

DISCUSSION PAPER

2004

Pure strategy and no-externalities with multiple agents: A comment

Andrea ATTAR¹, Eloisa CAMPIONI², Gwenaël PIASER³ and Uday RAJAN⁴

July 2004

Abstract

In this note we consider a basic property of common agency models: pure strategy equilibria of games where principals compete in direct mechanisms are robust to the possibility that principals might deviate and use more complex indirect mechanisms to design their contracts. We show that this property can be generalized to multi-principal multi-agent models.

JEL Classification: D82.

¹CORE Université Catholique de Louvain and DSE Università di Roma, La Sapienza; attar@core.ucl.ac.be

²CORE Université Catholique de Louvain and DIS Università di Roma, La Sapienza; campioni@core.ucl.ac.be

³Università Ca' Foscari di Venezia; piaser@core.ucl.ac.be

⁴University of Michigan Business School; urajan@bus.umich.edu

We thank Cl. d'Aspremont, E. Minelli, N. Porteiro and P. Siconolfi for helpful discussions.
This text presents research results of the Belgian Program on Interuniversity Poles of Attraction initiated by the Belgian State, Prime Minister's Office, Science Policy Programming. The scientific responsibility is assumed by the authors.

1 Introduction

In a recent paper, Peters (2004) provides two thought-provoking examples in multi-principal, multi-agent settings. His first example suggests that “*In a multiple agency environment [...] pure strategy equilibria are not robust against the possibility that principals might deviate to more complex indirect mechanisms*”¹. The second example shows that a “no externalities” assumption (see Peters, 2003) sufficient to imply the revelation principle in a multi-principal, single agent context fails to do so when there are many agents.

In this comment, we first highlight two critical features of the first example: (i) the timing of the interaction between principals and the agents has agents committing to an effort level, and (ii) the principals are restricted to deterministic strategies. Allowing for non-deterministic mechanisms, we then prove a general pure strategy theorem for multi-principal, multi-agent games, that shows that every equilibrium in direct mechanisms remains an equilibrium if principals choose more complicated indirect mechanisms.

We begin by describing example 1 in Peters (2004). There are two principals and two agents in an environment with pure moral hazard: agents take private actions that principals cannot observe. Hence, the contracts proposed by the latter cannot depend on the effort of the agents.

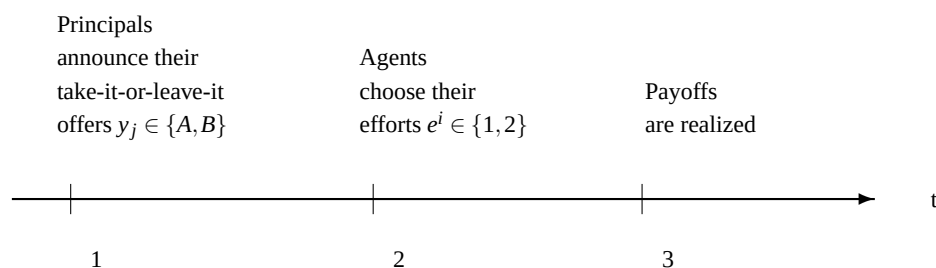


Figure 1: Timing 1, direct mechanism interaction

As shown in Figure 1, the direct mechanism game begins with Principals P_1 and P_2 simultaneously choosing actions in the space $\{A, B\}$. Then, Agents A_1 and A_2 observe the principals’ choices, and simultaneously choose a level of effort in the set $\{1, 2\}$. The normal form of the game is exhibited in Table 1. The principals’ choices affect which cell of the larger matrix is chosen. The agents then play the 2×2 subgame in that cell.

¹Peters, 2004, p. 184.

