

OVERCOMPENSATION AS A PARTIAL SOLUTION TO COMMITMENT AND RENEGOTIATION PROBLEMS: THE CASE OF *Ex Post* MORAL HAZARD

M. Martin Boyer

ABSTRACT

In a Costly State Verification world, an agent who has private information regarding the state of the world must report what state occurred to a principal, who can verify the state at a cost. An agent then has what is called *ex post moral hazard*: he has an incentive to misreport the true state to extract rents from the principal. Assuming the principal cannot commit to an auditing strategy, the optimal contract is such that: (1) the agent's expected marginal utility when there is an accident (high- and low-loss states) is equal to his marginal utility when there is no accident; (2) the lower loss is undercompensated, while the higher loss is overcompensated; and (3) the welfare of the agent is greater under commitment than under no-commitment. Result 2 is contrary to the results obtained if the principal can commit to an auditing strategy (higher losses underpaid and lower losses overpaid). The reason is that by increasing the difference between the high and the low indemnity payments, the probability of fraud is reduced.

INTRODUCTION

The traditional mechanism design literature has concentrated on designing contracts where it is optimal for all players to tell the truth. These mechanisms do not seem to apply in reality as we observe agents who sometimes do not tell the truth. The two most blatant examples are income-tax fraud and insurance fraud.

Martin Boyer is Associate Professor in the Department of Finance at HEC Montréal, Université de Montréal. The author can be contacted via e-mail: martin.boyer@hec.ca. This article draws extensively on a chapter of his doctoral dissertation at the Wharton School of the University of Pennsylvania. The author would especially like to thank Sharon Tennyson, my advisor, for her encouragement. The author would also like to thank Kodjovi Assoé, Marcel Boyer, Keith Crocker, Georges Dionne, Neil Doherty, Marc Duhamel, Nathalie Fombaron, Thomas Palfrey, Michel Poitevin, and seminar participants at the Université de Montréal and at the University of Minnesota for helpful discussions. The author would also like to thank the anonymous referees for their comments as well as the editor. Their investment in time and effort has greatly improved the quality of the article. This research benefited from the financial support of the S.S. Huebner Foundation and of the SSHRC of Canada. The continuing financial support of CIRANO is also acknowledged.

The mechanism design literature appropriately points out that the revelation principle may not hold when there is no commitment. More specifically, the commitment of the principal to a strategy is in many situations not a reasonable assumption since the players may be better off *a posteriori* if the principal deviates from the strategy agreed upon *a priori*. The case of *ex post* moral hazard is a clear example of the commitment problem. Suppose the agent's payoff depends on the state of the world he announces to the principal. Suppose also that the state of the world is private information to him. If it is costly for the principal to verify the state of the world, then the agent has an incentive to misrepresent it. Suppose the principal convinces the agent that she will audit every claim and detect and penalize any fraud. Since it is then optimal for the agent to tell the truth, the principal will not want to waste resources by verifying the state of the world, since she knows that the agent has told the truth about it. Thus the principal wants to renegotiate *a posteriori* to save resources.

The goal of this article is to address this commitment problem between the informed agent and the principal. An insurance-contract approach is used because the *ex post* moral hazard problem in insurance, namely insurance fraud, is very important. I shall restrict the analysis to the case of *build-up* because it appears to be the most costly to the economy.¹

Although a policyholder/insurer approach to the problem is used, it could just as well have been an entrepreneur/financier approach à la Gale and Hellwig (1985), or a polluter/regulator approach à la Laffont and Tirole (1993) or Lewis (1996). In the first case, the entrepreneur knows whether a project was a success or not, whereas the financier, who lent the money, does not. Thus there is an incentive for the entrepreneur to take the money and run.² In the second case, the polluter knows how much mess is left to be cleaned up, whereas the government does not unless it sends biochemists to the polluted site. In an insurance setting, the policyholder knows what loss he suffered in an accident, whereas the insurer knows only that an accident occurred.

The game is then for the agent to make a report to the principal concerning the severity of the accident (send a message) and for the principal to decide whether to verify the agent's report, at a cost. The results of the article are threefold. First, an agent's average marginal utility when there is an accident (high- and low-loss states) is equal to his marginal utility when he is not involved in an accident. The second result shows that an agent is overcompensated (undercompensated) in the event of an accident of high (low) severity. This means that in the event of a high (low) loss,³ the benefit the agent receives is greater (smaller) than his loss. This result stems from the principal's attempt

¹ Insurance fraud build-up is characterized by a policyholder who is involved in an accident and who, upon learning of his real loss, engages in claims padding (exaggerates his loss) to collect more money from his insurer. Build-up appears to be the most costly type of fraud. The Rand Corporation estimates that in 1993, more than \$18 billion was paid for apparently fraudulent medical claims arising from automobile accidents in the United States (*Research Brief*, 1995). Another case of fraud known as outright fraud relates to the case where the agent invents an accident instead of just exaggerating one. Outright fraud can be seen as a special case of build-up with $\pi = 1$ and $\lambda_L = 0$ in the model presented in this article. For more details on all types of insurance fraud, see Hoyt (1989).

² See, for example, Scheepens (1995), Persons (1997), and Khalil and Parigi (1998).

³ The terms *severity of accident* and *loss* are synonymous in this article.

to mitigate insurance fraud. As shown in the model, fraud is more (less) prevalent when the difference between the high-loss indemnity and the low-loss indemnity is small (large). The greater the difference between the two indemnity payments, the more the principal has to lose by not auditing a claim of high loss. Third, I show there is a welfare loss from not being able to commit to an auditing strategy.

The contribution of these results is easier to understand if one inspects the related literature on *ex post* moral hazard and on noncommitment. The difference between *ex post* and *ex ante* moral hazard depends on whether the agent plays after (*ex post*) or before (*ex ante*) Nature. Spence and Zeckhauser (1971) were the first to characterize these two forms of moral hazard. They show that there is no real difference between the two problems owing to their assumption that auditing is either costless or impossible.

Townsend (1979) was the first to present a model where the principal can verify the realized state of the world at a cost. This approach is known as the *costly state verification* approach. In Townsend's model, the optimal insurance contract between the principal and the agent stipulates a no-auditing region where a fixed payment is made, and an auditing region (where audits are always performed) where the agent has full insurance less a deductible. Mookherjee and Png (1989) showed that truth telling can be achieved using stochastic audits (see also Bond and Crocker, 1997). They also showed that a bonus should be paid to those agents who are audited and who are found to have told the truth. Moreover, they show that Townsend's debt contract is no longer optimal when agents are risk averse. Mookherjee and Png's optimal contract is Pareto superior to Townsend's, since less money is wasted on audits, but this is achieved only by allowing the principal to charge draconian fines.

Commitment is an important feature of contractual relationships between a principal and an agent. In a multiperiod setting, we know that an agent and a principal would be better off if they could sign a long-term, full-commitment contract. The problem, however, is that when the time comes to implement the contract, both players find it in their best interests to renegotiate the contract to achieve *ex post* efficiency. The commitment assumption has typically been challenged in multiperiod settings.⁴ In a one-period setting, the importance of commitment is also nontrivial in the presence of *ex post* moral hazard.⁵

The article that relates the most to this one is Khalil (1997). Using the same basic framework as Baron and Myerson (1982), Khalil develops a model where the principal is

⁴ See Cooper and Hayes (1987), Hosios and Peters (1989), Dionne and Doherty (1994), and Fombaron (1997).

⁵ See Reinganum and Wilde (1985), Graetz, Reinganum, and Wilde (1986), Beaudry and Poitevin (1993, 1994), Picard (1996, 1997), Jost (1996), Khalil (1997), and Boyer (2000). Melumad and Mookherjee (1989) attempt to address this commitment problem (see also Swierzbinski, 1994; Lewis and Sappington, 1995; Picard, 1997). By allowing the principal to delegate the auditing to a third party, they contend that the commitment problem may be avoided. Unfortunately, that is not necessarily the case because the third party itself may be faced with commitment and/or agency problems. Suppose the principal can commit to an auditing budget *ex ante*. This would yield the same truth-telling results as committing to an auditing strategy. The problem is then for the principal to determine whether the auditing has been conducted properly. In other words, *Quid custodiet ipsos custodes?* (Who audits the auditors?) Melumad and Mookherjee assume this problem away by arguing that it can be done at no cost.

unable to commit to an auditing strategy. This framework consists in a principal (regulator) who does not know the production cost function of the agent (firm) she seeks to regulate. I discuss the recent literature as it relates to my model more thoroughly in "Discussion and Conclusion."

