

Moral Hazard with Persistent Actions and Learning

Nolan H. Miller*

This draft: April 9, 1999

Abstract

This paper considers a two-period principal-agent model in which the agent's ex ante effort choice affects the distribution of outcomes throughout the game and the parties learn over time about the agent's risk classification. In this environment, the optimal long-term contract may involve overinsurance in the last period, giving the agent more utility following a loss than following no loss, even when the initial distribution satisfies the Monotone Likelihood Ratio Condition. In addition, it is shown that in the presence of persistent actions and learning, the optimal bilateral-commitment contract involves ex post inefficient payments in the second period that cannot be supported by short-term contracts. Because of this and the fact that bilateral commitment to long-term contracts allows the optimal contract to better allocate the agent's incentives over time, the bilateral-commitment contract Pareto dominates contracts where the parties are unable to commit.

*Managerial Economics and Decision Sciences, J.L. Kellogg Graduate School of Management. I thank Daniel Spulber for advice throughout the project as well as Nabil Al-Najjar, Sandeep Baliga, Matt Jackson, Elizabeth Lacey, Amit Pazgal, Bill Sandholm, and Sri Sridharan for helpful comments. All errors are my own.

1. Introduction

Models of repeated moral hazard typically focus on environments that are time-separable. Agents' actions affect only the current period, and each period's outcome distribution is history-independent. However, many real-world principal-agent relationships are not separable. For example, there are a number of natural ways in which the relationship between an insurance company and a new driver may fail to be time-separable. The first has to do with learning. Some people are simply better drivers than others. However, it may be difficult or impossible to identify *ex ante* how risky a particular driver is. As the driver amasses an experience record, both the insurer and the driver learn about the driver's risk classification, and this learning affects their beliefs about how likely particular outcomes are in the future. Poor performance in the past leads the parties to believe the agent is high-risk and consequently signals poor performance in the future. Thus the expected distribution of outcomes is history dependent, violating separability.

Another reason why the relationship between a driver and his insurer fails to be separable is that the probability that the new driver experiences a loss depends on the amount of effort that he puts forth during driving school. All else being equal, the more attention the driver pays in driving school, the lower the risk of loss *throughout the driver's lifetime*. Since the effects of effort are persistent, outcomes in future periods depend on the effort choice. Consequently these future outcomes provide information about which effort choice was made, again violating separability.

The non-separabilities in the principal-agent relationship induced by persistent actions and learning may alter the form of the optimal contract in a number of ways. This paper focuses on two. The first part of the paper investigates the way in which the optimal bilateral-commitment contract structures the agent's incentives within each period. The second part of the paper considers the manner

in which commitment to long-term contracts in moral hazard with persistent actions and learning affects the provision of incentives over time.

In this paper, learning and persistent actions are introduced into a two-period relationship between a risk-neutral, monopolistic principal and a risk-averse agent¹. There are two types of agents, good agents and bad agents, two levels of effort, high effort and low effort, and two outcomes, which, keeping with an insurance interpretation of the model, will be referred to as “loss” and “no loss”. Neither the principal nor the agent knows the agent’s type, but both share a common belief about the prior probability that the agent is the bad type and update this belief over time.

The first part of the paper investigates the provision of incentives within each period. The main result shows that in the presence of persistent actions and learning, the optimal long-term principal-agent contract may involve *overinsurance* in the second period. That is, the optimal contract may result in the agent receiving more utility following a loss than following no loss. In an insurance context, this is equivalent to the agent’s policy reimbursing her more in the event of a loss than the cash value of the loss. Thus overinsurance represents a failure of monotonicity in the agent’s second-period payoff function.

Surprisingly, overinsurance can arise in the second period even when the first-period distribution satisfies the Monotone Likelihood Ratio Condition, which is sufficient in the static principal-agent problem for the optimal incentive scheme to be monotonically increasing (see Grossman and Hart 1984, proposition 5 and section 4). This is due to the fact that the principal and agent learn over time about the probability of a loss. The effect of this learning is that the second-period distribution may not satisfy MLRC, even if the first period distribution does. Failure to satisfy MLRC in the second period means that increasing effort increases the likelihood of a loss occurring in the second period. And, since the

¹For clarity, throughout the paper the principal will be referred to as “he” and the agent as “she”.

principal provides the agent with incentives by rewarding her following actions whose likelihood increases when effort increases, this implies that the optimal contract will involve overinsurance.

The reason why the second-period distribution may fail to satisfy MLRC lies in the interaction between the agent's effort choice and the information about the agent's type contained in the first period outcome. Because effort is persistent and the players learn about the agent's type, the agent's effort choice gives rise to two separate effects. First, there is the *direct effect*, the idea that, all else being equal (i.e. for a given belief that the agent is the bad type), the agent is less likely to experience a loss in the second period if she chooses high effort rather than low effort.

The second effect is the *learning effect*. The learning effect arises from the fact that the parties use the first-period outcome to make inferences about whether the agent is the good or bad type, and a particular first-period outcome provides different information about the agent's type depending on whether effort is high or low. For example, experiencing a loss in the first period is always evidence in favor of the agent being the bad type. However, the strength of the evidence depends on whether the agent chose high or low effort. It may be that a loss results in a higher posterior probability that the agent is the bad type when effort is high than when it is low, or vice versa.

The direct effect and the learning effect combine to determine the expected probability of a loss occurring in the second period. To understand how, consider the case where a loss occurs in the first period². If a loss in the first period is stronger evidence that the agent is the bad type under low effort than under high effort, then the learning and direct effects tend to reinforce each other. For any posterior belief that the agent is the bad type, the direct effect implies that choosing high effort decreases the probability of a loss in the second period. And,

²Similar reasoning applies in the case where there is no loss in the first period.

since a first-period loss when effort is high yields a lower posterior probability that the agent is the bad type than a loss when effort is low, and bad types experience losses more often than good types, increasing effort also tends to decrease the expected probability of a loss in the second period via the learning effect.

If, on the other hand, a loss is stronger evidence that the agent is the bad type under high effort than under low effort, then the learning effect and direct effect work in opposite directions. The direct effect still tends to decrease the probability of a loss in the second period. However, since a loss in the first period when effort is high is stronger evidence that the agent is the bad type than a loss when effort is low, the learning effect tends to increase the probability of a loss in the second period. This is because agents who experience a loss when effort is high are believed more likely to be the bad type than agents who experience a loss when effort is low.

In cases where the effort and learning effects oppose each other, it is possible that increasing effort actually increases the history-conditional expected probability of a loss in the second period. This occurs when the first period outcome is *much* stronger evidence that the agent is the bad type when effort is high than when effort is low. In this case, the learning effect overwhelms the direct effect and the optimal contract features overinsurance.

Persistent actions, learning, and commitment to long-term contracts each play a role in the possibility of overinsurance. If actions do not have persistent effects, then the first period outcome is a sufficient statistic for effort, and consequently the optimal contract will not depend on the second period's outcome. Learning is needed because it is the learning effect that leads overinsurance, rather than partial insurance, to have positive incentive effects.

The role of commitment to long-term contracts by the principal and agent in moral hazard with persistent actions and learning is the subject of the second part of the paper. Commitment by the principal is shown to be necessary for the overinsurance result because overinsurance is *ex post* inefficient. In the ab-

sence of commitment, renegotiation will result in replacing a contract featuring overinsurance with one that fully insures the agent³.

Commitment to long-term contracts also affects the manner in which non-separabilities in the environment impact the provision of incentives over time. In any dynamic principal-agent problem, there are two ways in which the agent can be provided with incentives. First, the agent can be exposed to financial risk if a loss occurs in a period, such as through a deductible. Call this type of risk *outcome risk*. Second, there is the risk that the agent will acquire a bad record, and, as a result, decrease the utility she expects to receive in the future. Call this type of risk *classification risk*⁴.

Intuitively, outcome risk and classification risk each play a role in the real-life provision of incentives. Once again, consider an insured driver. The driver is likely to state two financial reasons why she attempts to prevent damage to her car⁵. First, if the driver's policy features a deductible, then any theft from or damage to her car will result in an out-of-pocket expense. Second, if she experiences a loss today she will likely face higher insurance rates in the future. The first cost represents outcome risk, and the second represents classification risk.

As in the overinsurance result, learning, persistent actions, and commitment are all important in determining whether incentives are given to the agent through outcome risk or classification risk. The main result of this section is to show that when both the principal and agent can commit to long-term contracts, the risk the agent must bear is divided equally between outcome risk and classification risk. However, if one or both of the parties cannot commit, then the additional constraints imposed on the problem prevent the principal from optimally allo-

³While commitment by the agent does play a role in the optimal contract, it is not necessary for the overinsurance result.

⁴The term "classification risk" is due to Palfrey and Spatt (1985).

⁵We ignore the fear of injury, since the incentives it generates are unlikely to be affected by the terms of the auto insurance contract.

cating the agent’s incentives. In this case, relative to the bilateral-commitment regime, the agent will bear too much classification risk when the cost of effort is low and too little classification risk when the cost of effort is high. As a result of this distortion and the fact that commitment by the principal allows the principal to implement contracts that make better use of the information provided by the second-period output, the bilateral-commitment contract Pareto dominates the other forms of contracting.

1.1. Related Literature

A number of papers have studied the impact of learning in models without moral hazard. Harris and Holmstrom (1982) investigate the role of learning in labor markets, where both employer and employee learn about the worker’s productivity over time, while Palfrey and Spatt (1985) consider the role of learning in repeated insurance contracts where the consumer makes a *contractible* effort choice at the beginning of each period, and the insurer and consumer share a common belief about the probability that the consumer is the high-risk type and update that belief after observing the first period outcome.

Palfrey and Spatt show that the principal insures the agent against classification risk through a series of subsidies, whereby old consumers subsidize young consumers and consumers whose record has revealed them to be low risk subsidize the other consumers. Harris and Holmstrom, on the other hand, show that if the agent is able to appeal to a competitive market for her services and the principal is able to commit to long-term contracts, then the optimal contract features wages that never fall and rise only when the market “bids up” the agent’s wage. Thus, as in Palfrey and Spatt, the agent is insured against being revealed to be the bad type.

Hirao (1993) considers a model of moral hazard with learning in which a firm employs a manager to implement a new project where both parties learn about

the quality of the project over time. It is shown that in the optimal long-term contract the level of effort the agent chooses is positively related to the information value of the first period output as a signal of the effort level⁶. If increasing effort makes it easier to distinguish between good and bad projects, the optimal contract will tend to call for a higher level of effort; if increasing effort makes it harder to distinguish between good and bad projects, the optimal contract will tend to call for a lower level of effort. Thus the optimal effort level is chosen based both on productivity and informational concerns⁷.

In the aforementioned papers, the ability to bind one or both parties to long-term contracts generally results in Pareto superior contracts. However, while long-term contracts permit the best trade-off between the provision of incentives and risk sharing, they suffer from significant drawbacks. First, there is the fact that fully-specified contingent contracts can be difficult and costly to write. Second, the law typically prohibits agents from entering into contracts that they cannot break, and it is rare that courts will order a worker to perform a job against her will.

The problems inherent in long-term contracting have spawned a number of papers that seek to determine when the benefits of bilateral commitment to long-term contracts can be replicated by spot contracting. Fudenberg, Holmstrom, and Milgrom (1990) show conditions under which the ability to commit to long-term contracts in repeated agency relationships provides no benefits to the parties. Malcomson and Spinnewyn (1988) and Chiappori et al. (1994) prove results similar to Fudenberg, Holmstrom, and Milgrom, although using different assumptions. Rey and Salanié (1990) consider a $T \geq 3$ period model and state conditions under which in separable environments long-term contracts can be replicated by a series

⁶Palfrey and Spatt come to a similar conclusion in the case where effort is contractible.

⁷Hirao's "information value" is related to what is being called the learning effect in this model, although due to differences in emphasis and technique he does not focus on the same questions as this paper.

of short-term contracts, each of which lasts more than one period but less than T periods.

