{\sqrt{2\pi(s^2 + \sigma^2)}} e^{-\frac{(x-s\omega+\frac{1}{2}s^2-\bar{\theta})^2}{2(s^2+\sigma^2)}},$$

which likewise is normal.

⁴That is to say, the shock ε is drawn from the standard normal distribution.

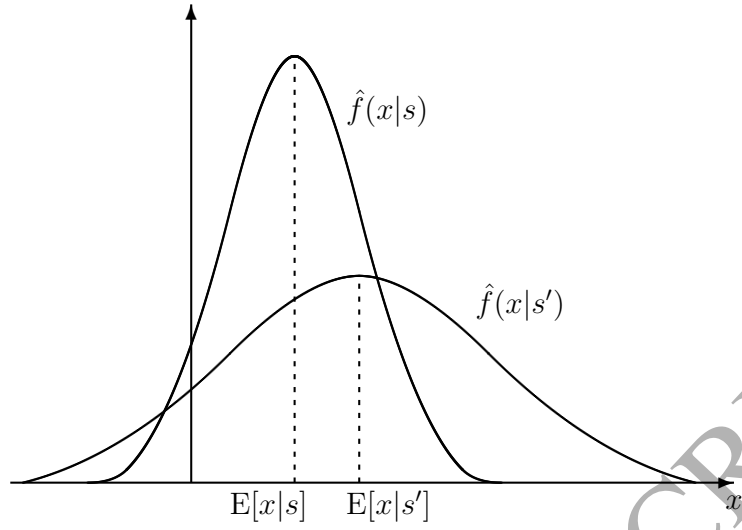


Figure 2: PDFs with $s < s' \leq \omega$

This production technology captures the two key features of strategic decision-making described in the introduction. The expected outcome,

$$E[x|s] = s\omega - \frac{1}{2}s^2 + \bar{\theta},$$

is strictly concave in s . While the strategy $s = \omega$, which maximises the expected outcome, lies in the interior of the choice set. At the same time, there exists a tradeoff between risk and expected return as the variance of the marginal distribution,

$$\text{Var}[x|s] = s^2 + \sigma^2,$$

is strictly increasing in s .

Figure 2 illustrates how the tradeoff between risk and expected return affects the shape of the marginal distribution's PDF. If the agent selects a low strategy $s < \omega$, the variance of the marginal distribution is low. The bulk of the PDF is tightly spaced around the expected outcome. If instead the agent selects a higher strategy s' , that lies closer to ω , both the expected outcome and the uncertainty surrounding that outcome will likewise be higher. Notice that the tails of the distribution are thicker with the strategy s' . Extreme outcomes, both high and low, become more likely as the strategy, and hence the project

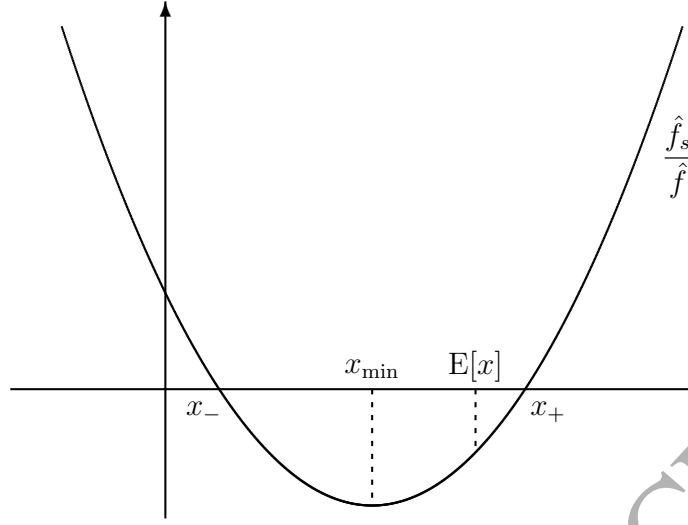
risk, increases; while intermediate outcomes become less likely. Not illustrated in figure 2 are the consequences of increasing the strategy past ω . In this range the variance continues to increase in s , while the expected outcome begins to decline.

The tradeoff between risk and expected outcome, captured in (2) and illustrated in figure 2, plays an important role in a number of principal-agent relationships. Some prominent examples include:

Firm Management: The owners of a firm grant a manager considerable autonomy to determine the strategic direction of the firm. For example, the manager may be responsible for choosing the capacity of a new factory, the quality of a new product, or the rate at which the firm expands its operations. In each of these problems the expected profit maximising strategy typically lies in the interior of the choice set. At the same time, the manager's decision is likely to affect the amount of risk carried by the firm. For example, in the face of uncertain demand, large scale expansion is likely to be riskier than small scale expansion.

Research and Development: When embarking on a project, a researcher must decide on how far to push beyond the existing technological/knowledge frontier. Research projects that lack ambition are unlikely to lead to interesting/profitable results, while overly ambitious projects are unlikely to be successful. Once again, the choice that maximises the expected outcome typically lies in the interior of the researcher's choice set. At the same time, the further the project ventures into the unknown, the riskier it becomes.

Creative Endeavours: The producers of creative works — writers, musicians, film makers, and the like — face the decision of how much novel content to include in new work. Lack of novel content is unlikely to lead to commercial success as consumers have “seen it all before.” However, too much novel content can make a work inaccessible to consumers. Because consumer reaction to novel content is uncertain, the more novel the project, the riskier it becomes.


