

DISCUSSION PAPER

2003

# Common Agency games with separable preferences

Andrea ATTAR<sup>1</sup>, Dipjyoti MAJUMDAR<sup>2</sup>, Gwenaël PIASER<sup>3</sup> and Nicolàs PORTEIRO<sup>4</sup>

November 2003

## Abstract

This paper examines the role of the revelation principle in common agency games. We show how the introduction of a separability condition on the preferences of the agent is sufficient for the revelation principle to hold. Therefore, it is still possible to restrict attention to direct mechanisms without any loss of generality even when competition over contracts is considered.

JEL Classification: D82.

---

<sup>1</sup>CORE Université catholique de Louvain and Università di Roma, La Sapienza.

<sup>2</sup>CORE Université catholique de Louvain.

<sup>3</sup>CORE Université catholique de Louvain and CREPP Université de Liège.

<sup>4</sup>CORE Université catholique de Louvain.

We are grateful to Catarina Goulão, Pierre Fleckinger, Enrico Minelli and seminar participants at the ARC seminar at CORE for helpful comments.

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 author.

# 1 Introduction

The present work tries to contribute to the analysis of common agency games. This class of games defines a situation where multiple principals compete over the contract offers they make to a single agent.<sup>1</sup>

Several important researches recently focused on such a theoretical structure with the aim of further developing the relationship between incentives and competition. Examples are given by: Biais, Martimort, and Rochet (2000) who analyzed a non-linear pricing environment, Parlour and Rajan (2001) who studied the competition among lenders offering non-exclusive credit contracts, and Kahn and Mookherjee (1998) who presented applications to the insurance literature. In the terminology introduced by Segal and Whinston (2003), these contributions can be seen as different *bidding* games. These games involve the following timing: principals act as mechanism designers and simultaneously offer contracts to the single agent. In a next step, the agent chooses her preferred alternative, contracts are executed and the equilibrium allocation is implemented.

A relevant consequence of allowing for a multi-principal interaction is that it is no longer possible to argue that we can restrict attention to the set of direct mechanisms with no loss of generality. In other words: “*When contracting situations become more complex, the applicability and the usefulness of the Revelation Principle comes into question*” (Martimort and Stole 2002, p. 1659). The main rationale behind the failure of the Revelation Principle can be summarized as follows: at the stage of contracting with any principal the agent’s private information involves both her type and the mechanisms simultaneously offered by all the other principals. Direct mechanisms are too simple as principals cannot profit from the extra information held by the agent.

Several independent researches developed examples of common agency games which equilibrium allocation cannot be supported when only direct mechanisms are involved.<sup>2</sup> The most recent literature offered two alternative ways to face this difficulty. On the one hand, Epstein and Peters (1999) showed that in multi-principal-multi-agent problems it is possible to construct a *universal* type space that embeds all the agents’ private informations: with respect to this space the Revelation Principle is still valid. In the same line, the recent work by Pavan and Calzolari (2003) introduces a *Markovian* Revelation Principle for common agency games. It is possible to restrict without loss of generality to mechanisms implying agent’s truthful revelation of her *extended* type, including agent’s private information, as well as all the payoff-relevant decisions contracted with the multiple principals. The main difficulty with these general solutions is the analytical complexity of the relevant mechanisms and types that limits their applicability to common contracting problems.

---

<sup>1</sup>The first attempt at modeling a common agency problem has been provided by Bernheim and Whinston (1986).

<sup>2</sup>See Peck (1997), Epstein and Peters (1999), Peters (2001) and Martimort and Stole (2002).

On the other hand, Peters (2001) and Martimort and Stole (2002), instead of looking at the relationship between indirect and direct mechanisms when multiple principals compete, emphasized how the payoffs supported by general indirect mechanisms can be as well supported as equilibria in the class of games where principals strategies are menus over the relevant alternatives. That is, they suggest an extension of the so-called taxation principle (Rochet 1986) to a multi-principal framework: this amounts to substituting a potentially rich set of mechanisms with a simpler one, where the agent is not required to strategically reveal her private information.

