

Collusion as an Indirect Signal

A Theory of Corrupt Delegation in Hierarchy

Yucheng Wang*

Abstract

In a principal-agent relationship with adverse selection, direct investigation from a third party (such as an outside supervisor) is a common practice, especially when agents lack effective and credible signals to reveal their types. Unlike traditional explanations such as expertise in investigation, this paper finds that corrupt (susceptible to side payments from the agents) but credible (truth-telling) delegation may be beneficial to the principal, even if the supervisor has the expertise of investigation no better than the principal. It is because under corrupt delegation, the collusion with the form of covert side payment can play as an effective ‘indirect signal’ revealing the type of agents. As a result, it improves the incentives offered by the principal and saves the cost of investigation. More importantly, this signal is ‘indirect’ because under some circumstances, it becomes ineffective without available direct investigation, or if it could be directly observed by the principal.

Keywords: Principle-Agent Model; Delegation; Collusion; Signaling

*Graduate student at Department of Economics, Duke University. E-mail: yucheng.wang124@duke.edu. My major research interests include micro theory and political economics. I would like to thank Charles Becker, Ed Tower, and John Guerard for their comments. All errors are my own.

Mundus vult decipi, ergo decipiatur. (The world wants to be deceived, so let it be deceived.) – a Latin phrase

1 Introduction

Inspired by studies of information asymmetry and contract theory, the basic model of the principal-agent relationship has been applied to various aspects of daily lives, providing insightful suggestions on how to design incentives to deal with adverse selection and/or moral hazard of the agent. In practice, a commonly considered method is to introduce a third-party supervisor; usually, his duty is to investigate the private information or monitor the agent's effort. Such delegation emerges in all business, administrative, and political contexts: the employer hires monitors to supervise employees, the government establishes special agencies to supervise firms, the society votes political candidates with the help of media, etc. Even though sometimes there is no explicit relationship of employment between the principal and supervisor (such as citizen and media), we still can regard some long-run relationships as 'employments' with implicit contracts.

However, it is natural to believe that since the supervisor cannot be completely controlled by the principal, the delegation creates a new principal-agent relationship and its subsequent possible problems: besides conventional ones such as the moral hazard, a major problem is the possibility of coalition between supervisor and agent through some covert collusion; in the simplest form, the agent may bribe the supervisor to align their interests.

Indeed, the possibility of collusion has been well recognized by scholars, and many studies focus on the optimal contract with additional coalition-proof constraints (Tirole, 1986; Baiman et al., 1991; Laffont and Tirole, 1991), or mechanisms to offset collusion (Kofman and Lawarrée, 1993; Laffont and Martimort, 1999; Celik, 2009). However, even in the pioneering study on collusion in organizations, Tirole (1986) admitted that the coalition between supervisor and agent may not be always harmful, and necessary to be eliminated:

In our model ... Coalitions and their enforcement mechanism, side transfers, ought to be fought. This conclusion is extreme. In practice, some side transfers exist because organizations do not want to (rather than cannot) curb them. ... preventing relationships between members of a hierarchy may result in efficiency losses. (Tirole, 1986, pp.207-208)

As Tirole (1986) pointed out, collusion may have positive effects on organizations, including motivating high-quality teamwork, reducing moral hazard issues within teams, and acting as mutual help. Aside from these side effects outside of the basic framework, even within the principal-supervisor-agent relationship, a strand of studies investigates the potential positive effects of collusive delegations under particular circumstances (Mookherjee, et al., 2020). For example, Strausz (1997b) points out two positive effects of delegation of monitoring by the use of identical technologies: the 'incentive-effect' such that the principal can better regulate incentives for the agent and the supervisor; and the 'commitment-effect' such that the principal can commit to the reward to the agent.

This paper argues a new positive effect of collusion: the collusion itself may perform as an ‘indirect’ signal from the agent when a third-party supervisor plays the role of costly direct investigation (or monitoring) of different types of agents. This signal is ‘indirect’ because it only involves the interaction between the supervisor and the agent, and cannot be observed (and credibly sent) by the principal, or it will become ineffective. In short, the side payments proposed by the agents will credibly reveal their types if they are ensured to be covert. Moreover, the effectiveness of this ‘indirect’ signal is also favorable to the principal, who is supposed to pay the supervisor to motivate his investigation.

To illustrate how collusion can perform as a signal, note that collusion with the form of side payment can be used to motivate or discourage the supervisor’s effort to investigate, and different agents may have opposite attitude toward reliable investigation. In particular, for the ‘good’-type agents, the investigation can help the principal acknowledge their performance and reward them (e.g. offer positions, pay higher wages, etc.).

As a result, there are two possible cases in which ‘good’ and ‘bad’ agents differ in their decision of collusion: First, if the supervisor is ‘reluctant’ in the sense that he will be motivated to investigate only by some sufficient bribe from an agent, then the ‘good’ agents may be willing to offer bribes since the investigation is beneficial to them, while the ‘bad’ ones realize that investigation does not help but only reveal their incompetence. Second, if the supervisor is ‘blackmailing’ in the sense that he will refrain from investigation only by some sufficient bribe from an agent, then the ‘bad’ agents may be willing to offer bribes to escape the investigation, while the ‘good’ ones are clearly not willing to imitate them. In both cases, the supervisor can distinguish agents from their bribes, and the collusion turns out to be an effective signal.

However, since the effectiveness of the collusion as a signal depends on direct investigation and its relationship with the principal’s action, it has to be ‘indirect’ such that the principal can distinguish the agents only through the supervisor’s investigation, instead of the collusion signal itself. This requires not only the supervisor to be credible (truth-telling), but also the collusion to be ‘covert’ and cannot be directly observed by the principal: once directly observed, the principal can simply skip the direct investigation and distinguish agents from collusion details if they are still effective signals, which will, in turn, makes them ineffective.

On the other hand, the principal is also willing to permit collusion in the delegation to save her cost paid to the supervisor: the aggregate cost of investigation can be reduced by narrowing the set of agents to be investigated, once the supervisor could identify the types of agents before direct investigation through the effective signal of collusion.

From the logic above, one can identify some key prerequisites of the effectiveness of collusion as an indirect signal. First, direct, albeit may not be perfect, investigation is possible and costly. Next, the supervisor’s report is credible, which can be ensured by a long-term relationship with the principal. Finally, the collusion is credibly covert, even though the principal can rationally anticipate their existence and reduce the cost by taking advantage of them.

Since the principal-agent model can be applied to various economic and social contexts, the conclusion of this paper can have some implications for some real-life issues. The very first implication is to explain delegation itself: in his pioneering work of collusion in hierarchies, Tirole (1986) discusses the necessity of delegation of supervision by assuming that

Axiom 1: The principal ... lacks either the time or the knowledge required to supervise

the agent. (Tirole, 1986, pp.183)

In this paper, this argument is formalized as the cost (or disutilities) of investigation; then a natural reason for delegation is that the supervisor is specialized with a lower cost of investigation, so that delegation is efficient. This argument definitely holds in most cases, but this paper points out why delegation is still beneficial even if the supervisor has the expertise of investigation no better than the principal: compared with self-investigation, the delegation permits covert collusion and creates an ‘indirect signal’ to improve the efficiency of investigation.

Related Literature. This paper by its content is deeply related to studies focusing on collusion in hierarchies, such as Tirole (1986), Laffont and Tirole (1991), Kofman and Lawarrée (1993), and Laffont and Martimont (1999). However, these studies generally hold a negative attitude toward collusion and focus on collusion/coalition-proof contracts or mechanisms.

