

Strausz, Roland

Working Paper

Timing of Verification Procedures: Monitoring versus Auditing

SFB/TR 15 Discussion Paper, No. 33

Provided in Cooperation with:

Free University of Berlin, Humboldt University of Berlin, University of Bonn, University of Mannheim, University of Munich, Collaborative Research Center Transregio 15: Governance and the Efficiency of Economic Systems

Suggested Citation: Strausz, Roland (2005) : Timing of Verification Procedures: Monitoring versus Auditing, SFB/TR 15 Discussion Paper, No. 33, Sonderforschungsbereich/Transregio 15 - Governance and the Efficiency of Economic Systems (GESY), München, <https://doi.org/10.5282/ubm/epub.13515>

This Version is available at:

<https://hdl.handle.net/10419/93845>

Standard-Nutzungsbedingungen:

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

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

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

Terms of use:

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

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

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



GOVERNANCE AND THE EFFICIENCY
OF ECONOMIC SYSTEMS
GESY

Discussion Paper No. 33

**Timing of Verification Procedures:
Monitoring versus Auditing**

Roland Strausz*

January 2005

*Roland Strausz, Free University Berlin, Department of Economics, Boltzmannstr. 20, D-14195 Berlin, strausz@fu-berlin.de

Financial support from the Deutsche Forschungsgemeinschaft through SFB/TR 15 is gratefully acknowledged.

Sonderforschungsbereich/Transregio 15 · www.gesy.uni-mannheim.de
Universität Mannheim · Freie Universität Berlin · Humboldt-Universität zu Berlin · Ludwig-Maximilians-Universität München
Rheinische Friedrich-Wilhelms-Universität Bonn · Zentrum für Europäische Wirtschaftsforschung Mannheim

Speaker: Prof. Konrad Stahl, Ph.D. · Department of Economics · University of Mannheim · D-68131 Mannheim,
Phone: +49(0621)1812786 · Fax: +49(0621)1812785

Timing of Verification Procedures: Monitoring versus Auditing

Roland Strausz*

January 3, 2005

Abstract

This paper studies the strategic effect of a difference in timing of verification in an agency model. A principal may choose between two equally efficient verification procedures: monitoring and auditing. Under auditing the principal receives additional information. Due to a double moral hazard problem, there exists a tension between incentives for effort and incentives for verification. Auditing exacerbates this tension and, consequently, requires steeper incentive schemes than monitoring. Hence, auditing is suboptimal if 1) steep incentives structures are costly to implement due to bounded transfers, or 2) steep incentive schemes induce higher rents due to limited liability.

Keywords: timing of verification, double moral hazard, monitoring, auditing

JEL Classification No.: D82

*Free University Berlin, Dept. of Economics, Boltzmannstr.20, D-14195 Berlin (Germany); email address: strausz@zedat.fu-berlin.de. I thank Jacques Lawarrée, Annamaria Menichini, two anonymous referees and the editor for helpful comments and suggestions. Financial support by the DFG (German Science Foundation) under SFB-TR 15 is gratefully acknowledged.

1 Introduction

One of the main internal problems of an organization is the existence of moral hazard. When an employee's effort or action cannot be observed, his remuneration cannot be linked to his actual decision, and room for moral hazard exists. As is well known, problems of moral hazard place a cost on the organization. Organizations will therefore have reasons to reduce the scope for moral hazard and obtain more accurate information through costly verification procedures (e.g. Townsend 1979). Many aspects of these procedures will lie in the hands of the organization itself. It must decide what kind of information it wants to acquire, when to acquire it, and how to use it. The purpose of this paper is to look into the "when" of information acquisition, its timing.¹

More specifically, in a standard agency setting this paper studies two alternative procedures of verification that I call monitoring and auditing. The difference between the two procedures is that monitoring takes place *while* the agent chooses his action whereas auditing occurs *after* he has taken his action. This difference in timing has a strategic consequence because after the agent's decision the principal receives supplementary information about the agent's actual behavior. Hence, with auditing the principal's decision to verify is taken on the basis of additional information that is not available under monitoring. It is this informational wedge that influences the principal's optimal verification procedure.

To focus on the effects that are due to this difference in timing I assume that monitoring and auditing concern the verification of identical variables and that both procedures are equally efficient.² This modelling allows for two different interpretations. First, the difference between monitoring and

¹The original work on moral hazard of Holmström (1979) and Shavell (1979) considered the question of how to use available information. Maskin and Riley (1985) and Khalil and Lawarrée (1995) address the first question of what kind of information the principal should gather.

²Moreover, I abstract from all other possible differences between auditing and monitoring, such as that under monitoring the agent may observe the outcome of verification and adjust his behavior accordingly, while under ex post auditing he cannot.

auditing may regard the exact point in time at which the decision to evaluate evidence is taken. For instance, the principal may decide to observe the agent with video cameras. Monitoring would then mean that the principal follows the agent’s behavior “live” on a video screen. Auditing, on the other hand, would mean that the principal collects the recordings and decides about reviewing them on the basis of some additional information (e.g., the success of the project the agent worked on) that comes available after the agent has completed his tasks. Second, the physical character of different procedures of verification may lead to a natural difference in timing. For instance, direct supervision of an agent necessarily implies monitoring whereas checking the agent’s reports about his actions involves auditing.

Clearly, if the principal can fully commit herself to a specific verification strategy *ex ante*, she can never be worse off under auditing. With auditing she can achieve any outcome under monitoring by simply mimicking the monitoring strategy (i.e., disregarding all intermediate information). The mimicking-strategy, however, requires that the principal’s verification strategy is verifiable such that her commitment to disregard additional information is credible. When such contractual commitment is not feasible, the weak optimality result of auditing may be overturned.

Indeed, if the principal cannot commit to a verification strategy, the principal’s verification behavior becomes a strategic variable that is chosen sequentially rationally. A non-commitment to verification seems reasonable if the effectiveness of verification depends on an unobservable scrutinizing effort by the principal. A second reason may be due to the difficulty of committing to random verification. As is well known (e.g. Mookherjee and Png 1989), optimal verification procedures often require a random use of verification, yet agents and outside courts may find it hard to verify whether the principal did indeed apply the correct random behavior as stipulated by some contract.³ This seems the most realistic reason why the assumption of non-verifiable verification makes sense: many real life contracts do stipulate the *possibility* that the agent is being verified, but do not determine the actual

³This argument is also used in Khalil (1997) and is further investigated in Strausz (2001).

frequency.⁴ Such contracts conform to the contractibility assumption in this paper: the principal binds herself contractually to a verification procedure, but its actual use is left at the principal's discretion.