Surprisingly, relative less attention has been devoted to the definition of domain conditions that guarantee the Revelation Principle to hold in common agency games.<sup>3</sup> The present paper explores such a research direction trying to identify conditions on agent's preferences that are sufficient for the Revelation Principle to remain valid when competing mechanisms are introduced. Our condition amounts to a separability requirement on the agent's preferences with respect to both the contract offers made by every principal and the actions taken by the agent herself. To be more precise, what we require is that the agent's preference ordering over the contract offers by one principal should not depend neither on those made by the others nor on the (uncontractible) action she takes. Under such a condition we are able to show that the Revelation Principle works. The remarkable result is that no restriction on principals' preferences is introduced. This allows us to deal with a large class of agency problems that did not fit into the most traditional works on common agency where no direct externality among principals is assumed.<sup>4</sup>

To our best knowledge, a similar approach has only been followed by Peters (2003), who developed an intuition close to ours. In spite of the analogies between these two works, they are aimed at studying different common agency interactions. Our research deals with situations in which the revelation of agent's exogenous private information is the crucial element, while Peters' concern is primarily the interaction between several principals and one agent who takes actions affecting principals' payoffs. The relationships between our research and Peters' analysis, will be discussed in detail in Section 5.

The paper is organized as follows: Section 2 presents some examples emphasizing the problem of the Revelation Principle in common agency games. Section 3 studies the basic framework of common agency games with incomplete information; Section 4 deals with a more general case where agents are allowed to take actions. Section 5 provides a general discussion of our results and Section 6 concludes.

---

<sup>3</sup>This line of research has originally been suggested by Martimort (1992).

<sup>4</sup>See for instance Bernheim and Whinston (1986) and Dixit, Grossman, and Helpman (1997).

## 2 Examples

We illustrate now the effects of assuming separability in agent's preferences. This section will be devoted to the discussion of some examples of common agency adverse selection games.

First consider the following game (introduced by Martimort and Stole, 2002). The agent's type space is degenerate and there are two principals: each of them can undertake the three actions  $d_{i \in \{1,2\}} \in \{A, B, C\}$ . The payoffs of principal 1, principal 2 and the agent,  $\{V^1, V^2, U\}$ , are given by the following matrix

|           | $d_2 = A$    | $d_2 = B$   | $d_2 = C$    |
|-----------|--------------|-------------|--------------|
| $d_1 = A$ | $(1, 1, a)$  | $(2, 0, 2)$ | $(-1, 5, b)$ |
| $d_1 = B$ | $(1, 2, 2)$  | $(1, 1, 1)$ | $(0, 0, c)$  |
| $d_1 = C$ | $(5, -1, b)$ | $(0, 0, c)$ | $(0, 0, 0)$  |

In the original example we have  $a = 1$ ,  $b = 10$  and  $c = 0$ . If we restrict principals to offer direct mechanisms, that correspond to take it or leave it offers over single actions  $d_i$ , the equilibrium allocation is  $\{d_1 = C, d_2 = C\}$ .<sup>5</sup> However, if principals are allowed to offer menus of simple actions, then the outcome  $\{B, B\}$  can also be supported at equilibrium. The menu offers  $\{B, C\}$  by both principals constitute a Nash equilibrium: given such offers the agent will choose  $B$  from both principals and the outcome  $(1, 1, 1)$  will be implemented. Principal 1 cannot profitably deviate by offering alternative menus because the agent is acting as a coordinating device. If he deviates by offering  $A$ , then the agent will choose  $C$  from Principal 2 making the deviation unprofitable. Martimort and Stole therefore emphasize that there exist outcomes implementable through menu offers<sup>6</sup> that are not implementable through direct mechanisms.

It is important to remark that under alternative assumptions on agent's preferences the outcome  $(1, 1, 1)$  cannot be implemented through menus neither. To show that it is enough to consider the case  $a > 2$ ,  $b < 2$  and  $c < 1$ . In such a case, allowing for menus does not enlarge the set of equilibria sustainable through direct mechanisms.

For  $(1, 1, 1)$  to be sustained as equilibrium payoffs, the utility of the agent should be such that, for every  $i \neq j$

$$U(d_i = B, d_j = B) \geq U(d_i = C, d_j = B)$$

$$U(d_i = B, d_j = A) < U(d_i = C, d_j = A)$$

The conditions above cannot be simultaneously fulfilled if the agent has separable preferences. This is true, because whenever the preferences of the agent are separable, if he gains when moving from  $d_i = B$  to  $d_i = C$  (given  $d_j = B$ ), then the

<sup>5</sup>Notice that in such a direct mechanism there is no strategic role for the agent and the outcomes that can be supported are just the Nash equilibrium outcomes of the choice-of-action game by the two principals.