In comparison, more recent papers discover some positive effects of supervisor-agent collusion. For example, Strausz (1997a) shows that the optimal collusion-proof contract may be ex-post inefficient and create the scope for re-negotiation; Strausz (1997b) elucidates two positive effects of delegation of monitoring, one of which remains even with collusion. Kessler (2000) demonstrates that collusion between the agent and the supervisor imposes no additional costs on the principal if the report is verifiable and cannot be forged. However, this paper discovers a new beneficial channel of collusion.

Similar to this paper, some studies directly conclude that the collusion may be beneficial to the principal under some circumstances. In a dynamic framework, Olsen and Torsvik (1998) argue that the prospect of collusion can make the principal better off by generating dynamic effects. Asseyer (2020) considers the design of signal by the principal and concludes that collusive delegation sometimes performs better than direct communication.

Finally, the discussion of ‘indirect signaling’ in this paper can be a novel and interesting contribution to the area of signaling games. Most signals are considered in two-player contexts and entailed in direct communication; however, this paper discovers an indirect but effective signal in a hierarchical context: the principal would not receive this signal directly from the agent (and it will become ineffective if so), but it becomes effective and beneficial once received by the supervisor since it can help the supervisor to investigate agents in a more efficient way, and thus help the principal to reduce the necessary cost to motivate investigation.

Structure of Paper. Section 2 introduces a simple and motivating example to show the basic ideas of this paper. Section 3 considers a general model of a principal-supervisor-agent relationship, characterizes the equilibrium side payments, and gives some sufficient conditions under which corrupt delegation will take place. Some real-life applications are proposed in Section 4. Section 5 concludes. All long proofs are in Appendix.

2 A Motivating Example

Suppose there are numerous N firms (as the agents) that either produce high or low-quality products, and for convenience they are called high or low-quality firms, respectively. To launch their product legally, firms need to be licensed by the government (as the principal) with discretion to grant licenses. Suppose the government is a social welfare maximizer, then it is natural to believe that the government prefers licensing high to low-quality firms. The proportion of high-quality firms is $\gamma \in (0, 1)$.

To specify the payoffs, suppose that each licensed firm could produce at most one unit of product, and the price of each product is exogenous given by p^1 . Here we temporarily assume that all products have zero marginal cost. Hence, a licensed firm has payoff p regardless of its quality.

As for the social welfare of products, suppose that a low-quality product generates negative externality (e.g. damage to health). Specifically, suppose the social welfare increases by $Q > 0$ if a high-quality firm is licensed, while drops by $D > 0$ if a low-quality firm and its product are permitted to the market.

Suppose the type of each firm is private information, and direct investigation to reveal the types is feasible, perfect but costly. It is ‘perfect’ because it will reveal the type of the firm for certain without any possible error, and the cost of the investigation is $V > 0$ for each firm regardless of its type. However, other than the direct investigation, there is no effective signal (similar to the education signal in a job market (Spence, 1973)) to credibly reveal the type of each firm in its direct communication with the government.

By simple calculation, it could be shown that the government will investigate all firms, and grant licenses to all high-quality firms if and only if

$$V \leq \gamma Q,$$

since the benefit from identifying firms is given by γNQ , and the cost is given by NV . If the condition above fails, no investigation will occur, and all firms will be licensed if the expected payoff of licensing a firm is non-negative, that is, $\gamma Q - (1 - \gamma)D \geq 0$.

Next, suppose the investigation is delegated to a third-party supervisor who is credible but susceptible to collusion. The supervisor shares the same investigation technology with the government. Before the direct investigation, the government proposes a reward schedule to the supervisor depending on his report, and then each firm i decides on its bribe c^i paid to the supervisor.

Now we claim that the government can save the cost of investigation by proposing a reward schedule to the supervisor such that he will be rewarded by V only if his report shows that a firm is of high quality. As a result, the cost of the government is reduced to γNV ; in addition, the government becomes willing to induce investigation for some higher cost V , as now the government will delegate investigation if and only if

$$V \leq Q.$$

Note that the benefit from identifying firms is still γNQ , while the cost is reduced to γNV . Therefore, the delegation improves social welfare as long as the cost of investigation V is not too high.

This schedule is effective because the supervisor could identify all firms before direct investigation through the ‘indirect signal’ of collusion. Specifically, in the interaction between the supervisor and firms, it could be shown that there exists a separating equilibrium, in which each high-quality firm will bribe the supervisor by some $r \in (0, p]$, while the low-quality firm is inactive to collusion. As a result, anticipating the identity of the firms, the supervisor will only investigate the firms with bribe r , since only these firms are high-quality and reporting them brings reward from the government.

To show the existence of separating equilibrium, note that it could be supported by the belief of the supervisor such that the firm only with bribe not less than r is of high-quality. Given this belief, a low-

¹Some readers may notice that price does not act as a signal to differentiate high and low-quality products. It may be because consumers are not informed of the quality difference in products.

quality firm is not willing to bribe, since inducing investigation and revealing its type make no benefit. On the other hand, a high-quality firm would like to bribe r to the supervisor to induce investigation revealing its type, and thus be licensed and obtain the profit p , as long as the cost of the bribe does not exceed its benefit (that is, $r \leq p$).

More interestingly, the signal of collusion may become ineffective once it could be observed by the government, even if collusion is permitted. It is because in this case, if the collusion still acts as an effective signal, the government could identify firms simply through their collusion proposals, which will in turn invalidate this signal.

To show this formally, now we suppose that the products have some marginal cost which is increasing in its quality. Let t_H, t_L be the unit cost of a high and low-quality product, respectively, and assume that $t_L < t_H < p$. Note that if the signal of collusion becomes effective, in equilibrium high and low-quality firms are differentiated in their collusion, denoted by c_L, c_H ; then the government will license all firms with collusion c_H , and refuse to license whom with side payment c_L . Then to prevent a high or low-quality firm from imitating the other type of firms, it is required that

$$\begin{cases} p - t_H - c_H \geq -c_L \\ -c_L \geq p - t_L - c_H \end{cases}$$

which implies $p - t_H \geq p - t_L$ and contradicts the assumption that $t_H > t_L$. Therefore, separating equilibrium does not exist and the signal of collusion is shown to be ineffective.

By discussion above, we conclude that the acquiescence to collusion may be cost-saving for the government (the principal), and it turns out to be effective only when it is ‘blinded’ to her. However, the government cannot commit to being ‘blinded’ in a credible manner: once the signal of collusion becomes effective, the government can benefit from direct observation of the details of collusion if it is possible; clearly, by doing so the government can identify firms at zero cost. Therefore, it has very explicit implications for the optimal structure of supervision: in this example, only an outside and independent supervisor can validate the indirect signal of collusion; a reliable commercial partner works better than a bureaucratic supervisory agency in the government, since the latter is more vastly controlled by the decision-maker. The reliability of the third-party supervisor may not be a problem if it has a long-term cooperation relationship with the government, and some reputation mechanism is introduced.

The next section will develop a model with two types of agents, the potential moral hazard of the supervisor, imperfect investigation technology, and general form of payoff functions. The results confirm the intuition that a separating signaling equilibrium generically exists, and corrupt delegation is beneficial to the principal under some moderate conditions.