The resulting non-trivial trade-off between monitoring and auditing is caused by a natural tension between incentives for working and incentives for verification. To induce high effort from an agent the principal must reward him when there is evidence that he worked. Such a payment structure implies that the principal has a relative strong incentive to audit when she receives bad news about the agent's behavior. In contrast, when she has an indication that the agent actually did work, the principal is less inclined to verify. This difference indicates that under auditing, the principal effectively chooses between two types of contracts: contracts that induce her to audit only after bad information and contracts that induce her to audit also after good news. I show that both types of contracts require relatively steep incentives. First, if the principal audits only after observing bad news, the auditing intensity is relatively low. As a consequence, the difference in the agent's payment after a positive verification and after non-verification must be large to induce the agent to work. In contrast, if the principal is to audit also after good news, the contract must give her an incentive to audit despite her information that it is relatively unlikely to catch the agent doing something wrong. Hence, the difference in the agent's payment after a negative verification and a non-verification must be large. Thus, both types of auditing contracts imply a steep incentive structure. An exclusive auditing after bad news requires steep incentives because of the agent's moral hazard problem, and an auditing after good news requires steep incentives due to the principal's moral hazard problem.

High powered incentive structures may render auditing suboptimal for two reasons. First, if transfers are bounded or the agent is risk averse, steeper incentive structures place a social cost on the organization and monitoring becomes preferable. Second, when the agent's liability is limited, his rent is increasing in the steepness of the incentive scheme. Auditing therefore leads

⁴In many countries employers are, by law, only allowed to use (stochastic) verification procedures if they inform their employees explicitly about their existence ex ante.

to higher rents for the agent. Hence, if the increase in rents is large, auditing is suboptimal.

The rest of the paper is organized as follows. The next section presents a simple model of verification. Section 3 derives the optimal contract under monitoring, and Section 4 analyzes the case of auditing. Section 5 compares the optimal contracts and derives conditions under which monitoring is superior to auditing. Some extensions are investigated in Section 6. Section 7 concludes. All proofs of the propositions are available on the JEBO website at [INSERT WEBADDRESS].

2 The Model

Consider a risk neutral principal who has a project that is run by a risk neutral agent. The agent chooses to work w or shirk s . If the agent works, the project is always successful. With shirking the project is only successful with probability p . A successful project yields the principal a payoff of y . An unsuccessful project is worthless. Hence, the productive gain when the agent works is $\Delta y \equiv (1 - p)y$. If the agent works he incurs a disutility of e . Shirking is costless to the agent. The difference $\Delta y - e$ measures therefore the potential social gain of working versus shirking. To have a non-trivial problem, I assume that this gain is strictly positive, $\Delta y > e$.

The agent's decision and the success of the project are not verifiable. Instead, the principal may, at a personal cost $c < e$, verify the agent's action to detect shirking. That is, there exists a verifiable signal $\sigma \in \{w, s\}$ about the agent's action whose informative content depends on a verification effort of the principal. The principal's effort is binary; she either verifies actively and incurs the cost c or she does not verify. Active verification reveals a shirking agent perfectly. If the principal does not verify actively, she cannot detect shirking and the signal σ always reports w .

The principal and agent write a contract t that stipulates transfers from the principal to the agent. Since only the signal σ about the action is verifiable, a general contract of transfers is a combination (t_w, t_s) . The agent's

liability is limited to zero. The maximum transfer that the principal can promise is bounded by $\bar{t} \geq e$. Hence, a feasible contract requires $t_w, t_s \in [0, \bar{t}]$. Due to the simple structure, the difference $\Delta t \equiv t_w - t_s$ measures the steepness of the agent's incentive structure. Limited liability implies that the steepness is at most $\bar{t}$.⁵ Finally, the agent's outside option is zero.

Before offering a contract to the agent the principal commits to one of the two verification procedures auditing or monitoring. If the principal adopts monitoring, she chooses her verification effort before knowing the agent's action, and the timing is as follows:

1. Principal offers a contract.
2. Agent decides whether to accept the contract.
3. Agent and principal decide simultaneously about action and verification effort respectively.

Hence, under monitoring the agent and principal play a simultaneous move game. In contrast, the principal and the agent play a sequential game under auditing. The agent chooses first his action, after which the principal observes the output. Only then she chooses whether to verify:

1. Principal offers a contract.
2. Agent decides whether to accept the contract.
3. Agent chooses his action.
4. Principal observes the output and chooses verification effort.

⁵The bounded transfers are intended to model a disadvantage to steep incentive structures (they make it impossible (i.e., infinitely costly) to implement incentive structures that are steeper than $\bar{t}$). Although risk aversion would be a more natural way to introduce such a disadvantage, the use of bounded transfers enables us to differ between two effects that may render auditing suboptimal. Moreover, risk aversion renders the analysis less tractable and makes it more difficult to separate the two effects that determine the optimal verification procedure.

The game with auditing is more complicated in that the principal takes her decision under asymmetric information. Hence, whereas with monitoring we may solve the subgame in stage 3 as a straightforward Nash equilibrium, the appropriate equilibrium concept in the game with auditing is Perfect Bayesian Equilibrium.

3 Monitoring

First suppose the principal uses monitoring as her procedure of verification. Clearly, if the principal does not monitor actively, the agent will shirk, since without verification his remuneration is independent of his actual action. Hence, if the principal wants to induce the agent to work, she must verify actively. Indeed, if the principal monitors with probability γ the agent receives a net utility of $\gamma t_w + (1 - \gamma)t_w - e$ if he works. Shirking, on the other hand, yields the agent $\gamma t_s + (1 - \gamma)t_w$. Hence, the agent has a weak incentive to work if

$$\Delta t \geq e/\gamma. \quad (1)$$

The inequality represents the agent's incentive constraint. It shows that the required steepness of the incentive structure, Δt , is inversely related to the principal's monitoring intensity γ . Indeed, if γ approaches zero, the required wedge Δt goes to infinity. It reflects the observation that at least some active verification has to occur to induce the agent to work.

Inducing the agent to work requires active verification by the principal, yet because verification is not contractible, the contract (t_w, t_s) must give the principal genuine incentives to monitor. Given that the agent works, the principal pays $t_w + c$ if she decides to monitor. If she, on the other hand, does not verify, she pays t_w . Hence, given that the agent worked, the principal will not monitor. Consequently, there is no equilibrium in which the agent works with probability one.