<sup>6</sup>Notice that menu offers belong to the class of indirect mechanisms.

same has to hold when  $d_j = A$ . That is to say, the marginal impact on the agent's utility of a change in principal  $i$ 's action has to be independent from principal  $j$ 's action. For this reason, the allocation  $\{B, B\}$  cannot be supported when preferences are separable, even if we allow the principals to offer menus to the agent. This first example already provides an important hint for the analysis: A requirement needed for menus to improve over direct mechanisms seems to be that the agent has a *preference reversal* over the principals' actions, that is not possible when preferences are separable. In the next section we will check how important is this insight.

Then, we focus on another example, still borrowed from Martimort and Stole (2002). There are two principals and one agent with type  $\theta \in \{-1, 1\}$ , drawn with equal probability. Each principal has two actions  $d_i \in \{0, 1\}$ . The utilities are  $V_1 = \theta(d_1 - d_2)$  for principal 1,  $V_2 = \theta(d_2 - d_1)$  for principal 2 and  $U = \theta(d_1 + d_2)$  for the agent. The aim of the original example was to show that any equilibrium that can be supported with arbitrarily complex indirect mechanisms can be also replicated through menus. Equilibrium strategies are given by the menus  $\{0, 1\}$  for both principals. The outcome will always be  $d_1 = d_2 = 0$  (when  $\theta = -1$ ) and  $d_1 = d_2 = 1$  (when  $\theta = 1$ ).

The example however assumes agent's preferences that are separable in the actions of the principals and, for this reason, the same outcome can be achieved as an equilibrium of a direct mechanism. Simply consider the following revelation contract: each principal offers the agent a direct mechanism  $T_i$  such that:

$$T_i = \begin{cases} d_i = 0 & \text{if } \theta = -1 \\ d_i = 1 & \text{if } \theta = 1. \end{cases}$$

This mechanism induces truthful revelation by the agent and implements the same outcome as the indirect mechanisms (or the menus) do.

These two examples considered together emphasize the relationship between the set of outcomes implementable by menus and by direct mechanisms if the agent has separable preferences. In the first example, the use of menus does not support outcomes that are not equilibria associated to direct revelation mechanisms. In the second one, that assumes separability, every equilibrium associated to a menu game can be as well supported by direct mechanisms.

### 3 The Incomplete Information Case

This section discusses how the assumption of separability in agents' preferences can restore the Revelation Principle in games with incomplete information. The natural reference is therefore given by Martimort and Stole (2002).

We consider a scenario where there are a number of principals (indexed by  $i \in \{1, \dots, N\}$ ) contracting with one agent (denoted by index, 0). The agent's type is drawn from a finite set  $\Theta$  having a probability distribution  $F(\cdot)$  that is common knowledge. Each principal can write a contract with the agent over an allocation  $y_i \in Y_i$ . Let  $Y = \times_{i \in N} Y_i$  be the set of all possible contract profiles available to the

principals. The agent's payoff is given by the vNM utility function  $U(y_1, \dots, y_N, \theta)$ , and for each principal  $i$  the payoff is denoted by  $V_i(y_1, \dots, y_N, \theta)$ . The set of message spaces available to principal  $i$  is  $\mathcal{M}_i$  with  $\mathcal{M} = \times_{i \in N} \mathcal{M}_i$ . Each principal's strategy is to choose a message space  $M_i \in \mathcal{M}_i$  and an outcome function  $\sigma_i : M_i \rightarrow Y_i$ . Let  $\Sigma_i$  be the collection of all maps  $\sigma_i$ , and denote by  $\Sigma = \times_{i \in N} \Sigma_i$  the Cartesian product of all those  $\Sigma_i$ .