3 Model

3.1 Basic Setup

Consider a principal-supervisor-agent model with adverse selection. There are numerous agents $i \in \mathcal{I}$ (as a group we denote it A) with two possible types $\theta : \mathcal{I} \rightarrow \Theta = \{0, 1\}$, where 0 and 1 can be regarded as ‘bad’ and ‘good’ agent, respectively. For convenience, the type of agent i is denoted $\theta^i \equiv \theta(i)$. From

now on we call agent i ‘good’ or ‘bad’ if $\theta^i = 1$ or 0 . The proportion of ‘good’ agents is $\gamma \in (0, 1)$, which is commonly known by all players.

At the beginning of the game, the principal P declares a reward schedule $\phi : \mathcal{R} \rightarrow \mathbb{R}_+$, which associates each possible report $r \in \mathcal{R}$ from the supervisor S with a monetary transfer $\phi(r) \geq 0$ to the supervisor. Suppose such a schedule can be credibly enforced and publicly known to all players.

Then concerning each agent, the supervisor can choose his effort level $p : \mathcal{I} \rightarrow [0, \bar{p}]$ to investigate this private information: with probability $p^i \equiv p(i)$ he can get the evidence of the agent i ’s true type θ^i , while with probability $1 - p^i$ no evidence is found; so that p also stands for the ‘quality of the signal’². Here $\bar{p} \in (0, 1]$ stands for the maximum opportunity to obtain the evidence. The effort level cannot be observed by both principal and the agent. The cost of collecting evidence from agent i takes a linear form $p^i V$ where $V > 0$ is constant to all agents.

To incorporate the moral hazard problem of the supervisor into this model, suppose that during the investigation the supervisor can obtain the evidence at no cost with an exogenous probability $\delta \in [0, 1)$ regardless and independent of the effort level p . Hence, given the effort level p , the probability that the supervisor can obtain the evidence is given by $\delta + (1 - \delta)p$.

If receiving any evidence, the supervisor is compelled to report it to the principal³; suppose that the report is evidentiary (Milgrom, 1981) so that (i) the supervisor cannot forge fake reports about the type of agents; (ii) the supervisor can convince the principal only if he reports the evidence in his signal.⁴ Suppose the report technology is perfect so that $\mathcal{R} = \Theta$ ⁵. Moreover, we assume that the communication between the principal and supervisor is covert and cannot be observed by the agent.

Based on the assumptions above, the principal will either know the true type of agents precisely or learn nothing new from the supervisor. After receiving reports (if any), concerning the agent i the principal makes the decision $a^i \in \mathcal{A}$ (where $\mathcal{A}$ is compact) which affects the payoff of both the principal and the agent.

Finally, suppose there could be potential collusion between the agent and supervisor: after the reward schedule $\phi(\cdot)$ is announced, the agent i can covertly make a payment $c^i \geq 0$ to the supervisor.

The payoffs of players are given as follows:

1. The supervisor’s payoff (concerning agent $i \in \mathcal{I}$) has three possible components, including the reward from principal $\phi(r^i)$ from evidentiary reports, the side payment from the agent c^i , and the cost of investigation $p^i V$.
2. The agents have homogeneous preferences varying only with their types. Agent i ’s payoff takes a general form $u^A(a^i, \theta^i) - c^i$, in which the first term u^A depends on the final decision a^i by the principal, and his true type.
3. The principal’s payoff (concerning agent $i \in \mathcal{I}$) also takes a general form $u^P(a^i, \theta^i) - \phi(r^i)$, which includes the gain depending on her action a and agent i ’s true type θ , and the wage paid to the

²This information structure is also seen in Laffont and Tirole (1991), while in this model the parameter p is flexible and endogenous.

³Indeed, this assumption can be relaxed to solely truth-telling: supervisor needs not to reveal all evidence, but cannot report a fake one. However, as one may see later, in a separating equilibrium it makes no difference to the supervisor’s decision: if he anticipates reporting evidence with some type is not beneficial to him, he will choose not to investigate in advance.

⁴These assumptions can be verified in a long-term relationship between the supervisor and the principal.

⁵Generally speaking, the report space $\mathcal{R}$ can be arbitrary; however, in the spirit of the well-known revelation principle (Myerson, 1981) the principal can design a direct report mechanism such that $\mathcal{R} = \Theta$. This assumption validates the feasibility of such a mechanism.

supervisor (if she receives report).

Here we need to say something about the ‘action’ $a \in \mathcal{A}$. In the motivating example, it stands for whether to grant the license to the firm, so that one can think of $\mathcal{A} = \{\text{grant}, \text{not grant}\}$. However, in general, the action space $\mathcal{A}$ can be arbitrary: in the motivating example, if the government is accessible to another (costless) signal $x \in \mathcal{X}$ relevant to the type of the firm, the action space will become a function space $\mathcal{A} = \mathcal{F}(\mathcal{X}, \{\text{grant}, \text{not grant}\})$.

Nevertheless, we can introduce some additional assumptions about optimal action a to reduce the possible complexity of the action space $\mathcal{A}$. Let $a^*(\mu)$ be the optimal action (from the perspective of the principal) given the belief that the agent has type $\theta = 1$ with probability $\mu \in [0, 1]$ (and has type $\theta = 0$ with probability $1 - \mu$), that is,

$$a^*(\mu) = \operatorname{argmax}[(1 - \mu)u^P(a, 0) + \mu u^P(a, 1)] \quad (1)$$

Moreover, the following two additional assumptions concerning $a^*(\mu)$ are introduced into this model.

Assumption 1. $a^*(\mu)$ exists for every $\mu \in [0, 1]$.

However, it is possible that $a^*(\mu)$ is not unique; in this case, we only consider the one most favored by the agent. For convenience, we abuse the notation $a^*(\mu)$ without causing confusion.

Assumption 2. $u^A(a^*(\mu), \theta)$ is increasing in μ for every $\theta \in \Theta$.

One can see that these assumptions hold in the motivating example, in which the action space is given by $\mathcal{A} = \{\text{grant}, \text{not grant}\}$. It is because a firm favors being granted regardless of its type (quality), and it will be granted if it is believed to be high-quality (the ‘good’ type) with probability higher than some threshold level:

$$a^*(\mu) = \begin{cases} \text{grant} & \mu Q - (1 - \mu)D \geq 0 \Leftrightarrow \mu \geq \frac{D}{Q+D} \\ \text{not grant} & \mu Q - (1 - \mu)D < 0 \Leftrightarrow \mu < \frac{D}{Q+D} \end{cases}$$

Intuitively, Assumption 2 suggests that the agent benefits from action $a^*(\mu)$ for higher $\mu \in [0, 1]$; that is, he would like to be regarded as a ‘good’ agent with probability as high as possible. From now on we always implicitly suppose that Assumption 1 and 2 hold.

In summary, the timing of the game is as follows:

1. The agents i privately know their type $\theta^i \in \Theta = \{0, 1\}$.
2. The principal sets the reward scheme $\phi : \Theta \rightarrow \mathbb{R}_+$ that depends on the content of the report.
3. Each agent i decides on the side payment $c^i \geq 0$ to the supervisor.
4. The supervisor chooses his effort level to investigate agent i , represented by the quality $p : \mathcal{I} \rightarrow [0, 1]$ of his evidence.
5. If any evidence is obtained, the supervisor covertly reports it to the principal.
6. For each agent $i \in \mathcal{I}$, the principal makes the final decision $a^i \in \mathcal{A}$ based on her belief about his type.
7. The payoffs are realized.