 Figure 3: Likelihood ratio with $s < \omega$

2.4 Strategy, talent and outcomes

The rate at which an outcome becomes relatively more likely as s increases is captured by the *likelihood ratio*. The likelihood ratio reveals additional characteristics of the relationship between the strategy chosen by the agent and the outcome of the project. It also plays an important role in the construction of incentives contracts (see sections 4 and 5). Formally, the likelihood ratio is found by taking the partial derivative of the marginal distribution with respect to s , and dividing through by $\hat{f}$. This yields,

$$\frac{\hat{f}_s(x|s)}{\hat{f}(x|s)} = \frac{s(x - s\omega + \frac{1}{2}s^2 - \bar{\theta})^2}{(s^2 + \sigma^2)^2} - \frac{(s - \omega)(x - s\omega + \frac{1}{2}s^2 - \bar{\theta})}{s^2 + \sigma^2} - \frac{s}{s^2 + \sigma^2}, \quad (3)$$

which is a convex quadratic in x for all $s, \omega \in S$. The likelihood ratio is illustrated in figure 3.

The convex quadratic shape of the likelihood ratio confirms that, for all parameter values, a marginal increase in s makes extreme outcomes more likely while intermediate outcomes become less likely. It follows that extreme outcomes, both high and low, are indicative of risk taking, while intermediate actions are indicative of a risk-averse strategy choice.

The non-monotonicity of the likelihood ratio also has implications for the way in which decision-makers rank the strategies within S . It follows from Milgrom (1981) that given any

two strategies $s, s' \in S$, the resultant marginal distributions $\hat{f}(x|s)$ and $\hat{f}(x|s')$, cannot be ranked in the sense of first-order stochastic dominance. As demonstrated in section 3, this means that the principal and agent may have different preferences over the choice of strategy, even where their respective payoffs depend only on the outcome, and are strictly increasing in the outcome.

The outcome of the agent's first-period project also carries information that all labour market participants can use to update their prior beliefs about the agent's talent. Given that both the prior beliefs and shock ε are normally distributed, the first-period posterior will likewise be normally distributed with mean,

$$E[\theta|x] = \frac{(s^*)^2 \sigma^2}{(s^*)^2 + \sigma^2} \left(\frac{\bar{\theta}}{\sigma^2} + \frac{x - s^* \omega + \frac{1}{2}(s^*)^2}{(s^*)^2} \right). \quad (4)$$

In this expression s^* is the solution to the agent's first-period optimisation problem. That is, the strategy the market infers that the agent selected. The expectation in (4) is then the average of the prior expectation, and that part of the observed outcome that cannot be explained by the agent's inferred action, weighted by the inverse of the variances of the prior and outcome respectively.

The expectation of the first-period posterior has two features which drive the agent's career concerns: First, using (2) to substitute for x into (4), and taking the expectation with respect to the marginal distribution, demonstrates that the posterior has the Martingale property,

$$E[E[\theta|x]|s^*] = \bar{\theta}.$$

Second, for all $s^*, \omega \in S$, the mean of the posterior is strictly increasing and linear in the observed outcome x . In other words, the higher the outcome the agent produces in the first period, the more talented the agent is regarded in the second-period labour market.

2.5 The principal

The final element of the career concerns model is the principal, the owner of the project. The principal's payoff is the project's outcome net of the wage paid to the agent. It is assumed

that the principal is risk-neutral as the principal is able to diversify risk across numerous projects.

3 Moral hazard

Moral hazard arises in the model when there exists a tension between the preferences of the principal and agent. In this section it is shown that no such tension exists in the second period, however career concerns create divergent preferences over risk in the first period.

3.1 First-best in the second period

The second period is the final period of the agent's (working) life and as such career concerns do not influence the agent's behaviour in this period. Because the agent has no innate preferences over the strategies available to a firm, it is possible to achieve a first-best outcome in the second period.

To see this, first note that the agent prefers a fixed-wage contract over an incentive contract with the same expected wage, as a fixed-wage contract insulates the risk-averse agent from risk. Prospective employers in the second-period labour market thus have no reason to offer the agent an outcome contingent contract, as the agent already weakly prefers to implement the strategy that maximises the project's expected outcome. This is first-best risk sharing, as the risk-neutral principal carries all the risk, and results in the first-best strategy choice.

The magnitude of the agent's second-period wage is of significance. In a competitive labour market the agent's second-period wage will be her (expected) marginal product; the expected outcome of her second-period project given that she implements the strategy ω . In turn, this expectation depends on the first-period posterior beliefs concerning the agent's talent. It then follows from (2) and (4) that the agent's second-period wage will be,

$$v = \frac{\omega^2}{2} + E[\theta|x] = \frac{\omega^2}{2} + \frac{(s^*)^2 \sigma^2}{(s^*)^2 + \sigma^2} \left(\frac{\bar{\theta}}{\sigma^2} + \frac{x - s^* \omega + \frac{1}{2}(s^*)^2}{(s^*)^2} \right), \quad (5)$$

where x is the outcome of the agent's first-period project, and s^* is the solution to the agent's first-period optimisation problem.

3.2 Moral hazard in the first period

The dependence of the agent's second-period wage on her first-period outcome creates career concerns for the agent in the first period. At first glance it may appear as though the agent's career concerns align her incentives with those of the agent's first-period employer. After all, (5) states that the agent's second-period wage is linear and strictly increasing in the project's outcome. However, this reasoning neglects the fact that the agent's career concerns expose her to the risk associated with x .