All of the above papers require time-separability as a condition for spot contracting to replicate the benefits of long-term contracting. Because of persistent effort and learning, however, the environment in this paper is not separable, and, in fact, long-term contracts will offer benefits over spot contracting. These benefits take two forms. First, since the outcome of the first period will not be a sufficient statistic for the full outcome vector with respect to the effort level, as Holmstrom (1979) shows, the optimal contract will be based on the first and second period outcomes. If the parties are able to commit to long-term contracts, then they can agree on a contract that depends on the second period outcome. However, contracts that depend on the second-period outcome cannot be implemented when the parties cannot commit. Hence the bilateral-commitment contract allows the parties to make better use of the information provided by the outcome vector. Second, in the presence of learning, the agent expects greater losses in the second period following a loss than following no loss. The optimal long-term contract insures the agent against this classification risk. Spot contracts cannot.

The paper proceeds as follows. Section 2 describes the basic model. Section 3 presents an example. In section 4 the overinsurance result is derived, and the roles of learning, persistent actions, and commitment in the provision of incentives within periods are examined. Section 5 considers the role of commitment in the moral hazard with persistent actions and learning and considers its impact on the provision of incentives over time. Section 6 shows that in the presence of persistent actions and learning the bilateral-commitment contract Pareto dominates contracts where one or both parties are unable to commit. Section 7 discusses applications and extensions. All proofs are in section 8.

2. The Model

2.1. Basic Structure

Consider a two period principal-agent model with moral hazard in which both the principal and agent learn over time about the probability of a loss occurring. In each period the wealth that the agent produces takes one of two values, w or $w - x$, where $w > w - x > 0$. When wealth equals $w - x$, it will be said that a “loss” has occurred. Throughout the paper, the subscript 1 will be used to refer to the initial period of the model, the subscript L to refer to the second period of the model following a loss in the first period, and the subscript N to refer to the second period of the model following no loss in the first period. These will frequently be called the three “states” or “histories” of the game.

The agent is risk averse. Her utility for wealth is given by the utility function $u(\cdot)$, a strictly increasing, strictly concave, differentiable function. The agent’s net utility for the two periods of the game is given by $u(w_1) + u(w_2) - \hat{c}$, where $\hat{c}$ is the level of effort the consumer chooses and w_r is the consumer’s final wealth in period $r \in \{1, 2\}$. The agent does not discount. All conclusions generalize in a straightforward manner to cases where the agent has a positive discount factor.

There are two types of agents, good agents (type G) and bad agents (type B). The agent’s type does not change between periods.

At the start of the game, the agent chooses effort $\hat{c} \in \{0, c\}$, where c represents high effort and 0 represents low effort. The agent’s effort choice remains constant for the entire game and is not observed by the principal.

If an agent of type t chooses effort level $\hat{c}$, the probability of a loss during period $r \in \{1, 2\}$ is given by $q_{tr}(\hat{c})$, where $0 < q_{tr}(c) \leq q_{tr}(0) < 1$ for $t \in \{G, B\}$, and $q_{Br}(\hat{c}) \geq q_{Gr}(\hat{c})$ for $\hat{c} \in \{0, c\}$. Thus for a given type, effort decreases the probability of a loss, and for a given effort level, good agents experience losses less often than bad agents.

Unless otherwise stated, throughout the paper it is assumed that all of the

inequalities in the previous paragraph are strict and that $q_{t1}(\hat{c}) = q_{t2}(\hat{c}) \equiv q_t(\hat{c})$. The added notational complexity is needed in order to be able to consider cases where actions are not persistent or there is no learning. If actions are not persistent, $q_{t2}(0) = q_{t2}(c)$ for both types. If there is no learning, $q_{Br}(\hat{c}) = q_{Gr}(\hat{c})$ for $r \in \{1, 2\}$ and $\hat{c} \in \{0, c\}$. When environments with either no learning or no persistent actions are being considered, these assumptions will be explicitly stated.

Neither the agent nor the principal knows the agent's type ex ante, but both share the common prior that the consumer is type B with probability p_1 . Thus the expected probability of a loss in period 1 is given by

$$q_1(\hat{c}) \equiv p_1 q_B(\hat{c}) + (1 - p_1) q_G(\hat{c}). \quad (2.1)$$

After observing the first period outcome, the principal and agent update their belief that the agent is the bad type according to Bayes' rule. If a loss occurred in the first period, the posterior probability that the consumer is the bad type if she chose effort $\hat{c}$ is given by

$$p_L(\hat{c}) \equiv \frac{p_1 q_B(\hat{c})}{q_1(\hat{c})}, \quad (2.2)$$

and the posterior probability of a loss occurring in period 2 is given by

$$q_L(\hat{c}) \equiv p_L(\hat{c}) q_B(\hat{c}) + (1 - p_L(\hat{c})) q_G(\hat{c}). \quad (2.3)$$

If no loss occurs in period 1, the posterior probability that the consumer is the bad type is given by

$$p_N(\hat{c}) \equiv \frac{p_1 (1 - q_B(\hat{c}))}{1 - q_1(\hat{c})}, \quad (2.4)$$

and the posterior probability of a loss occurring in period 2 is given by

$$q_N(\hat{c}) \equiv p_N(\hat{c}) q_B(\hat{c}) + (1 - p_N(\hat{c})) q_G(\hat{c}). \quad (2.5)$$

Throughout the paper, the following assumption is maintained. It ensures that under autarky the agent wishes to choose high effort for all relevant loss probabilities.

Assumption 1: $q_L(c)u(w-x) + (1-q_L(c))u(w) - \frac{c}{2} > q_L(0)u(w-x) + (1-q_L(0))u(w)$.

This paper follows Rogerson (1985) by assuming that the agent has no access to credit markets. This is done largely for expositional simplicity in order to focus on the provision of incentives over time in the presence of learning and action persistence. The role of credit in models without learning has received some attention, and the insights of these models will likely transfer to the present case⁸. The impact of giving the agent access to credit markets will be discussed in section 7.

The principal is a risk-neutral monopolist⁹. He maximizes expected profits over the course of the two periods. The principal may borrow or save at an interest rate of zero. The model easily adapts to the case where the principal has a positive interest rate.

2.2. The Contracting Problem

A contract between the principal and agent consists of a vector $(w_L, w_N, w_{LL}, w_{LN}, w_{NL}, w_{NN}) \in \mathbb{R}^6$ of wages w_H , $H \in \{L, N\}$ in the first period and w_{HJ} in the second period, where $H \in \{L, N\}$ is the first period outcome and $J \in \{L, N\}$ is the second period outcome. In the context of an insurance

⁸Specifically, Chiappori et al. show that the environment where the agent has no access to credit markets but the principal can save and borrow is equivalent to one where the agent has access to credit markets but this access can be monitored by the principal. Hence the conclusions in this paper will apply in “monitorable access” environments.

⁹The conclusions do not change if it is assumed that the principal is one firm in a perfectly competitive industry. However, the analysis is substantially more complicated since issues of asymmetric information at the renegotiation stage arise.

model, if π is the premium the agent pays to the principal and ρ is the reimbursement the agent pays to the principal in the event of a loss, then $w_N = w - \pi$ and $w_L = w - \pi - x + \rho$. Hence the formulation above can be used to represent either a labor model, as in Harris and Holmstrom (1982), or an insurance model, as in Palfrey and Spatt (1985). Denote a generic contract by Δ .

Following Grossman and Hart (1985), it is frequently useful to talk about the contracting problem using the utility offered to the consumer following a particular outcome as the control variables. Let $v_H \equiv u(w_H)$ be the agent's utility if the principal pays her w_H following outcome $H \in \{L, N\}$ in the first period. Similarly, define $v_{HJ} \equiv u(w_{HJ})$. Thus there is a one-to-one correspondence between contracts $(w_L, w_N, w_{LL}, w_{LN}, w_{NL}, w_{NN})$ in terms of wage payments and $(v_L, v_N, v_{LL}, v_{LN}, v_{NL}, v_{NN})$ in terms of utility. For the remainder of the paper, the latter representation of a contract will be employed.

Let $h(\cdot) \equiv u^{-1}(\cdot)$. Thus $h(v)$ is the cost to the principal of giving the agent utility v following a particular outcome. Since $u(\cdot)$ is strictly increasing and strictly concave, $h(\cdot)$ is strictly increasing and strictly convex, and $h'(\cdot) = \frac{1}{u'(\cdot)}$.

In order to make the analysis interesting, throughout the paper it is assumed that the principal wishes to implement high effort. However, if the rents generated by the contracting relationship are small enough, the principal will instead choose to implement low effort. The assumptions on the primitives needed to ensure that the principal wishes to implement high effort change depending on whether the principal and/or agent can commit to long-term contracts. Such considerations complicate the analysis but do not add significantly to the results. For this reason Assumption 2 does not refer to the primitives of the problem¹⁰.

Assumption 2: *In all of the contracting environments in this paper, assume*

¹⁰Discussion of optimal contracts that implement low effort is available from the author upon request. Ma (1991) investigates the question of when a principal may not want to implement high effort in a similar model and shows that in certain circumstances the principal may want to implement a mixed action.

that the principal wishes to implement high effort by the agent.

The principal maximizes expected profits. In terms of the utility v offered to the consumer following each of the outcomes of the game, the objective function is written as

$$\begin{aligned} & 2(w - q_1(c)x) \tag{OF} \\ & - q_1(c)(h(v_L) + q_L(c)h(v_{LL}) + (1 - q_L(c))h(v_{LN})) \\ & - (1 - q_1(c))(h(v_N) + q_N(c)h(v_{NL}) + (1 - q_N(c))h(v_{NN})) \end{aligned}$$

where $2(w - q_1(c)x)$ is the expected wealth generated by the consumer over the course of the relationship. The remainder of (OF) is the cost to the principal of giving the consumer utility vector $(v_N, v_L, v_{NN}, v_{NL}, v_{LN}, v_{LL})$.

When the principal wishes to implement high effort, the incentive compatibility constraint requires that the agent's expected utility under high effort exceed his expected utility under low effort by at least the cost of effort. For a fixed contract, Δ , define the consumer's expected utility over the course of the contract, $U(\hat{c}, \Delta)$, as follows

$$\begin{aligned} U(\hat{c}, \Delta) \equiv & q_1(\hat{c})(v_L + q_L(\hat{c})v_{LL} + (1 - q_L(\hat{c}))v_{LN}) \\ & + (1 - q_1(\hat{c}))(v_N + q_N(\hat{c})v_{NL} + (1 - q_N(\hat{c}))v_{NN}). \end{aligned}$$

The incentive compatibility constraint is therefore given by

$$U(c, \Delta) - U(0, \Delta) \geq c. \tag{IC}$$

Since the monopolistic principal is the only possible insurance provider, the alternative to the agent entering into a contract with the principal is autarky. Under autarky, the consumer expects utility

$$\begin{aligned} & q_1(c)(u(w - x) + q_L(c)u(w - x) + (1 - q_L(c))u(w)) \\ & + (1 - q_1(c))(u(w) + q_N(c)u(w - x) + (1 - q_N(c))u(w)) - c = 2U_1 - c \end{aligned}$$

over the two periods of the model. Here, $U_1 \equiv q_1(c) u(w) + (1 - q_1(c)) u(w - x)$ is the agent's unconditional expected utility before the cost of effort in each of the first and second periods. However, conditional on the outcome in the first period, the agent expects to earn

$$U_L \equiv q_L(c) u(w - x) + (1 - q_L(c)) u(w)$$

if there was a loss in the first period and

$$U_N \equiv q_N(c) u(w - x) + (1 - q_N(c)) u(w)$$

if there was no loss in the first period¹¹.

If the agent commits to long-term contracts, then she must expect to receive her reservation utility over the course of the entire relationship, although she may receive less than her reservation utility in one or more particular states. Hence the participation constraint is given by

$$U(c, \Delta) \geq 2U_1. \tag{P}$$

With this terminology in place, when both the principal and agent can commit to long-term contracts, the contracting problem consists of maximizing (OF) subject to the constraint that the agent prefer high effort to low effort, (IC) , and prefer contracting with the principal to autarky, (P) .