In this game with incomplete information, the concept of equilibrium applied here is the perfect Bayesian equilibrium, which is characterized by the strategy profile of players, the belief of the supervisor

depending on the side payments from agents, and the belief of the principal depending on if any report is received. It is required that the strategy of each player is optimal given others' strategies and the belief, and the belief is consistent with the strategy profile according to the Bayes rule. Of course, with the special interest we focus on the separating equilibrium in which the strategies of agents differ only by their types.

3.2 Self-Investigation and Clean Delegation

Now we first consider the self-investigation. In this case, we only need to focus on the payoff maximization problem of the principal alone; the only decision variables are the effort levels of investigation. Since there is no effective signal revealing the type of agents, all agents are homogeneous to the principal *a priori*, and the effort level is supposed to be identical among them. Specifically, this problem is given by

$$\max_p U^P = [\delta + (1 - \delta)p] \mathbb{E}u^P(a^*(\theta^i), \theta^i) + (1 - \delta)(1 - p) \mathbb{E}u^P(\bar{a}, \theta^i) - pV \quad (2)$$

Recall that the probability of being informed of the type of agents is given by $\delta + (1 - \delta)p$. Note that here $\bar{a} = a^*(\gamma)$ is supposed to be the optimal decision conditional on no evidence being found, since

$$\bar{a} = \operatorname{argmax}(1 - p) \mathbb{E}u^P(a, \theta^i) = \operatorname{argmax}[\gamma u^P(a, 1) + (1 - \gamma)u^P(a, 0)] = a^*(\gamma).$$

Since the objective function is linear in p , it is easy to see that the solution of (2) is given by

$$p = \begin{cases} \bar{p} & V \leq (1 - \delta)\Lambda \\ 0 & V > (1 - \delta)\Lambda \end{cases} \quad (3)$$

where $\Lambda = \mathbb{E}[u^P(a^*(\theta), \theta) - u^P(a^*(\gamma), \theta)] > 0$ denotes the expected loss of the principal from ignorance about the true type of all agents. The expected payoff of the principal is given by

$$U^P = \begin{cases} \mathbb{E}u^P(a^*(\theta), \theta) - (1 - \tilde{p})\Lambda - \bar{p}V & V \leq (1 - \delta)\Lambda \\ \mathbb{E}u^P(a^*(\theta), \theta) - (1 - \delta)\Lambda & V > (1 - \delta)\Lambda \end{cases} \quad (4)$$

where $\tilde{p} = \delta + (1 - \delta)\bar{p}$.

Next, we consider the clean delegation of investigation, in which the collusion between the supervisor and the agent can be prohibited. Note that because delegation does not generate effective signals, the supervisor still cannot identify agents in advance, and differentiate his effort level according to the type of agents. Hence, the optimal reward schedule $\{\phi(\theta)\}_{\theta \in \Theta}$ follows from the following maximization problem:

$$\begin{aligned} \max_{\phi(\cdot)} U^P &= [\delta + (1 - \delta)p] \mathbb{E}[u^P(a^*(\theta^i), \theta^i) - \phi(\theta^i)] + (1 - \delta)(1 - p) \mathbb{E}u^P(a^*(\gamma), \theta^i) \\ \text{s.t. } p &= \operatorname{argmax}[\delta + (1 - \delta)p] \mathbb{E}\phi(\theta^i) - pV \end{aligned} \quad (5)$$

Note that the constraint in (5) can be solved into

$$p = \begin{cases} \bar{p} & \mathbb{E}\phi(\theta^i) \geq V/(1-\delta) \\ 0 & \mathbb{E}\phi(\theta^i) < V/(1-\delta) \end{cases}$$

and thus the solution of (5) is given by $\mathbb{E}\phi(\theta^i) = \frac{V}{1-\delta}$ if V is moderate such that inducing investigation is profitable. The expected payoff of the principal is then given by

$$U^P = \begin{cases} \mathbb{E}u^P(a^*(\theta), \theta) - (1-\tilde{p})\Lambda - \tilde{p}V/(1-\delta) & V \leq (1-\delta)^2(\bar{p}/\tilde{p})\Lambda \\ \mathbb{E}u^P(a^*(\theta), \theta) - (1-\delta)\Lambda & V > (1-\delta)^2(\bar{p}/\tilde{p})\Lambda \end{cases} \quad (6)$$

It is easy to see that compared to self-investigation, the necessary cost to motivate the supervisor's investigation rises from V to $V/(1-\delta)$, which is increasing in δ . This result is intuitive in the context of the moral hazard: the uncertainty of investigation measured by δ creates an incentive for the supervisor to 'shirk' by taking effort level $p = 0$, thereby increasing the reward paid by the principal to offset this force. Therefore, by comparing (6) with (4) we conclude that

Proposition 1. *If the supervisor has identical investigation technology to the principal, then the principal prefers self-investigation to clean delegation.*

3.3 Corrupt Delegation

Now we turn to consider corrupt delegation, in which the supervisor is susceptible to collusion with the agent.

First, suppose the optimal decision of the principal is a^0 if no investigation and report are made; it is easy to see that $a^0 = a^*(\mu^0)$ for some $\mu^0 \in [0, 1]$. On the other hand, it is clear that the principal will take decision $a^*(\theta)$ if the report reveals the type of the agent as $\theta \in \Theta$ since the report is evidentiary and cannot be false.

We start by finding the best response of the supervisor given reward schedule $\phi : \Theta \rightarrow \mathbb{R}_+$ in a separating equilibrium. It is clear that if the supervisor can identify the type of agents in advance (through the signal of collusion), the payoff function of the supervisor is still linear in effort level p , so that he will investigate the agent i with maximum effort $\bar{p}$ if and only if

$$\phi(\theta^i) \geq V/(1-\delta) \quad (7)$$

and he will exert zero effort if (7) fails.

With this result in mind, we immediately know that (i) if the supervisor can identify agents through the signal of collusion, he will take identical decisions (the effort of investigation) concerning the same type of agents, and (ii) for agents of each type $\theta \in \Theta$, the principal can induce either maximum effort (and get evidence with probability $\tilde{p}$) or zero effort (and get evidence with probability δ). Moreover, to minimize the necessary reward, in these two cases the optimal rewards are given by $\phi(\theta) = V/(1-\delta)$ and $\phi(\theta) = 0$, respectively. Hence, in equilibrium for every θ we have either $\phi(\theta) = 0$ or $V/(1-\delta)$.

Next, we turn to discuss the interaction between the supervisor and agents. The following proposition characterizes the separating equilibria of this signaling game between them. Note that in a separating

equilibrium it is required that the side payment c^i of the agent i only depends on type $\theta^i \in \Theta$ (so that we denote $c(\theta) = c^i$ if $\theta^i = \theta$) since agents with same type are homogeneous, and $c(0) \neq c(1)$.