When the agent selects a strategy in the first period, the impact of this choice on the expectation of her second-period utility is given by the function,

$$\psi(s) \equiv \delta \int_{\mathbb{R}} u \left(\frac{\omega^2}{2} + \frac{(s^*)^2 \sigma^2}{(s^*)^2 + \sigma^2} \left(\frac{\bar{\theta}}{\sigma^2} + \frac{x - s^* \omega + \frac{1}{2}(s^*)^2}{(s^*)^2} \right) \right) \hat{f}(x|s) dx. \quad (6)$$

Two important characteristics of this career concerns function are established in the following proposition.

Proposition 1. *For all $s^* \in S$, the agent's career concerns have the following properties:*

- (a) *The function $\psi(s)$ is strictly concave in s .*
- (b) *The action $s_\psi \equiv \arg\max_{s \in S} \psi(s)$ is unique and strictly less than ω .*

Proof. The proof of (a) is adapted from the discussion in section 3 of Jewitt (1988). Define the function,

$$g(s, \varepsilon) = s\omega - \frac{s^2}{2} + \bar{\theta} + s\varepsilon.$$

Expressed in this way, $g(\cdot)$ is the outcome x , described as a function of the agent's action s , and a random variable $\varepsilon \sim \mathcal{N}(0, 1)$. Note that, for all $\varepsilon \in \mathbb{R}$, $g(\cdot)$ is strictly concave in s . Now define,

$$h(s, \varepsilon) = u \left(\frac{\omega^2}{2} + \frac{(s^*)^2 \sigma^2}{(s^*)^2 + \sigma^2} \left(\frac{\bar{\theta}}{\sigma^2} + \frac{g(s, \varepsilon) - s^* \omega + \frac{1}{2}(s^*)^2}{(s^*)^2} \right) \right).$$

Given that a strictly increasing concave transformation of a strictly concave function is, itself, strictly concave, $h(\cdot)$ is likewise strictly concave in s for all $\varepsilon \in \mathbb{R}$. Moreover,

$$\delta E[h(s, \varepsilon)] = \psi(s),$$

is strictly concave in s as the sum of strictly concave functions is strictly concave.

For (b), first note that the agent's second-period utility is strictly increasing, and strictly concave, in the first-period outcome x . Moreover,

$$\left. \frac{\partial}{\partial s} E[x|s] \right|_{s=\omega} = 0 \quad \text{and} \quad \left. \frac{\partial}{\partial s} \text{Var}[x|s] \right|_{s=\omega} = 2\omega > 0.$$

It follows that $\psi'(\omega) < 0$ as a marginal decrease in s , in the neighbourhood of $s = \omega$, reduces the variance of x without altering the expectation of x . The result then follows from the strict concavity of $\psi(s)$. $\square$

The intuition behind proposition 1 is straightforward. It has already been established that the agent's second-period wage is strictly increasing in the observed first-period outcome. This means that the agent is exposed to the risk associated with this outcome. But risk is costly for the risk-averse agent to bear. It follows that, to some degree, the agent is willing to sacrifice the expected outcome to reduce the project's risk. Given that, in the neighbourhood of ω , marginal changes in s affect the variance of the outcome but not the expectation, it must be the case that the agent prefers a strategy s_ψ that creates less project risk than ω .

Proposition 1 establishes the existence of a moral hazard problem in the first period. The principal would like the agent to implement the strategy ω that maximises the expected outcome, while the agent prefers the strategy s_ψ . Thus, from the concavity of $\psi(\cdot)$, it follows that there exists a tension between the preferences of the principal and agent on the interval $[s_\psi, \omega]$.⁵ Given that the agent's choice of strategy is hidden, the only means by which the principal can influence the agent's actions is an incentive contract. The next two sections derive features of the optimal contract given various constraints on the contracting problem.

⁵Of course, the preferences of the principal and agent are aligned outside of this interval.

4 Contracting

In section 3 it was established that in the first period, the principal wants the agent to take on more risk than the agent would prefer, given her career concerns. One solution to this problem is to offer the agent an outcome contingent contract in the first period. An incentive contract influences the agent's behaviour by rewarding outcomes that are indicative of the agent choosing the desired action. The outcomes that are indicative of risk-taking are those outcomes for which the likelihood ratio takes a high value.

The link between the likelihood ratio and the optimal contract is a common feature of moral hazard problems. Indeed, in many moral hazard problems the optimal contract is some strictly increasing transformation of the likelihood ratio (see for example [Mirrlees, 1999](#); [Holmstrom, 1979](#); [Rogerson, 1985](#); [Jewitt, 1988](#)). The likelihood ratio described in (3) and illustrated in figure 3, is convex quadratic in the project's outcome. It follows that any wage constructed as a strictly increasing transformation of the likelihood ratio, results in a wage that is strictly quasi-convex in the outcome, and in which extreme outcomes, both high and low, are rewarded relative to intermediate outcomes. This means paying the agent a higher wage for a project that is a failure, than for a project that performs more or less as expected. There are a number of practical reasons why the principal may be unwilling to write a contract that rewards failure in this way.

1. *Perverse strategy choice:* A wage that is decreasing in firm profits over some range may encourage a perverse choice of strategy. If the contract places a sufficiently high reward on poor outcomes, and the set of available strategies S is sufficiently broad, the agent has an incentive to implement a strategy that produces low outcomes with high probability. That is, there is an incentive for the agent to select a very high or low value of s . This issue arises because the standard approach to finding the optimal contract identifies an interior (local) maximum. However, if the agent's first-period choice problem is not globally concave, the global maximum may lie on the boundary

of the choice set.