The strategy for the agent is a map  $\sigma_0 : \Theta \times \Sigma \rightarrow M$  and  $\Sigma_0$  denotes the collection of such maps. For every collection of message spaces  $M \in \mathcal{M}$  chosen by principals, we can define an associated common agency game:

$$\Gamma_M = \{\Theta, (M_i, \Sigma_i)_{i \in N}, \Sigma_0, U(\cdot, \theta), (V_i(\cdot, \theta))_{i \in N}\}.$$

We consider the Perfect Bayesian Equilibrium (PBE) as the relevant equilibrium concept for the common agency game  $\Gamma_M$ .

**Definition 1** *Given the message space  $M$ , a strategy vector  $\sigma^* = ((\sigma_i^*)_{i \in N}, \sigma_0^*) \in \Sigma \times \Sigma_0$  is a PBE of  $\Gamma_M$  if:*

$$\forall \theta \in \Theta, \forall (\sigma_i)_{i \in N} \in \Sigma, \sigma_0^*((\sigma_i)_{i \in N}, \theta) = (m_1^*, \dots, m_N^*) \in \arg \max_{m \in M} U(\sigma_i(m_i)_{i \in N}, \theta)$$

$$\forall i \in N, \sigma_i^*(m_i^*((\sigma_i^*)_{i \in N})) \in \arg \max_{\sigma_i \in \Sigma_i} \int_{\theta \in \Theta} V_i(\sigma_i(m_i^*(\sigma_i, \sigma_{-i}^*(\cdot))), \sigma_{-i}^*(\cdot), \theta) dF(\theta)$$

We now consider direct mechanisms. In direct mechanisms, each principal restricts his message space to the agent's type space  $\Theta$ .

Given  $\Theta$ , considering direct mechanisms induces the following common agency game  $\Gamma_\Theta$ :

**Definition 2** *For the type space  $\Theta$ , a direct mechanism game  $\Gamma_\Theta$  is the array*

$$\Gamma_\Theta = \{\Theta, (\Theta, \tilde{\Sigma}_i)_{i \in N}, \tilde{\Sigma}_0, U(\cdot, \theta), (V_i(\cdot, \theta))_{i \in N}\}$$

We defined the strategy of principal  $i$  in the game  $\Gamma_\Theta$  by the map  $\tilde{\sigma}_i : \Theta \rightarrow Y_i$ ; we let  $\tilde{\Sigma}_i$  to be the strategy space for principal  $i$  and  $\tilde{\Sigma}$  be the collection of all such strategy profiles.

The strategy of the agent is then a map  $\tilde{\sigma}_0 : \Theta \times \tilde{\Sigma} \rightarrow \Theta^N$ , and  $\tilde{\Sigma}_0$  denotes the collection of all such maps.

We now introduce the condition of separable utility for an agent.

**Definition 3** *Weak Separability: the utility function  $U$  is separable in the principals' decisions for all  $\theta \in \Theta$ , if for all  $y_i, y'_i \in Y_i$  and for all  $y_{-i}, y'_{-i} \in Y_{-i}$*

$$[U(y_i, y_{-i}, \theta) > U(y'_i, y_{-i}, \theta)] \implies [U(y_i, y'_{-i}, \theta) > U(y'_i, y'_{-i}, \theta)] \quad (1)$$

Under weak separability the agent's ranking over principal  $i$ 's decision is preserved for every set of decisions taken by the other principals. Throughout this section we will only consider utility functions for the agent that satisfy the above requirement of separability.

Agent's preferences under separability can be represented by the notion of partial utility. Given the utility function  $U$ , the partial utility for the agent given that the offers made by the other principals are fixed at the level  $\hat{y}_{-i} \in Y_{-i}$ , is defined as:

$$\tilde{U}_{\hat{y}_{-i}}(y_i, \theta) = U(y_i, \hat{y}_{-i}, \theta)$$

We start by showing some properties satisfied by equilibrium strategies when the agent has separable preferences. In particular, we show that for each  $i$  and for any  $\hat{y}_{-i} \in Y_{-i}$ , the  $i$ -th component of the agent's optimal strategy maximizes her partial utility at the level  $\hat{y}_{-i}$ . More formally, denoting  $m_i$  as the projection of  $\sigma_0$  on  $M_i$ , we have the following:

**Claim 1** *If  $\sigma^* = (\sigma_0^*, (\sigma_i^*)_{i \in N})$  is a PBE then for all  $i \in N$ , for all  $\hat{y}_{-i} \in Y_{-i}$ :*

$$m_i^* \in \arg \max_{m_i \in M_i} \tilde{U}_{\hat{y}_{-i}}(\sigma_i^*(m_i), \theta)$$