Proposition 2. *In a separating equilibrium,*

1. *The agent believed (by the supervisor) to be of ‘bad’ type will be investigated with effort $p = 0$;*
2. *The agent believed (by the supervisor) to be of ‘good’ type will be investigated with effort $p = \bar{p}$;*
3. *The equilibrium strategy profile $\{c(0), c(1)\}$ is given by*

$$\begin{cases} c(0) = 0 \\ c(1) \leq (1 - \delta)\bar{p}[u^A(a^*(1), 1) - u^A(a^0, 1)] \end{cases} \quad (8)$$

or

$$\begin{cases} c(0) \leq (1 - \delta)\bar{p}[u^A(a^0, 0) - u^A(a^*(0), 0)] \\ c(1) = 0 \end{cases} \quad (9)$$

Proof. See Appendix. □

The intuition of this argument is as follows: because the investigation and report are truth-telling and reveal the type of the agents, it follows that ‘good’ agents are willing to be investigated, while ‘bad’ agents are not. Since collusion with the form of side payments is purely dissipative (similar to uninformative advertisement; see Milgrom and Roberts, 1986), a ‘bad’ agent would bribe less to the supervisor if the chance of being investigated could be reduced by doing so; on the other hand, a ‘good’ agent may also bribe less to the supervisor if the chance could increase by doing so. Hence, one may expect that either a ‘good’ or ‘bad’ agent will find it beneficial to deviate to zero collusion for any possible belief the supervisor holds in this situation.

As a result, in a separating equilibrium there must be either $c(1) = 0$ for the ‘good’ agents, or $c(0) = 0$ for the ‘bad’ agents, otherwise either a ‘good’ or ‘bad’ agent would find profitable deviation to zero collusion. For the agents from the other group, it could be shown that some positive collusion can be supported by some proper belief of the supervisor; the only constraint is that the side payment in equilibrium is bounded from above, in case the agent from the other group does not find it too costly to induce (or avoid) investigation and choose zero collusion instead.

By Proposition 2 we immediately have the following two conclusions. Note that by Bayes’ rule, when ‘good’ agents are investigated with effort $\bar{p}$ while ‘bad’ agents are with zero effort, an agent without evidence found is of ‘good’ type with probability $\mu^0 = \frac{\gamma(1-\bar{p})}{(1-\gamma)+\gamma(1-\bar{p})}$.

Corollary 1. *A separating equilibrium exists if*

$$u^A(a^*(1), 1) > u^A(a^0, 1) \quad \text{or} \quad u^A(a^0, 0) > u^A(a^*(0), 0)$$

where $a^0 = a^*(\mu^0) = a^*\left(\frac{\gamma(1-\bar{p})}{(1-\gamma)+\gamma(1-\bar{p})}\right)$.

Corollary 2. *Suppose investigation technology is perfect such that $\bar{p} = \bar{p} = 1$. Then a separating equilibrium exists if*

$$u^A(a^*(1), 1) > u^A(a^*(0), 1).$$

In a separating equilibrium the strategy profile $\{c(0), c(1)\}$ is given by

$$\begin{cases} c(0) = 0 \\ c(1) \leq (1 - \delta)[u^A(a^*(1), 1) - u^A(a^*(0), 1)] \end{cases}$$

Corollary 2 could refer to the motivating example in Section 2, in which no uncertainty of investigation is involved (i.e. $\delta = 0$), and high-quality firms benefit from license (note that $a^*(1) = \text{grant}$, $a^*(0) = \text{not grant}$).

Indeed, one can find some real-life correspondence to these separating equilibria. Note that the $c(1)$ in (8) refers to a ‘threshold’ level of bribe to be investigated (with effort), while $c(0)$ in (9) resembles a ‘threshold’ level of bribe in order not to be investigated (with effort). In this way, the former separating equilibrium can be characterized by a ‘reluctant’ supervisor: he will investigate only the agent with the bribe not lower than a threshold level. When this rule prevails among agents, the ‘good’ agents find it optimal to bribe at this level, while ‘bad’ agents do not.

On the other hand, the separating equilibrium in (9) depicts a ‘blackmailing’ supervisor: he would like to investigate all agents, except the ones with bribes not lower than a threshold level. As a result, only the ‘bad’ agents would like to pay the bribe to escape the investigation.

Even though these separating equilibria are appealing to our analysis, some pooling equilibria may also exist in this signaling game, in which $c(0) = c(1)$ and the signal of collusion becomes ineffective, so that the supervisor cannot distinguish agents through their collusion decisions. Fortunately, the following proposition makes them trivial and could be reasonably neglected.

Proposition 3. *In a pooling equilibrium we have $c(0) = c(1) = 0$.*

Proof. See Appendix. □

The intuition of this argument is similar to one in Proposition 2: the situation in which all agents propose strictly positive side payments cannot sustain a pooling equilibrium, since either a ‘good’ or ‘bad’ agent is willing to deviate to zero collusion for any possible belief of the supervisor concerning those agents.

As we have seen, in this signaling game the complexity of equilibria is involved, and a natural question that which equilibrium (or equilibria) will be observed emerges. Indeed, it is reasonable to believe that the separating equilibria will realize instead of the pooling equilibria since only the former ones involve positive collusion which is beneficial to the supervisor. As a result, the supervisor is willing to set a belief supporting a separating equilibrium, and such belief will be self-confirmed. Hence, from now on we only consider separating equilibria in which the supervisor is capable to identify agents before any direct investigation.

Finally, we turn to discuss the optimal reward schedule proposed by the principal. By Proposition 2 and discussion above, to sustain a separating equilibrium she needs to set $\phi(1) = V/(1 - \delta)$, $\phi(0) = 0$; otherwise, we will observe a pooling equilibrium and the result will be identical to the case of clean delegation. The expected payoff of the principal in the former case is thus given by

$$U^P = \gamma \tilde{p}[u^P(a^*(1), 1) - V/(1 - \delta)] + \gamma(1 - \tilde{p})u^P(a^0, 1) + (1 - \gamma)u^P(a^0, 0) \quad (10)$$

where $a^0 = a^*(\mu^0) = a^*\left(\frac{\gamma(1-\bar{p})}{(1-\gamma)+\gamma(1-\bar{p})}\right)$. By comparing (10) with (4) we have

Proposition 4. *The principal prefers corrupt delegation to self-investigation if*

$$\frac{(1-\gamma)(1-\delta)\bar{p}-\gamma\delta}{1-\delta}V \geq (1-\tilde{p})\mathbb{E}[u^P(a^*(\gamma), \theta) - u^P(a^*(\mu^0), \theta)] + (1-\gamma)\tilde{p}[u^P(a^*(0), 0) - u^P(a^*(\mu^0), 0)] \quad (11)$$

where the expectation is taken over the whole population, and $\mu^0 = \frac{\gamma(1-\bar{p})}{(1-\gamma)+\gamma(1-\bar{p})}$.

As shown in (11), whether the principal prefers corrupt delegation depends on two main effects of the introduction of collusion as a signal: first, if this signal is effective and the supervisor can identify the type of each agent before any direct investigation, the principal can improve her reward schedule and induce the supervisor to find evidence of the ‘good’ agents only. In other words, in the equilibrium the set of agents being investigated can be narrowed. However, since the moral hazard elevates the unit cost of the principal to motivate effort, the direction of this effect may be mixed: it becomes negative if δ is high and so is the unit cost, while it can be a positive effect if γ is small and so is the set of agents being investigated. The direction and size of this effect are shown on the left-hand side of (11).

The second effect of corrupt delegation shown on the right-hand side of (11) reflects the change in optimal action a^0 compared to the self-investigation. Recall that in the case of self-investigation in which the effort cannot depend on the type of agents, the distribution of types among agents without evidence is identical to the whole population, so that the $a^0 = a^*(\gamma)$. However, under corrupt delegation only ‘good’ agents will be investigated, and there will be changes in action taken toward two groups of agents: for all agents from whom no evidence to be found (with share $1 - \tilde{p}$ of the population), the action taken toward them changes from $a^*(\gamma)$ to $a^*(\mu)$; for ‘bad’ agents from whom the evidence to be found in the case of self-investigation (with share $(1-\gamma)\tilde{p}$ of the population), the action changes from $a^*(0)$ to $a^*(\mu)$. This effect is certainly non-positive to the principal, since now the actions taken toward these two groups of agents cannot be differentiated to their first-best level, but have to be identical and equal to the second-best level.