2. *Sabotage*: In each of the principal-agent relationships that motivate this model, it is likely that the agent will learn of the project's outcome before it is revealed to the principal and the market as a whole. Moreover, the agent may be able to take hidden actions, such as generating additional costs or engage in other 'money burning' activities, that reduce the outcome before it is revealed. Once again, the agent has an incentive to damage the project if the rewards for doing so are sufficiently large.
3. *Cultural prohibition*: Many societies hold to the principal that success should be rewarded. Both the punishment of success and rewarding of failure are regarded as perverse outcomes. Social attitudes are particularly relevant in the context of publicly traded firms. Shareholders are unlikely to condone executive compensation that rewards low profits and losses, relative to higher profit outcomes.

The three issues outlined above are only of concern if the agent's contract specifies a wage that decreases in the outcome over some range. A straightforward solution to these problems is, therefore, to impose a constraint on the contracting problem requiring that the agent's wage is non-decreasing in the outcome.

4.1 The principal's problem

The structure of the contracting problem is familiar. The principal selects an outcome contingent wage $w(x)$, and strategy $s^* \in S$, in order to maximise the net return from the project,

$$\max_{w(x), s^*} \int_{\mathbb{R}} (x - w(x)) \hat{f}(x|s^*) dx.$$

The principal's choices of $w(x)$ and s^* are subject to three constraints. First, the *participation constraint* states that the expected lifetime utility the agent receives by taking the contract must be at least as great as her outside option $\bar{U}$,

$$\int_{\mathbb{R}} u(w(x)) \hat{f}(x|s^*) dx + \psi(s^*) \geq \bar{U}. \quad (\text{PC})$$

Note that the agent's lifetime utility comprises the utility the agent receives from her first-period wage, and her discounted expected second-period utility, captured by the career concerns term $\psi(s^*)$.

Second, the *incentive compatibility constraint* states that the strategy s^* must be optimal for the agent given the structure of the wage and the agent's career concerns. Because the principal wants to implement a strategy that lies between the agent's preferred strategy s_ψ , and the expected outcome maximising strategy ω , s^* will lie in the interior of the agent's choice set S . It follows that s^* must satisfy the first-order condition,

$$\int_{\mathbb{R}} u(w(x)) \hat{f}_s(x|s^*) dx + \psi'(s^*) = 0. \quad (\text{IC})$$

Finally, the wage must be selected from the family of non-decreasing functions on $\mathbb{R}$. That is,

$$w'(x) \geq 0. \quad (\text{MW})$$

This constraint is henceforth referred to as the *monotone-wage constraint*.

4.2 The optimal contract

Let λ and μ be the Lagrange multipliers for (PC) and (IC) respectively. Consider the contracting problem with (MW) removed. The solution is for $w(\cdot)$ to be defined implicitly by the familiar equation,

$$\frac{1}{u'(w(x))} = \lambda + \mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)}.$$

Given that the lefthand-side is increasing in $w(\cdot)$, it follows that $w(\cdot)$ is itself increasing (resp. decreasing) in x wherever the likelihood-ratio is increasing (resp. decreasing) in x .

Now consider the problem with (MW) included. Let $x_{\min}$ represent the outcome that minimises the likelihood ratio; the turning-point of the quadratic equation described in (3). (This outcome is illustrated in figure 3.) Given that the likelihood ratio is decreasing in x for all $x \in (-\infty, x_{\min})$, it follows that (MW) binds on this interval, and is slack elsewhere.

Therefore, the optimal contract can be characterised by the equation,

$$\frac{1}{u'(w(x))} = \begin{cases} \lambda + \mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)} & x \geq x_{\min} \\ \lambda + \mu \frac{\hat{f}_s(x_{\min}|s^*)}{\hat{f}(x_{\min}|s^*)} & x < x_{\min} \end{cases}. \quad (7)$$

Notice that the effect of (MW) is to censor the optimal contract.⁶ The resulting contract is option-like; delivering a fixed payment for outcomes below $x_{\min}$, and strictly increasing in x past this point.

Proposition 2. *The optimal non-decreasing wage contract is option-like with a ‘strike-price’ at,*

$$x_{\min} = E[x|s^*] - \frac{1}{2} ((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}. \quad (8)$$

The optimal strategy s^* is strictly less than ω .

Proof. It is first necessary to establish that λ and μ are strictly positive. For $\lambda > 0$: Taking expectations of both sides of (7) yields,

$$\int_{\mathbb{R}} \frac{1}{u'(w(x))} \hat{f}(x|s^*) dx = \lambda + \int_{-\infty}^{x_{\min}} \mu \frac{\hat{f}_s(x_{\min}|s^*)}{\hat{f}(x_{\min}|s^*)} \hat{f}(x|s^*) dx + \int_{x_{\min}}^{\infty} \mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)} \hat{f}(x|s^*) dx.$$

Using the fact that $\int_{\mathbb{R}} \mu \hat{f}_s dx = 0$, and solving for λ ,

$$\lambda = \int_{\mathbb{R}} \frac{1}{u'(w(x))} \hat{f}(x|s^*) dx + \int_{-\infty}^{x_{\min}} \left(\mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)} - \mu \frac{\hat{f}_s(x_{\min}|s^*)}{\hat{f}(x_{\min}|s^*)} \right) \hat{f}(x|s^*) dx > 0.$$

The inequality follows because $x_{\min}$ is the minimum point of $\mu \hat{f}_s / \hat{f}$.