¹¹In defining the history-conditional reservation utility in this manner, we are implicitly assuming that the model refers to an insurance relationship rather than a labor relationship, since in a labor relationship the agent may not be able to produce at all under autarky if she has no access to productive assets. However, this model can be interpreted as a labor model if the agent is seen as a contractor who may either work for a firm or work on her own. Alternatively, in a labor context one can simply think of U_L and U_N as the agent's reservation utility if she leaves the principal's employment after the first period without linking them to the primitives of the problem. As long as $U_L < U_N$, the general conclusions of the paper will continue to hold. In particular, Propositions 1 and 2 do not depend on the definition of U_L and U_N .

3. Example

This example illustrates the interaction of learning and persistent actions in two-period insurance contracts where the principal and agent can commit to contracts that bind for the entire relationship.

Consider the example of the new driver discussed in the introduction. The probability that this driver experiences a loss depends on two factors, whether she is a good driver or a bad driver, and whether she puts forth low effort or high effort. Assume that the driver lives for two periods and that she experiences one of two outcomes each period, loss or no loss.

Suppose that $q_B(0) = 0.9$, $q_B(c) = 0.7$, $q_G(0) = 0.2$, and $q_G(c) = 0.1$. Thus for a given type, effort reduces the probability of a loss, and for a given effort level, good agents experience losses less often than bad agents.

Suppose the common prior that the agent is the bad type is 0.8. In this case, the expected loss probabilities in the first period given low effort and high effort as computed using (2.1) are $q_1(0) = 0.76$ and $q_1(c) = 0.58$.

The outcome of the first period provides information to the parties about the driver's type. The Bayesian posterior probabilities, $p_H(\hat{c})$, that the agent is the bad type given the first period outcome $H \in \{L, N\}$ and the level of effort $\hat{c} \in \{0, c\}$ as computed using (2.2) and (2.4) are (approximately) $p_L(0) = 0.947$, $p_L(c) = 0.966$, $p_N(0) = 0.333$, $p_N(c) = 0.571$. According to formulas (2.3) and (2.5), this yields history-conditional expected loss probabilities $q_L(0) = 0.863$, $q_L(c) = 0.679$, $q_N(0) = 0.433$ and $q_N(c) = 0.443$.

Interestingly, increasing effort decreases the expected probability of a second period loss following a first period loss from 0.863 to 0.679. However, increasing effort increases the expected probability of a second period loss following no loss in the first period from 0.433 to 0.443. The fact that increasing effort makes the L state better but the N state worse will be the driving force behind the form of the optimal contract.

Suppose the agent's utility function is $u(w) = \sqrt{2w}$, the cost of effort is $c = 10$, and $U_1 = 100$. If the principal wishes to implement high effort and both parties can commit to long-term contracts, the optimal principal-agent contract maximizes the insurer's profits (OF) subject to the constraint that the driver expect to earn at least her reservation utility (P), and that she prefer choosing high effort to low effort, (IC).

Given the above assumptions, the optimal contract can be explicitly computed by evaluating the Lagrangian of the constrained problem. The payments are given by:

$$\begin{aligned} w_N^* &= 7245.7 & w_L^* &= 3633.1 \\ w_{LN}^* &= 7316.1 & w_{LL}^* &= 2337.8 \\ w_{NN}^* &= 7189.9 & w_{NL}^* &= 7316.1 \end{aligned}$$

This contract exhibits a number of interesting properties. To begin, the optimal contract features memory, since payments in the second period depend on the first period outcome. In addition, since w_N^* is greater than w_L^* , the optimal contract exhibits partial insurance in the first period; the agent's net financial position is worse following a loss than following no loss. This conforms with the intuition that optimal insurance contracts in the presence of moral hazard typically involve partial insurance. In addition, since the first period distribution satisfies MLRC, one would expect the contract to be monotonic in the first-period output.

In the second period, the L state offers partial insurance as is expected, but since $w_{NN}^* < w_{NL}^*$, the N state features *overinsurance*. That is, if the agent has no loss in the first period, her net financial position is greater if she experiences a loss in the second period than if she experiences no loss. This reversal in the payments arises from the fact that $q_N(c) > q_N(0)$, a phenomenon that occurs when experiencing no loss in the first period when effort is high is a much weaker signal that the agent is the good type than experiencing no loss in the first period

when effort is low. In this case, the learning effect overwhelms the direct effect and leads $q_N(c)$ to be greater than $q_N(0)$. And, since the optimal contract rewards the agent for outcomes that become more likely under high effort, this leads to overinsurance in the N state.

4. The Overinsurance Result

In this section, the issues raised by the example are developed in a more general setting, and the roles of persistent actions, learning, and commitment in shaping the optimal contract are discussed.

Consider the case where the principal and agent enter into agreements that bind both parties for the entire two-period game. Before period 1 the principal offers the agent a contract Δ , which she may either accept or reject. If the agent rejects Δ , then she proceeds under autarky. If the agent accepts the contract, it binds both parties for the entire game. In this case, the agent chooses an effort level $\hat{c} \in \{0, c\}$ and the first-period output is realized. The principal pays the agent according to Δ , and then the second-period outcome is realized, following which the principal once again pays the agent according to Δ . Since both parties are bound by their original contract, there is no opportunity for renegotiation.

Although the exact form of the contract cannot be determined without making further assumptions about the agent's utility function and the technology, several interesting properties emerge by considering the optimality conditions of the constrained optimization problem. Under bilateral commitment contracting, the principal maximizes (OF) subject to (IC) and (P) as defined above. Evaluating the Lagrangian of this problem yields the following proposition:

Proposition 1: *Consider the bilateral long-term contracting problem. The solution¹² to this problem exhibits the following properties (optimized values of the variables are denoted by asterisks):*

¹²Throughout the paper, assume an interior solution.

Property 1.1: *At the optimal solution, (P) and (IC) bind and their respective shadow prices are strictly positive.*

Property 1.2: *$v_L^* \neq v_N^*$ and $v_{LJ}^* \neq v_{NJ}^*$ for $J \in \{L, N\}$. Hence, the optimal bilateral long-term commitment contract features memory.*

Property 1.3: *$v_N^* > v_L^*$. That is, partial insurance is offered in the first period.*

Property 1.4: For $H \in \{L, N\}$,

If $q_H(0) > q_H(c)$, then $v_{HN}^ > v_{HL}^*$*

If $q_H(0) = q_H(c)$, then $v_{HN}^ = v_{HL}^*$.*

If $q_H(0) < q_H(c)$, then $v_{HN}^ < v_{HL}^*$.*

Property 1.5: *Either $v_{LL}^* < v_{LN}^*$ or $v_{NL}^* < v_{NN}^*$, i.e. at least one of the second period states features partial insurance.*

All proofs are in Appendix A.

Property 1.1 establishes that the constraints bind in this problem and that the shadow prices of the constraints are strictly positive. This arises from the fact that the interests of the principal and agent are conflicting and any slack in the constraints could be used by the principal to improve expected profits.

Rogerson (1985) and Lambert (1983) consider the role of memory in repeated moral hazard models where the technology is separable over time and show that memory plays a role in any Pareto optimal long-term contract. A contract exhibits memory whenever payments in the second period depend on the payments made in the first period. In the context of the present model, a contract exhibits memory if $w_L \neq w_N$ implies that $w_{LJ} \neq w_{NJ}$ for some $J \in \{L, N\}$. Since intertemporal linkage in the technology due to persistent actions or learning only provides additional reasons why the optimal long-term contract should exhibit memory, their results hold *a fortiori* in the present model¹³. Property 1.2 con-

¹³In fact, Rogerson's Propositions 1 and 2, which prove the memory result, are valid in the present environment as well.

firms that the optimal bilateral commitment contract features memory. In fact, it exhibits a strong form of memory, where $w_L \neq w_N$ and $w_{LJ} \neq w_{NJ}$ for both $J \in \{L, N\}$.

One of the most general properties of moral hazard problems is that in order to induce the agent to choose high effort the contract must reward the agent following outcomes that are more likely under high effort than under low and punish the agent following outcomes that are more likely under low effort than high. Since $q_1(c) < q_1(0)$, increasing effort increases the likelihood of no loss in the first period, and thus one would expect the optimal contract to punish the occurrence of a loss in the first period and reward the occurrence of no loss. In an insurance context, this is accomplished via a partial insurance contract, such as one featuring a deductible or copayment, where the occurrence of a loss reduces the agent's net utility. Property 1.3 states that the optimal contract does, in fact, offer partial insurance in the first period.

The technical condition that guarantees that the first-period incentive scheme offers partial insurance (i.e. is monotonically increasing in the output) is the Monotone Likelihood Ratio Condition. In words, if a distribution satisfies MLRC, then high outcomes are relatively more likely to arise from high effort than low outcomes are. Mathematically, MLRC arises as a sufficient condition for monotonicity in the incentive scheme because the payment the agent receives following a particular outcome y can be shown to be decreasing in the “likelihood ratio,”

$$\frac{\text{Prob}(y \mid \text{low effort})}{\text{Prob}(y \mid \text{high effort})},$$

the ratio of the probability of outcome y occurring under low effort to the probability of that outcome occurring under high effort (see Grossman and Hart (1985) for example). As this ratio decreases, outcome y becomes stronger and stronger evidence that high effort was undertaken, and thus rewarding the agent following y becomes more effective at inducing the agent to exert high effort.

Since $\frac{q_1(0)}{q_1(c)} > 1 > \frac{(1-q_1(0))}{(1-q_1(c))}$, the distribution of outcomes in the first period

satisfies MLRC. When there are only two actions, MLRC is sufficient for the optimal incentive scheme to be monotonically increasing in the output in the static principal agent problem (see Grossman and Hart, Proposition 5 and section 4), which yields Property 1.3.

Property 1.4 shows that even though the first-period distribution satisfying MLRC is sufficient for the first-period incentive scheme to be increasing, it does not guarantee that the compensation scheme will be increasing in the second-period as well. Partial insurance arises in second period state $H \in \{L, N\}$ if $q_H(0) > q_H(c)$. This is in accordance with the standard intuition in moral hazard problems. However, if $q_H(c) > q_H(0)$, then the optimal contract is one that offers the agent more utility in the event of a loss in the second period than in the event of no loss. That is, the optimal contract features *overinsurance*. This “contradiction” of the Grossman and Hart result occurs because in the presence of persistent actions and learning, the history-conditional distribution of outcomes in the second period may not satisfy MLRC even if the first period distribution does.

The analysis of Property 1.4 is broken into two parts. First, it is shown how the sign of $q_H(0) - q_H(c)$ determines whether partial insurance or overinsurance is optimal. Then, the conditions under which increasing effort increases the expected probability of a loss after a particular first period outcome are identified.

The reason why the sign of $q_H(0) - q_H(c)$ determines whether the optimal contract features partial insurance or overinsurance begins with the intuition that inducing the agent to choose high effort requires giving her more utility following outcomes that become more likely when effort increases. When $q_H(0) > q_H(c)$, then increasing effort tends to decrease the probability of a loss in the second period following outcome $H \in \{L, N\}$ in the first period. In this case, the optimal incentive scheme rewards the occurrence of no loss relative to the occurrence of a loss in the second period.

However, when $q_H(c) > q_H(0)$, increasing effort tends to increase the prob-

ability of a loss. In fact, since $q_H(c) > q_H(0)$, it follows that $\frac{q_H(0)}{q_H(c)} < \frac{1-q_H(0)}{1-q_H(c)}$, violating MLRC. Thus a loss is relatively more likely to occur under high effort than no loss. Since the principal wants to reward the agent for outcomes that become more likely under high effort, and the probability of a loss in the second period following outcome H in the first period increases when effort is increased, it follows that the principal will want to reward the agent following a loss in the second period. And, in an insurance context, rewarding the agent for a loss relative to no loss amounts to reimbursing her in excess of the cash value of the loss, i.e. overinsurance.