Now suppose the agent chooses to work with a probability α less than one. This requires that the agent must be indifferent between working and

shirking. That is, the agent's incentive constraint (1) holds in equality:

$$\Delta t = e/\gamma. \quad (2)$$

If the principal monitors, she expects to pay $\alpha t_w + (1 - \alpha)t_s + c$, whereas she pays t_w if she does not monitor. Hence, the principal has an incentive to monitor if

$$\Delta t \geq \frac{c}{1 - \alpha}. \quad (3)$$

Inequality (3) represents the principal's incentive constraint to monitor actively. It confirms the former observation that the principal cannot be given incentives to monitor if the agent works with probability one, as the required wedge Δt goes to infinity when α goes to one.

The incentive constraints (2) and (3) describe the implementation restrictions due to the double moral hazard problem. In addition to these constraints, the contract must ensure the participation of the agent by yielding the agent a non-negative utility. Yet, because the agent is protected by a limited liability of zero and shirking is costless, any admissible contract yields the agent a non-negative payoff if he chooses to shirk. Hence, any incentive compatible contract that satisfies limited liability ensures the agent at least his reservation utility. It follows that one may disregard the agent's individual rationality constraint and that the optimal contract solves the following problem:

$$\begin{aligned} P^m : \max_{t_w, t_s, \gamma, \alpha} \quad & V_m = (\alpha + (1 - \alpha)p)y - \alpha t_w - (1 - \alpha)[\gamma t_s + (1 - \gamma)t_w] - \gamma c \\ \text{s.t.} \quad & (1 - \gamma)(\Delta t - c/(1 - \alpha)) = 0 \\ & (2) \text{ and } (3), \end{aligned} \quad (4)$$

where the constraint (4) guarantees that the principal is indifferent about monitoring if she monitors with a probability less than one.

Proposition 1 *Under monitoring the optimal contract is $(t_w, t_s) = (t_w^*, 0)$ and yields the principal V_m^* . It induces the agent to work with probability*

$\alpha = 1 - c/t_w^*$ and the principal to monitor with probability $\gamma = e/t_w^*$, where

$$t_w^* = \begin{cases} e & \text{if } c\Delta y < e^2 \\ \sqrt{c\Delta y} & \text{if } c\Delta y \in [e^2, \bar{t}^2] \\ \bar{t} & \text{if } c\Delta y > \bar{t}^2 \end{cases} \quad V_m^* = \begin{cases} y - e - c\Delta y/e & \text{if } c\Delta y < e^2 \\ y - 2\sqrt{c\Delta y} & \text{if } c\Delta y \in [e^2, \bar{t}^2] \\ y - \bar{t} - c\Delta y/\bar{t} & \text{if } c\Delta y > \bar{t}^2. \end{cases}$$

The proposition shows that the maximum punishment principle holds so that t_s is set to its minimum of zero.⁶ The optimal level of t_w depends on the cost of verification c . If monitoring costs are relatively small, the principal chooses t_w such that she monitors with probability 1. For larger monitoring costs it is optimal for the principal to monitor with a probability less than one. Since t_w^* is increasing in c , we obtain the intuitive result that the monitoring intensity is decreasing in the cost of verification c . Finally, the maximum allowable transfer $\bar{t}$ restricts the principal only if it is relatively small.

4 Auditing

Now suppose the principal chooses auditing as her procedure of verification. In this case the principal decides about active verification after observing the project's outcome. Hence, she may audit failed and successful projects with different intensities. Suppose the principal audits successful projects with probability γ_σ and failures with probability γ_f . To induce the agent to work with positive probability, the decision to work must yield the agent at least as much as shirking. Given the principal's auditing intensities γ_f and γ_σ , the agent receives a utility of $p(\gamma_\sigma t_s + (1 - \gamma_\sigma)t_w) + (1 - p)(\gamma_f t_s + (1 - \gamma_f)t_w)$ when he shirks. Working, on the other hand, yields a net utility of $t_w - e$. Hence, the agent has a weak incentive to work if

$$\Delta t \geq \frac{e}{\gamma_\sigma p + \gamma_f(1 - p)}. \quad (5)$$

⁶Proofs of the propositions are available on the JEBO website at [INSERT WEBADDRESS].

Constraint (5) represents the agent's incentive constraint under auditing. It shows that at least some auditing must take place if the agent is to work with positive probability.

The principal's auditing behavior is guided by the contract t and her belief about the agent's behavior. More precisely, given that the principal believes that the agent worked with probability ω , she has a weak incentive to audit if

$$\omega t_w + (1 - \omega)t_s + c \leq t_w.$$

The principal's belief ω depends on the outcome of the project. If the principal observes a failure, this can only have come because the agent shirked. Hence, $\omega_f = 0$ and the principal has a (weak) incentive to audit a failed project if

$$\Delta t \geq c. \tag{6}$$

On the other hand, if the principal observes a successful project, the agent either worked or shirked but was lucky. Given that the agent works with probability α the probability that the agent worked follows from Bayes' rule:

$$\omega_\sigma = \frac{\alpha}{\alpha + (1 - \alpha)p}.$$

Consequently, the principal has a (weak) incentive to audit after observing a successful project if

$$\Delta t \geq c \left(1 + \frac{\alpha}{(1 - \alpha)p} \right). \tag{7}$$

Effectively, constraints (6) and (7) imply that the principal may choose between two basic auditing strategies. Either she audits only after observing a failure, or she audits also after observing a success. Quite intuitively, the principal cannot induce herself to audit only successful projects because constraint (7) is stricter than (6). That is, the auditing intensities γ_f and γ_σ are interdependent. If the principal audits successful projects with a positive probability, then she must audit failed projects with probability

one. Alternatively, if the principal audits failed projects with a probability less than one, she does not audit successful projects.

One may use the interdependence to simplify the agent's incentive constraint (5). If the principal audits failures with a probability less than one ($\gamma_f < 1$), she does not audit successes ($\gamma_\sigma = 0$), and the agent's incentive constraint (5) reduces to

$$\Delta t \geq \frac{e}{\gamma_f(1-p)}. \quad (8)$$

In contrast, if the principal audits successful projects with a probability $\gamma_\sigma > 0$, the principal audits failed project with certainty, $\gamma_f = 1$. Moreover, it requires that inequality (7) must be strictly satisfied such that $\alpha < 1$. That is, if in equilibrium the principal audits also successful projects, the agent has to shirk with positive probability and must, therefore, be indifferent between working and shirking. Consequently, the incentive constraint (5) rewrites as

$$\Delta t = \frac{e}{1 - (1 - \gamma_\sigma)p}. \quad (9)$$

Whether the principal chooses a contract that induces her to audit only failures or also successes depends on which type of contract yields the highest utility. The optimal contract when the principal audits both successful and failed projects solves the following problem:⁷

$$\begin{aligned} P^b : \max_{t_w, t_s, \gamma_\sigma, \alpha} \quad & V_b = (\alpha + (1 - \alpha)p)y - \alpha(t_w + \gamma_\sigma c) \\ & - (1 - \alpha)[p((1 - \gamma_\sigma)t_w + \gamma_\sigma t_s + \gamma_\sigma c) - (1 - p)(t_s + c)] \\ \text{s.t.} \quad & (1 - \gamma_\sigma)c[1 + \alpha/((1 - \alpha)p)] = 0 \\ & (7) \text{ and } (9), \end{aligned} \quad (10)$$

where equality (10) guarantees that the principal is indifferent about auditing a successful project if she is to audit such projects with a probability less than one.