Although I am interested in designing an optimal contract in a costly state *verification* world where the principal cannot commit, it is appropriate to mention the contributions of Lacker and Weinberg (1989) and Crocker and Morgan (1998) in a costly state *falsification* framework. Crocker and Morgan (1998) construct a model where an agent is compensated not on his loss directly, but on his falsified loss. In the falsification literature, an agent may spend resources to conceal his loss so that his claimed loss never equals his actual loss (except at the two end-points of the loss distribution). The principal can exploit the fact that the cost of falsification increases in the amount falsified. It then becomes increasingly costly for an agent who suffered a very small loss to imitate an agent who suffered a larger loss; there is therefore truth in the messages, although there is no truth in the action.⁶ The optimal contract found is such that harder to falsify claims (smaller losses) are overcompensated whereas easier to falsify claims are undercompensated (larger losses). This contrasts with the results of my model where smaller losses are undercompensated and larger losses are overcompensated.

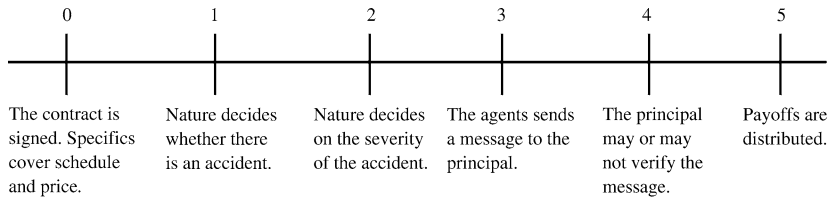
The remainder of the article is organized as follows. In the next section, the claiming game between the agent and the principal is presented. In this game the informed agent must report his type (severity of accident) to the principal, who must then decide whether to audit. "The Optimal Contract" section of the article presents the optimal contract in such a setting as well as its implications. Finally, "Discussion and Conclusion" discusses the results and presents a conclusion.

THE CLAIMING GAME

Before the model is examined in detail, it seems appropriate to state some basic assumptions. The agent is risk-averse ($U'(W) > 0$, $U''(W) < 0$), where $U(W)$ is the agent's von Neumann–Morgenstern utility function over final wealth. All agents have exogenous wealth given by Y . The principal is risk-neutral and is making zero expected profits. Any message sent by the agent to the principal is costless, whether the message is truthful or not. The distribution of losses given by a quadruplet $\Xi = \{\pi, \rho, \lambda_L, \lambda_H\}$ is common knowledge. If an accident occurs (with probability π), then only two possible losses can occur, λ_H and λ_L , with $\lambda_H > \lambda_L > 0$. Conditional on an accident occurring, the agent suffers loss λ_H with probability ρ . Although the agent suffers loss $\lambda \in \{\lambda_H, \lambda_L\}$, the payment he receives is given by $\beta \in \{\beta_H, \beta_L\}$. Therefore, when the claiming game is played, the players must take into account not only the values of λ_H and λ_L , but also those of β_H and β_L . The only information that is private to the agent is the severity of the accident (the loss he suffered). The insurer does not know the severity of the accident unless she conducts a costly audit of the claim at a cost of c which is not too large (we shall define *too large* in "The Other Two Cases" section of the article). Auditing is perfect if conducted, and the audit cost c is common knowledge.

⁶ The principal is, however, assumed to know the falsification cost function of the agent. If this cost function is unknown, we may be faced with a problem à la Baron and Myerson (1982) of regulating agents with unknown costs.

FIGURE 1
Sequence of Play



The sequence of the game is displayed in Figure 1. The claiming game is played in stages 2–5. In stage 0 the contract is signed. Nature plays in stages 1 and 2. In stage 1, Nature reveals to all players whether there is an accident⁷ or not, while in stage 2 it reveals the severity of the accident only to the agent. The agent then makes a report to the principal, who then decides whether to audit the report. Since the occurrence of an accident is common knowledge, the only thing the principal does not know is whether the loss is equal to λ_H or to λ_L . I shall assume that if the principal verifies the state of the world, then she must compensate the agent according to the benefit corresponding to his true loss. Thus an agent caught sending a false message receives compensation equal to what he would have received had he told the truth.

An agent caught committing fraud must incur two kinds of penalties, k_1 and k_2 . Penalty k_1 , which is denominated in utility terms, is a *deadweight loss* in the sense that it is only incurred by the agent and does not profit the principal.⁸ The amount k_2 , on the other hand, is *paid to the principal*.

Following Karpoff and Lott (1993) and Waldfogel (1994) who find that the reputation loss (k_1 in my model) associated with a criminal conviction far outweighs any direct penalty (k_2 in my model) inflicted on criminal elements, one should think of k_2 as being small compared to k_1 . Waldfogel (1994) finds that fraud conviction reduces the offender's future income by as much as 30 percent. In the case of corporations, Karpoff and Lott (1993) find that the reputation loss is associated with an abnormal stock return between -1.34 and -5.05 percent. This market value loss represents

⁷ The term *accident* is used as a generic term to represent an action of Nature that is common knowledge. It can be viewed as an automobile accident where there are witnesses, a flood, a hurricane, etc. Everyone knows if a flood occurs or if a hurricane hits. Another example is the case of police reports that are often needed when claiming a burglary or some type of physical injury in an automobile accident. As shown in Picard (1997), the optimal contract characterized by Bond and Crocker (1997) when the accident state is observable is not optimal when the principal cannot observe whether an accident truly occurred. If the accident state is not observable, the optimal insurance contract will surely be quite different as it must now deter outright fraud. See also Crocker and Snow (1985).

⁸ Another way to view this deadweight penalty is as an opportunity cost of being in autarchy. Suppose an agent who is caught cheating is shut out of the insurance market, and thus receives the expected utility of remaining in autarchy. If $\bar{V}$ represent the agent's expected utility of participating in the market, and if $\underline{V}$ represents the agent's expected utility of remaining in autarchy, then $k_1 = (\frac{\delta}{1-\delta})(\bar{V} - \underline{V})$ is the implicit penalty for being shut out of the market (where δ is the discount factor).

more than ten times the sum of any fine, penalty, or court-imposed costs. Because this direct loss is small compared to the reputation loss, I shall assume that it is smaller than the auditing cost ($k_2 < c$). This arguably strong assumption is based on a study by the U.S. Sentencing Commission (1988) where it is shown that the total monetary penalty paid for fraud in terms of criminal restitution plus any fine is less than the damage produced by fraud. The two penalties are fixed before the start of the game and common knowledge.^{9,10}

There are three choice variables for the principal in this model: the coverage in case of a high loss (β_H), the coverage in the case of a low loss (β_L), and the premium (α). The variables are chosen by the principal to maximize the agent's expected utility before the claiming game starts. They are thus fixed for the duration of the claiming game and can be manipulated as parameters. The game can end in two ways. First, there may be no accident, in which case there is no opportunity for the agent to exaggerate his claim. The payoff to the agent and the principal are then given by $U(Y - \alpha)$ and α . Second, there could be an accident, and the claiming game is played. Table 1 summarizes all the variables used in this article, while Table 2 lists all the possible payoffs to the players contingent on their actions. Figure 2 is an extensive-form representation of the claiming game.

Given that β_L , β_H , and α are chosen in period 0, the solution to the claiming game played in periods 1–5 gives us a perfect Bayesian Nash equilibrium (PBNE). By definition the PBNE of this game is a sextuplet.

Definition 1: A PBNE is defined in this game as

$$PBNE = \left(\begin{array}{l} \text{Agent's strategy if Nature chose } \lambda = \lambda_L, \\ \text{Agent's strategy if Nature chose } \lambda = \lambda_H; \\ \text{Principal's strategy if the Agent reported } \lambda' = \lambda_L, \\ \text{Principal's strategy if the Agent reported } \lambda' = \lambda_H; \\ \text{Principal's beliefs in the information set resulting from } \lambda' = \lambda_L, \\ \text{Principal's beliefs in the information set resulting from } \lambda' = \lambda_H \end{array} \right).$$

I shall use η to represent the probability that an agent announces a loss λ_H when his loss is in fact λ_L . Then, ν represents the probability that an agent who reported a loss

⁹ Alternatively, I could have let the principal choose the penalties subject to some upper bounds. The resulting penalties would have been equal to the maximum allowable, which are set outside the principal's jurisdiction (see Mookhejee and Png, 1989). There is therefore no loss in generality in letting the penalty be given exogenously. Furthermore, given that the optimal contract I find is independent of the deadweight penalty, the size and form of the deadweight penalty has no bearing on the optimal allocation in the economy (except the equilibrium probability of audit).