For another approach to the intuition, recall that the payment the agent receives under the optimal contract is decreasing in the ratio

$$\frac{\text{Prob}(y \mid \text{low effort})}{\text{Prob}(y \mid \text{high effort})}. \quad (4.1)$$

If a loss occurs in the first period¹⁴, then (4.1) is equal to

$$\frac{q_1(0) q_L(0)}{q_1(c) q_L(c)} \quad (4.2)$$

for the outcome LL and

$$\frac{q_1(0) (1 - q_L(0))}{q_1(c) (1 - q_L(c))} \quad (4.3)$$

for the outcome LN . Note that the ratio $\frac{q_1(0)}{q_1(c)}$ is common to both (4.2) and (4.3). Thus (4.2) is larger than (4.3) whenever

$$\frac{q_L(0)}{q_L(c)} > \frac{(1 - q_L(0))}{(1 - q_L(c))}$$

or

$$q_L(0) > q_L(c).$$

Thus even though the size of the payments in the LL and LN states are determined by the entire likelihood ratios (4.2) and (4.3), the fact that $\frac{q_1(0)}{q_1(c)}$ is common

¹⁴A similar argument holds if no loss occurs in the first period.

to both implies that the order of these two ratios depends only on the sign of $q_L(0) - q_L(c)$. It is as if we condition out the effect of effort in the first period, $\frac{q_1(0)}{q_1(c)}$, and look only at the effect of effort in the second period in determining the relative size of v_{LL} and v_{LN} .

The possibility that the history-conditional loss probability is greater under high effort than under low effort is due to learning. In the presence of learning, the effect of increasing effort on the probability of a loss in the second period consists of two parts, the *direct effect* and the *learning effect*. The direct effect captures the idea that putting forth high effort reduces the probability of a loss, regardless of whether the agent is the good or bad type. The learning effect, on the other hand, has to do with the fact the parties update their belief about whether the agent is the good or bad type based on the first period's outcome and whether the agent chose high effort or low.

To isolate the two effects, suppose the first period outcome is $H \in \{L, N\}$ and rewrite $q_H(c) - q_H(0)$ as

$$\{p_H(c) q_B(c) + (1 - p_H(c)) q_G(c)\} - \{p_H(0) q_B(0) + (1 - p_H(0)) q_G(0)\}.$$

Adding and subtracting $p_H(c) q_B(0) + (1 - p_H(c)) q_G(0)$ and rearranging the terms yields

$$\begin{aligned} & p_H(c) (q_B(c) - q_B(0)) + (1 - p_H(c)) (q_G(c) - q_G(0)) \\ & + (p_H(c) - p_H(0)) (q_B(0) - q_G(0)). \end{aligned} \tag{4.4}$$

The first line of (4.4) represents the direct effect. Since $q_t(c) - q_t(0)$ is always negative for $t \in \{G, B\}$, the direct effect is always negative; for a given belief that the consumer is the bad type, increasing effort decreases the probability of a loss in the second period.

The second line of (4.4) represents the learning effect. It arises from the fact that a loss conveys different information about the agent's type depending on whether she chose high or low effort. Since $q_B(0) - q_G(0) > 0$, the sign of the

learning effect depends on the sign of $(p_H(c) - p_H(0))$, the difference between the posterior belief that the agent is the bad type if outcome H occurs and effort is high and if effort is low.

The conditions under which $p_H(c) - p_H(0)$ is positive or negative depend on whether $H = N$ or $H = L$. The difference $p_N(c) - p_N(0)$ has the same sign as

$$\frac{(1 - q_B(c))}{(1 - q_G(c))} - \frac{(1 - q_B(0))}{(1 - q_G(0))}. \quad (4.5)$$

Since the occurrence of no loss in the first period is always evidence in favor of the agent being the good type, if (4.5) is positive, then no loss is stronger evidence that the agent is the good type when effort is low than when effort is high. Conversely, if (4.5) is negative, then no loss is stronger evidence that the agent is the good type when effort is high than when effort is low.

If $H = L$, then $p_L(c) - p_L(0)$ has the same sign as

$$\frac{q_B(c)}{q_G(c)} - \frac{q_B(0)}{q_G(0)}. \quad (4.6)$$

In this case, since the occurrence of a loss in the first period is always evidence in favor of the agent being the bad type, if (4.6) is positive then a loss is stronger evidence that the agent is the bad type when effort is high than when effort is low. On the other hand, if (4.6) is negative, then a loss is stronger evidence that the agent is the bad type when effort is low than when effort is high.

Thus regardless of whether $H = L$ or $H = N$, if $p_H(0) > p_H(c)$ then outcome H is a weaker signal that the agent is the bad type (or equivalently, a stronger signal that the agent is the good type) when effort is low than when effort is high. All else being equal, the lower the posterior belief that the agent is the bad type, the less the expected loss probability. Thus by choosing high effort instead of low effort the agent tends to decrease the expected probability of a loss. In this case, the direct effect and learning effect work in the same direction. Both effects tend to decrease the expected probability of a loss in the second period following outcome H in the first period.

If, on the other hand, $p_H(c) > p_H(0)$, then increasing effort tends to increase the posterior belief that the agent is the bad type following outcome H in the first period. And, since the agent is more likely to be the bad type following outcome H when effort is high, the impact of the learning effect is that choosing high effort instead of low effort tends to increase the expected probability of a loss in the second period. Thus the direct effect and learning effect work in opposite directions. All else being equal, effort tends to decrease expected probability of a loss, but the mere occurrence of outcome H when effort is high is strong evidence that the agent is the bad type, which tends to increase the expected probability of a loss in the second period.

If the learning effect is small relative to the direct effect, then the overall effect of effort will be to decrease the expected loss probability in the second period. However, if $p_H(c) > p_H(0)$ and the learning effect is large relative to the direct effect, then the learning effect may overwhelm the direct effect, resulting in the expected loss probability in the second period being *larger* under high effort than under low effort. Since in this case, as was discussed earlier, high effort makes a loss following outcome H more likely and incentives are provided to the agent by rewarding outcomes that become more likely under high effort, the optimal contract will overinsure the agent.

As indicated in the previous paragraph, a positive learning effect alone is not sufficient for $q_H(c) - q_H(0) > 0$ and thus for overinsurance. For overinsurance to arise, the learning effect must be sufficiently positive to outweigh the direct effect. Thus while (4.5) and (4.6) being negative is a sufficient condition for $v_{HL}^* < v_{HN}^*$, it is by no means necessary. It may be that the learning effect is positive, but not so positive that it outweighs the direct effect.

4.1. Overinsurance and Stochastic Structure

While Property 1.4 shows that overinsurance can arise in the optimal contract, it does not necessarily arise in the optimal contract. The situations under which overinsurance will arise in the N and L states are slightly different. We will consider each state in turn.

The conditions under which the optimal contract may feature overinsurance in the second period following no loss in the first period include:

- $q_B(0) \gg q_G(0)$
- $q_B(0) > q_B(c)$
- $q_G(0) \sim q_G(c)$

That is, the optimal contract may¹⁵ feature overinsurance in the second period following no loss in the first period when bad agents are much more likely to experience losses than good agents, but effort on the part of bad agents is more effective at reducing the probability of a loss than effort on the part of good agents. Thus one can interpret a contract featuring overinsurance as the principal telling the agent “I want you to choose high effort. I realize that by doing so you decrease your ability to distinguish yourself as a good agent by producing no loss, and as a result you expect to do worse in the N state. I will compensate you for this by overinsuring you in the N state.”

The conditions under which overinsurance may¹⁶ arise in the optimal contract in the L state differ from the conditions for the N state in that they involve effort by the good type of agent having a larger loss-reducing effect than effort by the bad type of agent. Thus the conditions are given by:

¹⁵When these conditions hold, the optimal contract will feature overinsurance for an open set of beliefs about the prior probability that the agent is the bad type, but not necessarily all beliefs.

¹⁶See the previous footnote.

- $q_B(0) \gg q_G(0)$
- $q_B(0) \sim q_B(c)$
- $q_G(0) > q_G(c)$

As the previous discussion indicates and Property 1.5 proves, it is impossible for the optimal long-term contract to feature overinsurance in both the L and N states.

4.2. Robustness of the Learning Effect

As the discussion in the previous section shows, overinsurance arises only in special circumstances. While there is a reasonably large set of specifications of the primitives that can lead to overinsurance, it is by no means the norm. However, the overinsurance result is in some sense a pedagogical device designed to highlight the learning effect, which is a very robust phenomenon.

The learning effect arises from the fact that the inferences the parties make about how likely it is that the agent is the bad type depends on the level of effort she has chosen. Since in general, $p_H(c) \neq p_H(0)$, the learning effect will play a role in every optimal contract. Even when it is not so strong that it overwhelms the direct effect and leads to overinsurance, it will play a role in fine-tuning the way in which the agent is given incentives to choose high effort. When the learning effect is positive, $p_H(c) > p_H(0)$, it works against the direct effect, decreasing the power of the incentives given to the agent; all else being equal, a positive learning effect in state H narrows the gap between v_{HL} and v_{HN} . In extreme cases a positive learning effect even leads to overinsurance. Conversely, when the learning effect is negative, $p_H(c) < p_H(0)$, the learning effect tends to reinforce the direct effect, increasing the power of the incentives offered to the agent; all else equal, a negative learning effect increases the gap between v_{HL} and v_{HN} .

The observations concerning the learning effect made in the previous paragraph point to some interesting empirical implications of the analysis in Proposition 1.

When the learning effect is positive, we should expect the power of incentives to shrink over time. On the other hand, when the learning effect is negative, we should expect the power of incentives to shrink over time.

In addition, the analysis of Proposition 1 shows that changes in the power of incentives can be traced to differences between $p_H(c)$ and $p_H(0)$ arising from the learning effect. However, consider extending the model presented here from a two period model to a many period model. In such a model, eventually the parties would learn whether the agent is good or bad. Thus after a number of periods differences between the posterior belief that the agent is bad given low or high effort must eventually disappear. Thus we would expect the power of incentives to stabilize after a time.

4.3. Simpson's Paradox

The overinsurance result is an example of a statistical reversal phenomenon known as Simpson's Paradox¹⁷. Simpson's Paradox arises when the first of two alternatives performs better than the second on several different measures, but the second alternative performs better on average. For example, consider the following situation. Two hospitals, Good Hospital and Bad Hospital, serve a small town. In the town, there are two types of patients, low-risk patients and high-risk patients, and all patients require the identical operation. Low-risk patients survive with probability 0.4 if they choose Bad Hospital. They survive the operation with probability 0.6 if they choose Good Hospital. High-risk patients, on the other hand, survive with probability 0.1 if they choose Bad Hospital and 0.3 if they choose Good Hospital. Thus Good Hospital is better than Bad Hospital in the sense that any particular patient has a better probability of survival if it goes to Good Hospital than if it goes to Bad Hospital.

Now suppose that 90% of the patients at Bad Hospital are low-risk and 10%

¹⁷I thank Bill Sandholm for introducing me to Simpson's Paradox.

are high-risk. On the other hand, 10% of the patients at Good Hospital are low-risk and 90% are high-risk. In this case, the expected survival probability at Bad Hospital is $0.9 * 0.4 + 0.1 * 0.1 = 0.37$, while the expected survival probability at Good Hospital is $0.1 * 0.4 + 0.9 * 0.3 = 0.31$. Thus even though Good Hospital has a higher probability of survival for any particular patient, the fact that a particular patient has chosen to go to Good Hospital is such a strong signal that he is the high-risk type that his expected survival probability is lower than a randomly chosen patient at Bad Hospital. This “statistical reversal” is Simpson’s Paradox.

In the hospital example, it is the fact that a particular patient who shows up at Good Hospital is so much more likely to be the high-risk type than a patient at Bad Hospital that leads to the statistical reversal. Thus by choosing Good Hospital the patient “signals” that she is the bad type of patient. Similarly, in the model considered in this paper, the statistical reversal occurs when a particular first period outcome becomes a much stronger signal that the agent is the bad type under high effort than low. Thus by choosing high effort, the agent makes that outcome into a much worse signal about her type. Consequently, the expected posterior probability of a loss is larger under high effort than under low effort.