	A	B
A	$ \begin{array}{cc} & e=1 & e=2 \\ e=1 & * & * \\ e=2 & * & * \end{array} $	$ \begin{array}{cc} & e=1 & e=2 \\ e=1 & \left(\frac{7}{8}, \frac{7}{8}, -\frac{3}{2}, \frac{5}{4}\right) & \left(\frac{7}{8}, \frac{7}{8}, \frac{5}{4}, -\frac{3}{2}\right) \\ e=2 & \left(\frac{7}{8}, \frac{7}{8}, \frac{5}{4}, -\frac{3}{2}\right) & \left(\frac{7}{8}, \frac{7}{8}, -\frac{3}{2}, \frac{5}{4}\right) \end{array} $
B	$ \begin{array}{cc} & e=1 & e=2 \\ e=1 & \left(0, 0, -\frac{3}{2}, \frac{5}{4}\right) & \left(0, 0, \frac{5}{4}, -\frac{3}{2}\right) \\ e=2 & \left(0, 0, \frac{5}{4}, -\frac{3}{2}\right) & \left(0, 0, -\frac{3}{2}, \frac{5}{4}\right) \end{array} $	$ \begin{array}{cc} & e=1 & e=2 \\ e=1 & (3, 0, -1, -1) & (1, 1, 1, 1) \\ e=2 & (1, 1, 1, 1) & (0, 3, -1, -1) \end{array} $

Table 1: Normal form of the example in Peters (2004)

Observe that principals' strategies are restricted to deterministic offers. That is, principals are not allowed to use lotteries over the actions $\{A, B\}$. Given this restriction, in the direct mechanism game in Figure 1, there exist three subgame perfect equilibria in which both principals play the pure strategy B . These equilibria are:

- P_1 and P_2 both play B ; A_1 plays 1 and A_2 plays 2; each player gets a payoff of 1.
- P_1 and P_2 both play B ; A_1 plays 2 and A_2 plays 1; each player gets a payoff of 1.
- P_1 and P_2 both play B ; A_1 and A_2 randomize and play 1 with probability $1/2$; P_1 and P_2 both get a payoff of $5/4$, A_1 and A_2 get a payoff of 0.

If principal 1 can negotiate with the agents before choosing his mechanism, these pure strategy equilibria do not survive. Suppose P_1 asks each agent $i = 1, 2$ to send a message $m_1^i \in \{A, B\}$. Consider the following mechanism:

$$\sigma_1(m_1^1, m_1^2) = \begin{cases} B & \text{if } m_1^1 = m_1^2, \\ A & \text{otherwise.} \end{cases}$$

Whenever the agents do not coordinate on their messages, P_1 chooses A . There is full commitment in this game from both parties. In addition communication and effort are chosen simultaneously. That is, when the agents communicate a message to P_1 , they *commit* to an effort choice.

If principal 2 continues to play B , then the deviation σ_1 is profitable for principal 1. The continuation game associated with the strategies $\{\sigma_1, B\}$ has three subgame-perfect equilibria. In the first one (referred to as E_1), agents both report B and equally randomize over efforts, inducing for P_1 a payoff equal to $5/4$. In the second and third ones (E_2 and E_3 respectively), agents randomize over *both* messages and actions. These respectively induce payoffs of $17/16$ and $19/16$ for P_1 . Hence the strategy σ_1 is a profitable deviation.

Hence, when principals are restricted to offer deterministic direct mechanisms and we want to represent a moral hazard framework, the sequential nature of the revelation of information implies that the deviation σ_1 is not profitable.

1.2 The role of non-deterministic direct mechanisms

Ignoring timing issues, let us consider the original set-up of example 1 and assume that principals can use non-deterministic direct mechanisms. We show that, in such a scenario, all payoffs associated with principal 1's deviation to σ_1 can be supported using direct mechanisms.

To develop this argument, we start with a simple remark: In the equilibria discussed above, principal 2 always chooses the strategy B . Thus, we can restrict our analysis to the second column of table 1. We can therefore focus on the optimal action of principal 1, and interpret this example as a single principal, multi-agent game, without loss of generality. This makes the equilibrium outcomes associated with E_2 and E_3 particularly puzzling, since they contradict the Revelation Principle.²

For such games, with one principal contracting with several agents, Strausz (2003) has shown that it is no longer true that any payoff implementable by deterministic indirect mechanisms can at least be matched with a *deterministic* direct mechanism. However, the Revelation Principle retains its power if we allow principals to use non-deterministic direct mechanisms.