⁷Again any incentive compatible contract satisfying limited liability is automatically individual rational to the agent.

In order to derive the solution to the problem P^b define

$$\hat{t}_b \equiv \frac{\sqrt{(1-p)c(c+py)} - (1-p)c}{p}.$$

Proposition 2 *The optimal contract that induces the principal to audit both failed and successful projects is $(t_w, t_s) = (t_b^*, 0)$, where*

$$t_b^* = \begin{cases} \min\{\hat{t}_b, e/(1-p), \bar{t}\} & \text{if } \hat{t}_b \geq e \\ e & \text{if } \hat{t}_b < e. \end{cases}$$

It yields the principal the payoff V_b^ , where*

$$V_b^* = \begin{cases} \frac{\bar{t}p(y-\bar{t})-c^2(1-p)}{(1-p)c+\bar{t}p} & \text{if } \hat{t}_b > \bar{t} \text{ and } \bar{t} < e/(1-p) \\ \frac{ep(y(1-p)-e)-c^2(1-p)^3}{(1-p)^2c+ep} & \text{if } \hat{t}_b > e/(1-p) \text{ and } \bar{t} \geq e/(1-p) \\ y - \frac{2(1-p)c-2\sqrt{(1-p)c(c+py)}}{p} & \text{if } \hat{t}_b \in [e, \bar{t}] \\ \frac{ep(y-e)-c^2(1-p)}{(1-p)c+ep} & \text{if } \hat{t}_b < e. \end{cases}$$

In the associated equilibrium the agent works with probability $\alpha = [t_w^ - c]/[t_w^* - c + c/p]$. The principal audits failed projects with probability one and successful projects with probability $\gamma_f = 1 - [t_w^* - e]/(t_w^*p)$.*

The proposition shows that the optimal contract that leads to auditing both failed and successful projects induces a similar equilibrium under monitoring. In both equilibria the principal and agent use a mixed strategy. Also the intuition behind the result is the same. If the agent would work with probability one, the principal will not monitor or verify a successful project.

As an alternative to auditing both failed and successful projects, the principal may choose a contract that induces her to audit only failures. An optimal contract of this kind solves

$$\begin{aligned} P^f : \max_{t_w, t_s, \gamma_f, \alpha} \quad & V_f = (\alpha + (1-\alpha)p)(y - t_w) - (1-\alpha)(1-p)(\gamma_f t_s + (1-\gamma_f)t_w + c\gamma_f) \\ \text{s.t.} \quad & (1-\gamma_f)[\Delta t - c] = 0; \quad (1-\alpha)[\Delta t - e/(\gamma_f(1-p))] = 0 \\ & (6) \text{ and } (8), \end{aligned} \tag{11}$$

where the constraints in (11) guarantee that the principal or agent is indifferent if she or he uses a mixed strategy.

Proposition 3 *An optimal contract that induces the principal to audit only failed projects exists only if $\bar{t} \geq e/(1-p)$. It exhibits $t_w^* = e/(1-p)$ and $t_s^* = 0$ and induces the principal to audit failed projects with probability one and the agent to work with probability one. It yields the principal $V_f^* = y - e/(1-p)$.*

The nature of the optimal contract differs from the other type of contract under auditing and monitoring. First, in equilibrium the principal's incentives are strict.⁸ That is, the principal's incentive constraint is not binding. At first sight, this is surprising as it implies that the principal's inability to commit does not constrain the equilibrium. The result is nevertheless intuitive. If the principal observes a failure, she is sure that the agent did not work. A failure, therefore, leads to the lower payment t_s . The previously used argument that the principal will not verify if she is sure that the agent works does therefore not hold.⁹ In fact, the principal uses auditing only as a threat to withhold the agent from shirking and auditing does not take place in equilibrium.

This type of contract has therefore two attractive features. First it enables the principal to induce the agent to work with probability one without incurring any verification costs. Second, the contract only exists if the maximum possible payment $\bar{t}$ is large enough. This observation follows directly from the agent's incentive constraint (8). Indeed, given that the principal audits failed projects exclusively, the agent's shirking is identified only if the project fails. Hence, the detection probability of shirking is at most $1 - p$. This reveals the disadvantage that an exclusive auditing of failures implies a low detection probability and, therefore, requires high powered incentives for the agent. Such high powered incentives are possible only if the maximum allowable payment $\bar{t}$ is large enough.

⁸This follows from $c < e < e/(1-p) = \Delta t$.

⁹Technically, the agent's decision shifts the support of the principal's observation. If a failure could also occur with some (small) positive probability when the agent works then working with probability one is not sustainable in equilibrium. Since this would make auditing less attractive, the shifting support gives auditing an extra advantage over monitoring. It will be shown that it may nevertheless be suboptimal.

5 Monitoring versus Auditing

The previous two sections derived the optimal contracts under monitoring and auditing. This section compares the three different types of contracts. Comparing the optimal contract under monitoring to the optimal contract that gives the principal incentives to audit both failed and successful projects yields the following result.

Proposition 4 *It holds $V_m^* \geq V_b^*$.*

The proposition establishes the superiority of monitoring over an auditing of failed and successful projects. To understand the result, note that in both equilibria the principal verifies probabilistically. That is, she is, in equilibrium, indifferent concerning verification. Therefore, the difference in her payoffs in the two equilibria depends only on the contract t and is independent of the intensity of verification. It then follows that Proposition 4 is entirely driven by the principal's commitment problem: for a given working intensity α the principal believes that, if she monitors, she saves on the wage payment t_w with probability $1 - \alpha$. Her consideration is different under auditing where the principal receives additional information before she decides to verify. In particular observing a success raises the principal's belief from α to ω_σ . Hence, after a success she considers it less likely to save on the payment t_w if she verifies. To induce her to verify nevertheless, the saved amount t_w must be larger. This implies a steeper incentive structure, which leads to more risk on the agent and a higher wage bill for the principal.