4.4. Institutional Constraints

The possibility of overinsurance raises serious problems relating to fraud, since an agent that is rewarded following a loss in the second period will take every effort to produce one. Hence sabotage of the outcome becomes a real concern. Typically, in order to address the possibility of fraud, the insurance industry will limit reimbursement to be no larger than the size of the loss. Thus there is no possibility for overinsurance. These “institutional constraints” against overinsurance are incorporated into the model considered in this paper by adding the following constraints to the problem:

$$v_N \geq v_L, \quad v_{NN} \geq v_{NL}, \quad \text{and} \quad v_{LN} \geq v_{LL}.$$

Proposition 2 characterizes the optimal contract when both parties can commit to long-term contracts with the addition of the institutional constraints.

Proposition 2: Institutional Constraints. *If for some $H \in \{L, N\}$, $q_H(c) > q_H(0)$, then the optimal solution to the bilateral-commitment contracting problem with the addition of the institutional constraints is such that $v_{HN}^* = v_{HL}^*$ in that state. If $q_H(0) \geq q_H(c)$ for both H , then the imposition of the institutional constraints leaves the optimal solution unchanged.*

Proposition 2 says that if there is overinsurance in some state before the institutional constraint is imposed, then there is full insurance in that state after it is imposed, with other payments adjusting in order to induce the agent to choose high effort. Since when $q_H(c) > q_H(0)$, setting $v_{HL} > v_{HN}$ punishes the agent for choosing high effort and increases the cost to the principal, the principal has nothing to gain by imposing risk on the agent in state H . If, on the other hand, $q_H(0) \geq q_H(c)$, then there is partial insurance in the optimal contract before the imposition of the institutional constraint. Since any contract featuring partial insurance satisfies the institutional constraints, their imposition has no effect on the optimal contract.

Proposition 2 has an interesting empirical interpretation. In the case where $q_L(c) - q_L(0) > 0$, Proposition 2 states that the optimal contract will feature partial insurance in the first period, partial insurance in the second period following no loss, and full insurance in the second period following a first-period loss. Thus Proposition 2 captures the insight that many types of insurance policies exhibit aggregate spending limits. Health insurance policies, for example, often feature annual limits on the policy holder's out-of-pocket expenses; once an insured consumer pays a certain amount in health care expenses during a year, the insurer pays all future expenses.

Such contracts make sense in health insurance because “bad” consumers tend to incur large health care costs throughout their lives. Thus consumers are very concerned with insurance against being the bad type. In terms of the present

model, $q_L(c) - q_L(0) > 0$ arises when the difference between the good and bad types of consumers is large relative to the impact of effort. Thus the model agrees with the intuition that limits on out-of-pocket spending by agents should arise in such a setting. In summary, Proposition 2 can be interpreted as predicting that annual limits on out-of-pocket spending will arise in environments where the difference between types of agents is more important than the difference between low-effort and high-effort agents. When effort is more important than the agent's type, then the insurance contract will exhibit a deductible each period.

The opposite extreme is the case where $q_N(c) - q_N(0) > 0$. This case will arise when bad agents are much more likely to experience losses than good agents, and increasing effort is more effective at decreasing the probability of a loss for bad agents than for good agents. In this case, under the institutional constraints the optimal contract will feature partial insurance in the first period, partial insurance in the L state, and full insurance in the N state. Thus this is a case where successful agents experience less risky wages than unsuccessful agents. This may correspond to situations such as the tenure of professors, where good performance leads to a situation where wages are riskless.

The next several subsections investigate the role of persistent effort, learning, and commitment in producing the overinsurance result. Specifically, we will seek to determine what gives rise to memory, partial insurance in the first period, and the possibility of overinsurance in the second period.

4.5. The Role of Persistent Actions

In order to isolate the role of persistent effort in Proposition 1, consider the case where the effort choice affects only the distribution of outcomes in the first period. This is accomplished by assuming that $q_{t2}(c) = q_{t2}(0)$ for $t \in \{G, B\}$. Hence we allow the probability of a loss in the second period to depend on the history but not on the level of effort. The following proposition characterizes the optimal

long-term contract if effort is not persistent.

Proposition 3: Effort is not Persistent. *Suppose $q_{t2}(c) = q_{t2}(0)$ for $t \in \{G, B\}$, i.e. effort is not persistent. The optimal contract under bilateral long-term commitment is such that*

$$\begin{aligned} v_N^* &= v_{NN}^* = v_{NL}^* = U_1 + \frac{cq_1(c)}{2(q_1(0) - q_1(c))} \\ v_L^* &= v_{LN}^* = v_{LL}^* = U_1 - \frac{c(1 - q_1(c))}{2(q_1(0) - q_1(c))}. \end{aligned}$$

The key difference between the optimal bilateral commitment contract if actions are persistent and if actions are not persistent is that when actions are not persistent, the optimal contract imposes no outcome risk on the agent in the second period. Since the agent's action choice is sunk at the beginning of the second period, the only reason why the principal would want to impose risk on the agent in the second period is if the second-period outcome were informative about the agent's original action choice. However, if actions are not persistent, then the first period outcome is a sufficient statistic for the joint outcome of the two periods with respect to the agent's effort choice. Thus it follows from Holmstrom (1979) that basing the contract on the outcome of both periods offers no advantage over a contract based on only the first-period outcome since the second-period outcome provides no additional information beyond what is contained in the first-period outcome. Since imposing outcome risk on the agent tends to increase the principal's cost and has no beneficial incentive effects, there is no reason for the principal to do so.

In addition to the lack of second-period outcome risk, the contract in Proposition 3 also sets the agent's utility constant along any path through the game. The reason for this is best understood as an application of a result due to Rogerson (1985, Proposition 1). Rogerson proves that if both parties can commit to long-term contracts, then the utility the agent receives in the first and second period

following first-period outcome H must satisfy the following relationship

$$h'(v_H) = q_H(c)h'(v_{HL}) + (1 - q_H(c))h'(v_{HN}). \quad (4.7)$$

The reason for this is that when both parties can commit, the principal is free to shift utility along the path between the first and second periods as long as the contract satisfies the overall participation constraint (P), and providing the agent with slightly more utility in period 1 by increasing v_H by k has the same incentive effects as increasing both v_{HL} and v_{HN} by k . From this it follows that at the optimum the marginal cost of increasing utility in period 1 must be the same as the marginal cost of increasing utility in all states in period 2. This is exactly what (4.7) says. Since the cost to the principal is convex, (4.7) implies that when there is no outcome risk in the second period, the optimal allocation of utility along paths through the game will set $v_H = v_{HL} = v_{HN}$.

4.6. The Role of Learning

In order to isolate the role of learning, assume that $q_{Gr}(\hat{c}) = q_{Br}(\hat{c})$ for $r \in \{1, 2\}$ and $\hat{c} \in \{0, c\}$. This implies that the loss probability in the second period is the same as the loss probability in the first period and depends only on the agent's effort choice. Thus there is no learning. The optimal bilateral commitment contract is derived in the following proposition.

Proposition 4: No Learning. *Assume that $q_{Gr}(\hat{c}) = q_{Br}(\hat{c})$ for $r \in \{1, 2\}$, i.e. there is no learning. In the optimal bilateral commitment contract, $v_N^* > v_L^*$, $v_{NN}^* > v_{NL}^*$, $v_{LN}^* > v_{LL}^*$.*

In the absence of learning, the optimal contract to the agent offers partial insurance in all states. Hence it is the fact that $q_L(\hat{c})$ and $q_N(\hat{c})$ differ from $q_1(\hat{c})$ that leads to the possibility of overinsurance in the optimal contract. Seen another way, recall that, all else being equal, effort decreases the posterior probability of an accident. In the presence of learning, overinsurance arose when the learning

effect was sufficiently positive. However, with no learning there is no learning effect, and thus nothing to overcome the direct effect.

5. The Role of Commitment

This section examines the role of commitment in moral hazard problems with persistent actions and learning. It has three broad goals. First, the additional constraints imposed on the contracting environment when the parties are unable to commit are explained. Second, the optimal contracts are derived when neither party can commit to long-term contracts and when only the principal can commit, and these contracts are compared to the bilateral-commitment contract. Third, the manner in which commitment by the principal and agent affects the distribution of the agent's incentives over time is examined.

If the principal and agent are not bound by long-term commitments, then the parties will have an opportunity to renegotiate following the first period, during which the principal may revise his original offer and the agent decide whether she wishes to continue the contract or proceed under autarky. Since the nature of the renegotiation will impact the form of the optimal contract, it is necessary to carefully specify the timing of the game at the renegotiation stage.

The renegotiation stage takes place after the first-period outcome has been observed and the payment has been made to the agent, but before the second-period outcome is realized. At the start of the renegotiation stage, the principal offers the agent a new continuation contract $(\hat{v}_{HN}, \hat{v}_{HL})$. If the principal is bound by commitments, then he cannot revise the contract. He must offer the original contract terms again at the renegotiation stage. If the principal is not bound by commitments, then he is free to offer any continuation contract he wishes.

The agent may either accept or reject the principal's offer. If the agent accepts $(\hat{v}_{HN}, \hat{v}_{HL})$, then play proceeds to period 2 with $(\hat{v}_{HN}, \hat{v}_{HL})$ as the current contract. The alternative to $(\hat{v}_{HN}, \hat{v}_{HL})$ in the event that the agent rejects it depends

on whether or not she is bound by commitments.

The idea of commitment by the agent is somewhat more subtle than commitment by the principal. Since the principal makes the agent a take-it-or-leave-it offer, commitment by the agent to a specified contract in the second period is meaningless if the principal does not offer her that contract at the renegotiation stage. Because of this, commitment by the agent is taken to consist of a binding promise by the agent to accept any renegotiated contract $(\hat{v}_{HN}, \hat{v}_{HL})$ in the second period provided that this contract offers at least as much expected utility in the H state as did the original contract, Δ .

The effect of commitment by the agent is to determine the status quo in the event that the agent rejects the principal's offer $(\hat{v}_{HN}, \hat{v}_{HL})$ at the renegotiation stage. If the agent is not bound by commitments, then she reverts to autarky following a rejection of $(\hat{v}_{HN}, \hat{v}_{HL})$; if the agent is bound by commitments, then Δ continues to be in effect after a rejection and period 2 begins.

Since the principal always wants to contract with the agent, the effect of the renegotiation opportunity is to limit the contracts that the principal will offer originally. Since (as will be shown) the optimal renegotiation contracts when the principal cannot commit do not depend on the first period contract, the principal will take these into account when making his original contract offer, and, in equilibrium, no renegotiation will take place.

5.1. Commitment and the Contracting Problem

Commitment by the principal and commitment by the agent each play a role in determining the form of the optimal contract. If the agent cannot commit to long-term contracts, then she will elect to leave the relationship with the principal and proceed under autarky if the contract does not offer her at least her history-conditional reservation utility in the second period. Thus if the agent cannot commit, in order for the agent to “credibly” accept the original contract Δ , it must

satisfy the following two history-conditional participation constraints in addition to the overall participation constraint (P) :

$$q_L(c) v_{LL} + (1 - q_L(c)) v_{LN} \geq U_L \quad (P_L)$$

$$q_N(c) v_{NL} + (1 - q_N(c)) v_{NN} \geq U_N. \quad (P_N)$$

When the principal cannot commit to long-term contracts he cannot be held to any ex ante commitment to non-profit maximizing behavior in the second period. Consequently he will offer the contract that maximizes expected profits in state $H \in \{L, N\}$ subject to the constraint that he give the agent the least utility possible. If the agent cannot commit either, then the “least utility possible” is the agent’s history-conditional reservation utility¹⁸, U_H . Consequently, when neither the principal nor the agent can commit and the principal believes that the agent chose high effort, contingent on outcome H in the first period, the contract the principal offers at the renegotiation stage solves

$$\begin{aligned} (v_{HL}, v_{HN}) \in \arg \max w - q_H(c) x - q_H(c) h(v_{HL}) - (1 - q_H(c)) h(v_{HN}) \\ \text{subject to } (P_H). \end{aligned}$$

Because the agent’s action choice is sunk, this is a simple risk-sharing problem. Since the principal maximizes a concave function subject to a linear constraint, standard arguments show that the solution to this problem sets $v_{HL}^* = v_{HN}^* = U_H$.