¹⁰ With part of the total penalty collected by the principal ($k_2 > 0$), we are then faced with the problem that auditing becomes a rent-extracting device for the principal as in Khalil (1997). If no part of the penalty is paid to the principal (i.e., $k_2 = 0$), the principal would not profit from audits (he can only reduce his payout at best), which would only become a means of reducing the number of false messages sent in the economy. Picard (1996) studies the outright fraud case where a proportion of the penalty is paid to the principal.

TABLE 1
Summary of All the Variables Used in This Article

Variable Name	Signification
Y	Agent's initial wealth
β_L	Benefit in case of a low loss
β_H	Benefit in case of a high loss
α	Premium
λ_L	Size of the low loss
λ_H	Size of the high loss
π	Probability of an accident occurring
ρ	Probability that the accident results in a high loss
c	Cost of auditing
k_1	Deadweight penalty paid by agent for committing fraud
k_2	Penalty paid by agent to principal for committing fraud
η	Agent's probability of committing fraud given the loss is low
ν	Principal's probability of auditing given that the reported loss is high

TABLE 2
Payoffs to the Agent and the Principal Contingent on Their Actions and the State of the World

State of the World	Action of Agent	Action of Principal	Payoff to Agent	Payoff to Principal
No accident	–	–	$U(Y - \alpha)$	α
<i>Low loss</i>	<i>Tell truth</i>	<i>Audits</i>	$U(Y - \alpha - \lambda_L + \beta_L)$	$\alpha - \beta_L - c$
Low loss	Tell truth	Does not audit	$U(Y - \alpha - \lambda_L + \beta_L)$	$\alpha - \beta_L$
Low loss	Lie	Audits	$U(Y - \alpha - \lambda_L + \beta_L - k_2) - k_1$	$\alpha - \beta_L + k_2 - c$
Low loss	Lie	Does not audit	$U(Y - \alpha - \lambda_L + \beta_H)$	$\alpha - \beta_H$
High loss	Tell truth	Audits	$U(Y - \alpha - \lambda_H + \beta_H)$	$\alpha - \beta_H - c$
High loss	Tell truth	Does not audit	$U(Y - \alpha - \lambda_H + \beta_H)$	$\alpha - \beta_H$
<i>High loss</i>	<i>Lie</i>	<i>Audits</i>	$U(Y - \alpha - \lambda_H + \beta_H - k_2) - k_1$	$\alpha - \beta_H + k_2 - c$
<i>High loss</i>	<i>Lie</i>	<i>Does not audit</i>	$U(Y - \alpha - \lambda_H + \beta_L)$	$\alpha - \beta_L$

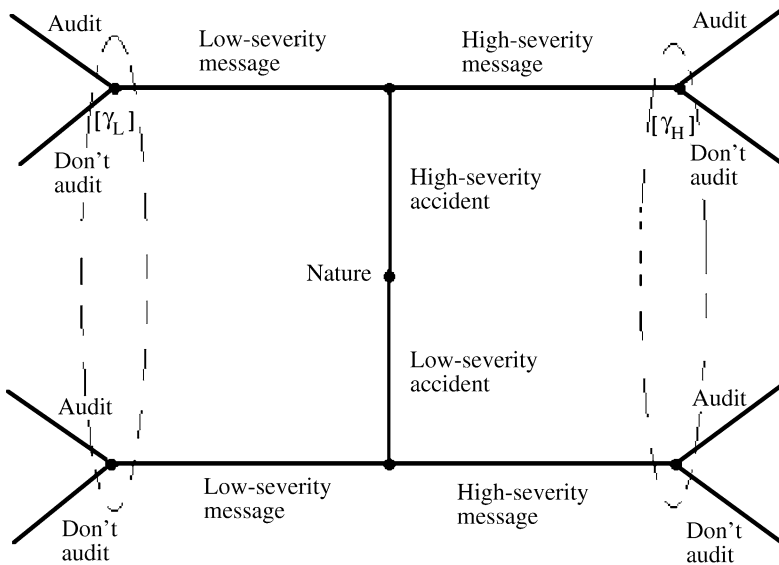
Note: The contingent states in italics never occur in equilibrium: they represent actions that are off the equilibrium path.

of λ_H is audited. I shall also assume for now that higher losses are compensated more than lower losses (i.e., $\beta_H > \beta_L$). This brings me to the first result of the article.

Lemma 1: *The unique PBNE in mixed strategy of this game¹¹ is such that: (1) the agent who suffered a high loss always tells the truth; (2) the agent who suffered a low loss tells the*

¹¹ A necessary condition for the mixed strategy to exist is that $\eta \in [0, 1]$. This occurs if and only if $c < (1 - \rho)(\beta_H - \beta_L + k_2)$. As for uniqueness, Gibbons (1992) and Myerson (1991) state that in a 2×2 game (two players each with two possible actions) there will be at most one mixed-strategy equilibrium.

FIGURE 2
Extensive Form of the Claiming Game



truth with probability $1 - \eta$ and reports a high loss with probability η ; (3) the principal never audits a low-loss claim; and (4) the principal audits a high-loss claim with probability v and does not audit with probability $1 - v$, where

$$\eta = \left(\frac{c}{\beta_H - \beta_L + k_2 - c} \right) \left(\frac{\rho}{1 - \rho} \right), \quad (1)$$

$$v = \frac{U(Y - \alpha - \lambda_L + \beta_H) - U(Y - \alpha - \lambda_L + \beta_L)}{U(Y - \alpha - \lambda_L + \beta_H) - U(Y - \alpha - \lambda_L + \beta_L - k_2) + k_1}. \quad (2)$$

Proof: All the proofs are given in the Appendix.

Q.E.D.

What Lemma 1 does is model the strategic behavior of each player depending on the relationship between the contingent payoff in the case of a high-severity loss and in the case of a low-severity loss. Although there are other equilibria, I shall concentrate on this one because it is the most interesting and the most plausible.¹²

¹² The other two possibilities occur when $\beta_L = \beta_H$ or when $\beta_H < \beta_L$. The flat payment is not interesting because there will be no fraud in the economy, although it might be preferable to the case where $\beta_H > \beta_L$ if the cost of auditing is high; and the case where $\beta_H < \beta_L$, is always dominated by the case where $\beta_H = \beta_L$. A brief discussion of these other two cases is found in "The Other Two Cases" section.

Given the players' optimal strategies, it is possible to find the price that gives zero expected profits to the principal. This price is implicitly given by

$$\alpha = \underbrace{\pi[\rho\beta_H + (1 - \rho)\beta_L]}_{\text{Benefit when truth}} + \underbrace{\pi(1 - \rho)\eta(1 - \nu)(\beta_H - \beta_L)}_{\text{Rent}} + \underbrace{\pi c \nu[\rho + (1 - \rho)\eta]}_{\text{Cost of audit}} + \underbrace{\pi(1 - \rho)\eta \nu k_2}_{\text{Monetary penalty}}. \quad (3)$$

On the right, $\pi[\rho\beta_H + (1 - \rho)\beta_L]$ represents the expected benefits paid when the agent tells the truth. The second term, $\pi(1 - \rho)\eta(1 - \nu)(\beta_H - \beta_L)$, is the rent an agent can expect to extract from the principal. The third term, $\pi c \nu[\rho + (1 - \rho)\eta]$, is the expected cost of the auditing strategy. Finally, $\pi(1 - \rho)\eta \nu k_2$, is the expected amount collected by the principal when the agent is caught committing fraud.¹³ By substituting in (3) for the equilibrium value of η given in (1), the price of the contract becomes

$$\alpha = \pi[\rho\beta_H + (1 - \rho)\beta_L] + \pi \rho c \left(\frac{\beta_H - \beta_L}{\beta_H - \beta_L + k_2 - c} \right). \quad (4)$$

This zero-profit constraint has the characteristic that the principal's auditing strategy does not directly affect the price of the contract. Therefore, the price function is not related to the agent's utility function, the agent's initial wealth (Y), or the deadweight penalty (k_1). The reason is that Y and k_1 only appear in the price of the contract by way of the audit probability as we can see in (2) and (3). But since the audit probability washes out once we substitute for η , all dependence of α on Y and k_1 is discarded.

The term $\pi \rho c \left(\frac{\beta_H - \beta_L}{\beta_H - \beta_L + k_2 - c} \right)$ represents the implicit loading factor of the insurance contract.¹⁴ This loading increases in the probability of accident, the probability of a high severity loss and in the cost of auditing. It decreases, however, in the difference in the indemnity payments and in the penalty paid to the principal, provided that $k_2 < c$, which we assume to be true. Because of this *ex post* moral hazard induced loading, it will not be possible to have a full insurance contract. The only way in which the implicit loading would disappear in this model is if the agent never committed fraud ($\eta = 0$).