Indeed, as the principal's incentive constraints (3) and (7) capture the principal's inability to commit, a direct comparison of these two constraints confirms the reasoning. The constraint (3) is weaker than (7), implying that for a given working intensity α , auditing requires a higher powered incentive structure than monitoring.

Because for $\bar{t} < e/(1-p)$ a contract that induces the agent to work requires the principal to audit both failed and successful projects, the proposition has a straightforward corollary.

Corollary 1 *Monitoring is strictly better than auditing if $\bar{t} < e/(1-p)$.*

Hence, if the maximum payment $\bar{t}$ is relatively low, such that the principal is unable to induce the agent to work by auditing only failures, the principal prefers monitoring over auditing. The result follows from the exogenously bounded transfers, but is ultimately due to the need for high powered incentives. That is, auditing can only be better than monitoring if the implementation of steep incentives is not too costly. Since a steeper incentive structure implies more risk, steep incentives are costly to implement if the agent is risk averse. That is, if transfers are unlimited but the agent is risk averse, then auditing is inferior to monitoring because the cost of compensating the agent for his increased risk will outweigh the gain from a selective auditing of failed projects. Indeed, one may see a boundedness of transfers as an extreme form of risk aversion. Instead of restricting wages to the interval $[0, \bar{t}]$, one may assume that the agent has a utility of $u(t) = \min\{t, \bar{t}\}$ for positive transfers $t \geq 0$. That is, the agent is risk neutral in the interval $[0, \bar{t}]$ and infinitely risk averse for wages lower than 0 and exceeding $\bar{t}$.

The corollary shows that one factor that influences the optimal verification procedure is the social costs associated with steep incentive schemes due to risk aversion or limited funds. The remainder of this section shows that even if steep incentive schemes do not involve a social cost, auditing may still be suboptimal due to the principal's concern over rents.

As discussed earlier the advantage of auditing is that the principal may circumvent her commitment problem and induce the agent to work with probability one without incurring verification costs. However, under an exclusive auditing of failed projects, the verification intensity is at most $1 - p$ and implies that the agent's incentive structure must be rather steep to induce him to work. As the agent is protected by limited liability, a steep incentive scheme results in a large rent to the agent. Hence, apart from the increased risk a second drawback of the auditing strategy is that it may require a relatively large rent to the agent. As an alternative namely the principal may use monitoring. In this case, the principal is able to verify with a larger probability, which translates into a smaller rent to the agent. A drawback, however, is that the principal is unable to induce the agent to work with probability one and, in addition, incurs positive verification costs, yet when

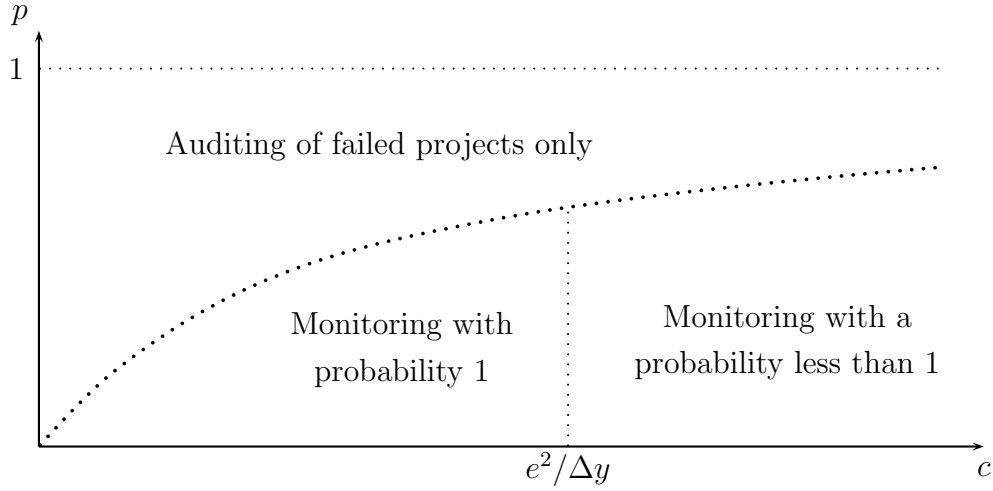


Figure 1: Optimal verification procedures due to rent extraction.

the principal has to give up too many rents under auditing, monitoring will be optimal.¹⁰

The following proposition identifies the optimal verification strategy due to the problem of rent extraction.

Proposition 5 *Suppose the maximum transfer is unbounded; then monitoring is strictly optimal if and only if*

$$c\Delta y < S(p, e)$$

with

$$S(p, e) = \begin{cases} \frac{e^2 p}{(1-p)} & \text{for } p \leq 1/2 \\ \frac{e^2}{4(1-p)^2} & \text{for } p > 1/2. \end{cases}$$

The result shows that when c and Δy are relatively large in comparison to $e/(1-p)$, auditing is optimal. Indeed, if the cost of verification c and Δy are high, the advantage of auditing is large because it saves on verification costs and allows a working intensity of one. On the other hand $e/(1-p)$

¹⁰Note that without a social cost to steep incentive schemes, this implies that the principal does not maximize the overall surplus and takes a socially inefficient decision. See Strausz (2005) for a more careful analysis of this effect.

represents the agent's rent under auditing and renders auditing suboptimal if it is too large.

Figure 1 illustrates the proposition graphically. When the verification procedure is accurate (high p), verification is relatively effective in extracting the agent's rent. Hence, the fact that auditing requires higher rents plays only a minor role. Consequently, auditing is optimal. For less accurate signals, the rent argument gains in relative importance, and monitoring becomes optimal. In this case, the principal always monitors when the cost of monitoring is low. When costs increase, the principal decreases the monitoring frequency, which implies that she uses it randomly.

6 An Extension: Verifiable Output

Until now we used an extremely stylized model that allowed us to calculate explicitly the optimal contracts under monitoring and auditing and compare them accordingly. The analysis becomes less tractable if one considers more standard contracting settings, in which the agent's contract conditions on additional verifiable variables such as output. In this section we confirm, however, that the tension between incentives for working and verification also exists in these more elaborate models. Again, assume that the agent chooses between two effort levels e_h and e_l with costs $e > 0$ and zero, respectively. As is more standard, let effort result in a distribution $f(y|e_i)$ over the possible output levels $Y \subset \mathbb{R}$. As before, the principal chooses between two intensities of verification v_1 and v_0 , where the higher intensity v_1 has a cost of $c > 0$ and the lower intensity has costs zero. Verification at intensity level v_j leads to a verifiable result $a \in A$ according to the distribution function $g_j(a|e_i)$. Verification at the intensity level v_0 is uninformative. That is, $g_0(a|e_h) = g_0(a|e_l)$ for all $a \in A$. In contrast, the verifiable result a is informative if the principal verifies with the higher intensity level v_1 . This means that there exists some subset $A_1 \subset A$ with positive measure such that $g_1(a|e_h) \neq g_1(a|e_l)$ for all $a \in A_1$.