The previous paragraph illustrates that any time the principal is unable to commit, the agent will be fully insured in the second period. This is because in the absence of the power to commit, the principal will always offer a profit maximizing contract in the second period. Since the agent’s action choice is sunk at the renegotiation stage, the principal has nothing to gain by imposing risk on the agent. Consequently any profit maximizing contract will fully insure the agent.

¹⁸If the agent can commit, then she may promise to accept a contract that offers less than U_H in state H . In this case, the principal will offer the lower utility.

When the principal cannot commit, there is an additional technical issue that must be addressed. At the time that the principal makes his renegotiation offer, the agent knows whether she has chosen high effort or low, but the principal does not. Thus there appears to be an asymmetric information problem. However, this problem is easily addressed. Consider the case where neither the principal nor the agent commits. If the agent has chosen high effort, the second-period contract is given by $v_{HL}^* = v_{HN}^* = U_H$, for $H \in \{L, N\}$. The natural assumption with respect to the principal's beliefs is to assume that if the first period contract (v_L, v_N) proposed by the principal coupled with this second-period contract satisfies (IC) , then the principal believes the agent has chosen high effort at the recontracting date. Otherwise, the principal assumes the agent has chosen low effort. Since proposing such a contract actually induces the agent to choose high effort, this is sufficient to address the asymmetric information problem.

5.2. Spot Contracting and Unilateral Commitment by the Principal

Beside the bilateral-commitment regime, there are two other natural environments in which to investigate the role of commitment. The first is the case where neither the principal nor the agent can commit. This corresponds to principal-agent relationships where the parties sign a new short-term contract each period. The second is the case where only the principal can commit. Contracts that are enforceable against the principal but not the agent frequently arise in employment relationships, where courts will force employers to fulfill their end of the contract but will rarely force agents to go to work against their will.

When neither the principal nor the agent can be bound in the second period by a contract that was negotiated before the first period, the relationship is governed by a series of “spot” contracts which are negotiated at the beginning of each period and last only for that period. In this case, the contracting problem consists of maximizing (OF) subject to (IC) , (P) , and the requirement that $v_{HL}^* = v_{HN}^* =$

U_H for $H \in \{L, N\}$. The solution to this problem is described in Proposition 5.

Proposition 5: Let $A \equiv \frac{c}{q_1(0) - q_1(c)} - (U_N - U_L)$. If $A \leq 0$, the optimal series of spot contracts is given by

$$v_L^* = v_N^* = U_1 \qquad v_{LL}^* = v_{LN}^* = U_L \qquad v_{NL}^* = v_{NN}^* = U_N.$$

If $A > 0$, the optimal series of spot contracts is given by

$$\begin{aligned} v_L^* &= U_1 - (1 - q_1(c)) A & v_N^* &= U_1 + q_1(c) A \\ v_{LL}^* &= v_{LN}^* = U_L & v_{NL}^* &= v_{NN}^* = U_N. \end{aligned}$$

Learning plays an important role in the form of the optimal series of spot contracts since it determines that the agent receives U_N in the N state and U_L in the L state. Since increasing effort increases the likelihood of no loss in the first period, and $U_N > U_L$, the agent expects to do better in the second period if she experiences no loss in the first period. This provides her with an incentive to choose high effort. When the cost of effort is small, i.e. $A \leq 0$, the difference in the agent's expected utility in the N and L states alone is sufficient to induce the agent to choose high effort. Consequently the principal does not need to impose any additional risk on the agent in the first period, and so he sets $v_L^* = v_N^* = U_1$ in order to satisfy (P).

However, when $A > 0$, the risk imposed on the agent due to the fact that $v_{NL}^* = v_{NN}^* = U_N > U_L = v_{LL}^* = v_{LN}^*$ is not sufficient to induce the agent to choose high effort. Consequently, the principal must set $v_N > v_L$ in order to expose the agent to enough risk that she prefers high effort to low. Hence the optimal series of spot contracts partially insures the agent in the first period.

The optimal series of spot contracts differs from the optimal bilateral-commitment contract (Proposition 1) in two ways. First, regardless of c , the spot-contracting regime features full insurance in the second period. Second, for small levels of c , the spot-contracting regime fully insures the agent in the first period. If the principal can commit, these differences can often be eliminated, although due to

the additional constraints imposed by the agent's inability to commit, there will be cases where the optimal contract resembles the optimal series of spot contracts more than the optimal bilateral-commitment contract.

When the principal can commit to long-term contracts but the agent cannot, the contracting problem involves maximizing (OF) subject to (IC) , (P) , (P_N) , and (P_L) . Proposition 6 characterizes the solution to the contracting problem when only the principal can commit.

Proposition 6: *Consider the case where the principal commits to long-term contracts but the agent does not.*

If $c > 2(q_1(0) - q_1(c)) \frac{U_N - U_1}{q_1(c) + 1}$, then the solution exhibits Properties 1.1-1.5 as stated in Proposition 1.

If $c \leq 2(q_1(0) - q_1(c)) \frac{U_N - U_1}{q_1(c) + 1}$, then the optimal contract sets $v_{NL}^ = v_{NN}^* = U_N$, and $v_L^* = v_N^* = v_{NL}^* = v_{NN}^* = \left(\frac{2U_1 - U_N(1 - q_1(c))}{q_1(c) + 1} \right) < U_1$.*

Proposition 6 establishes that because commitment by the principal allows the contract to depend on the second period outcome, the optimal contract when only the principal can commit resembles the bilateral-commitment contract as long as the cost of effort is sufficiently large. In this case, overinsurance is possible and arises under the same circumstances as in the bilateral-commitment contract¹⁹. In the event that overinsurance arises in the optimal contract the effect of imposing the institutional constraints is the same as in bilateral-commitment contracting. That is, Proposition 2 applies in this case as well.

The requirement that $c > 2(q_1(0) - q_1(c)) \frac{U_N - U_1}{q_1(c) + 1}$ is necessary for the contract to resemble the bilateral-commitment contract because for very small values of the cost, (IC) may fail to bind. The reason for this is that due to the constraint

¹⁹Note, however, that it is not necessarily the case that (P_N) or (P_L) bind at the optimum. If either (P_N) or (P_L) binds, then it can be shown that the optimal contract when only the principal can commit differs from the optimal contract in bilateral long-term commitment contracting. But, it may be the case that efficient provision of incentives demands that the agent expect more than her reservation utility in each of the second periods.

(P_N) , the minimum amount of risk imposed on the agent by any feasible contract is bounded away from zero. Since the principal cannot promise the agent less than U_N in the second period, the lowest risk contract that satisfies (P) is given by $v_{NL}^* = v_{NN}^* = U_N$, and $v_L^* = v_N^* = v_{NL}^* = v_{NN}^* = \left(\frac{2U_1 - U_N(1 - q_1(c))}{q_1(c) + 1} \right) < U_1$. This is also the lowest cost contract that satisfies (P) , (P_N) , and (P_L) . For c sufficiently small, this contract will satisfy (IC) , and consequently the optimal contract will exhibit no outcome risk at all. The agent will be fully insured in both the first and second periods.

Propositions 5 and 6 illustrate that failure to commit by either the principal or the agent will have an impact on the allocation of incentives over time. When the agent cannot commit, she cannot credibly commit to a contract that offers less than her history-conditional reservation utility U_H in state $H \in \{L, N\}$. If the principal cannot commit, then he cannot credibly offer more than U_H in period 2. In either case, the result is that the amount of classification risk imposed on the agent by any feasible contract will be bounded away from zero. Hence, for small values of c , the entire incentive will be provided by classification risk and the agent will be fully insured in the second period. The intertemporal provision of incentives is the subject of next section of this paper, where the full impact of commitment by the principal and commitment by the agent on the provision of incentives over time will be discussed.

5.3. Outcome Risk vs. Classification Risk

In a dynamic model of insurance with moral hazard and learning, the agent faces two types of risk, outcome risk and classification risk. To illustrate, consider the agent under autarky, and assume that she chooses high effort. In this case, the agent's wealth in state $S \in \{1, L, N\}$ is equal to $w - x$ with probability $q_S(c)$ and equal to w otherwise. Thus the agent faces risk of a loss, or outcome risk, in each period. In addition, the agent expects utility U_S in state S , where $U_L < U_1 < U_N$.

Thus the agent expects to do worse in the second period if she experiences a loss in the first period. This is classification risk, the risk in expected future utility due to learning about the agent's type.

Both outcome risk and classification risk can be used to induce the agent to choose high effort. As illustrated by Propositions 5 and 6, the contracting environment affects the extent to which the principal can make use of each of these. Thus the task of the principal is to find for the given contracting environment the least costly way to give the agent her reservation utility while exposing her to enough outcome and classification risk to induce her to choose high effort. This section considers the question of what determines the optimal balance between outcome risk and classification risk in providing the agent with incentives.

When actions are persistent, the optimal bilateral-commitment contract (Proposition 1) and the optimal contract when only the principal can commit (Proposition 3) are quite complicated. Because the contracts exhibit outcome risk in both the first and second periods, it is impossible to explicitly solve for the optimal contracts without making additional assumptions on the form of the utility functions and technology. However, making such assumptions does not yield significant insight into the interaction between outcome risk and classification risk in the provision of incentives over time.

An approach that is useful in characterizing the relative roles of outcome risk and classification risk is to assume that actions are not persistent. When actions are not persistent, the first period output is a sufficient statistic for the effort choice. Consequently, the optimal contract will not depend on the second period outcome, and the agent will be fully insured in each of the second period states. This facilitates the analysis since it leaves only two sources of risk for the agent, outcome risk due to partial insurance in the first period and classification risk due to the difference between the utility offered to the agent in the second period following no loss and following a loss.

The optimal contract when actions are persistent will differ from the optimal

contract when actions are not persistent in that when actions are persistent the principal will have at his disposal another instrument, outcome risk in the second period, which he will generally use to impose at least some outcome risk on the agent in the second period, adjusting the levels of first period outcome risk and classification risk accordingly. However, the general insights into the interaction between outcome risk and classification risk will continue to hold. The remainder of this subsection assumes that actions are not persistent.

The natural benchmark to begin the analysis is with the bilateral-commitment regime. The optimal contract when both parties can commit to long-term contracts and actions are not persistent was derived in Proposition 3 and is given by:

$$v_N^* = v_{NN}^* = v_{NL}^* = U_1 + \frac{cq_1(c)}{2(q_1(0) - q_1(c))} \quad (5.1)$$

$$v_L^* = v_{LN}^* = v_{LL}^* = U_1 - \frac{c(1 - q_1(c))}{2(q_1(0) - q_1(c))}. \quad (5.2)$$

Since the incentive compatibility constraint binds in this problem, $U(c, \Delta) - U(0, \Delta) = c$. The left-hand side of this expression can be decomposed into outcome risk and classification risk. Since there is no outcome risk in the second period when actions are not persistent, the total outcome risk under contract Δ is given by

$$(q_1(0) - q_1(c))(v_N - v_L). \quad (5.3)$$

Substituting (5.1) and (5.2) into (5.3) and simplifying, the total outcome risk in the optimal bilateral commitment contract when actions are not persistent is equal to $\frac{1}{2}c$. The total amount of risk needed to induce the agent to choose high effort is c . Hence the optimal bilateral commitment contract “evenly” divides the risk between outcome risk and classification risk.

In addition to the even balance between outcome risk and classification risk, this contract also exhibits perfect consumption smoothing along all paths through the game. That is, along any path the agent’s consumption in the first and second

periods is equal. As discussed earlier, this arises from the fact that utility in the first period following outcome H and utility in all second period outcomes following outcome H are perfect substitutes in terms of the incentives they provide; increasing v_H by k and increasing v_{HL} and v_{HN} by k have the same incentive effect. Due to the convexity of $h(\cdot)$, in the absence of any constraints to shifting utility payments forward or backward in time this implies that the marginal cost of increasing utility in the first period must equal the marginal cost of increasing utility in the second period. Again, this is a simple application of Rogerson (1985) Proposition 1.

Seen in this light, the even balance between outcome risk and classification risk and the presence of perfect consumption smoothing along paths through the game are really two aspects of the same phenomenon. Namely, that when the cost of increasing the agent's utility is convex, the principal wishes to spread the risk the agent must face as evenly as possible.