For $\mu > 0$: Multiplying (IC) by its Lagrange multiplier yields,

$$\int_{\mathbb{R}} u(w(x)) \mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)} \hat{f}(x|s^*) dx = -\mu \psi'(s^*).$$

Given that the expectation of the likelihood ratio is zero, this equation can be restated as,

$$\text{Cov} \left[u(w(x)), \mu \frac{\hat{f}_s(x|s^*)}{\hat{f}(x|s^*)} \middle| s^* \right] = -\mu \psi'(s^*). \quad (9)$$

⁶In this respect, (MW) functions in much the same way as the upper- and lower-bounds in Jewitt, Kadan and Swinkels (2008).

From (7) it follows that $u(w(x))$ and $\mu f_s/f$ covary in the same direction on $[x_{\min}, \infty)$, and are unrelated on $(-\infty, x_{\min})$ as $u(w(x))$ is constant in this range. Therefore, the lefthand-side of (9) is strictly positive. Given that $s^* > s_\psi$, it follows from proposition 1 that $\psi'(s^*) < 0$, and therefore $\mu > 0$.

The value of $x_{\min}$ can be established by taking the partial derivative of (3) with respect to x ,

$$\frac{\partial}{\partial x} \left(\frac{\hat{f}_s(x|s)}{\hat{f}(x|s)} \right) = \frac{2s(x - s\omega + \frac{1}{2}s^2 - \bar{\theta})}{(s^2 + \sigma^2)^2} - \frac{(s - \omega)}{s^2 + \sigma^2}.$$

Setting this equation equal to zero, and solving for x yields (8).

For $s^* < \omega$, first note that for each marginal increase in s^* , there must be a marginal increasing in the risk carried by the agent. Given that (PC) binds, and that risk is costly for the agent to bear, each marginal increase in s^* , marginally increases the expected cost to the principal of the agent's contract. It then follows from the symmetry of $E[x|s^*]$ around ω , that any expected outcome that can be implemented by a strategy greater than ω , can be implemented by a strategy less than ω at lower cost to the principal. The case $s^* = \omega$ can be excluded as, in this neighbourhood, a marginal decrease in s^* will marginally decrease the cost of the agent's contract without affecting the expected outcome. $\square$

Imposing the monotone wage constraint on the contracting problem prevents the principal from exploiting all the information present in the project's outcome. The optimal non-decreasing wage contract's incentive structure is tied only to those outcomes at which the likelihood ratio is non-decreasing. The information carried by all remaining outcomes must be discarded.

The outcome that serves as the 'strike-price' for the contract is $x_{\min}$. Intuitively, $x_{\min}$ is the outcome that is most indicative of the agent implementing a risk-averse strategy. The likelihood ratio is monotone increasing on the interval $[x_{\min}, \infty)$ and, therefore, on this interval higher outcomes are rewarded with higher wages.

The location of $x_{\min}$ is of interest. Rearranging (8) reveals that $x_{\min}$ is strictly less than the expected outcome, and that the amount by which the strike-price is discounted relative

to expectations is,

$$E[x|s^*] - x_{\min} = \frac{1}{2} ((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}.$$

This discount is increasing in both the variance of the marginal distribution, and the proportion by which the expected profit maximising strategy ω exceeds the strategy implemented by the agent s^* . The second term can be interpreted as a measure of the severity of the moral hazard problem.

4.3 In-the-money executive stock options

Suppose that the owners of a firm (the principal) want to construct an executive compensation scheme to address the moral hazard in strategic decision-making by the firm's CEO (the agent). Suppose further that x represents the share price of the firm at the end of the first period. Stock options allow the owners to link executive compensation to the share price, while discarding the information associated with realisations of x below the strike-price.

Proposition 2 demonstrates that the optimal strike-price lies below the ex-ante expectation of the firm's vesting date share price. This implies that, on the date of issue, it is optimal for managers to be 'in-the-money,' with the strike-price below the prevailing share price. In-the-money stock options allow the owners of the firm to utilise the information present in share prices that lie between $x_{\min}$ and $E[x|s^*]$. Conversely, this information, and the incentives associated with its exploitation, are lost if the options are priced 'at-the-money.'

It is also possible to say something about how the magnitude of the optimal discount varies across firms and industries. In industries in which there is relatively little share price volatility, or in which the actions of management are relatively transparent (limiting the moral hazard problem), proposition 2 states that the optimal strike-price is relatively close to the share price on the date of issue. Conversely, if a firm's share price is highly volatile, or if the actions of management within the firm are opaque, the strike-price should be more generous.⁷ In the presence of very high risk, proposition 2 implies that instead of options,

⁷Morgan (2002) calculates a measure of firm 'opacity' for a range of industries.

managers should be vested with stock, effectively options with a strike-price of zero.

Finally, proposition 2 highlights the costs of taxation policies that penalise in-the-money options relative to at-the-money (or out-of-the-money) options. Suppose that the tax treatment of in-the-money stock options is so unfavourable that no manager is willing to accept a contract with a strike-price below $E[x|s^*]$. For any given number of stock options, raising the strike-price from $x_{\min}$ to $E[x|s^*]$ reduces the incentive power of the compensation scheme. With weaker incentives, the manager will implement a more risk-averse strategy than would be the case under the optimal contract.