THE OPTIMAL CONTRACT

The problem for the principal is to choose a coverage pair and a premium that maximize the agent's expected utility over final wealth. The principal anticipates rationally the agent's optimal strategies in the event of an accident. The principal must include those strategies in her contract design.

¹³ It is interesting to note that if $\eta = 0$ and $\nu = 0$, then the price is given only by the first term of (3). Thus if the agent never lies and the principal never audits, then the price of the contract collapses to the actuarially fair price of the contract. In such a setting, the agent would choose coverage equal to his loss in each state, i.e., $\beta_L = \lambda_L$ and $\beta_H = \lambda_H$.

¹⁴ By loading factor we mean the amount paid above and beyond the expected loss, $\pi[\rho\beta_H + (1 - \rho)\beta_L]$.

Results

The problem faced by the principal is

$$\begin{aligned}
 \max_{\alpha, \beta_L, \beta_H, v, \eta} & (1 - \pi)U(Y - \alpha) + \pi\rho U(Y - \alpha - \lambda_H + \beta_H) \\
 & + \pi(1 - \rho)(1 - \eta)U(Y - \alpha - \lambda_L + \beta_L) \\
 & + \pi(1 - \rho)\eta v[U(Y - \alpha - \lambda_L + \beta_L - k_2) - k_1] \\
 & + \pi(1 - \rho)\eta(1 - v)U(Y - \alpha - \lambda_L + \beta_H), \quad (\text{MP})
 \end{aligned}$$

subject to the break-even constraint, (4), the equilibrium probability of filing a fraudulent claim, (1), and the probability of auditing a high-loss claim, (2). The participation constraint can be ignored because it is not binding at the optimum. By substituting (1), (2), and (4) into (MP), the maximization problem simplifies to

$$\begin{aligned}
 \max_{\beta_L, \beta_H} V = & (1 - \pi)U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L - c + k_2}\right]\right) \\
 & + \pi\rho U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L - c + k_2}\right] - \lambda_H + \beta_H\right) \\
 & + \pi(1 - \rho)U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L - c + k_2}\right] - \lambda_L + \beta_L\right). \quad (\text{SP})
 \end{aligned}$$

It is interesting to note that the deadweight penalty (k_1) is never a factor in the design of the optimal contract as it is nowhere to be found in (SP). Therefore it cannot be a factor when the optimal contract is chosen. A related point is that it will not matter whether the deadweight penalty imposed on the agent caught committing fraud is a decision variable in the model. As long as the decision to audit comes last in the sequence of play, the deadweight penalty parameter will disappear from the simplified problem.¹⁵

This is not true, however, for the penalty paid to the principal. As we can see, the penalty paid to the principal by the agent in case he is found to have committed fraud reduces the insurance premium paid. This penalty k_2 enters the maximization problem of the agent through the Nash equilibrium strategies and the zero-profit constraint. Since the Nash equilibrium constraints disappear, the agent is left with considering the impact of the penalty k_2 only through the premium he must pay. It is clear that the greater the penalty, the smaller the premium the agent must pay.

Another interesting point to mention is that the audit probability v as well as the fraud probability η drop off completely from the optimization problem. In the first case, the reason is that the principal is indifferent between auditing and not auditing, which means that changes in the audit probability do not affect the zero-profit constraint.

¹⁵ There is no loss in generality here to use a fixed deadweight penalty rather than a deadweight penalty that would be a function of the reported loss, or the difference between the reported loss and the real loss because there is only one possible report that may be fraudulent. In other words, letting the deadweight penalty equal $k_1(\lambda_H)$, $k_1(\lambda_H, \lambda_L)$, or k_1 has no impact on the contracting problem.

In the second case, it is the agent who is indifferent between committing fraud and telling the truth.

I am then able to show the following lemma.

Lemma 2: *Letting*

$$\alpha'_H = \frac{\partial \alpha}{\partial \beta_H} = \pi \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L - 2c + 2k_2) + (k_2 - c)k_2}{(\beta_H - \beta_L - c + k_2)^2}, \quad (5)$$

$$\alpha'_L = \frac{\partial \alpha}{\partial \beta_L} = \pi \left[1 - \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L - 2c + 2k_2) + (k_2 - c)k_2}{(\beta_H - \beta_L - c + k_2)^2} \right], \quad (6)$$

and

$$V' = (1 - \pi)U'(Y - \alpha) + \pi \rho U'(Y - \alpha - \lambda_H + \beta_H) + \pi(1 - \rho)U'(Y - \alpha - \lambda_L + \beta_L), \quad (7)$$

where V' is defined as the expected marginal utility with respect to wealth, the optimal (β_L^*, β_H^*) solves

$$\pi \rho \frac{U'(Y - \alpha - \lambda_H + \beta_H)}{V'} = \alpha'_H \quad (\text{NC1})$$

$$\pi(1 - \rho) \frac{U'(Y - \alpha - \lambda_L + \beta_L)}{V'} = \alpha'_L = \pi - \alpha'_H \quad (\text{NC2})$$

It is clear that the left-hand sides of (NC1) and (NC2) are positive. Therefore, α'_H and α'_L must also be positive. It then becomes clear that the participation constraint does not bind because it would be equivalent to choosing $\beta_H = \beta_L = 0$, which cannot be optimal because of risk aversion. What do these two necessary conditions imply for the economy? With some manipulation, I can prove my first theorem.

Theorem 1: *At the optimum, the agent's expected marginal utility with respect to wealth is the same in the accident state as in the no-accident state.*

The statement of Theorem 1 is not surprising; this is a classic result in the literature (see Winter, 1992; Bond and Crocker, 1997). The theorem only tells us that the marginal utilities across accident states are the same; it says nothing about the marginal utility in a given loss state conditional on an accident occurring. More specifically, it does not tell us whether there is perfect income smoothing (marginal utilities equal in every state), nor whether the agent's marginal utility in the high-loss state is greater (or smaller) than in the low-loss state. The reason this result is obtained is that there is no information asymmetry between the principal and the agent regarding whether there was an accident or not. We know from classical results that under perfect information an agent will choose to be fully insured. In my case, there is perfect information concerning whether an accident occurred or not. As in the case of perfect information, the optimal contract is such that full insurance (in the sense of equal marginal utilities) is provided. The same occurs here with the exception that it is the *expected* marginal utilities that are equal rather than the marginal utilities *per se*.

Solely on the basis of Theorem 1, it seems possible to envision a contract where perfect income smoothing is obtained ($\beta_H = \lambda_H$ and $\beta_L = \lambda_L$). Such a contract is optimal only when $k_2 = c$. When $k_2 < c$ as I assumed, the optimal coverage for the agent will be such that coverage is better (i.e., the agent's out-of-pocket amount is smaller) in the event of a high-severity accident than in the event of a low-severity accident. This is proven as Corollary 1.

Corollary 1: *The optimal contract is such that $\beta_H - \lambda_H > \beta_L - \lambda_L$ if and only if $k_2 < c$. In other words, the agent does not smooth his income perfectly.*

This corollary not only shows that the optimal contract does not give the agent the same final wealth in every state of the world, it also tells us that the agent gets greater utility in the high-loss state than in the low-loss state. This means that the agent is better covered if he suffered a high-loss accident, which has the strange implication that the agent would rather be involved in a high-severity accident than in a low-severity one. How can this be explained?

The reason why larger losses are better covered is that they give the insurer greater incentive to audit, and thus reduce the agent's probability of lying. We know from the construction of mixed equilibria that one player's strategy must take into account only the parameters the other player cares about. Thus the agent must consider what effect a greater difference between β_H and β_L has on the principal's payoff. It is straightforward to see that the greater it is, the more the principal has to lose by not auditing. Knowing this, the agent must reduce his probability of sending a false signal.

The question then becomes whether the compensation in the high-loss state may be so large that the agent prefers to be in it rather than in the no-accident state. In other words, is the agent underinsured? Or does he choose coverage that is greater than his loss? This question is answered in the following theorem, which is the most striking result of the article.

Theorem 2: *The agent's higher loss is overinsured while his lower loss is underinsured. In other words, $\beta_H > \lambda_H$ and $\beta_L < \lambda_L$.*

In the proof of the theorem, it appears that neither penalty (k_1 or k_2) enters the decision to overcompensate the agent for a high loss or to undercompensate him for a low loss. Theorem 2 shows that overcompensation of the high loss always occurs, no matter what the penalty is (provided, of course, that $c > k_2$). This means that as much as the principal can use auditing as an extracting device, she will still need to offer a contract that provides insurance to the agent across accident states. Relating to Picard (1996), I would obtain the same results if $\pi = 1$, that is, if the no-accident state never occurs. Incidentally, it is interesting to note that the probability that an accident occurs (π) has no impact on the decision to overcompensate the high loss and to undercompensate the low loss.