Let output $y \in Y$ be verifiable such that transfers between the principal and agent may now depend directly on y and a . We write $t(y, a)$. Feasibility

requires $t(y, a) \in [0, \bar{t}]$. Given some contract $t(y, a)$ we may define

$$t_{ij} \equiv \int \int t(y, a) f(y|e_i) g_j(a|e_i) dy da \quad (12)$$

as the expected payment to the agent if the agent chooses effort level e_i and the principal verifies with intensity level v_j .¹¹

If the principal chooses the verification intensity v_1 with probability γ , the agent receives $\gamma t_{h1} + (1 - \gamma)t_{h0} - e$ when he works and $\gamma t_{l1} + (1 - \gamma)t_{l0}$ when he shirks. The agent's incentive constraint is therefore

$$\gamma \Delta t \geq e + t_{l0} - t_{h0}, \quad (13)$$

with

$$\Delta t \equiv t_{h1} - t_{l1} - t_{h0} + t_{l0}.$$

If the right hand side of (13) is negative, the contract is able to induce the agent to work while verifying at the low intensity v_0 . However, depending on the difference $f(y|e_h) - f(y|e_l)$ this may require transfers $t(y, a)$ that are not feasible or imply large rents. Hence, suppose that $t_{h0} - t_{l0} < e$ so that the principal must verify with the higher intensity v_1 in order to induce the agent to work. In this case, the agent's incentive constraint (13) implies that Δt must be positive.

If the principal believes the agent to work with probability α^e , she expects to pay $\alpha^e t_{h1} + (1 - \alpha^e)t_{l1} + c$ when she verifies and $\alpha^e t_{h0} + (1 - \alpha^e)t_{l0}$ when she does not verify. Consequently, she verifies if

$$\alpha^e \Delta t \leq t_{l0} - t_{l1} - c. \quad (14)$$

As before, the incentive constraints (13) and (14) reveal a tension between incentives for working and incentives for verification. On the one hand, Δt must be large enough to induce the agent to work. On the other hand, Δt must be small enough to induce the principal to verify. In particular, the larger α^e the stronger the tension between the two constraints. This relation

¹¹In the more stylized model we had $t_{h0} = t_{l0} = t_{h1} = t_w$ and $t_{l1} = t_s$.

drives the result that under auditing it is more difficult to induce an intensive verification of the agent. To see this, note that under monitoring the principal's belief α^e coincides in equilibrium with the agent's actual randomization α^* .

Under auditing the principal's belief depends on some additional information s . To make this more precise, let $h(s_j|e_i) \in [0, 1]$ represent the probability that the principal receives the signal $s_j \in S = \{s_1, \dots, s_N\}$ when the agent chooses action e_i . Suppose the set S is ordered such that a higher signal indicates that the agent worked, $h(s_j|e_h)/h(s_j|e_l) > h(s_{j-1}|e_h)/h(s_{j-1}|e_l)$. Now, if the agent randomizes with probability α^* then, in equilibrium, the principal's belief after receiving the signal s_i satisfies

$$\alpha^e(s_i) = \frac{\alpha^* h(s_i|e_h)}{\alpha^* h(s_i|e_h) + (1 - \alpha^*) h(s_i|e_l)}.$$

Due to the ordering of the signal s , the belief $\alpha^e(s_i)$ is increasing in i .

It now follows that it is harder for the principal to verify with a high probability under auditing than under monitoring. For instance, if the principal wants to audit with probability one, she must have an incentive to audit for any s_i and in particular for s_N , but since, necessarily, $h(s_N|e_h) > h(s_N|e_l)$, it holds $\alpha^e(s_N) > \alpha^*$ so that the principal's incentive constraint is stricter under auditing than under monitoring. This implies that with monitoring the principal is able to verify more easily with a high frequency than under auditing. Consequently, the optimal frequency under auditing tends to be lower than under monitoring, but if the agent is less frequently verified under auditing, steeper incentives are required to induce the agent to take a high effort level. We therefore obtain the same results as in the simpler model in which output was not contractible.

7 Discussion and Conclusion

This paper studied the strategic effect of a difference in the timing of verification. It showed that when the principal's verification behavior is non-contractible, monitoring may be optimal. The non-contractibility creates a

double moral hazard problem and, thereby, requires steeper incentive structures under auditing. The agent's moral hazard problem asks for steeper incentives if the verification intensity is low. The principal's moral hazard problem demands a steeper incentive structure if the principal is to use a high verification intensity. For two reasons steeper incentive schemes may render auditing suboptimal. First, they may be too costly to implement due to reasons of risk aversion or limited funds. Second, the incentive structure affects the division of rents to create an additional trade-off. At the expense of a higher rent to the agent, auditing enables a higher working intensity at lower verification costs. Depending on the outcome of this trade-off either monitoring or auditing is optimal. Ultimately, the optimality of monitoring is due to a natural tension between the principal's incentives to verify and the incentives for the agent which is stronger under auditing. Indeed, by switching to monitoring as her procedure of verification the principal relaxes the tension. Hence, monitoring may be seen as a commitment device not to act on the additional information.

Although the paper used a simple model the main result that auditing (i.e., additional information) requires steeper incentives is robust. As shown in the previous section an extension to multiple output levels and different sources of information do not affect qualitative results. Crucial and driving the result is the double moral hazard problem. Additional information about the agent's action intensifies the tension between the two problems and requires higher stakes and steeper incentives. If such incentive structures are costly to implement, preventing oneself from receiving additional information (i.e., choosing monitoring rather than auditing) may be an optimal strategy.

The limited liability on the part of the agent is responsible for the positive rent to the agent. Hence, without limited liability the second disadvantage of auditing, that steeper incentives require higher rents, disappears. In this case, auditing also imposes more risk on the agent than monitoring and may therefore be suboptimal due to risk aversion.

Since the difference between monitoring and auditing is only the additional information, our results also shed light on the value of information to the principal. More precisely, only if steep incentives are not too costly to

implement can the information have a positive value to the principal. In this case, the difference $V_f^* - V_m^*$ expresses the value of information and represents the principal's maximum willingness to pay for the information.