Suppose, in particular, that principal 1 randomizes equally between A and B (with principal 2 playing B). Then, the agents are effectively faced with the following subgame, where the first number in each cell is the principal's expected payoff, and the remaining two numbers are the expected payoffs to agents 1 and 2, respectively.³

	$e = 1$	$e = 2$
$e = 1$	$(31/16, -5/4, 1/8)$	$(15/16, 9/8, -1/4)$
$e = 2$	$(15/16, 9/8, -1/4)$	$(7/16, -5/4, 1/8)$

The agents' subgame exhibits only a mixed strategy equilibrium, with both agents equally randomizing over actions. This randomization yields the principal a payoff of $\frac{17}{16}$, and the agents a payoff $-\frac{1}{8}$.⁴ Assigning a higher probability to the choice of B in the direct mechanism increases the principal's payoff. In general, if the principal plays A with probability p and B with probability $(1 - p)$, where $p \in (0, 1)$, the unique equilibrium in the agents' subgame still has the agents randomizing equally over actions, and the principal gets an expected payoff $\frac{5}{4} - \frac{8}{3}p$. If $p = \frac{1}{6}$, the principal's payoff is $\frac{19}{16}$, the same payoff that was achievable with the indirect mechanism σ_1 . If $p \in (0, \frac{5}{16})$, the principal's payoff is even higher.

Hence, if the principal is allowed to randomize over contracts, the outcome $(1, 1, 1)$ cannot be an equilibrium. The same argument can be reproduced if we consider another principal, say principal 2,

²Whenever either E_2 or E_3 is implemented, the corresponding principal's payoff cannot be achieved in the direct mechanism game.

³Since principal 2's strategy is fixed at B , his payoffs are ignored. As mentioned above, this is now equivalent to a single principal, two-agent game.

⁴If participation constraints (e.g., a reservation of utility of zero for the agents) are a factor, note that the agents payoffs in each cell of the original game can be increased by $\frac{1}{8}$ without affecting the equilibria.

playing a fixed strategy B as in Peters' original example, and examine the best response of principal 1. This suggests that the restriction to deterministic contracts in the direct mechanism game is critical.

2 A general result

We now provide a general argument to show that in multi-principal multi-agent games every equilibrium outcome of a direct mechanism game remains an equilibrium if we allow principals to design more complex communication mechanisms. Our structure is essentially similar to Peters (2004). There are n principals dealing with k agents. Each principal j chooses an allocation $y_j \in Y_j$. The allocation y_j can be monetary transfers, or tax rates, prices, or quantities, depending on the particular interpretation of the model. Each agent i has a type drawn from a set Ω^i . Let $F(\cdot) : \times_{i=1}^k \Omega^i \rightarrow [0, 1]$ denote the joint probability distribution of the array $\omega = (\omega^1, \dots, \omega^k) \in \Omega = \times_{i=1}^k \Omega^i$. $F(\cdot)$ is assumed to be common knowledge.

We assume that agent i can take two types of unobservable actions: $e^i \in E^i$ and $a^i \in A^i$. The two actions are different in that they are chosen at two distinct times (more on this below). We denote the vectors of unobservable actions as $e = (e^1, e^2, \dots, e^k) \in E = \times_{i=1}^k E^i$ and $a = (a^1, a^2, \dots, a^k) \in A = \times_{i=1}^k A^i$. For convenience (rather than any intrinsic interpretation) we refer to a^i as the action of agent i , and e^i as his effort.

Each principal chooses an allocation $y_j^i \in Y_j^i$ for each agent i . Let $Y_j = \times_{i=1}^k Y_j^i$ denote the allocation set of principal j .

We consider a "communication game". First, each principal j chooses a message space M_j^i for each agent i , and an allocation strategy σ_j . Principal j can make the allocation y_j contingent on all messages he receives. Let $M_j = \times_{i=1}^k M_j^i$ denote the set of feasible messages principal j can receive. His allocation strategy is represented by $\sigma_j : M_j \rightarrow \Delta(Y_j)$, so that it specifies a probability distribution over allocations, depending on the message array m_j chosen by the agents. We define (M_j, σ_j) to be the mechanism offered by principal j . Mechanisms are public: they are known by all agents.