If either the principal or the agent is unable to commit, then the principal will not be free to shift utility between the first and second periods, and consequently the optimal contract will fail to exhibit the even division of incentives between outcome risk and classification risk that the bilateral-commitment contract does. This is most simply demonstrated by the case where neither the principal nor the agent can commit.

When neither party can commit, the equilibrium contract is as described in Proposition 5. In the spot contracting regime, the fact that neither party can commit forces the contract to be such that $v_{HL}^* = v_{HN}^* = U_H$. In this case, the total outcome risk to which the agent is exposed is computed by substituting the optimized values for v_N and v_L from Proposition 5 into (5.3), and is equal to:

$$c - (q_1(0) - q_1(c))(U_N - U_L). \quad (5.4)$$

Since the second term of (5.4) does not depend on c , the total outcome risk imposed on the agent is smaller than the total outcome risk in the bilateral-

commitment regime when $c < 2(q_1(0) - q_1(c))(U_N - U_L)$ and larger when $c > 2(q_1(0) - q_1(c))(U_N - U_L)$. Only when $c = 2(q_1(0) - q_1(c))(U_N - U_L)$ does the spot contracting regime exhibit the same balance of incentives between outcome and classification risk as the bilateral-commitment contract.

The behavior of the optimal contract when either only the principal or only the agent can commit is similar to the behavior when neither part can commit, although the fact that one of the parties can commit permits incentives to be more evenly divided between outcome risk and classification risk. Still, the additional constraints imposed on the problem by the lack of commitment result in the optimal contract involving too much classification risk when the cost of effort is small and too little classification risk when the cost of effort is large.

6. Comparison of the Contracts

Thus far, this paper has shown that in moral hazard problems with persistent actions and learning, commitment plays two roles. First, commitment by the principal allow the principal and agent to enter into contracts that depend on the second period outcome, thus making use of the information it provides to more effectively provide the agent with incentives. Second, commitment by both the principal and the agent permit the optimal contract to more efficiently divide the agent's incentives between outcome risk and classification risk. This section shows that these two factors together imply that the bilateral-commitment contract Pareto dominates other forms of contracts.

Call the bilateral long-term contracting regime (BC), the case where only the principal can commit (PC), and the spot contracting regime (SC). The key insight in comparing the contracts is that the feasible regions of the various problems are nested. The (BC) contracting problem maximizes (OF) subject to (IC) and (P). The (PC) contracting problem adds the second period participation constraints (P_L) and (P_N). The (SC) contracting problem adds the additional

constraint that the principal's second period contracts must be profit maximizing. Thus the feasible sets in the (BC) contains the feasible set in (PC) which in turn contains the feasible set in (SC) .

Let Π_{BC}^* , Π_{PC}^* , and Π_{SC}^* be defined as the optimal value of the principal's profits in the (BC) , (PC) , and (SC) problems. Proposition 7 ranks the principal's profits under the various regimes.

Proposition 7: *The profits to the principal are ranked as follows*

$$\Pi_{BC}^* \geq \Pi_{PC}^* > \Pi_{SC}^*$$

The consumer is indifferent between all types of contracts.

As discussed following Proposition 6, the optimal contract in the (BC) problem will differ from the optimal contract in the (PC) problem only if (P_L) or (P_N) binds at the optimum in the (PC) problem. While the principal's desire to smooth the agent's consumption along paths through the game implies that, generally, at least one of these constraints will bind, it is not necessarily the case that one must bind. Thus the optimal (BC) contract need not strictly dominate the optimal (PC) contract.

The most important aspect of Proposition 9 is the fact that the bilateral-commitment contract strictly outperforms the optimal series of spot contracts. There are two reasons for this. The first is that in the bilateral-commitment regime the principal faces no constraints on shifting utility back and forth along paths through the game. Consequently the bilateral-commitment contract permits better allocation of the agent's incentives over time. While this does contribute to the superiority of the bilateral-commitment contract, it may not be robust to extending the model by giving the agent free access to credit markets. This is because the agent can use the credit markets to provide the benefits due to intertemporal consumption smoothing that the bilateral-commitment contract does. Hence bilateral-commitment may not be necessary. However, to the extent that the agent is limited in her access to credit markets or her access is monitored by

the principal there are benefits to be gained via superior intertemporal provision of incentives in the bilateral-commitment contract.

A more fundamental reason why the optimal bilateral-commitment contract out-performs the optimal series of spot contracts is that the most efficient provision of incentives to the agent involves outcome risk in the second period. This is because when actions are persistent the second-period outcome offers information beyond that contained in the first-period outcome which can be used to construct better incentives if the contract can be conditioned on the second-period outcome. Since spot contracts cannot impose outcome risk on the agent in the second period, they will necessarily depart from the optimal provision of incentives. Consequently the bilateral-commitment contract will out-perform spot contracts. Further, the idea that the bilateral-commitment contract makes better use of the information contained in the second-period outcome is robust to all natural extensions of the model, including giving the agent free access to credit markets.

One implication of proposition 9 is that in the presence of transaction costs, the principal and agent will enter into contracts more often when long-term contracts are enforceable than when they are not. To illustrate, suppose that the principal is an insurer and has a cost of administering an insurance contract of k per period. Since the principal's profits are ranked as in proposition 9, there will be some values of the transaction cost where the principal cannot break even under spot contracting but can break even if one or more of the parties are able to commit to long-term contracts. In this context, if bilateral-commitment is not possible, some agents run the risk of becoming uninsurable following particularly bad histories while they would not if long-term contracts could be enforced. In such circumstances, the ability to commit will not only increase the principal's expected profits, but also affect whether or not the principal and agent contract at all.

The subject of uninsurability is a public policy issue of great concern. Proposition 9 implies that one way to remedy the problem of uninsurability is to create

laws that help to promote commitment to long-term contracts. Interestingly, the trend toward passing legislation that prohibits denying an agent insurance based on a particular history can be seen as one such law.

In addition, since the ability to commit leads to Pareto improvements, one would expect that the principal and agent would attempt to find ways to commit beyond the creation of new laws. This may include lengthening the duration of a “period” by paying an insurance period once a year instead of once every six months. By signing the longer-term contract, the principal and agent are able to garner some of the benefits of commitment.

Similarly, the desire to increase the level of commitment in insurance relationships may also provides insight into why health-insurance benefits are provided by employers. When health insurance is provided by the agent’s employer, the fact that the agent is committed to her employment relationship also acts to commit her to her insurance relationship, and the fact that the insurance provider is committed to its relationship with the employer helps it to commit to relationships with individual consumers. In addition, through the employment contract the principal is able to implement policies such as permitting the agent to change insurance providers only during specified “employee benefits choice” periods, which also plays a role in increasing the level of commitment in the relationship.

7. Conclusion

7.1. Extensions

This paper considers the simplest model that captures both learning and persistence of action in a moral hazard framework. However, it is robust to the usual extensions. The conclusions presented here would translate in a straightforward manner to models with larger finite action or outcome spaces, as well as increasing the duration of the game to any finite number of periods. The key insight is

that in more general models the overinsurance results in this paper will manifest themselves as non-monotonicities of the optimal contract.

One extension that could impact the conclusions would be to allow the agent free access to credit markets. As Chiappori et al. (1994) note, an environment where the agent's savings can be monitored is equivalent to a situation where the agent has no access to credit but the principal does. In this paper, the principal is assumed to be able to save or borrow at zero interest rate. Thus the conclusions of this paper should remain robust to the assumption that the agent has access to credit, but that her savings can be monitored by the principal.

The case where the agent's savings behavior cannot be monitored is more complicated. If the agent's savings cannot be monitored, then there is asymmetric information at the renegotiation stage; the agent knows how much she saved but the principal does not. Due to this, there is an adverse selection problem, and the analysis becomes much more complicated. The interested reader is referred to Chiappori et al. (1994) for an introduction to the subject in the stationary environment.

However, while access to credit markets complicates the analysis, it can only eliminate the differences between long-term contracts and spot contracts that are due to superior intertemporal consumption smoothing under bilateral commitment contracting. However, while long-term contracts do help to spread incentives across time, long-term contracting also allows for more efficient provision of incentives to the agent within periods, due to the fact that commitment allows the principal to use the second-period outcome in structuring incentives. Consequently one would not expect access to credit markets to completely eliminate the benefits associated with long-term contracting.

This paper considers a simple model where the agent makes a single ex ante effort choice and the principal and agent learn about the agent's type over time. However, the general insights will transfer to more complex environments.

For example, suppose we add to the current model an effort choice made at

the beginning of the second period so that the first period's outcome depends on the ex ante effort choice but the second period's outcome depends on both effort choices. Assume that the rest of the model remains the same, and that the principal wishes to implement high effort in both the first and second periods²⁰.

In order to implement high effort in the second period, the contract must impose risk on the agent in the second period. If effort increases the likelihood of no loss occurring, such a contract will necessarily give the agent more utility following no loss in the second period than following a loss. Thus if there are only two outcomes, any second-period contract that implements high effort will always feature partial insurance, even when the learning effect is positive. The need to implement high effort in the second period will override the fact that overinsurance in the second period helps to induce the agent to choose high effort ex ante.

While the second period contract will always be monotonic if there are only two outcomes, it may fail to be monotonic if there are more than two outcomes. When there are many outcomes, the learning effect may lead to non-monotonicities in the second period incentive scheme, since in such an environment the optimal contract can implement high effort in the second period by rewarding the agent for high outcomes while not being monotonically increasing in the outcome. Such a scheme will tend to reward the agent for choosing high effort, but the desire to provide the agent with incentives in the first period may lead to non-monotonicities in the second period incentive scheme, especially in the middle of the range of outcomes.

²⁰If the principal wishes to implement low effort then the addition of the second effort choice does not impact the problem. Since low effort is the agent's lowest cost effort, the requirement that the contract implement low effort imposes no additional constraints on the contracting environment.

7.2. Discussion

The main result of this paper is to show that in the presence of learning, persistent actions, and commitment, overinsurance may arise in the second period of a repeated moral hazard model. This is true regardless of the outcome in the first period. This paper considers the simplest possible model containing moral hazard, learning, and persistent actions. In addition, the stochastic structure is also quite well behaved. The first period distribution satisfies MLRC, which is sufficient for monotonicity in the static problem, and the second period distribution differs from the first period distribution only through the difference in the posterior belief that the agent is the bad type. In short, the model considered here is as well behaved as possible, and even in this model the second period history-conditional distribution may fail to be increasing in the outcome.

The non-monotonicity identified in this paper will *a fortiori* appear in models that are less well behaved. Thus it provides a strong argument against monotonic contracts. In environments with larger action and effort spaces, the conditions needed to guarantee monotonicity in the static problem include MLRC and the Concavity of the Distribution Function Condition (see Grossman and Hart 1984), which already impose great restrictions on the distribution. However, this model shows that these conditions are not sufficient to guarantee monotonicity beyond the first period when there is learning and persistent action choice, and so one is left wondering if it is ever reasonable to expect dynamic agency contracts to exhibit monotonicity. More troubling, since most real-world insurance contracts are monotonic, relying on a fixed deductible, what causes this discrepancy between the theory and reality?

The second part of the paper considers the role of outcome risk and classification risk in providing the agent with incentives over time. It is shown that the bilateral-commitment contract exhibits an even division between outcome risk and classification risk, while environments where one or both parties can-

not commit involve too much classification risk for small values of the effort cost and too little classification risk for larger values of the effort cost. The benefits due to superior intertemporal provision of incentives combined with the fact that the bilateral-commitment contract makes better use of the information provided by the second-period outcome implies that the bilateral-commitment contract is Pareto superior to the other forms of contracting.