Proof. Suppose the claim is false. This implies that there exists a  $m'_i \in M_i$  such that  $\tilde{U}_{\hat{y}_{-i}}(\sigma_i^*(m'_i), \theta) > \tilde{U}_{\hat{y}_{-i}}(\sigma_i^*(m_i^*), \theta)$ . Now, from separability it follows that:  $U(\sigma_i^*(m'_i), \sigma_{-i}^*(\sigma_0^*), \theta) > U(\sigma^*, \theta)$ .

This contradicts the assertion that  $\sigma^*$  is a PBE. ■

A direct consequence of the Claim is that:

$$\sigma_i^*(m_i^*(\sigma^*, \theta)) = \sigma_i^*(m_i^*(\sigma_i^*, \theta)) \quad (2)$$

That is, every PBE exhibits the following property: principal  $i$ 's optimal strategy is independent of principal  $j$ 's mechanism. It is worth noting that this property on principals' behavior is just an implication of the restriction to separable preferences by the agent.

We also remark that given the mechanism offered by principal  $i$ , then the message he receives from the agent depends on the agent's type  $\theta$  only. At this point we are in the condition to state our main theorem that reestablishes the Revelation Principle under separability.

**Theorem 1** *If the preferences of the agent satisfy the Weak Separability condition, then for every outcome that can be supported as a PBE of an arbitrary indirect mechanism game  $\Gamma_M$ , there exists a game  $\Gamma_\Theta$  associated to direct mechanisms that implements the same outcome.*

Proof. Let us consider a PBE  $\sigma^*$  of the indirect mechanism game  $\Gamma_M$ . We will now demonstrate that the distribution over outcomes generated by the strategy profile  $\sigma^*$  can be replicated in a direct mechanism game.

Consider the direct mechanism game  $\Gamma_\Theta$  where principal  $i$ 's strategy  $\tilde{\sigma}_i : \Theta \rightarrow Y_i$  is defined as  $\tilde{\sigma}_i(\theta) = \sigma_i^*(m_i^*(\sigma_i^*, \theta))$ . The agent's strategy is a map  $\tilde{\sigma}_0 : \Theta \times \tilde{\Sigma} \rightarrow \Theta^N$ . Notice that the agent's best reply coincides with the one defined in  $\Gamma_M$  whenever principals were using direct mechanisms.

We can then rewrite principal  $i$ 's payoff as

$$\int_{\theta \in \Theta} V_i(\tilde{\sigma}_i(\theta), \tilde{\sigma}_{-i}(\theta), \theta) dF(\theta) = \int_{\theta \in \Theta} V_i(\sigma_i^*(m_i^*(., \theta)), \sigma_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta)$$

Specifically,

$$\int_{\theta \in \Theta} V_i(\tilde{\sigma}_i(\theta), \tilde{\sigma}_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta) = \int_{\theta \in \Theta} V_i(\sigma_i^*(m_i^*(., \theta)), \sigma_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta)$$

We now show that no principal has a unilateral profitable deviation. Suppose that this is not true. Then, let us suppose that there exists a principal  $i$  and a strategy  $\sigma'_i(.)$  such that:

$$\int_{\theta \in \Theta} V_i(\sigma'_i(\tilde{\theta}), \tilde{\sigma}_{-i}^*(\theta), \theta) dF(\theta) > \int_{\theta \in \Theta} V_i(\tilde{\sigma}_i^*(\theta), \tilde{\sigma}_{-i}^*(\theta), \theta) dF(\theta) \quad (3)$$

where  $\tilde{\theta}$  is the agent's best reply to the deviating principal. Obviously, if principal  $i$  changes his offered mechanism, then the agent's reply to this deviation will be different. But separability implies that the agent's reply to the other principals will be unchanged. Therefore, from (2) it follows that:

$$\begin{aligned} \int_{\theta \in \Theta} V_i(\sigma'_i(\tilde{\theta}), \sigma_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta) &> \int_{\theta \in \Theta} V_i(\tilde{\sigma}_i^*(\theta), \tilde{\sigma}_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta) \\ \implies \\ \int_{\theta \in \Theta} V_i(\sigma'_i(\tilde{\theta}), \sigma_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta) &> \int_{\theta \in \Theta} V_i(\sigma_i^*(m_i^*(., \theta)), \sigma_{-i}^*(m_{-i}^*(., \theta)), \theta) dF(\theta) \end{aligned}$$