Let $\Sigma_j(M_j)$ be the collection of all maps σ_j for a given message set M_j and denote by $\Sigma(M) = \times_{j=1}^n \Sigma_j(M_j)$ the Cartesian product of these collections. Let $\Delta(Y) = \times_{j=1}^n \Delta(Y_j)$ be the set of possible outcomes. Finally, let M^i denote the message space for agent i in all his transactions with principals, so that $M^i = \times_{j=1}^n M_j^i$. Let $M = \times_{i=1}^k \times_{j=1}^n M_j^i$ denote the collection of message spaces in the game. We let $\mathcal{M}_j^i$ denote the set of feasible message spaces for principal j in his interaction with agent i , and $\mathcal{M} = \times_{j=1}^n \times_{i=1}^k \mathcal{M}_j^i$.

For an arbitrary message space M , the multi-principal multi-agent game may be intractable. The main purpose of the present paper is to consider specific message spaces which make the problem tractable. In particular, we consider direct mechanisms, in which $M_j^i = \Omega^i$ for each agent i and principal j , and see how results obtained with this particular set can be extended to more complex messages spaces.

We assume here that each principal communicates with all agents. In the Myerson (1982) model, principals are allowed to communicate only with subsets of agents. Both structures are restrictive but we do not explicitly model these potential restrictions on the communication schemes.

The action a^i is taken after the principals announce their mechanisms, but before agent i observes the (possibly randomized) outcomes $\psi \in \Delta(Y)$. Since there are multiple agents, and principals' allocations depend on all messages, agent i need not know the allocations offered to him at this stage. The effort e^i is only chosen after observing the outcome of the mechanisms. This representation is useful to highlight the difference in the timing structures (Timing 2 and 3) modelled in the previous section.⁵ A strategy for agent i , for a given message space M^i , consists of the two mappings:

$$\begin{aligned}\lambda^i &: \Omega^i \times \Sigma \rightarrow \Delta(A^i \times M^i) \\ \delta^i &: \Omega^i \times \Delta(Y) \times A^i \times M^i \rightarrow \Delta(E^i).\end{aligned}$$

We denote by Λ^i and Δ^i the sets of all such maps. Though we suppress this dependence in the notation, it should be understood that these maps vary with the message spaces chosen by the principals.

Agent i 's payoff is given by the von Neumann–Morgenstern utility function $U^i(y, a, e, \omega)$, and for each principal j the payoff is given by $V_j(y, a, e, \omega)$. With a slight abuse of notation, we sometimes refer to $U^i(\sigma, \lambda, \delta, \omega)$ and $V_j(\sigma, \lambda, \delta, \omega)$.

For any collection of message spaces $\mathcal{M}$, we can define an associated multi-principal, multi-agent game:

$$\Gamma_{\mathcal{M}} = \left\{ \Omega, (M_j, \Sigma_j)_{j=1}^n, (\Lambda^i, \Delta^i)_{i=1}^k, (V_j(\cdot))_{j=1}^n, (U^i(\cdot))_{i=1}^k, F(\cdot) \right\}.$$

We consider Perfect Bayesian Equilibrium (PBE) as the relevant equilibrium concept for $\Gamma_{\mathcal{M}}$.

We now consider direct mechanism games. In direct mechanisms, each principal is restricted to the message space Ω^i for each agent i . That is, $M_j^i = \Omega^i$ for every $j = 1, \dots, n$ and for every $i = 1, \dots, k$. Given Ω , considering direct mechanisms induces the following multi-principal multi-agent game Γ_{Ω} :

Definition 1 For a given type space Ω , a direct mechanism game Γ_{Ω} is the array

$$\Gamma_{\Omega} = \left\{ \Omega, (\tilde{\Sigma}_j)_{j=1}^n, (\tilde{\Lambda}^i, \tilde{\Delta}^i)_{i=1}^k, (V_j(\cdot))_{j=1}^n, (U^i(\cdot))_{i=1}^k, F(\cdot) \right\}.$$

The main difference between the indirect game $\Gamma_{\mathcal{M}}$ and the direct mechanism game Γ_{Ω} is that the message spaces are not choices for the principals in the latter case.