As discussed in the introduction it is trivial that auditing is (weakly) superior if the principal can commit to a verification strategy. Starting from a moral hazard problem, this paper therefore studied non-contractible verification. This transforms the model in a double moral hazard problem, which is more severe under auditing. Subsequently, conditions were studied under which monitoring is optimal. It remains an open question as to how far the results and the provided intuition of this model extend for a setting in which the principal's basic problem is an adverse selection rather than a moral hazard problem. As is well known, the analysis of an adverse selection problem with limited commitment is rather involved (for a possible approach see Bester and Strausz 2001) and fundamentally different from a moral hazard problem. Nevertheless, also in this class of models, the principal is often hurt by more information so that one may expect to find conditions under which auditing may be suboptimal as well. How far these conditions relate to the ones arrived under moral hazard remains a question for future research.

8 Appendix

Proof of Proposition 1:

Proof: We solve problem P^m by first disregarding (4) since the solution of the relaxed problem automatically satisfies the constraint. Substitution of (2) yields the simplified problem

$$\begin{aligned} \max_{t_s, \gamma, \alpha} \quad & V = (\alpha + (1 - \alpha)p)y - t_s - (1 - (1 - \alpha)\gamma)e/\gamma - \gamma c \\ \text{s.t.} \quad & \alpha \leq 1 - c\gamma/e \end{aligned} \tag{15}$$

with $t_w = t_s + e/\gamma$. Since (15) is independent of t_s and the objective constraint is decreasing in t_s , optimality requires $t_s = 0$. Moreover, assuming that (15) is not binding leads to a contradiction: if it does not bind, the objective function is linear in α and, since, by assumption, $\alpha > 0$ is optimal, linearity

implies that $\alpha = 1$ is optimal, yet this violates the constraint. That is, (15) is binding so that (4) is indeed satisfied. By substitution of (2) and (4) and using $t_s = 0$, the original maximization problem P^m may be rewritten as

$$\max_{t_w} y - t_w - \frac{c\Delta y}{t_w},$$

which is concave in t_w as the 2nd derivative with respect to t_w is $-2c\Delta y/t_w^3$. Hence, the first order condition is sufficient and yields

$$\hat{t}_w = \sqrt{c\Delta y}.$$

If $\hat{t}_w > \bar{t}$, then optimally $t_w^* = \bar{t}$. Otherwise, $t_w^* = \max\{e, \hat{t}_w\}$, where $\hat{t}_w \geq e$ if and only if $c\Delta y \geq e^2$.

Q.E.D.

Proof of Proposition 2:

Proof: Substituting out t_w by using (9) yields

$$V_b = (\alpha + (1-\alpha)p)y - c(1 - (1-\alpha)(1-\gamma_\sigma)p - \alpha(1-\gamma_\sigma)) - \frac{(\alpha + (1-\gamma_\sigma)(1-\alpha)p)e}{1 - (1-\gamma_\sigma)p} - t_s,$$

and shows that V_b is decreasing in t_s (a lower t_s also relaxes the remaining constraint (7)). Hence, optimality requires $t_s = 0$. Moreover, since V_b is linear in α and, by assumption, $\alpha > 0$ is optimal, it follows that (7) must bind at the optimum. This implies that (10) is automatically satisfied. By using (7) and (9) to substitute out α and γ_s the problem P^b reduces to

$$\max_{t_w \in [e, \min\{e/(1-p), \bar{t}\}]} V_b(t_w) \equiv \frac{pt_w(y - t_w) - c^2(1-p)}{pt_w + (1-p)c}, \quad (16)$$

where the domain restriction $t_w \in [e, e/(1-p)]$ guarantees that $\gamma \in [0, 1]$. The first order condition is

$$\hat{t}_b = \frac{\sqrt{(1-p)c(c+py)} - (1-p)c}{p}.$$

It leads to a utility of

$$\hat{V}_b^* \equiv y - \frac{2\sqrt{(1-p)c(c+py)} - 2(1-p)c}{p}. \quad (17)$$

Since the 2nd derivative with respect to t_w at $\hat{t}_{2w}$ is $-2p/\sqrt{(1-p)c(c+py)} < 0$ the first order condition is sufficient if it satisfies the domain restrictions, that is, $t_w^* = \hat{t}_b$ is optimal if $\hat{t}_b \in [e, \min\{e/(1-p), \bar{t}\}]$. Now if $\hat{t}_b < e$ then $t_w^* = e$ is optimal. On the other hand, if $\hat{t}_b > \min\{e/(1-p), \bar{t}\}$, then optimality requires $t_w^* = \min\{e/(1-p), \bar{t}\}$. For the principal's utility it follows

$$V_b^* = \begin{cases} \frac{\bar{t}p(y-\bar{t})-c^2(1-p)}{(1-p)c+\bar{t}p} & \text{if } \hat{t}_b > \bar{t} \text{ and } \bar{t} < e/(1-p) \\ \frac{ep(y(1-p)-e)-c^2(1-p)^3}{(1-p)[(1-p)^2c+ep]} & \text{if } \hat{t}_b > e/(1-p) \text{ and } \bar{t} \geq e/(1-p) \\ \hat{V}_b^* & \text{if } \hat{t}_b \in [e, \min\{e/(1-p), \bar{t}\}] \\ \frac{ep(y-e)-c^2(1-p)}{(1-p)c+ep} & \text{if } \hat{t}_b < e. \end{cases}$$

Q.E.D.

Proof of Proposition 3:

Proof: From (8), $\gamma_f \leq 1$ and $t_s \geq 0$, it follows that $t_w \geq e/(1-p)$, but since $t_w \leq \bar{t}$, a necessary condition for the existence of a contract that induces the agent to work is $\bar{t} \geq e/(1-p)$.

If $\bar{t} \geq e/(1-p)$ then constraint (11) implies that either $\alpha = 1$ or (8) is binding. First suppose (8) is slack such that $\alpha = 1$. The maximization problem reduces then to

$$\begin{aligned} \max_{t_w, t_s, \gamma_f} \quad & y - t_w \\ \text{s.t.} \quad & (1 - \gamma_f)(\Delta t - c) = 0; \Delta t \leq c \end{aligned}$$

and yields $t_w = c$, $t_s = 0$ as an optimum. Since $c < e$, this violates (8) for any $\gamma_f \in [0, 1]$. Consequently, (8) must be binding at the optimum, which implies that (11) is automatically satisfied. Disregarding (11) implies that α only enters linearly in the objective function V_f . Hence, $\alpha = 1$ must be optimal, while (8) is binding. This reduces the maximization problem to

$$\begin{aligned} \max_{\gamma_f \in [0, 1], t_s \geq 0} \quad & y - t_s - e/((1-p)\gamma_f) \\ \text{s.t.} \quad & (1 - \gamma_f)[e/(\gamma_f(1-p)) - c] = 0 \\ & e/(\gamma_f(1-p)) \geq c. \end{aligned} \tag{18}$$

It follows that $\gamma_f = 1$ and $t_s = 0$ is optimal such that $t_w^* = e/(1-p)$ and $V_f^* = y - e/(1-p)$; apart from necessary, the condition $\bar{t} \geq e/(1-p)$ is also sufficient for existence.