This is a contradiction. ■

The Theorem guarantees us that the set of equilibria of the indirect mechanism game  $\Gamma_M$  is preserved in the corresponding direct mechanism game  $\Gamma_\Theta$ . To provide a general statement of the Revelation Principle we also have to show that the reverse holds, that is, the set of equilibria of  $\Gamma_\Theta$  is not enlarged if we consider  $\Gamma_M$ . This is implied by the following

**Theorem 2** *Consider a message space  $M = \times_{i \in N} M_i$  such that  $\Theta \subseteq M_i$  for all  $i \in N$ . If  $\sigma^*$  is a PBE in the common agency game  $\Gamma_\Theta$ , then  $\sigma^*$  is a PBE in the common agency game associated with message space  $M$ , that is  $\Gamma_M$ , then every outcome that can be supported as a PBE of the direct mechanism game  $\Gamma_\Theta$  can be as well implemented in the game  $\Gamma_M$ .*

Proof. Let  $\sigma^*$  a *PBE* in the game  $\Gamma_\Theta$  and contrary to our assumption, let us suppose that in the game  $\Gamma_M$ , there exists a player  $i \in N \cup \{0\}$  who has a profitable deviation. Let us first consider the case where  $i \in N$ . This implies there exists a strategy  $\sigma'_i : M_i \rightarrow Y_i$  s.t. the following are true:

$$\int_{\Theta} V_i [\sigma'_i(\hat{m}_i), \sigma_{-i}^*(\hat{\theta}), \theta] dF(\theta) > \int_{\Theta} V_i [\sigma^*(\hat{\theta}), \theta] dF(\theta) \quad (4)$$

and

$$(\hat{m}_i(\theta), \hat{\theta}(\theta)) \in \arg \max_{(m_i, \theta)} U [\sigma'_i(\hat{m}_i), \sigma_{-i}^*(\hat{\theta}), \theta]$$

Let us consider the direct mechanism  $\tilde{\theta}_i : \Theta \rightarrow Y_i$  for principal  $i$  in the following manner: for any  $\theta \in \Theta$ ,  $\tilde{\sigma}_i = \sigma'_i(\hat{m}_i(\theta))$ .

From (4) it follows that

$$\int_{\Theta} V_i [\tilde{\sigma}_i(\theta), \sigma_{-i}^*(\hat{\theta}), \theta] dF(\theta) > \int_{\Theta} V_i [\sigma^*(\hat{\theta}), \theta] dF(\theta).$$

This is a contradiction.

We have shown so far that no principal has unilateral deviation in the game  $\Gamma_M$ . Obviously if the principals use direct mechanisms, and since  $\sigma^*$  is an equilibrium in the direct mechanism game  $\Gamma_\Theta$ , the agent does not have an unilateral deviation.

■

From Claim 1 it follows that at the stage of considering the contract between the agent and principal  $i$  we can restrict our attention to the agent's partial utility for any fixed  $\hat{y}_{-i}$ , and that principal  $i$ 's optimizing behavior is independent of the agent's contract with other principals. This implies that for any fixed level  $\hat{y}_{-i}$ , we can break the general common agency game into  $N$  games, each of which is a single principal-agent problem where the Revelation Principle applies. Since the optimizing behavior for the two conceived partners remain unchanged in the original problem, we can apply the Revelation Principle taking one principal at a time and obtain the Revelation Principle from the original problem.

## 4 Moral hazard and contractible actions

We generalize here our approach to the context of incomplete information common agency games where the agent is also allowed to take an action  $e$  belonging to a measurable set  $E$ . Our set-up, therefore, encompasses both the traditional hidden action moral hazard framework (when the action of the agent is non-contractible), and that of contractible actions.

The main result of this section is that we can extend our Theorem 1 to the more general setting of moral hazard, if the separability assumption is properly specified.

At the same time, we provide a counter-example illustrating the difficulties arising when  $e$  is assumed to be contractible.