We define the strategy of principal j in the game Γ_{Ω} as the map:

$$\tilde{\sigma}_j : \Omega \rightarrow \Delta(Y_j).$$

Let $\tilde{\Sigma}_j$ be the strategy space for principal j and $\tilde{\Sigma}$ the collection of all such strategy profiles. The strategy of agent i is then defined by the two maps:

$$\begin{aligned}\tilde{\lambda}^i &: \Omega^i \times \tilde{\Sigma} \rightarrow \Delta(\Omega^i \times A^i) \\ \tilde{\delta}^i &: \Omega^i \times \Delta(Y) \times A^i \rightarrow \Delta(E^i).\end{aligned}$$

⁵Note that Peters (2004) defines the contract space of a principal to be mappings from E to Y_j . However, in Example 1 in his paper, contracts are not conditional on effort, so our definition is sufficient.

We denote the respective strategy spaces as $\tilde{\Lambda}^i$ and $\tilde{\Delta}^i$.

The time structure of the interaction is provided in Figure 4. Remember that both action and effort, a and e are unobservable to the principals. Notice that the first one is taken before the realizations of principals decisions, while the second one after observing them hence with additional and relevant information.

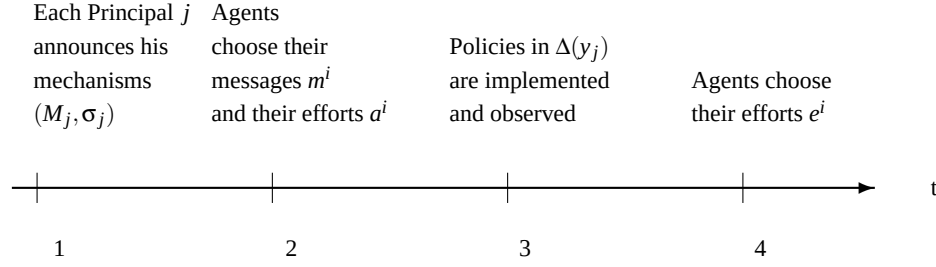


Figure 4: Timing with both action and effort

Our main aim is to remark that even in multi-principal multi-agent games, we do not lose a basic property of contract design problems: every pure strategy equilibrium of a direct mechanism game remains such when we allow principals to use more complex indirect mechanisms. That is, in the more complex game (in which principals can choose message spaces) it is an equilibrium for each principal j to choose a message space $M_j^i = \Omega^i$ for each agent i , and to play the same strategy conditional on messages ω that he plays in the direct mechanism game. It follows, of course, that it is a best response for the agents to behave exactly as they would in the direct mechanism game.

We formally state this intuition in the following theorem.

Theorem 1 *Consider a collection of message spaces $\mathcal{M}$ such that $\Omega^i \subseteq \mathcal{M}_j^i$ for all $i = 1, \dots, k$ and $j = 1, \dots, n$. If $\left((\sigma_j^*)_{j=1}^n, (\lambda^{i*}, \delta^{i*})_{i=1}^k \right)$ is a pure strategy PBE in the game Γ_Ω , there exists a PBE in the multi-principal multi-agent game $\Gamma_{\mathcal{M}}$ that generates the same outcome.*

Proof. Let $\left((\sigma_j^*)_{j=1}^n, (\lambda^{i*}, \delta^{i*})_{i=1}^k \right)$ be a pure strategy PBE in the game Γ_Ω , where $\lambda^{i*} = (a^{i*}, \omega^{i*})$. Now consider the game $\Gamma_{\mathcal{M}}$. Suppose that every principal j in the game $\Gamma_{\mathcal{M}}$ offers the mechanisms (Ω, σ_j^*) , and every agent i chooses actions according to the maps $(\lambda^{i*}, \delta^{i*})_{i=1}^k$. Notice that this leads to the same outcome as the equilibrium $\left((\sigma_j^*)_{j=1}^n, (\lambda^{i*}, \delta^{i*})_{i=1}^k \right)$ in Γ_Ω .

Contrary to the statement of the theorem, suppose these strategies are not an equilibrium in the game $\Gamma_{\mathcal{M}}$. Then, there must be at least one principal j who has a profitable deviation. That is, suppose all other principals $\tilde{j} \neq j$ offer the mechanisms $(\Omega, \sigma_{\tilde{j}}^*)$. Then, principal j has a strategy (M'_j, σ'_j) , where $M'_j \neq \Omega$ and $\sigma'_j : M'_j \rightarrow \Delta(Y_j)$, that leads to a higher payoff than the strategy (Ω, σ_j^*) .