In particular, the second effect is nullified in two cases: first, the investigation technology is perfect in the sense that he could perfectly identify all agents if he exerts the maximum effort; second, the action space is relatively simple such that the optimal action taken toward the agent is constant if he is believed to be of ‘good’ type with low probability (e.g. in the motivating example with only two possible actions, if the expected payoff of licensing a firm is negative according to the prior, that is, $\gamma Q - (1-\gamma)D < 0$, then a firm will not be granted whenever it is believed to be high-quality with probability not greater than γ). Under these circumstances, we only need to focus on the trade-off between higher unit cost due to the moral hazard, and a narrower set of agents to be investigated with effort.

Corollary 3. *Suppose investigation technology is perfect such that $\bar{p} = \tilde{p} = 1$. Then the principal prefers corrupt delegation to self-investigation if*

$$u^A(a^*(1), 1) > u^A(a^*(0), 1) \quad (12)$$

and

$$\gamma + \delta \leq 1. \quad (13)$$

Proof. Note that when $\tilde{p} = 1$ the right-hand side in (11) becomes zero since now $\mu = 0$. (12) immediately follows from Corollary 1 by letting $a^0 = a^*(\mu^0) = a^*(0)$, and (13) follows from $(1-\gamma)(1-\delta) - \gamma\delta \geq 0$. $\square$

Corollary 4. *Suppose $a^*(\mu)$ is constant for every $\mu \in [0, \gamma]$. Then the principal prefers corrupt delegation to self-investigation if (12) and (13) hold.*

Proof. The proof is same as in Corollary 3, because the right-hand side in (11) also becomes zero as $a^*(\gamma) = a^*(\mu^0) = a^*(0)$. $\square$

3.4 Corrupt Self-Investigation?

In previous subsections, we have shown that collusion between supervisor and agents can be viewed as an ‘indirect signal’ as agents will propose side payments depending on their types, so that in the most ideal case the supervisor can distinguish agents before any direct investigation.

It is important to note that a key underlying assumption is that the collusion itself is neither directly observed by the principal, nor revealed by the supervisor through any report. This assumption not only validates the ‘secrecy’ of collusion but also sustains the functioning of bribery as the ‘indirect signal’. This subsection will show that this signal will become ineffective once the supervisor is compelled to report the detail of collusion.

Indeed, the case of transparent collusion is essentially equivalent to so-called ‘corrupt self-investigation’: that is, there is no delegation of investigation, and agents can bribe the principal directly by proposing the same side payments as in corrupt delegation. Compared with the corrupt delegation with observable collusion, the principal has the same information structure but there will be no need to pay the supervisor.

What will happen under corrupt self-investigation? The most significant difference is that since the principal is additionally informed of the detail of collusion, she can simply distinguish agents by side payments instead of direct investigation. If they differ by agents’ types in equilibrium, then we have a separating equilibrium, in which the collusion becomes a ‘direct signal’ like those in typical signaling games.

To formally characterize this situation, consider a modified version of the model setup without a third-party supervisor, and the possible collusion occurs between the principal and agents. After being informed of his true type, the agent i decides on the side payment $c^i \geq 0$ directly flows to the principal, and then the principal decides whether to investigate and finally makes decision $a^i \in \mathcal{A}$. Moreover, an additional assumption about the preference of agents is given below.

The following proposition gives a moderate condition under which the signal of collusion becomes ineffective.

Proposition 5. *If $u^A(a^*(1), 1) + u^A(a^*(0), 0) < u^A(a^*(0), 1) + u^A(a^*(1), 0)$, then this direct signaling game only exists pooling equilibria in which $c(0) = c(1) = 0$.*

Proof. See Appendix. $\square$

One may notice that the assumption in Proposition 5 can be regarded as the inverse of a discrete version of the well-known Spence-Mirrlees condition, so that a separating equilibrium cannot be sustained. It is easy to see that the condition in Proposition 5 holds in the motivating example, in which

the profit of being licensed is greater for low-quality firms as they have a lower marginal cost. Indeed, the intuition of Proposition 5 is as same as the reasoning in the motivating example: it is impossible to sustain a differential in side payments between ‘good’ and ‘bad’ agents, since they may find it beneficial to deviate by imitating the strategy of agents with the other type.

In a pooling equilibrium, the principal cannot distinguish agents through the signal of collusion, so that the principal’s investigation decision and expected payoff still follow from (3) and (4), respectively. In other words, when the collusion can be directly observed by the principal, the situation will degenerate into the case of self-investigation. As a result, Proposition 4 is also applicable to the comparison between corrupt delegation and corrupt self-investigation.

3.5 Contractible Collusion

In this model, the collusion takes the monetary form of side payments before investigation. However, one may notice that these side payments are purely dissipative: their only role is to credibly reveal if the agents are ‘good’ or ‘bad’; from an outsider’s perspective, it also acts as the necessary cost to trigger or prevent investigation. However, readers may worry that it does not capture the essence of ‘collusion’ in actual relationships between the supervisor and agents: side payments are a part of the ‘implicit contract’ between them, and the agent can commit to rewarding the supervisor depending on his performance.

To incorporate contractible collusion into this model, now we suppose that each agent i can propose a ‘collusion contract’ $c^i : \mathcal{A} \rightarrow \mathbb{R}_+$ which associates each possible action $a \in \mathcal{A}$ taken by the principal to a non-negative side payment $c^i(a)$ to the supervisor. To simplify the analysis, suppose the investigation technology is perfect; then only the separating equilibrium in (8) remains in the original model. Moreover, in this case, the action taken by the principal could only be either $a^*(0)$ or $a^*(1)$, and a collusion contract only needs to cover these two contingencies. The following proposition shows a similar result to Proposition 2.

Proposition 6. *Suppose the investigation technology is perfect. Then the following strategy profile yields a separating equilibrium: for a ‘good’ agent i ,*

$$\begin{cases} c^i(a^*(0)) = 0 \\ c^i(a^*(1)) = c(1) \leq (1 - \delta)[u^A(a^*(1), 1) - u^A(a^*(0), 1)] \end{cases}, \quad \theta^i = 1$$

and for a ‘bad’ agent i ,

$$\begin{cases} c^i(a^*(0)) = 0 \\ c^i(a^*(1)) = 0 \end{cases}, \quad \theta^i = 0$$

Similar to the separating equilibrium characterized by (8), one can see that this strategy profile can be supported by the belief that all agent i with collusion $c^i(a^*(1)) < c(1)$ are of ‘bad’ type, and otherwise ‘good’ type. However, in this case, both $c^i(a^*(0))$ and $c^i(a^*(1))$ can act as signals, and this fact enlarges the set of separating equilibria: for example, in an equilibrium all ‘good’ agents propose strictly positive side payments if he is treated as a ‘bad’ agent by the principal; then a ‘bad’ agent will not imitate since he knows that this side payment will be realized for him, but not for those ‘good’ agents. However, one can see that these separating equilibria bring zero benefit to the supervisor, as

the contingencies with positive side payments will never be realized. Therefore, by the similar logic to exclude pooling equilibria in the original model, it is reasonable to consider the separating equilibrium characterized in Proposition 6 only.