The overcompensation of large losses and the undercompensation of small losses is due to the fact that the principal cannot commit. In Boyer (2003), it is shown that overcompensation of large losses is optimal even when the accident state is not observable. When the no-accident state never occurs (say $\pi = 1$), Boyer (2001) shows that there would still be overcompensation of the larger losses.

Commitment

In this section, the commitment concern is addressed. What I want to do is compare the welfare of the agent when the principal can commit to an auditing strategy to his welfare when the principal cannot commit.

To analyze the problem under commitment, a new decision variable is needed: the principal's auditing strategy, v . There is no Nash equilibrium constraint, but there is an incentive constraint that stipulates that the agent is (weakly) better off telling the truth. It is interesting to see that the no-commitment case is much easier to analyze than the commitment case. Without commitment the simplified maximization problem (SP) was unconstrained with two variables. In the commitment case, there are four variables and two constraints. This adds two Lagrange multipliers, which means that there are six first-order conditions instead of only two. Theorem 3 follows.

Theorem 3: *If the principal can commit to an auditing strategy, then the optimal contract is such that $\beta_H < \lambda_H$ and $\beta_L > \lambda_L$.*

If the principal can commit to an auditing strategy, then the insurance contract is such that the agent is overcompensated for his loss in the low-loss state, and undercompensated for his loss in the high-loss state. This result is the complete opposite of the one under no-commitment. The reason is that the principal does not need to signal *ex ante* an *ex post* auditing strategy if she can fully commit credibly *ex ante*. Without commitment, the principal's signal was achieved by overcompensation of the high loss and undercompensation of the low loss. If, on the other hand, the principal can commit to an auditing strategy *ex ante*, then she will design a contract that minimizes her audit costs. The principal does this by reducing her probability of auditing, which is achieved by increasing the payment in the case of a low-severity accident and reducing the payment in the case of a high-severity accident. It is then clear that the premium paid by the agent is greater under no-commitment, and that there is a welfare loss if the principal is unable to commit. This is shown in the following corollary.

Corollary 2: *There is a welfare loss from being unable to commit to an auditing strategy.*

Corollary 2 shows that the agent's utility is in fact strictly greater if the principal can commit than it will be if she cannot commit. The fact that the problem under commitment Pareto dominates the problem under no-commitment can also be achieved by noting that the problem under commitment is nothing but the original problem (MP), but with one less constraint: the Nash equilibrium strategy of the policyholder. It must therefore be true that the equilibrium allocation under commitment dominates the equilibrium allocation under no-commitment.

From Theorem 2, we know that the optimal contract entails overpayment of the higher loss and underpayment of the lower loss, which means that the marginal utility of wealth is greatest in the low-loss state. The question is then to compare the expected wealth in the accident state with the wealth in the no-accident state. The optimal contract is such that the agent agrees to receive lower expected wealth when there is no accident than when there is an accident. That is the same as saying that the expected benefit is greater than the expected loss. As Theorem 2 shows, the agent bears more risk in the accident state because he does not know *a priori* what will be his

final wealth when he is involved in an accident. It then becomes logical to see that the agent's expected wealth is greater in the accident state than in the no-accident state.

Overcompensation

It seems odd that the optimal contract in the presence of *ex post* moral hazard entails overpayment of the higher loss. If fraud is so common in the economy, why is it that we do not systematically observe contracts similar to the one we found? Two reasons come to mind: insurer expenses and *ex ante* moral hazard.

Suppose that the insurer's expenses, such as underwriting, marketing, and taxes, can be modeled as a proportional loading factor τ on the premium. In accordance with standard results in the literature, proportional premium loading reduces insurance coverage in the sense that the agent receives greater expected marginal utility in the accident state than in the no-accident state. It is then possible to show that β_H decreases as τ increases, as in Boyer (2000). It is left to the reader to show that there exists a $\hat{\tau} > 0$ so that $\beta_H = \lambda_H$ and $\beta_L < \lambda_L$. Therefore, for any proportional loadings τ greater than $\hat{\tau}$ overinsurance should not be observed.

The other reason that overinsurance is not observed more often in reality is the presence of *ex ante* moral hazard. Although I will not address this concern here, *ex ante* moral hazard has important implications given that the agent is overcompensated in the event of a high loss, which may induce him to increase the probability that he will find himself in that state of the world.

The Other Two Cases

I have concentrated in this article on the case that was the most interesting and the most plausible. The other two cases, $\beta_L = \beta_H$ and $\beta_L > \beta_H$, were put aside to make the article more readable. This section of the article addresses those two cases briefly and presents the optimal contracts in the last theorem of the article.

Theorem 4: *A contract where $\beta_L = \beta_H = \beta^*$ gives greater utility to the agent than a contract where $\beta_H > \beta_L$ if the cost of auditing is larger than some $\tilde{c}$ that solves*

$$U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - \tilde{c}} \right] \right) = U'(Y - \pi\beta^*). \quad (8)$$

Furthermore, the contract where $\beta_L = \beta_H$ always gives greater utility to the agent than a contract where $\beta_L > \beta_H$.

I see two reasons that a flat payment may be preferable to a variable payment where the higher loss is overcompensated. First, when audit costs are high, the amount of money wasted through auditing increases because the greater the cost of auditing, the greater the probability of committing fraud. This invariably increases the number of audits in the economy (if the probability of auditing remains constant), which increases the total amount of money wasted. It is therefore logical to expect that there is a limit to the cost of auditing so that a variable payment is preferable to the fixed payment. The second reason has to do with the insurance concept of the contract. If the cost of auditing increases, the distance between β_H and β_L increases. The difference

may become so large that the payment schedule no longer gives any insurance to the agent; it becomes a lottery.

Under the fixed indemnity payment scheme the implicit loading factor presented in Equation (4) goes to zero. The reason is that there is no longer any reason for the agent to commit fraud because he has nothing to gain; whether he reports a high loss or a low loss he receives the same indemnity payment. When the insurance contract offers a fixed indemnity payment, the agent gains by not having to incur the implicit loading imbedded in the premium to control *ex post* moral hazard.

Another interesting aspect of this theorem is that it allows us to find under which conditions on c is the probability of fraud the highest. In fact, the last result of the article shows that $\tilde{c}$ found in Theorem 4 is not only the highest cost such that insurance fraud still occurs, but also the cost associated with the greatest probability of fraud.

Corollary 3: *The highest probability with which an agent will commit fraud is given by*

$$\tilde{\eta} = \left(\frac{\tilde{c}}{\beta_H - \beta_L + k_2 - \tilde{c}} \right) \left(\frac{\rho}{1 - \rho} \right), \quad (9)$$

which is always strictly smaller than 1.

Because agents are better off with a fixed payment contract in case of an accident than a contract that indemnified agents differently based upon the loss when the cost of auditing is large, the condition under which the mixed-strategy equilibrium is sustained always holds. This gives us the interesting result that insurance fraud increases as the cost of auditing increases (also, $\eta = 0$ when $c = 0$), but stops abruptly when $c = \tilde{c}$. For $c \geq \tilde{c}$, there is no fraud because agents gain nothing from it. I also have that the agent's expected utility decreases as the cost of auditing increases, up until $c = \tilde{c}$. For $c > \tilde{c}$ the agent's expected utility remains the same as when $c = \tilde{c}$.

DISCUSSION AND CONCLUSION

This article examines from the insurance standpoint what contract a principal, who cannot commit credibly to an auditing strategy, would offer an informed agent who is faced with *ex post* moral hazard. In a two-point distribution of losses conditional on the occurrence of an accident, the optimal contract entails overpayment of the larger loss and underpayment of the lower loss. Furthermore, the optimal contract is such that the agent receives the same expected marginal utility when an accident occurs as when no accidents occur, even though the agent's expected final wealth is greater in the accident state. Finally, there is a loss of welfare resulting from the principal's being unable to commit.