Given the mechanisms offered in the game $\Gamma_{\mathcal{M}}$, agents choose strategies λ' and δ' that constitute a Perfect Bayesian Equilibrium in the continuation game. Since the mechanisms are already chosen at this stage, for each agent i we let $\lambda^{i'}$ and $\delta^{i'}$ be functions only of ω^i . Further, let $\lambda_j^{i'} = (\alpha_j^{i'}, \mu_j^{i'})$ where $\alpha_j^{i'} : \Omega^i \rightarrow \Delta(A^i)$, and $\mu_j^{i'} : \Omega^i \rightarrow \Delta(M_j^i)$.

Now, if principal j has a profitable deviation in the game $\Gamma_{\mathcal{M}}$, it must be that

$$\int_{\omega \in \Omega} V_j [\sigma'_j(\mu'_j(\omega)), \sigma_{-j}^*(\mu'_{-j}(\omega)), \alpha', \delta', \omega] dF(\omega) > \int_{\omega \in \Omega} V_j [\sigma_j^*(m_j^*(\omega)), \sigma_{-j}^*(m_{-j}^*(\omega)), a^*, \delta^*, \omega] dF(\omega). \quad (1)$$

Note that the deviation by principal j need not induce a unique equilibrium in the continuation game. That is, there may exist multiple (λ', δ') pairs that constitute Perfect Bayesian Equilibria of the agents' game, given the offered mechanisms. However, at least one such pair must lead to a strictly higher payoff for principal j .

Following the deviation by principal j , the message sent by each agent to each principal is a function of his type. We now construct a mapping $\tilde{\sigma}_j(\omega)$ in the following manner:

$$\forall \omega \in \Omega, \quad \tilde{\sigma}_j(\omega) = \sigma'_j \left((\mu_j^i(\omega^i))_{i=1}^k \right).$$

Notice that since $\lambda^{i'}, \delta^{i'}$ are strategies played after a deviation by principal j , these can be mixed strategies, even though the original equilibrium is a pure strategy equilibrium. This is why we need the mechanisms $\tilde{\sigma}_j(\omega)$ to be non-deterministic.

>From (1) it follows that

$$\int_{\omega \in \Omega} V_j [\tilde{\sigma}_j(\omega), \sigma_{-j}^*(\mu'_{-j}(\omega)), \alpha', \delta', \omega] dF(\omega) > \int_{\omega \in \Omega} V_j [\sigma_j^*(m_j^*(\omega)), \sigma_{-j}^*(m_{-j}^*(\omega)), a^*, \delta^*, \omega] dF(\omega) \quad (2)$$

By definition, since principals $\tilde{j} \neq j$ are offering the message space Ω , the mapping μ'_{-j} maps Ω^i to Ω^i .

For each agent i , the message strategy $\mu_j^i(\omega^i)$ is part of his best reply when principal j offers the mechanism (M'_j, σ'_j) . By construction, the following message strategy is then part of his best reply when principal j offers the direct mechanism $\tilde{\sigma}_j(\omega)$: report ω^i to principal j , and follow the message strategy $\mu_{-j}^i(\omega^i)$ in making reports to all other principals. Then, equation (2) leads to a contradiction—the mechanism σ_j^* is not a best reply for principal j in the direct mechanism game.

We have shown that no principal has any profitable deviation from (Ω, σ_j^*) in the game $\Gamma_{\mathcal{M}}$. Since $\left((\sigma_j^*)_{j=1}^n, (\lambda^{i*}, \delta^{i*})_{i=1}^k \right)$ is an equilibrium in the direct mechanism game Γ_{Ω} , the agents' strategies in $\Gamma_{\mathcal{M}}$ remain the same as in Γ_{Ω} . ■

This theorem emphasizes that in a general multi-principal multi-agent interaction every best reply for any principal in the game Γ_{Ω} remains a best reply even if the message space M_i is enlarged with respect

to the type space Ω . The theorem is an extension to Peters (2003)'s pure strategy theorem.⁶ The intuition is simple: at equilibrium, principal i plays in a context where other principals' behavior is taken as given. He designs his best mechanism considering this environment and what he knows about agents. Then, when considering profitable deviations principal i can restrict himself to direct mechanisms without loss of generality. In this specific context, we are simply applying the basic argument given by Myerson (1982).