4 Applications

Outsourcing of Monitoring. As in the motivating example, under some circumstances the ‘outsourcing’ of monitoring may perform better by reducing the cost from the principal. Except for the reliability of the third-party supervisor, from (11) the following determinants also matter: the loss from inaccurate action (the right-hand side of (11)) and moral hazard (characterized by δ) should be moderate, while the cost of investigation V is supposed to be sufficiently high.

These determinants suggest that whether corrupt delegation is efficient remains contingent, and depends on the nature and effects of investigation. To show this, it is helpful to compare two cases relevant to the monitoring structure of the government: as the motivating example suggests, the ‘outsourcing’ of monitoring via business consultants or ‘think tanks’ may be a common practice with respect to granting licenses to firms, but it is hardly the case when the government manages its officials; instead, collusion in the bureaucracy is usually not tolerable and to be fought. Besides moral concerns, corrupt delegation may turn out not to be efficient in the latter case. It may be because it is easier to monitor bureaucrats than commercial corporations, and the cost of investigation V is relatively lower; alternatively, the inaccuracy in the arrangement of officials brings greater loss to social welfare.

Executive Compensations. Another interesting application of this model is to explain executive compensations within a business organization. Most discussions account this phenomenon as either a remedy to potential ‘agency problem’ from the managers, or simply the result of their rent extraction (Bebchuk and Fried, 2003). If one incorporates the principal-supervisor-agent relationship into the interaction among board, manager, and branches (or business projects), the executive compensation is associated with both explanations: the board designs contracts to overcome moral hazard and save the cost of investigation, while the CEOs are also receiving bribes from employees.

The correspondence between the model and their relationship is as follows: the board (principal) decides on the resource allocation (for example, budgets) among various branches (agents) within the corporation, and the manager (supervisor) is responsible to identify (‘good’) branches with better performance, and report them to the board. According to the model, one of the following may be the case: only ‘good’ branches will try to bribe the CEO to be appreciated (say, raised budget) by the board, or only ‘bad’ branches will bribe him in order not to be punished (say, cut budget or even dismissed) by the board.

Contrary to the traditional view, the model suggests that such collusion between the manager and employees might not be detrimental to the corporation, as long as the report from the manager remains reliable. As a result, this gives an additional explanation for delegation of management in modern corporations, since it reduces the necessary cost from adverse selection, and it has been a long theoretical tradition that the firm acts as an effective structure to minimize the ‘transaction cost’.

Moreover, we can find some other evidence in accordance with the predictions of this model: for example, the model shows that the signal of collusion will become ineffective once observed by the principal. In the context of executive compensation, this fact accounts for the autonomy of CEOs: from

the perspective of the board, it is an effective way to commit not to observe and take the advantage of the signal of collusion. Indeed, in reality the rent extraction of the managers will be usually tolerated and seldom exposed in the business world.

5 Conclusion

This paper manages to construct a general model within the basic framework of the principal-supervisor-agent relationship, in which the supervisor manages the collection and flow of the agent's private information, while both the principal and agent may be capable to influence the supervisor's incentives.

The most prominent conclusion of this paper is that corrupt but credible delegation could create effective signals through collusion, which is usually thought to be detrimental to the interest of the principal. However, in an adverse selection setting our model has verified how preventing collusion 'may result in efficiency losses' as Tirole (1986) worries. Note that this is not to say that collusion and bribery are morally good and acceptable; this conclusion provides a more solid explanation of the prevalence of delegation without reference to the positive effects outside of the framework of the principal-supervisor-agent relationship.

Undoubtedly, this paper is far from the end of studies on collusion in a general principal-supervisor-agent relationship; the model in this paper also calls for completion and extensions. The bargaining process of the reward schedules is omitted in this model; the production structure of information through direct investigation can be further generalized; the moral hazard of the agent can also be introduced into this model's setup. Most importantly, it might be true that the main results of this model could hold even if the supervisor has imperfect credibility with his report. I believe that some of these extensions can inspire new insights into how credible signals can be created and work in a hierarchical relationship.

Appendix

Proof of Proposition 2.

1. By discussion at the beginning of Section 3.3, we know that the effort level p can be either 0 or $\bar{p}$. We first show that if this argument does not hold, then in a separating equilibrium we have $c(0) = 0$. If $c(0) > 0$, then given the default action a^0 , to ensure that a ‘bad’ agent i will not deviate to $c^i = 0$, it is required that

$$\tilde{p}u^A(a^*(0), 0) + (1 - \tilde{p})u^A(a^0, 0) - c(0) \geq [\delta + (1 - \delta)\hat{p}(0)]u^A(a^*(0), 0) + (1 - \delta)(1 - \hat{p}(0))u^A(a^0, 0)$$

where $\hat{p}(c) \in \{0, \bar{p}\}$ ⁶ is the effort level of the supervisor concerning agent i given $c^i = c$ in equilibrium. However, the inequality above fails for every $\hat{p}(0) \in [0, \bar{p}]$ and $a^0 = a^*(\mu^0)$ for $\mu^0 \in [0, 1]$, since $u^A(a^*(0), 0) \leq u^A(a^0, 0)$ and $\tilde{p} \geq \delta + (1 - \delta)\hat{p}(0)$.

Next, if agent believed to be of ‘bad’ type will be investigated, and $c(0) = 0$ in equilibrium, then an ‘good’ agent i ($\theta^i = 1$) will find it beneficial to deviate to $c^i = 0$ and be investigated with effort $\bar{p}$, since

$$[\delta + (1 - \delta)\hat{p}(c(1))]u^A(a^*(1), 1) + (1 - \delta)(1 - \hat{p}(c(1)))u^A(a^0, 1) - c(1) < \tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1).$$

This contradicts the definition of separating equilibrium. Therefore we complete the proof of (1).

2. If not, by (1) all agents will be investigated with effort $p = 0$ in a separating equilibrium. In this case, the effort decision of the supervisor is invariant to any collusion from the agents, and it is natural to believe that the only Nash equilibrium strategy profile is that $c(0) = c(1) = 0$, which contradicts the definition of separating equilibrium.
3. Suppose $c(0) < c(1)$; if $c(0) > 0$, then to ensure that a ‘bad’ agent i not to deviate to $c^i = 0$, we have

$$\delta u^A(a^*(0), 0) + (1 - \delta)u^A(a^0, 0) - c(0) \geq [\delta + (1 - \delta)\hat{p}(0)]u^A(a^*(0), 0) + (1 - \delta)(1 - \hat{p}(0))u^A(a^0, 0)$$

which implies $\hat{p}(0) \neq 0$ and $\hat{p}(0) = \bar{p}$; however, then a ‘good’ agent i will find it beneficial to deviate to $c^i = 0$ since

$$\tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1) - c(1) < \tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1).$$

Hence, in a separating equilibrium we have $c(0) = 0$ and $\hat{p}(0) = 0$. To ensure a ‘good’ agent i not deviate to $c^i = 0$, it is required that

$$\tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1) - c(1) \geq \delta u^A(a^*(1), 1) + (1 - \delta)u^A(a^0, 1).$$

By simple calculation we have $c(1) \leq (1 - \delta)\bar{p}[u^A(a^*(1), 1) - u^A(a^0, 1)]$. This strategy profile can be supported by the belief of the supervisor in which all agents i with collusion $c^i < c(1)$ are of ‘bad’ type, and otherwise ‘good’ type.