To counter this effect, the firm can increase the incentive power of the compensation scheme by issuing a greater number of options. It is straightforward to see that this solution is not optimal. If it were, $w(x) = w(E[x|s^*])$ for all $x \in (-\infty, E[x|s^*])$ would be a solution to the contracting problem, and it is not. Intuitively, the optimal contract exploits the information present in share prices in the interval $[x_{\min}, E[x|s^*]]$. These outcomes are relatively likely as they are adjacent to the mean of the marginal distribution. It follows that they receive significant weight in a manager's optimisation problem. In contrast, an incentive scheme with a strike-price of $E[x|s^*]$ relies more heavily on outcomes in the upper tail of the marginal distribution. Because these outcomes are less likely they receive lower weight in a manager's optimisation problem.

5 Tenure

Another principal-agent relationship in which strategic decision-making may give rise to a moral hazard, is the relationship between university and academic. A researcher must decide how ambitious to make her research agenda. Significant contributions are unlikely without significant risk-taking. However, early in her career, an academic may prefer to 'play-it-safe', pursuing a marginal contribution that has a high probability of success. In academia, prizes such as tenure and promotion are used to create incentives over an agent's choice of research strategy.

5.1 A tenure-track agent

Suppose that in the first-period, the agent holds a tenure-track position. The principal in this problem is the institution employing the agent, and the outcome is the quality of the agent's first-period research portfolio. The agent's first-period wage w is fixed and negotiated prior to her commencing employment at the institution. The agent's second-period wage is determined by her value in the second-period labour market, given the first-period posterior beliefs regarding her talent, as described in section 3.

In order to give the agent an incentive to engage in risky research, with a higher expected return, the principal awards the agent with a prize (tenure) if the quality of her research portfolio exceeds a threshold level x_{crit} . The prize is valuable to the agent, increasing her utility by an amount $\Delta u > 0$. It follows that the agent's first-period optimisation problem is to,

$$\max_{s \in \mathcal{S}} \left[u(w) + \left(1 - \hat{F}(x_{\text{crit}}|s) \right) \Delta u + \psi(s) \right], \quad (10)$$

where $\hat{F}(x_{\text{crit}}|s)$ is the cumulative distribution function (CDF) for the marginal distribution, evaluated at x_{crit} . The mechanism by which the principal influences the agent's behaviour is the choice of the threshold value x_{crit} . In words, x_{crit} is the institution's 'tenure standard.'

The significance of x_{crit} for the agent's optimisation problem can be seen in the first-order condition implied by (10),

$$\hat{F}_s(x_{\text{crit}}|s^*) = \frac{\psi'(s^*)}{\Delta u}. \quad (11)$$

From proposition 1 it follows that the term on the right-hand-side of (11) is negative and decreasing for all $s^* > s_\psi$. Therefore, in order to motivate risk-taking, the principal must select x_{crit} such that the term on the left-hand-side of (11) is likewise negative. The relationship between x_{crit} and the partial derivative $\hat{F}_s(x_{\text{crit}}|s^*)$ is described in the following lemma.

Lemma 1. *Let,*

$$x_0 = \mathbb{E}[x|s^*] - ((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}. \quad (12)$$

For all $s^* \in S$, the partial derivative $\hat{F}_s(x|s^*)$ has the following properties:

- (a) On the domain $(-\infty, x_0)$, $\hat{F}_s(x|s^*)$ is strictly positive, strictly quasi-concave, with a unique maximum at,

$$x_- = x_{min} - \sqrt{\left(\frac{1}{2}((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}\right)^2 + ((s^*)^2 + \sigma^2)}, \quad (13)$$

where x_{min} is defined in (8).

- (b) On the domain (x_0, ∞) , $\hat{F}_s(x|s^*)$ is strictly negative, strictly quasi-convex, with a unique minimum at,

$$x_+ = x_{min} + \sqrt{\left(\frac{1}{2}((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}\right)^2 + ((s^*)^2 + \sigma^2)}. \quad (14)$$

- (c) $\hat{F}_s(x_0|s^*) = 0$.

Proof. First note that $\hat{F}(-\infty|s^*) = 0$ and $\hat{F}(\infty|s^*) = 1$ for all $s^* \in S$, implies $\hat{F}_s(-\infty|s^*) = \hat{F}_s(\infty|s^*) = 0$. The partial derivative $\hat{F}_s(x|s^*)$ can be written,

$$\hat{F}_s(x|s^*) = \int_{-\infty}^x \hat{f}_s(\chi|s^*) d\chi = \int_{-\infty}^x \left(\frac{\hat{f}_s(\chi|s^*)}{\hat{f}(\chi|s^*)} \right) \hat{f}(\chi|s^*) d\chi.$$

Given that $\hat{f}(x|s^*) > 0$, $\hat{F}_s(x|s^*)$ is strictly increasing (resp. decreasing) where the likelihood ratio is strictly positive (resp. negative). Recall from (3), that the likelihood ratio is a convex quadratic in x . It follows that $\hat{F}_s(x|s^*)$ is strictly positive at low values of x , and strictly negative for high values of x , with turning points corresponding to the zeros of the likelihood ratio. From (3) it follows that the upper and lower zeros of the likelihood ratio are given by (14) and (13) respectively.