We therefore start by briefly defining the basics of the relevant communication games when the strategy space of the agent is enlarged allowing for her action choice. Players' payoffs are given by the vNM utility  $U(y_1, \dots, y_N, e, \theta)$  and  $V_i(y_1, \dots, y_N, e, \theta)$  for the agent and for principal  $i \in N$ , respectively.

If the action  $e$  is contractible, then every principal  $i$  decision is given by the mapping  $y_i(e) : E \rightarrow Y_i$ . That is, for every observed  $e$  the principal can commit to offer the allocation  $y_i$ ; notice that non contractible actions (i.e. the moral hazard case) can be represented by considering the degenerate situation where the mapping  $y_i(\cdot)$  is independent of  $e$  for every  $i$ . Thus, principal  $i$ 's strategy is given by the choice of a message space  $M_i$  and of a map  $\sigma_i : M_i \rightarrow Y_i(e)$ , where  $Y_i(e)$  is the set of all the  $y_i$  mappings. As in the previous section,  $\Sigma_i$  denotes the collection of all the maps  $\sigma_i$  and  $\Sigma = \times_{i \in N} \Sigma_i$  is the Cartesian product of all the  $\Sigma_i$ .

The strategy for the agent is the map  $\sigma_0 : \Theta \times \Sigma \rightarrow M \times E$  and  $\Sigma_0$  is the collection of all the  $\sigma_0$ . For notational convenience, we define  $\sigma_0 = (m, e)$ , with  $m \in M$ . Finally, let  $m_i$  denote the projection of  $m$  on  $M_i$ .

The relevant game  $\Gamma_M^e$  can therefore be defined as follows:

$$\Gamma_M^e = \{\Theta, F(\cdot), (M_i, \Sigma_i)_{i \in N}, \Sigma_0, U(\cdot, e, \theta), (V_i(\cdot, e, \theta))_{i \in N}\}$$

As a consequence, we can define the associated Perfect Bayesian Equilibrium (PBE) notion:

**Definition 4** *Given a message space  $M$ , a strategy vector  $\sigma^* = ((\sigma_i^*)_{i=1}^n, \sigma_0^*) \in \Sigma \times \Sigma_0$  is a PBE of  $\Gamma_M$  if:*

$$\forall \theta \in \Theta, \forall (\sigma_i)_{i \in N} \in \Sigma, \sigma_0^* = (m^*, e^*) \in \arg \max_{(m, e) \in M \times E} U((\sigma_i(m_i))_{i \in N}, e, \theta)$$

$$\forall i \in N, \sigma_i^* \in \arg \max_{\sigma_i \in \Sigma_i} \int_{\theta \in \Theta} V_i(\sigma_i(m_i^*(\sigma_i, \sigma_{-i}^*(\cdot))), \sigma_{-i}^*(\cdot), e^*(\sigma_i, \sigma_{-i}^*), \theta) dF(\theta)$$

## 4.1 The moral hazard case

We start by analyzing the case in which the adverse selection problem co-exists with one of pure moral hazard. In the corresponding communication games, therefore, principal  $i$ 's strategies are not contingent on  $e$ .

The notion of separable preferences for the agent can be extended to this case in a direct way:

**Definition 5** *Strong Separability: the utility function  $U$  is strongly separable in the principals' decisions  $y_i$  for all  $\theta \in \Theta$ , if for all  $y_i, y'_i \in Y_i$ , for all  $y_{-i}, y'_{-i} \in Y_{-i}$  and for all  $e', e'' \in E$*

$$[U(y_i, y_{-i}, e', \theta) > U(y'_i, y_{-i}, e', \theta)] \implies [U(y_i, y'_{-i}, e'', \theta) > U(y'_i, y'_{-i}, e'', \theta)]$$

If preferences are strongly separable, then the preference ordering over principal  $i$ 's offers is independent of the other principals' offers, but also of the chosen effort.