The contribution of this article is that it is the first where the optimal contract stipulates that the higher loss must be overcompensated, whereas the lower loss must be undercompensated (as opposed to Bond and Crocker (1997) or Crocker and Morgan (1998)). Also, overcompensation does not depend on an agents being audited (as in Mookherjee and Png, 1989). The reason that the players agree to this counterintuitive scheme is that it reduces the cost of fraud in the economy. The reduction in the cost of fraud does not come primarily from the reduction of false messages that are not detected, which are merely transfers from the principal to the agent. It comes from

the smaller number of messages that are audited. Since the agent is less likely to send a false message when there is overcompensation of large losses, the total number of high-loss claims is reduced. This means that the total audit cost may go down. And since the cost of auditing is a deadweight cost to the economy, a reduction in this deadweight can only make the agent better off. Another gain from overcompensation is that agents are less likely to pay the deadweight penalty. Because the agent sends fewer false messages, he is less likely to be caught having sent a false message. This means that less money is wasted in deadweight penalties.

Fraud is reduced when the difference between the high-loss and the low-loss indemnity payment is increased because it gives greater monetary incentives to the insurer to make sure the claim is not fraudulent. Given that the insurer has more incentives to audit, the policyholder must reduce his likelihood of committing fraud to make the insurer indifferent again between auditing and not auditing.

To what extent can this contract be applied to the real world? Although there is little evidence that large losses are overcompensated,¹⁶ there are many instances where smaller losses are undercompensated; the preponderance of deductibles is a blatant example. With respect to overcompensation, anecdotal evidence suggests that it may occur. For example, in France some hospital patients are not only treated for whatever ails them, they also receive some monetary compensation when they are discharged from the hospital.¹⁷ Another example may be the case of policyholders who may be overcompensated as a result of a natural disaster loss.

Khalil (1997) treats a problem similar to the one considered here, but with some major differences. Khalil studies the case of an informed risk-neutral agent who must make a report to the risk-averse principal. The basic framework he uses is that of Baron and Myerson (1982), i.e., the information the agent has concerns his production cost. Khalil finds that the agent overproduces in comparison with the case where there is full information. This is similar to my result of overinsurance in the high-loss state. Surprisingly Khalil obtains these results using a completely different set of assumptions. The major differences between the Khalil article and the present one are as follows: (1) Khalil has the penalty paid by the agent to the principal (as in Picard, 1996), so that auditing is used by the principal to extract rents from the agent; (2) the agent is risk-neutral and the principal is risk-averse; and (3) there is no state of the world where the game is not played.

It is interesting to note in Khalil (1997), Boyer (2000), and Khalil and Parigi (1998), and in this article as well that the traditional predictions of the contract literature where the principal can perfectly commit are reversed. Khalil (1997) finds that the agent overproduces when the principal cannot commit, whereas Baron and Myerson (1982) find the opposite. The same applies to Khalil and Parigi (1998) who find that the entrepreneur should overinvest (or overborrow) when the financier cannot commit to an auditing strategy, compared with underinvestment in Gale and Hellwig (1985). Here, I obtain that the agent is overinsured in the high-loss state when the principal

¹⁶ The lack of evidence may be due in part to some of the assumptions used in the article, such as the assumption that the loading factor is zero. The discussion in "Overcompensation" section addresses this concern.

¹⁷ I am indebted to Nathalie Fombaron for this anecdote.

cannot commit. This contradicts the general literature on commitment, which shows that higher losses should be undercompensated (see Winter, 1992).

Can these results be generalized? In other words, will noncommitment of the principal to an action always lead to a prediction that is the opposite of the prediction under full commitment: overinsurance vs. underinsurance, overproduction vs. underproduction, overinvestment vs. underinvestment, etc. Although there appears to be a pattern to that effect, no one has shown any direct relationship or the conditions for a direct relationship. Another avenue for generalization would be to look at T possible losses given the occurrence of an accident. In such a setting, Boyer (2003) indeed finds overcompensation of higher losses and equal expected marginal utilities across accident states.

APPENDIX: PROOFS

Proof of Lemma 1: From the left-hand side of Figure 2, it is clear that γ_L , the principal's posterior belief that the true loss is high given that the agent sent message $\lambda' = \lambda_L$, is zero. Suppose Nature chooses the severity of the accident to be high (λ_H). Sending message $\lambda' = \lambda_H$ then always dominates sending message $\lambda' = \lambda_L$, whatever the principal does. Also, if the principal hears message $\lambda' = \lambda_L$, then she knows for sure that she is not playing at the upper node of left information set. Consequently, the only meaningful strategy for the principal is never to audit because she collects $-c - \beta_L$ if she does, and $-\beta_L$ if she does not. Let's now move to the right-hand side of Figure 2.

Let η be the probability (in the mixed-strategy sense) that the agent sends message $\lambda' = \lambda_H$ when Nature chooses $\lambda = \lambda_L$. By Bayes' rule we can find γ_H , the principal's posterior belief that the true loss is high given that the agent sent message $\lambda' = \lambda_H$:

$$\gamma_H = \frac{\rho}{\rho + (1 - \rho)\eta}. \quad (A1)$$

Only one strategy on the part of the agent makes the principal indifferent as to whether to audit or not. That strategy must be such that

$$\gamma_H = \frac{(\beta_H - \beta_L) + k_2 - c}{(\beta_H - \beta_L) + k_2}.$$

Substituting γ_H in (A1) yields

$$\eta = \left(\frac{c}{\beta_H - \beta_L + k_2 - c} \right) \left(\frac{\rho}{1 - \rho} \right).$$

All that is left to calculate is the principal's strategy at the information set associated with the high loss. Her strategy must be such that the agent is indifferent between telling the truth and lying, given that $\lambda = \lambda_L$. Let ν be the probability (in a mixed-strategy sense) of auditing a λ'_H message. ν must then solve

$$U(Y - \alpha - \lambda_L + \beta_L) = \nu[U(Y - \alpha - \lambda_L + \beta_L - k_2) - k_1] + (1 - \nu)U(Y - \alpha - \lambda_L + \beta_H),$$

which means that

$$\nu = \frac{U(Y - \alpha - \lambda_L + \beta_H) - U(Y - \alpha - \lambda_L + \beta_L)}{U(Y - \alpha - \lambda_L + \beta_H) - U(Y - \alpha - \lambda_L + \beta_L - k_2) + k_1}.$$

Since all six elements of the PBNE have been found, the proof is done. Q.E.D.

Proof of Lemma 2: The first-order conditions of (SP) are

$$\begin{aligned} \frac{\partial V}{\partial \beta_H} = 0 = & -\alpha'_H(1 - \pi)U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right]\right) \\ & + (1 - \alpha'_H)\pi\rho U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right] - \lambda_H + \beta_H\right) \\ & - \alpha'_H\pi(1 - \rho)U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right] - \lambda_L + \beta_L\right) \end{aligned} \quad (\text{FOC}_H)$$

and

$$\begin{aligned} \frac{\partial V}{\partial \beta_L} = 0 = & -\alpha'_L(1 - \pi)U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right]\right) \\ & - \alpha'_L\pi\rho U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right] - \lambda_H + \beta_H\right) \\ & + (1 - \alpha'_L)\pi(1 - \rho)U'\left(Y - \pi \left[\beta_L + \rho \frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c} \right] - \lambda_L + \beta_L\right). \end{aligned} \quad (\text{FOC}_L)$$

From (FOC_H), I obtain

$$\begin{aligned} \pi\rho U'(Y - \alpha - \lambda_H + \beta_H) = & \alpha'_H[(1 - \pi)U'(Y - \alpha) + \pi\rho U'(Y - \alpha - \lambda_H + \beta_H) \\ & + \pi(1 - \rho)U'(Y - \alpha - \lambda_L + \beta_L)]. \end{aligned}$$

Substituting and dividing everywhere by V' yields the desired result. The same technique is used with (FOC_L), except that I find that

$$\begin{aligned} \pi(1 - \rho)U'(Y - \alpha - \lambda_L + \beta_L) = & \alpha'_L[(1 - \pi)U'(Y - \alpha) + \pi\rho U'(Y - \alpha - \lambda_H + \beta_H) \\ & + \pi(1 - \rho)U'(Y - \alpha - \lambda_L + \beta_L)]. \end{aligned}$$

Substituting and dividing everywhere by V' yields the desired result. Q.E.D.

Proof of Theorem 1: By substituting (NC1) into (NC2),

$$\pi (1 - \rho) \frac{U'(Y - \alpha - \lambda_L + \beta_L)}{V'} = \pi - \pi \rho \frac{U'(Y - \alpha - \lambda_H + \beta_H)}{V'}.$$

Expanding V' , regrouping terms and simplifying yields

$$(1 - \rho) U'(Y - \alpha - \lambda_L + \beta_L) + \rho U'(Y - \alpha - \lambda_H + \beta_H) = U'(Y - \alpha).$$

The left-hand side is the expected marginal utility in the accident state, while the right-hand side is the marginal utility in the no-accident state. Q.E.D.