It only remains to locate the central zero of $\hat{F}_s(x|s^*)$. Given that the marginal distribution is normal, the CDF can be written,

$$\hat{F}(x|s^*) = \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{x - s^*\omega + \frac{1}{2}(s^*)^2 - \bar{\theta}}{\sqrt{2((s^*)^2 + \sigma^2)}} \right) \right],$$

where $\text{erf}(\cdot)$ is the error function. Taking the partial derivative of $\hat{F}(x|s^*)$ with respect to s yields,

$$\hat{F}_s(x|s^*) = \left((s^* - \omega) - \frac{s^* (x - s^* \omega + \frac{1}{2}(s^*)^2 - \bar{\theta})}{(s^*)^2 + \sigma^2} \right) \hat{f}(x|s^*).$$

Setting this equation equal to zero and solving for x yields (12). $\square$

5.2 The optimal tenure standard

The institution's goal is to maximise the quality of the research produced by the agents it employs. At any given time the institution employs numerous tenure-track faculty, allowing the institution to diversify its risk across a range of research projects. For this reason it is assumed that the institution is risk-neutral with regards to any given agent's research agenda. Finally, it is assumed that the institution has a weak preference to grant tenure. Thus, if there are multiple tenure standards that implement the institution's preferred strategy, the institution will always opt for the lowest standard.

At first glance it may seem odd to assume that the institution has only a weak preference over whether or not it grants the agent tenure. After all, tenure lines are scarce and granting tenure to one agent typically means that the institution forgoes the opportunity to seek a more promising replacement in the labour market. However, as demonstrated in the following proposition, the optimal tenure standard requires an agent to exceed expectations. This means that an agent is only offered tenure if the institution's beliefs regarding the agent's talent improve over the course of their tenure-track employment.

Proposition 3. *The optimal tenure standard x_{crit} is strictly greater than the expected outcome, and lies within one standard deviation of the expected outcome.*

Proof. There are two cases to consider:

Case 1: There exists $x_{\text{crit}} \in (x_0, x_+]$ such that $s^* = \omega$ solves (11). In this case the institution can implement the strategy that maximises the expected outcome. From lemma 1 it follows that $x_{\text{crit}} \in (x_0, x_+]$ where x_0 and x_+ are defined in (12) and (14) respectively. Substituting

$s^* = \omega$ into (12) and (14) yields $x_{\text{crit}} \in (E[x|\omega], E[x|\omega] + \sqrt{\omega^2 + \sigma^2})$. Here, the term $\sqrt{\omega^2 + \sigma^2}$ is the standard deviation of the marginal distribution.

Case 2: The highest strategy the institution can implement is $s^ \in (s_\psi, \omega)$.* From (11) and the concavity of $\psi(\cdot)$ it follows that in order to maximise s^* , the institution must select x_{crit} to minimise $F_s(x_{\text{crit}}|s^*)$. From lemma 1 it follows that $F_s(x_{\text{crit}}|s^*)$ achieves its minimum at $x_{\text{crit}} = x_+$. Substituting (8) into (14) yields,

$$x_{\text{crit}} = E[x|s^*] - \frac{1}{2} ((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*} + \sqrt{\left(\frac{1}{2} ((s^*)^2 + \sigma^2) \frac{\omega - s^*}{s^*}\right)^2 + ((s^*)^2 + \sigma^2)},$$

from which it follows $x_{\text{crit}} \in (E[x|s^*], E[x|s^*] + \sqrt{(s^*)^2 + \sigma^2})$. Here, the term $\sqrt{(s^*)^2 + \sigma^2}$ is the standard deviation of the marginal distribution. $\square$

The intuition behind proposition 3 is straightforward. The tenure-standard that maximises the incentive power of the contract rewards only those outcomes in the upper tail of the distribution that become more likely with a marginal increase in s . That is, those outcomes that lie above x_+ , the upper zero of the likelihood ratio (see figure 3). In the event that such a contract would cause the agent to overshoot the principal's desired strategy ω , the incentive power can be reduced by lowering the tenure-standard. Lowering the standard extends the reward to a range of outcomes in the centre of the distribution, that become less likely with a marginal increase in s .

Proposition 3 establishes two bounds on the optimal tenure standard. First, the optimal standard is set in excess of the conditional expectation of the marginal distribution. This means that the agent does not expect to be granted tenure, even if she implements the institution's preferred strategy. Put another way, an ambitious research agenda is not, in and of itself, sufficient to get across the line. Second, the optimal tenure-standard exceeds the conditional expectation by less than one standard deviation. Thus while difficult to achieve, tenure is 'within reach.'

An immediate corollary of proposition 3 is that a tenure-standard that is easy to achieve, has low incentive power, and may even encourage a perverse choice of action. To see this,

consider what happens as x_{crit} decreases from the optimum described in proposition 3. From lemma 1 it follows that the left-hand-side of the first-order condition in (11) increases. In turn, from the strict concavity of $\psi(\cdot)$ (proposition 1) it follows that the agent's optimal action decreases.

So long as $\hat{F}_s(x_{\text{crit}}|s^*)$ remains negative, the tenure-standard encourages risk taking to some degree. Once x_{crit} falls so far that $\hat{F}_s(x_{\text{crit}}|s^*)$ becomes positive, the incentive created by the standard *discourages* risk taking, and the agent's optimal strategy falls below s_ψ . That is, the strategy implemented by the agent is worse for the institution than would be the case if tenure were either certain or impossible. From lemma 1 it follows that the outcome at which this transition occurs is,

$$x_{\text{crit}} = s_\psi\omega - \frac{1}{2}s_\psi^2 + \bar{\theta} - (s_\psi^2 + \sigma^2) \frac{\omega - s_\psi}{s_\psi},$$

which is x_0 from lemma 1 evaluated at $s^* = s_\psi$. The worst possible tenure-standard, from the institution's perspective, is $x_{\text{crit}} = x_-$. Again from lemma 1, this is the standard at which $\hat{F}_s(x_{\text{crit}}|s^*)$ is maximised.