⁶It is because for any belief of the supervisor concerning agent i ’s type θ^i given his collusion $c^i \geq 0$, his optimal investigation strategy is given by $\hat{p}(c^i) = 0$ if $\mathbb{E}\phi(\theta^i) < V/(1 - \delta)$, and $\hat{p}(c^i) = \bar{p}$ if $\mathbb{E}\phi(\theta^i) \geq V/(1 - \delta)$, where the expectation is taken over the belief.

By similar method, if $c(1) < c(0)$ it could be shown that $c(1) = 0$, and to ensure that a ‘bad’ agent i not to deviate to $c^i = 0$, it is required that

$$\delta u^A(a^*(0), 0) + (1 - \delta)u^A(a^0, 0) - c(0) \geq \tilde{p}u^A(a^*(0), 0) + (1 - \tilde{p})u^A(a^0, 0).$$

By simple calculation we have $c(0) \leq (1 - \delta)\bar{p}[u^A(a^0, 0) - u^A(a^*(0), 0)]$. This strategy profile can be supported by the belief of the supervisor in which all agents i with collusion $c^i \leq c(0)$ are of ‘good’ type, and otherwise ‘bad’ type. □

Proof of Proposition 3. It is because if $c(0) = c(1) = c > 0$, then either a ‘bad’ or a ‘good’ agent i will find it beneficial to deviate to $c^i = 0$. To show this formally, first suppose that all agents with collusion $c > 0$ will be investigated with effort $p = \bar{p}$. To ensure that a ‘good’ agent i will not deviate to $c^i = 0$, it is required that

$$\tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1) - c \geq [\delta + (1 - \delta)\hat{p}(0)]u^A(a^*(1), 1) + (1 - \delta)(1 - \hat{p}(0))u^A(a^0, 1).$$

The inequality above holds only if $\hat{p}(0) = 0$ (recall that $\hat{p}(c) \in \{0, \bar{p}\}$ for every $c \geq 0$). However, in this case, a ‘bad’ agent i will find it beneficial to deviate to $c^i = 0$ since

$$\tilde{p}u^A(a^*(0), 0) + (1 - \tilde{p})u^A(a^0, 0) - c < \delta u^A(a^*(0), 0) + (1 - \delta)u^A(a^0, 0)$$

as $\tilde{p} \geq \delta$ and $u^A(a^*(0), 0) \leq u^A(a^0, 0)$.

Next, suppose that all agents with collusion $c > 0$ will be investigated with effort $p = 0$. To ensure that a ‘bad’ agent i will not deviate to $c^i = 0$, it is required that

$$\delta u^A(a^*(0), 0) + (1 - \delta)u^A(a^0, 0) - c \geq [\delta + (1 - \delta)\hat{p}(0)]u^A(a^*(0), 0) + (1 - \delta)(1 - \hat{p}(0))u^A(a^0, 0).$$

The inequality above holds only if $\hat{p}(0) = \bar{p}$. However, in this case, a ‘good’ agent i will find it beneficial to deviate to $c^i = 0$ since

$$\delta u^A(a^*(1), 1) + (1 - \delta)u^A(a^0, 1) - c < \tilde{p}u^A(a^*(1), 1) + (1 - \tilde{p})u^A(a^0, 1)$$

as $\tilde{p} \geq \delta$ and $u^A(a^*(1), 1) \geq u^A(a^0, 1)$.

Finally, this pooling equilibrium can be supported by the strategy that all agents with collusion $c > 0$ are investigated with effort $p = 0$; this strategy can be supported by, for example, the belief that all agents with positive collusion are of ‘bad’ type, and report revealing ‘bad’ type will not be rewarded by the principal. □

Proof of Proposition 5. First, note that in an (either separate or pooling) equilibrium there must be $c(0) = 0$; if not, then a ‘bad’ agent i will find it beneficial to deviate to $c^i = 0$ since

$$u^A(a^*(0), 0) - c(0) < p(0)u^A(a^*(0), 0) + (1 - p(0))u^A(a^*(1), 0)$$

for all $p(0) \in [0, 1]$, where $p(c)$ denotes the probability that the supervisor believes that an agent with

$c^i = c$ is of ‘good’ type.

Suppose $c(1) > 0$; then to ensure that each type of agent will not find it beneficial to imitate each other, we have

$$u^A(a^*(1), 1) - c(1) \geq u^A(a^*(0), 1)$$

$$u^A(a^*(0), 0) \geq u^A(a^*(1), 0) - c(1)$$

which implies $u^A(a^*(1), 1) + u^A(a^*(0), 0) \geq u^A(a^*(0), 1) + u^A(a^*(1), 0)$, where has a contradiction. Therefore, there is no separating equilibrium in this signaling game.

Finally, the pooling equilibrium in which $c(0) = c(1) = 0$ can be supported by the belief of the principal that any agent i with $c^i > 0$ is of ‘bad’ type with probability 1. $\square$

References

- [1] Asseoyer, Andreas. “Collusion and delegation under information control.” *Theoretical Economics* 15.4 (2020): 1547-1586.
- [2] Baiman, Stanley, John H. Evans, and Nandu J. Nagarajan. “Collusion in auditing.” *Journal of Accounting Research* 29.1 (1991): 1-18.
- [3] Bebchuk, Lucian Arye, and Jesse M. Fried. “Executive compensation as an agency problem.” *Journal of Economic Perspectives* 17.3 (2003): 71-92.
- [4] Celik, Gorkem. “Mechanism design with collusive supervision.” *Journal of Economic Theory* 144.1 (2009): 69-95.
- [5] Kessler, Anke S. “On monitoring and collusion in hierarchies.” *Journal of Economic Theory* 91.2 (2000): 280-291.
- [6] Kofman, Fred, and Jacques Lawarrée. “Collusion in hierarchical agency.” *Econometrica: Journal of the Econometric Society* (1993): 629-656.
- [7] Laffont, Jean-Jacques, and David Martimort. “Separation of regulators against collusive behavior.” *The Rand Journal of Economics* (1999): 232-262.
- [8] Laffont, Jean-Jacques, and Jean Tirole. “The politics of government decision-making: A theory of regulatory capture.” *The Quarterly Journal of Economics* 106.4 (1991): 1089-1127.
- [9] Milgrom, Paul R. “Good news and bad news: Representation theorems and applications.” *The Bell Journal of Economics* (1981): 380-391.
- [10] Milgrom, Paul, and John Roberts. “Price and advertising signals of product quality.” *Journal of Political Economy* 94.4 (1986): 796-821.
- [11] Myerson, Roger B. “Optimal auction design.” *Mathematics of Operations Research* 6.1 (1981): 58-73.
- [12] Olsen, Trond E., and Gaute Torsvik. “Collusion and renegotiation in hierarchies: A case of beneficial corruption.” *International Economic Review* (1998): 413-438.
- [13] Tirole, Jean. “Hierarchies and bureaucracies: On the role of collusion in organizations.” *Journal of Law, Economics, & Organization*, Vol.2, No.2 (1986): 181-214.
- [14] Spence, Michael. “Job Market Signaling.” *The Quarterly Journal of Economics* 87.3 (1973): 355-374.
- [15] Strausz, Roland. “Collusion and renegotiation in a principal-supervisor-agent relationship.” *Scandinavian Journal of Economics* 99.4 (1997): 497-518.
- [16] Strausz, Roland. “Delegation of monitoring in a principal-agent relationship.” *The Review of Economic Studies* 64.3 (1997): 337-357.