Proof of Corollary 1: From (NC1) and (5), we find that $\frac{\alpha'_H}{\pi \rho} < 1$ if and only if $k_2 < c$, which we assumed to be true. Thus

$$\frac{U'(Y - \alpha - \lambda_H + \beta_H)}{V'} < 1. \quad (A2)$$

Similarly, from (NC2) and (A2) we find that $\frac{\pi - \alpha'_H}{\pi(1 - \rho)} > 1$ if and only if $k_2 < c$. Thus

$$\frac{U'(Y - \alpha - \lambda_L + \beta_L)}{V'} > 1. \quad (A3)$$

Combining (A2) and (A3), we find that $\beta_H - \lambda_H > \beta_L - \lambda_L$. Q.E.D.

Proof of Theorem 2: From Theorem 1, we know that expected marginal utilities are equalized across accident states. Thus

$$\rho = \frac{U'(Y - \alpha) - U'(Y - \alpha - \lambda_L + \beta_L)}{U'(Y - \alpha - \lambda_H + \beta_H) - U'(Y - \alpha - \lambda_L + \beta_L)}.$$

It is clear that the denominator on the right is negative from Corollary 1. Thus the numerator must also be negative. This occurs if and only if $\beta_L < \lambda_L$. Similarly, from Theorem 1, we have

$$1 - \rho = \frac{U'(Y - \alpha) - U'(Y - \alpha - \lambda_H + \beta_H)}{U'(Y - \alpha - \lambda_L + \beta_L) - U'(Y - \alpha - \lambda_H + \beta_H)}.$$

Here, it is clear that the denominator on the right is positive. Thus the numerator must also be positive. This occurs if and only if $\beta_H > \lambda_H$. Q.E.D.

Proof of Theorem 3: See Mookherjee and Png (1989). Q.E.D.

Proof of Corollary 2: The revelation principal (see Myerson, 1979) states that amongst all feasible optimal contracts, at least one induces the agent to tell the truth. Therefore, the contract under no-commitment cannot be better than the contract under commitment. Q.E.D.

Proof of Theorem 4: The proof has three parts. First, I show for some value of the auditing cost that the agent is better off with the variable payment. I then show for some other value that the agent is better off with the fixed payment. I complete the proof by showing that the utility function is continuous in the cost of auditing, which means, from the mean value theorem, that there is a critical value where the agent's preferences change from one system to the next.

Part 1: I want to show that there exists a c such that

$$\begin{aligned} & (1 - \pi)U(Y - \pi\beta^*) + \pi\rho U(Y + (1 - \pi)\beta^* - \lambda_H) + \pi(1 - \rho)U(Y + (1 - \pi)\beta^* - \lambda_L) \\ & < (1 - \pi)U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c}\right]\right) \\ & \quad + \pi\rho U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c}\right] - \lambda_H + \beta_H\right) \\ & \quad + \pi(1 - \rho)U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c}\right] - \lambda_L + \beta_L\right). \end{aligned}$$

It is clear that the inequality holds if c is very small; specifically,

$$\lim_{c \rightarrow 0} \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c}\right] = \pi[(1 - \rho)\beta_L + \rho\beta_H],$$

which is obviously the premium paid under full information. Clearly, the agent is then better off with a variable compensation because the optimal contract entails full insurance ($\beta_H = \lambda_H$ and $\beta_L = \lambda_L$).

Part 2: I want to show that there exist a value of c such that the agent prefers a fixed payment to a variable payment.

Instead of going directly through the expected utility, I will use marginal utilities. If the expected utility when $\beta_L = \beta_H$ is greater than the expected utility when $\beta_L < \beta_H$, then the expected marginal utility with $\beta_L = \beta_H$ is smaller than the expected marginal utility when $\beta_L < \beta_H$. Using Theorem 1 (i.e., the expected marginal utility is the same in the accident state as in the no-accident state), I want to find the conditions on c such that the following inequality holds:

$$U'(Y - \pi\beta^*) < U'\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - c}\right]\right).$$

Since $\lim_{c \rightarrow \infty} (\beta_H - \beta_L) = \frac{c}{1 - \rho}$, taking the limit on the right-hand side yields that the agent will prefer the fixed payment to the variable payment if and only if

$$U'(Y - \pi\beta^*) < U'(-\infty),$$

which obviously holds. Thus as $c \rightarrow \infty$, the expected utility of the fixed payment is greater than the expected utility of the variable payment.

Part 3: To conclude the proof, the expected utility must be shown to be continuous and monotone in c : that is, $\text{sign}(\frac{\partial EU(c)}{\partial c}) = \text{cst}$. Let

$$\Omega = (1 - \pi)U(Y - \alpha) + \pi\rho U(Y - \alpha - \lambda_H + \beta_H) + \pi(1 - \rho)U(Y - \alpha - \lambda_L + \beta_L).$$

The total derivative with respect to the cost of auditing is

$$\begin{aligned} \frac{d\Omega}{dc} = & -(1 - \pi)\frac{d\alpha}{dc}U'(Y - \alpha) + \pi\rho\left(\frac{\partial\beta_H}{\partial c} - \frac{d\alpha}{dc}\right)U'(Y - \alpha - \lambda_H + \beta_H) \\ & + \pi(1 - \rho)\left(\frac{\partial\beta_L}{\partial c} - \frac{d\alpha}{dc}\right)U'(Y - \alpha - \lambda_L + \beta_L). \end{aligned}$$

Expanding $\frac{d\alpha}{dc}$, letting

$$\Phi = \frac{\partial\alpha}{\partial c} + \frac{\partial\alpha}{\partial\beta_H}\frac{\partial\beta_H}{\partial c} + \frac{\partial\alpha}{\partial\beta_L}\frac{\partial\beta_L}{\partial c},$$

and combining terms yield

$$\begin{aligned} \frac{d\Omega}{dc} = & -\Phi(1 - \pi)U'(Y - \alpha) \\ & - \Phi\pi[\rho U'(Y - \alpha - \lambda_H + \beta_H) + (1 - \rho)U'(Y - \alpha - \lambda_L + \beta_L)] \\ & + \pi\rho\frac{\partial\beta_H}{\partial c}U'(Y - \alpha - \lambda_H + \beta_H) + \pi(1 - \rho)\frac{\partial\beta_L}{\partial c}U'(Y - \alpha - \lambda_L + \beta_L). \end{aligned}$$

Using Theorem 1, (NC1), and (NC2), the above equation simplifies to

$$\frac{d\Omega}{dc} = -\Phi U'(Y - \alpha) + \frac{\partial\beta_H}{\partial c}U'(Y - \alpha)\frac{\partial\alpha}{\partial\beta_H} + \frac{\partial\beta_L}{\partial c}U'(Y - \alpha)\frac{\partial\alpha}{\partial\beta_L}.$$

Substituting back for Φ and using the fact that $\frac{\partial\alpha}{\partial\beta_L} = \pi - \frac{\partial\alpha}{\partial\beta_H}$ yields

$$\frac{d\Omega}{dc} = U'(Y - \alpha)\left[-\frac{\partial\alpha}{\partial c} + \frac{\partial\alpha}{\partial\beta_H}\left(\frac{\partial\beta_L}{\partial c} - \frac{\partial\beta_H}{\partial c}\right) - \pi\frac{\partial\beta_L}{\partial c} + \left(\frac{\partial\beta_H}{\partial c} - \frac{\partial\beta_L}{\partial c}\right)\frac{\partial\alpha}{\partial\beta_H} + \pi\frac{\partial\beta_L}{\partial c}\right].$$

Simplifying finally yields

$$\frac{d\Omega}{dc} = -U'(Y - \alpha)\frac{\partial\alpha}{\partial c}.$$

Since $\frac{\partial\alpha}{\partial c}$ is positive for all c , the agent's expected utility is monotone and decreasing for all c . There must therefore by a $\tilde{c}$ such that

$$U\left(Y - \pi\left[\beta_L + \rho\frac{(\beta_H - \beta_L)(\beta_H - \beta_L + k_2)}{\beta_H - \beta_L + k_2 - \tilde{c}}\right]\right) = U(Y - \pi\beta^*).$$

This completes the proof.

Q.E.D.

Proof of Corollary 3: I want to show that the necessary and sufficient for an agent's probability of committing fraud is never binding because it becomes optimal to offer a nondifferentiated contract before that. In other words, I want to show that the necessary condition for $\eta \in [0, 1]$, $c < \bar{c} = (1 - \rho)(\beta_H - \beta_L + k_2)$, is never binding because the cost of auditing after which it is better to offer a fixed payment contract (i.e., $\bar{c}$), no matter the loss, is such that $\tilde{c} < \bar{c}$.