6 Discussion

This paper develops a theory of moral hazard in which the agent takes the role of strategic decision-maker. While the agent has no innate preferences over the set of available strategies, her career concerns give rise to preferences over risk, which in turn create an incentive for her to manipulate the project's risk-return tradeoff to the detriment of the principal.

Two methods for ameliorating the moral hazard problem are investigated. The optimal non-decreasing wage contract involves granting the agent 'in-the-money' options. The discount at which these options are offered is increasing in both the risk associated with the project and the magnitude of the moral hazard problem. If, instead, incentives are created for the agent through the award of academic tenure, the optimal tenure standard requires the agent to exceed expectations, and lies within one standard deviation of the expected outcome.

6.1 Career Structure

One strong assumption utilised in this paper is that the projects into which the agent may be hired are identical. The effect of this assumption within the model is to ensure that the agent's second-period wage is linear in the observed outcome. Of course this need not be the case.

Suppose that projects vary in the marginal product of talent. This can be captured by modifying the production technology such that,

$$x = s\omega - \frac{s^2}{2} + \alpha\theta + s\varepsilon,$$

an expression that differs from (2) only in the addition of the α term. Here $\alpha > 0$ is the marginal return to talent for the project in question. Now suppose that the labour market sorts agents into projects such that those agents with the highest expected talent are assigned to projects in which talent is most productive.

A marginal increase in the observed outcome now increases the agent's second-period wage in two ways: First, as above, it increases the expectation of the agent's productivity on any given project. Second, it results in the agent being assigned to a project on which her talent has a higher marginal product (a project with marginally higher α). These effects compound one and other, resulting in a second-period wage $v(x)$ that is convex increasing in x .⁸ If $v(x)$ is sufficiently convex to overcome the concavity of the agent's utility function, then the composition $u(v(x))$ will likewise be increasing convex.

It is worthwhile briefly discussing what happens in the model if $u(v(x))$ is a convex transformation of x . This convexity means that the agent's career concerns, her expected second-period utility, now benefit from an increase in the project's risk. From the proof of proposition 1 it then follows that the agent's preferred strategy s_ψ will be greater than ω as the agent will be willing to tradeoff the expected return in order to increase the probability of an outcome in the upper tail of the distribution.

⁸A similar convexity would be observed if the production technology gave rise to superstars as described by Rosen (1981).

In this case the purpose of an incentive contract would be to discourage risk-taking. It remains the case that extreme outcomes are indicative of risk-taking while intermediate outcomes are indicative of a risk-averse strategy choice. Thus, from proposition 2 it follows that the optimal non-decreasing wage contract involves ‘capped incentives’. The agent’s wage should monotonically increase in x until it reaches the turning point of the likelihood ratio, and thereafter $w'(x) = 0$.⁹

References

- Chevalier, J. and G. Ellison (1999), Career Concerns of Mutual Fund Managers, *Quarterly Journal of Economics* 114, 389–432.
- de Meza, D. and D. Webb (2007), Incentive Design Under Loss Aversion, *Journal of the European Economic Association* 5, 66–92.
- Dewatripont, M., Jewitt, I. and Tirole, J. (1999a), The Economics of Career Concerns, Part I: Comparing Information Structures, *Review of Economic Studies* 66, 183–198.
- Dewatripont, M., Jewitt, I. and Tirole, J. (1999b), The Economics of Career Concerns, Part II: Application to Missions and Accountability of Government Agencies, *Review of Economic Studies* 66, 199–217.
- Harris, M. and Holmstrom, B. (1982), A Theory of Wage Dynamics, *Review of Economic Studies* 49, 315–333.
- Holmstrom, B. (1979), Moral Hazard and Observability, *Bell Journal of Economics* 10, 74–91.
- Holmstrom, B. (1999), Managerial Incentive Problems: A Dynamic Perspective, *Review of Economic Studies* 66, 169–182.

⁹To see this, note that $\omega < s^* < s_\psi$ and (9) together imply that $\mu < 0$ and, therefore, the wage should be increasing in x wherever the likelihood ratio is decreasing.

- Holmstrom, B. and Ricart i Costa, J. (1986), Managerial Incentives and Capital Management, *Quarterly Journal of Economics* 101, 835–860.
- Jewitt, I. (1988), Justifying the First-Order Approach to Principal-Agent Problems. *Econometrica* 56, 1177–1190.
- Jewitt, I., O. Kadan and J. Swinkels (2008), Moral Hazard with Bounded Payments, *Journal of Economic Theory* 143, 59–82.
- Manso, G. (2011), Motivating Innovation, *The Journal of Finance* 66, 1823–1860.
- Mirrlees, J. (1976), The Optimal Structure of Incentives and Authority within an Organization, *Bell Journal of Economics* 7, 105–131.
- Mirrlees, J. (1999), The Theory of Moral Hazard and Unobservable Behavior: Part I, *Review of Economic Studies* 66, 3–21.
- Milgrom, P. (1981), Good News and Bad News: Representation Theorems and Applications, *Bell Journal of Economics* 12, 380–391.
- Morgan, D. (2002), Rating Banks: Risk and Uncertainty in an Opaque Industry, *American Economic Review* 92, 874–88.
- Rogerson, W. (1985), The First-Order Approach to Principal-Agent Problems. *Econometrica* 53, 1357–1367.
- Rosen, S. (1981), The Economics of Superstars, *American Economic Review* 71, 845–58.