References

- [1] P. Chiappori, I. Macho, P. Rey, and B. Salanié, Repeated moral hazard: The role of memory, commitment, and the access to credit markets, *European Economic Review* **38** (1994) 1527-1553.
- [2] N. Falletta, *The Paradoxicon*, John Wiley & Sons, New York (1990) pp. 136-143.
- [3] D. Fudenberg, B. Holmstrom, and P. Milgrom, Short-term contracts and long-term agency relationships, *Journal of Economic Theory* **51** (1990) 1-31.
- [4] S. Grossman and O. Hart, An analysis of the principal-agent problem, *Econometrica* **51** (1985) 7-54.
- [5] M. Harris and B. Holmstrom, A theory of wage dynamics, *Review of Economic Studies* **49** (1982), 315-333
- [6] B. Holmstrom, Moral hazard and observability, *Bell Journal of Economics* **10** (1979), 74-91.
- [7] Y. Hirao, Learning and incentive problems in repeated partnerships, *International Economic Review* **34** (1993), 101-119.
- [8] R. Lambert, Long-term contracts and moral hazard, *Bell Journal of Economics* **14** (1983) 441-452.

- [9] C. Ma, Adverse selection in dynamic moral hazard, *Quarterly Journal of Economics* **106** (1991), 225-275.
- [10] J. Malcomson and F. Spinnewyn, The multiperiod principal-agent problem, *Review of Economic Studies* **55** (1988) 391-408.
- [11] T. Palfrey and C. Spatt, Repeated insurance contracts and learning, *Rand Journal of Economics* **16** (1985), 356-367.
- [12] P. Rey and B. Salainé, Long-term, short-term and renegotiation: on the value of commitment in contracting, *Econometrica* **58** (1990) 597-619.
- [13] W. Rogerson, Repeated moral hazard, *Econometrica* **53** (1985) 67-76.

8. Proofs

Proof of Proposition 1: Let λ be the Lagrange multiplier on (P) and μ be the Lagrange multiplier on (IC) . The optimal solution is determined by the complementary slackness conditions that (P) binds or $\lambda^* = 0$ and (IC) binds or $\mu^* = 0$ and

$$\begin{array}{ll} \text{the first order conditions} \\ h'(v_L^*) = \lambda^* + \mu^* \left(1 - \frac{q_1(0)}{q_1(c)}\right) & h'(v_N^*) = \lambda^* + \mu^* \left(1 - \frac{(1-q_1(0))}{(1-q_1(c))}\right) \\ h'(v_{LL}^*) = \lambda^* + \mu^* \left(1 - \frac{q_1(0)q_L(0)}{q_1(c)q_L(c)}\right) & h'(v_{LN}^*) = \lambda^* + \mu^* \left(1 - \frac{q_1(0)(1-q_L(0))}{q_1(c)(1-q_L(c))}\right) \\ h'(v_{NL}^*) = \lambda^* + \mu^* \left(1 - \frac{(1-q_1(0))q_N(0)}{(1-q_1(c))q_N(c)}\right) & h'(v_{NN}^*) = \lambda^* + \mu^* \left(1 - \frac{(1-q_1(0))(1-q_N(0))}{(1-q_1(c))(1-q_N(c))}\right). \end{array}$$

Proof of property 1.1: To show that (P) binds, assume $\lambda^* = 0$. Since $\frac{(1-q_1(0))}{(1-q_1(c))} > 1$, $\mu \left(1 - \frac{(1-q_1(0))}{(1-q_1(c))}\right) < 0$. Since $h'(\cdot) > 0$, this contradicts the assumption that an optimum exists if $\lambda^* = 0$.

To show that (IC) binds and its shadow price is positive, suppose $\mu^* = 0$. Then all payments are such that $h'(v) = \lambda^*$, which cannot implement high effort. ■

Proof of Property 1.2: Follows from the fact that $q_S(0) \neq q_S(c)$ for $S \in \{1, L, N\}$ and the fact that $h(\cdot)$ is strictly convex. ■

Proof of Property 1.3:

$h'(v_N^*) - h'(v_L^*) = \left(\lambda + \mu \left(1 - \frac{(1-q_1(0))}{(1-q_1(c))} \right) \right) - \left(\lambda + \mu \left(1 - \frac{q_1(0)}{q_1(c)} \right) \right) = \mu \frac{q_1(0) - q_1(c)}{(1-q_1(c))q_1(c)} > 0$. Therefore $v_N^* > v_L^*$ by the fact that $h(\cdot)$ is strictly convex and increasing. ■

Proof of Property 1.4: Begin with the L state. The sign of $v_{LN}^* - v_{LL}^*$ is the same as the sign of $h'(v_{LN}^*) - h'(v_{LL}^*) = \mu q_1(0) \frac{q_L(0) - q_L(c)}{q_1(c)(1-q_L(c))q_L(c)}$, which has the same sign as $q_L(0) - q_L(c)$. The proof for the N state is the same. ■

Proof of Property 1.5: First, show that $v_{LN}^* = v_{NL}^*$.

Claim: $q_1(\hat{c})(1 - q_L(\hat{c})) = (1 - q_1(\hat{c}))q_N(\hat{c})$

Proof of claim: $q_1(\hat{c})(1 - q_L(\hat{c}))$

$$\begin{aligned} &= (p_1 q_B(\hat{c}) + (1 - p_1) q_G(\hat{c})) \frac{(p_1 q_B(\hat{c})(1 - q_B(\hat{c})) + (1 - p_1) q_G(\hat{c})(1 - q_G(\hat{c})))}{p_1 q_B(\hat{c}) + (1 - p_1) q_G(\hat{c})} \\ &= (1 - q_1(\hat{c})) \frac{(p_1(1 - q_B(\hat{c}))q_B(\hat{c}) + (1 - p_1)(1 - q_G(\hat{c}))q_G(\hat{c}))}{1 - q_1(\hat{c})} = (1 - q_1(\hat{c}))q_N(\hat{c}). \end{aligned}$$

Proof of Property: Suppose $v_{NN}^* < v_{NL}^*$ and $v_{LN}^* < v_{LL}^*$. By Property 1.5, $v_{LL}^* > v_{LN}^* = v_{NL}^* > v_{NN}^*$. According to Rogerson (1985) Proposition 1, $h'(v_H) = q_H(c)h'(v_{HL}) + (1 - q_H(c))h'(v_{HN})$. Using this fact, $h'(v_N) = q_N(c)h'(v_{NL}) + (1 - q_N(c))h'(v_{NN}) < q_L(c)h'(v_{LL}) + (1 - q_L(c))h'(v_{LN}) = h'(v_L)$, which implies $h'(v_N) < h'(v_L)$, a contradiction of property 1.4. A similar argument shows that it cannot be that $v_{LL}^* = v_{LN}^*$ and $v_{NL}^* = v_{NN}^*$, and that it cannot be that one state features full insurance and one state features overinsurance. ■

Proof of Proposition 2: Suppose $q_H(c) > q_H(0)$ and $v_{HL} < v_{HN}$. Decrease v_{NN} by $q_N(c)k$ and increase v_{NL} by $(1 - q_N(c))k$. This does not violate the participation constraint and decreases cost. The effect on the left hand side of (IC) is given by

$$\begin{aligned} &((1 - q_1(c))(1 - q_N(c)) - (1 - q_1(0))(1 - q_N(0)))(-q_N(c)k) \\ &+ ((1 - q_1(c))q_N(c) - (1 - q_1(0))q_N(0))(1 - q_N(c))k \\ &= -k(q_N(0) - q_N(c))(1 - q_1(0)) > 0. \end{aligned}$$

Thus this also relaxes (IC) . Since the feasible change decreases cost, it cannot be that $v_{HL} < v_{HN}$ at the optimum. Hence $v_{NL}^* = v_{NN}^*$. ■

Proof of Proposition 3: Substitute $q_L(0) = q_L(c)$ and $q_N(0) = q_N(c)$ into the solution to the bilateral long-term commitment problem as in proposition 1 above. (IC) becomes $2(q_1(0) - q_1(c))(v_N - v_L) = c$, and (P) becomes $q_1(c)v_L + (1 - q_1(c))v_N = U_1$. The solution to these equations implies the result. ■

Proof of Proposition 4: $q_B(\hat{c}) = q_G(\hat{c})$ implies that $q_L(\hat{c}) = q_N(\hat{c}) = q_1(\hat{c})$. The optimal contract is as in proposition 1 with the appropriate substitutions. The same arguments as in properties 1.1 and 1.2 prove that λ and μ are strictly positive. The results follow from the fact that $h'()$ is strictly increasing. ■

Proof of Proposition 5: Suppose $A > 0$. First, show that (IC) and (P) bind at the optimum. To show (IC) binds at the optimum, suppose it doesn't. $A > 0$ implies $v_N > v_L$ in any feasible contract. Perform the following adjustment. Decrease v_N by $q_1(c)k$ and increase v_L by $(1 - q_1(c))k$. Since (IC) does not bind, it is not violated. This keeps expected utility constant and decreases the expected cost to the principal, violating the assumption of optimality. To show (P) binds, suppose it doesn't. Decreasing v_L and v_N by k does not violate (P) since it is assumed not to bind. The LHS of (IC) remains constant, so it is not violated, and the cost to the principal decreases, contradicting the optimality assumption. Since (IC) and (P) bind, the optimum is found by solving for the v_N and v_L that satisfy the constraints.

If $A \leq 0$, then the (IC) constraint is satisfied when $v_N = v_L$. Since setting $v_N = v_L = U_1$ is the cost minimizing way to satisfy (P) and also satisfies (IC) , it comprises the solution to the short-term contracting problem. The values of the other variables are carried over from the above discussion. ■

Proof of Proposition 6: Letting λ be the multiplier on P , λ_L be the multiplier on P_L , λ_N be the multiplier on P_N , and μ be the multiplier on IC , the solution to the problem is given by

$$h'(v_L^*) = \lambda^* + \mu^* \left(1 - \frac{q_1(0)}{q_1(c)} \right)$$

$$\begin{aligned}
h'(v_N^*) &= \lambda^* + \mu^* \left(1 - \frac{1 - q_1(0)}{1 - q_1(c)} \right) \\
h'(v_{LL}^*) &= \lambda^* + \frac{\lambda_L^*}{q_1(c)} + \mu^* \left(1 - \frac{q_1(0) q_L(0)}{q_1(c) q_L(c)} \right) \\
h'(v_{LN}) &= \lambda^* + \frac{\lambda_L^*}{q_1(c)} + \mu^* \left(1 - \frac{q_1(0) (1 - q_L(0))}{q_1(c) (1 - q_L(c))} \right) \\
h'(v_{NL}) &= \lambda^* + \frac{\lambda_N^*}{1 - q_1(c)} + \mu^* \left(1 - \frac{(1 - q_1(0)) q_N(0)}{(1 - q_1(c)) q_N(c)} \right) \\
h'(v_{NN}^*) &= \lambda^* + \frac{\lambda_N^*}{1 - q_1(c)} + \mu^* \left(1 - \frac{(1 - q_1(0)) (1 - q_N(0))}{(1 - q_1(c)) (1 - q_N(c))} \right)
\end{aligned}$$

and the complementary slackness conditions that either the multipliers λ^* , λ_L^* , λ_N^* , and μ^* are equal to zero or their respective constraint binds. If (IC) binds, then the proofs of the individual properties are similar to the proofs of the corresponding property of proposition 1. The proof of proposition 8.5 follows from the fact that the proof of property 1.5 also implies that it is not the case that $q_L(c) > q_L(0)$ and $q_N(c) > q_N(0)$.

If (IC) doesn't bind then the agent is fully insured in the second period. Further, the principal gains no benefit from decreasing the amount of risk to which the agent is exposed. Since the principal can always benefit from decreasing the risk to which the agent is exposed, this means that if (IC) doesn't bind then the contract must exhibit the minimum amount of risk possible. The minimum risk contract that satisfies (P) , (P_N) , and (P_L) is given by $v_{NL}^* = v_{NN}^* = U_N$, and $v_L^* = v_N^* = v_{NL}^* = v_{NN}^* = \left(\frac{2U_1 - U_N(1 - q_1(c))}{q_1(c) + 1} \right)$. ■

Proof of Proposition 7: The ranking of profits follows from the nesting of the feasible sets. Since the optimal short-term contract is feasible but not optimal in the case where only the principal can commit, the inequality is strict. The consumer's indifference derives from the fact that the participation constraints always bind. ■