Using Equation (8), a contract pays different amounts to agents who suffered loss λ_H and λ_L if and only if

$$c < \bar{c} = (\beta_H - \beta_L + k_2) \left(\rho \frac{\beta^* - \beta_H}{\beta^* - \beta_L} + 1 - \rho \right).$$

I want to show that $\tilde{c} < \bar{c}$ so that

$$(\beta_H - \beta_L + k_2) \left(\rho \frac{\beta^* - \beta_H}{\beta^* - \beta_L} + 1 - \rho \right) < (1 - \rho)(\beta_H - \beta_L + k_2).$$

This occurs if and only if

$$\rho \frac{\beta^* - \beta_H}{\beta^* - \beta_L} < 0,$$

which is always true because $\beta_L < \beta^* < \beta_H$.¹⁸ Since $\tilde{c} < \bar{c}$, it follows that the highest probability of fraud for an agent in this economy is

$$\tilde{\eta} = \left(\frac{\tilde{c}}{\beta_H - \beta_L + k_2 - \tilde{c}} \right) \left(\frac{\rho}{1 - \rho} \right) < 1.$$

This completes the last proof of the article.

Q.E.D.

REFERENCES

- Baron, D. P., and R. B. Myerson, 1982, Regulating a Monopolist With Unknown Costs, *Econometrica*, 50: 911-930.
- Beaudry, P., and M. Poitevin, 1993, Signalling and Renegotiation in Contractual Relationships, *Econometrica*, 61: 745-782.
- Beaudry, P., and M. Poitevin, 1994, The Commitment Value of Contracts under Dynamic Renegotiation, *Rand Journal of Economics*, 25: 501-517.
- Bond, E. W., and K. J. Crocker, 1997, Hardball and the Soft Touch: The Economics of Optimal Insurance Contracts With Costly State Verification and Endogenous Monitoring, *Journal of Public Economics*, 63: 239-254.

¹⁸ To see why $\beta_L < \beta^* < \beta_H$, recall that when $c = 0$, $\beta_L = \lambda_L$ and $\beta_H = \lambda_H$. Since $\lambda_L < \lambda_H$, it follows that $\beta_L < \beta_H$. In such a case, it is easily proven that $\beta_L < \beta^* < \beta_H$. When c increases, the difference between β_H and β_L must also increase because $\beta_H - \beta_L > \frac{c}{1-\rho} - k_2$ (see Equation (1)).

- Boyer, M. M., 2000, Insurance Taxation and Insurance Fraud, *Journal of Public Economic Theory*, 2: 101-134.
- Boyer, M. M., 2001, Les clauses de valeur-à-neuf sont-elles optimales?, *L'Actualité économique*, 77: 53-74.
- Boyer, M. M., 2003, Contracting Under Ex Post Moral Hazard, Costly Auditing and Principal Non-Commitment, *Review of Economic Design*, 8: 1-38.
- Crocker, K. J., and J. Morgan, 1998, Is Honesty the Best Policy? Curtailing Insurance Fraud Through Optimal Incentive Contracts, *Journal of Political Economy*, 106: 355-375.
- Crocker, K. J., and A. Snow, 1985, The Efficiency Effect of Categorical Discrimination in the Insurance Industry, *Journal of Political Economy*, 94: 321-344.
- Cooper, R., and B. Hayes, 1987, Multi-Period Insurance Policies, *International Journal of Industrial Organization*, 5: 211-231.
- Dionne, G., and N. A. Doherty, 1994, Adverse Selection, Commitment, and Renegotiation: Extension to and Evidence From the Insurance Markets, *Journal of Political Economy*, 102: 209-235.
- Fombaron, N., 1997, No-Commitment and Dynamic Contracts in Competitive Insurance Markets With Adverse Selection, Mimeo, THEMA-Université de Paris-X.
- Gale, D., and M. Hellwig, 1985, Incentive-Compatible Debt Contracts: The One-Period Problem, *Review of Economic Studies*, 52: 647-663.
- Gibbons, R., 1992, *Game Theory for Applied Economists* (Princeton University Press).
- Graetz, M. J., J. F. Reinganum, and L. L. Wilde, 1986, The Tax-Compliance Game: Toward an Interactive Theory of Law Enforcement, *Journal of Law, Economics and Organization*, 2: 1-32.
- Hosios, A. J., and M. Peters, 1989, Repeated Insurance Contracts With Adverse Selection and Limited Commitment, *Quarterly Journal of Economics*, 104: 229-253.
- Hoyt, R. E., 1989, The Effect of Insurance Fraud on the Economic System, *Journal of Insurance Regulation*, 8: 304-315.
- Jost, P.-J., 1996, On the Role of Commitment in a Principal-Agent Relationship With an Informed Principal, *Journal of Economic Theory*, 68: 510-530.
- Karpoff, J. M., and J. R. Lott, Jr., 1993, The Reputational Penalty Firms Bear From Committing Criminal Fraud, *Journal of Law and Economics*, 36: 757-802.
- Khalil, F., 1997, Auditing Without Commitment, *Rand Journal of Economics*, 28: 629-640.
- Khalil, F., and B. M. Parigi, 1998, Loan Size as a Commitment Device, *International Economic Review*, 39: 135-150.
- Lacker, J. M., and J. A. Weinberg, 1989, Optimal Contracts Under Costly State Falsification, *Journal of Political Economy*, 97: 1345-1363.
- Laffont, J.-J., and J. Tirole, 1993, *A Theory of Incentives in Procurement and Regulation* (Cambridge, MA: MIT Press).
- Lewis, T. R., 1996, Protecting the Environment When Costs and Benefits Are Privately Known, *Rand Journal of Economics*, 27: 819-847.

- Lewis, T. R., and D. E. M. Sappington, 1995, Using Markets to Allocate Pollution Permits and Other Scarce Resource Rights Under Limited Information, *Journal of Public Economics*, 57: 431-455.
- Melumad, N. D., and D. Mookherjee, 1989, Delegation as Commitment: The Case of Income Tax Audits, *Rand Journal of Economics*, 20: 139-163.
- Mookherjee, D., and I. Png, 1989, Optimal Auditing, Insurance and Redistribution, *Quarterly Journal of Economics*, 104: 205-228.
- Myerson, R. B., 1979, Incentive Compatibility and the Bargaining Problem, *Econometrica*, 47: 61-73.
- Myerson, R. B., 1991, *Game Theory* (Cambridge, MA: Harvard University Press).
- Persons, J. C., 1997, Liars Never Prosper? How Management Misrepresentation Reduces Monitoring Costs, *Journal of Financial Intermediation*, 6: 269-306.
- Picard, P., 1996, Auditing Claims in the Insurance Market With Fraud: The Credibility Issue, *Journal of Public Economics*, 63: 27-56.
- Picard, P., 1997, On the Design of Optimal Insurance Policies Under Manipulation of Audit Cost, Mimeo, Université de Paris-X.
- Research Brief*, May 1995, Rand Corporation Institute for Civil Justice.
- Reinganum, J. F., and L. L. Wilde, 1985, Income Tax Compliance in a Principal-Agent Framework, *Journal of Public Economics*, 26: 1-18.
- Scheepens, J., 1995, Bankruptcy Litigation and Optimal Debt Contract, *European Journal of Political Economy*, 11: 535-556.
- Spence, A. M., and R. Zeckhauser, 1971, Insurance, Information, and Individual Action, *American Economic Review*, 61: 380-387.
- Swierzbinski, J. E., 1994, Guilty Until Proven Innocent—Regulation With Costly and Limited Enforcement, *Journal of Environmental Economics and Management*, 27: 127-146.
- Townsend, R. M., 1979, Optimal Contracts and Competitive Markets With Costly State Verification, *Journal of Economic Theory*, 21: 265-293.
- U.S. Sentencing Commission, 1988, Department of Justice Corporate Fraud Survey, Memorandum, U.S. Sentencing Commission. As reported in Karpoff, J. M. and J. R. Lott, Jr., 1993, The Reputational Penalty Firms Bear From Committing Criminal Fraud, *Journal of Law and Economics*, 36: 757-802.
- Waldfoegel, J., 1994, The Effect of Criminal Conviction on Income and the Trust "Reposed in the Workmen," *Journal of Human Resources*, 29: 62-81.
- Winter, R. A., 1992, Moral Hazard and Insurance Contracts, in: G. Dionne, ed., *Contributions to Insurance Economics* (Boston: Kluwer Academic Publishers).

Copyright of Journal of Risk & Insurance is the property of Blackwell Publishing Limited and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.