Q.E.D.

Proof of Proposition 4:

Proof: We must show that $V_m^* \geq V_b^*$. Since $\bar{t} > e$ it suffices to distinguish between three cases:

1) If $c\Delta y \in [e^2, \bar{t}^2]$ then, according to Proposition 1, $V_m^* = y - 2\sqrt{c\Delta y}$. From the proof of Proposition 2 it follows $V_b^* \leq \hat{V}_2^*$. We now show that

$$V_m^* - \hat{V}_b^* = \frac{2 \left[\sqrt{(1-p)c(c+py)} - (1-p)c - p\sqrt{(1-p)cy} \right]}{p} \quad (19)$$

is non-negative because the term in the square bracket is non-negative. This follows from

$$\begin{aligned} p^2(c - (1-p)y)^2 \geq 0 &\Rightarrow (c + py + (1-p)c - p^2y)^2 \geq 4(1-p)c(c + py) \\ &\Rightarrow c + py + (1-p)c - p^2y \geq 2\sqrt{(1-p)c(c + py)} \\ &\Rightarrow c + py + (1-p)c - 2\sqrt{(1-p)c(c + py)} \geq p^2y \\ \Rightarrow (1-p)c \left[(c + py) + (1-p)c - 2\sqrt{(1-p)c(c + py)} \right] &\geq (1-p)cp^2y \\ &\Rightarrow \sqrt{(1-p)c(c + py)} - (1-p)c \geq p\sqrt{(1-p)cy}; \end{aligned}$$

2) for $c\Delta y < e^2$ it holds $V_m^* = y - e - c\Delta y/e$ and $V_b^* \in \{V_b(t_w) | t_w \in [e, e/(1-p)]\}$ with $V_b(t_w)$ as defined by (16), but for any $t_w \in [e, e/(1-p)]$ it holds

$$V_m^* - V_b(t_w) = \frac{\overbrace{p(t_w - e)}^{\geq 0} \overbrace{[et_w - c\Delta y]}^{\geq 0} + c(1-p) \overbrace{(e - c)}^{> 0} \overbrace{[\Delta y - e]}^{\geq 0}}{\underbrace{(1-p)ce + pt_w e}_{> 0}} \geq 0,$$

where the non-negativity of the first bracketed term follows from $c\Delta y < e^2 < et_w$;

3) if $c\Delta y > \bar{t}^2$ then $V_m^* = y - \bar{t} - c\Delta y/\bar{t}$. Due to $c < e < \Delta y$ it holds $\hat{t}_b > \sqrt{c\Delta y} > \bar{t}$. Therefore, $\hat{t}_b$ exceeds $\bar{t}$, which implies $V_b^* \leq \hat{V}_b^* = (\bar{t}p(y - \bar{t}) - c^2(1-p))/((1-p)c + \bar{t}p)$. It follows that

$$V_m^* - \hat{V}_b^* = \frac{c(1-p)(\bar{t} - c)[\Delta y - \bar{t}]}{\bar{t}(c + p(\bar{t} - c))} > 0.$$

The inequality holds because from $\sqrt{c\Delta y} > \bar{t}$ and $\bar{t} > c$, it follows $\Delta y > \bar{t}^2/c > \bar{t}$.

Q.E.D.

Proof of Proposition 5:

Proof: If the maximum transfer $\bar{t}$ is unbounded, then $V_f^* = y - e/(1 - p)$. It follows from Proposition 4 that auditing is superior if and only if $V_m^* < V_f^*$.

For $c\Delta y < e^2$ it holds $V_m^* = y - 2\sqrt{c\Delta y}$. Hence, $V_m^* < V_f^*$ if and only if $c\Delta y < e^2p/(1 - p)$. This yields a critical threshold $p_1(c) = c\Delta y/(c\Delta y + e^2)$ for the range $c < e^2/\Delta y$.

For $c\Delta y \geq e^2$ it holds $V_m^* = y - e - c\Delta y/e$. Hence, $V_m^* < V_f^*$ if and only if $4c\Delta y < e^2/(1 - p)^2$. This yields a critical threshold $p_2(c) = 1 - e/(2\sqrt{c\Delta y})$ for the range $c \geq e^2/\Delta y$.

Note that $p_1(e^2/\Delta y) = p_2(e^2/\Delta y) = 1/2$ so that the combined function

$$P(c) = \begin{cases} p_1(c) & \text{for } c < e^2/\Delta y \\ p_2(c) & \text{for } c \geq e^2/\Delta y \end{cases}$$

is continuous in c . Moreover, $P(c)$ is also differentiable at $c = e^2/\Delta y$ since $p'_1(e^2/\Delta y) = p'_2(e^2/\Delta y)$. Finally, the function $S(p, e)$ obtains from inverting $P(c)$.

Q.E.D.

References

- Bester, H., Strausz, R., 2001. Contracting with imperfect commitment and the revelation principle: the single agent case. *Econometrica* 69, 1077–1098.
- Holmström, B., 1979. Moral hazard and observability. *Bell Journal of Economics* 10, 74–91.
- Khalil, F., 1997. Auditing without commitment. *RAND Journal of Economics* 28, 629–640.
- Khalil, F., Lawarrée, J., 1995. Input versus output monitoring: who is the residual claimant? *Journal of Economic Theory* 66, 139–157.
- Maskin, E., Riley, J., 1985. Input versus output incentive schemes. *Journal of Public Economics* 28, 1–23.
- Mookherjee, D., Png, I., 1989. Optimal auditing, insurance, and redistribution. *Quarterly Journal of Economics* 104, 399–415.

- Shavell, S., 1979. Risk sharing and incentives in the principal and agent relationship. *Bell Journal of Economics* 10, 55–73.
- Strausz, R., 2001. Mitigating non-contractibility with interim randomization. *Journal of Institutional and Theoretical Economics* 157, 231–245.
- Strausz, R., 2005. Buried in paperwork: excessive reporting in organizations. *Journal of Economic Behavior and Organization*, in press.
- Townsend, R., 1979. Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory* 21, 265–293.