Theorem 1 cannot be generalized if one allows principals to play mixed strategies.⁷ To see this, consider a game with two principals (principal 1 and principal 2) and only one agent. If principal 2 randomizes over mechanisms, the array $(m'_j, a', \omega', \delta')$ will depend on the type of the agent denoted ω and on the mechanisms σ_2 randomly chosen by the principal 2. As the mechanism σ_2 is ex-ante unknown, one cannot construct a direct mechanism $\tilde{\sigma}_1$ as in the proof, because it would depend on a σ_2 which may take different values. If the theorem cannot be extended, it is mainly because the mechanism σ_2 is a random variable for principal 1, but it is observed by the agent.

In addition, it is now a well-known finding that in multi-principal games the Revelation Principle may fail to hold.⁸ This does not contradict our theorem.

To see this, consider a simple setting in which there are two principals (1 and 2) only, and some fixed number of agents. Principal 1 can always respond to principal 2 using a direct mechanism, without loss of generality. If principal 2 uses an indirect mechanism $\sigma_2(\cdot)$, then if principal 1's best reply is the indirect mechanism $\sigma_1(\cdot)$, we know that there exists an equivalent direct mechanism $\tilde{\sigma}_1(\cdot)$. The mechanism $\sigma_2(\cdot)$ can be a best reply to the mechanism $\sigma_1(\cdot)$, but it is not necessarily a best reply to the direct mechanism $\tilde{\sigma}_1(\cdot)$. Thus, the equilibrium involving the strategies $(\sigma_1(\cdot), \sigma_2(\cdot))$ can disappear if one restricts the principals' strategy spaces.

3 Conclusion

The literature on competing mechanisms with multiple agents makes ad hoc assumptions about the set of mechanisms that are feasible for principals: models mainly focus on 'direct mechanisms'. For example Doğan (2003), McAfee (1993), Prat and Rustichini (2002) assume that principals use direct mechanisms. Theorem 1 offers some support for this approach—equilibria defined in the direct mechanism game are also equilibria in the most general game in which principals use richer indirect mechanisms, if one correctly defines the strategy spaces. Nevertheless, as Peters (2004) shows with his second example, if we focus on direct mechanisms, we may lose some interesting equilibria.

⁶The same intuition is also discussed in Martimort and Stole (2002) in a less formal way.

⁷In our scenario, mixed strategies correspond to randomization over non-deterministic mechanisms, i.e. randomization over lotteries.

⁸See Martimort and Stole (2002), Peters (2001). Peck (1997) provides an example in a multi-principal multi-agent model.

References

- [1] Doğan, P. (2003) “Vertical Integration in the Internet Industry,” *mimeo*, Koç University.
- [2] McAfee, P. (1993) “Mechanism Design by Competing Sellers,” *Econometrica* 61, 1281-1312.
- [3] Martimort, D., and L. Stole (2002) “The Revelation and Taxation Principles in Common Agency Games,” *Econometrica* 70, 1659-1673.
- [4] Myerson, R. (1982) “Optimal Coordination Mechanisms in Generalized Principal-Agent Problems,” *Journal of Mathematical Economics* 10, 67-81.
- [5] Peters, M. (2001) “Common Agency and the Revelation Principle,” *Econometrica* 69, 1349-1372.
- [6] Peters, M. (2003) “Negotiation and take it or leave it in common agency,” *Journal of Economic Theory*, 111(1), 88-109.
- [7] Peters, M. (2004) “Pure strategy and no-externalities with multiple agents,” *Economic Theory*, 23, 183-194.
- [8] Peck, J. (1997) “Competing Mechanisms and the Revelation Principle,” *mimeo*, Ohio State University.
- [9] Prat, A., and A. Rustichini (2003) “Games Played Through Agents,” *Econometrica* 71, 989-1026.
- [10] Strausz, R. (2003) “Deterministic Mechanisms and the Revelation Principle,” *Economic Letters* 79, 333-37.