In order to restate Theorem 1, we start by representing agent's utility as a function of the equilibrium strategies

$\sigma_0^* = (m^*, e^*)$  and  $\sigma_i^* = y_i^*: U(y_i^*(m_i^*), y_{-i}^*(m_{-i}^*), e^*(y_i, y_{-i}, \theta), \theta)$ . Considering the pair  $\hat{y}_{-i} \in Y_{-i}$ ,  $\hat{e} \in E$ , then the corresponding partial utility for the agent is:

$$\tilde{U}'_{\hat{y}_{-i}, \hat{e}}(y_i, \theta) = U(y_i, \hat{y}_{-i}, \hat{e}, \theta)$$

This notion allows to parallel our Claim 1 in the following way:

**Claim 2** *If  $\sigma^* = (\sigma_0^*, (\sigma_i^*)_{i \in N})$ , with  $\sigma_0^* = (m^*, e^*)$  is a PBE then for all  $i \in N$ , for all  $\hat{y}_{-i} \in Y_{-i}$  and for all  $\hat{e} \in E$ :*

$$m_i^* \in \arg \max_{m_i \in M_i} \tilde{U}'_{\hat{y}_{-i}, \hat{e}}(\sigma_i^*(m_i), \theta)$$

Proof. Suppose not. Then, there must exists a  $m'_i \in M_i$  such that:  $\tilde{U}'_{\hat{y}_{-i}, \hat{e}}(\sigma_i^*(m'_i), \theta) > \tilde{U}'_{\hat{y}_{-i}, \hat{e}}(\sigma_i^*(m_i^*), \theta)$ . It follows that:  $U(\sigma_i^*(m'_i), \hat{y}_{-i}, \hat{e}, \theta) > U(\sigma_i^*(m_i^*), \hat{y}_{-i}, \hat{e}, \theta)$ . Then, given the strong separability assumption, we get:

$$U(\sigma_i^*(m'_i), \sigma_{-i}^*, e^*, \theta) > U(\sigma_i^*(m_i^*), \sigma_{-i}^*, e^*, \theta)$$

contradicting the assumption that  $\sigma^*$  is a PBE. ■

This Claim shows how in an environment with strongly separable preferences for the agent, in every PBE the message sent to principal  $i$  will not, therefore, be affected either by the mechanism designed by other principals or by the effort choice. Hence, this claim mimics the result obtained in the pure adverse selection case.

The next step is to provide a general argument:

**Theorem 3** *If the preferences of the agent satisfy the Strong Separability condition, then for every outcome that can be supported as a PBE of an arbitrary indirect mechanism game  $\Gamma_M^e$ , there exists a game  $\Gamma_\Theta^e$  associated to direct mechanisms that implements the same outcome.*

Proof. We have to show that every payoff profile associated to the equilibrium strategies  $\sigma^*$  of the game  $\Gamma_M^e$  can be replicated in the game  $\Gamma_\Theta^e$ .

Let the mechanism game  $\Gamma_\Theta^e$  to be:

$$\Gamma_\Theta^e = \{\Theta, F(\cdot), (\Theta, \tilde{\Sigma}_i)_{i \in N}, \tilde{\Sigma}_0, U(\cdot, e, \theta), (V_i(\cdot, e, \theta))_{i \in N}\}$$

Players' strategies are:  $\tilde{\sigma}_i : \Theta \rightarrow Y$  for principal  $i \in N$  and  $\tilde{\sigma}_0 : \Theta \times \tilde{\Sigma} \rightarrow \Theta \times E$  for the agent. Now, we can define the function  $\tilde{V}_i$  such that:

$$\tilde{V}_i(\tilde{\sigma}_i(\cdot), \tilde{\sigma}_{-i}(\cdot), \theta) = V_i(\tilde{\sigma}_i(\cdot), \tilde{\sigma}_{-i}(\cdot), e^*(\tilde{\sigma}_i(\cdot), \tilde{\sigma}_{-i}(\cdot)), \theta)$$

We consider the following strategies:  $\tilde{\sigma}_i(\theta) = \sigma_i^*(m_i^*(\sigma_i^*, \theta))$  for principal  $i \in N$  and  $\tilde{\sigma}_0 = [m_i^*(\theta, \tilde{\sigma}) = \theta, e^*(\sigma_i^*(m_i^*(\sigma_i^*, \theta)), \sigma_{-i}^*(m_i^*(\sigma_{-i}^*, \theta))]$  for the agent.

Then, considering  $\tilde{V}_i$  for  $i \in N$ , we can replicate the argument already given in Theorem 1. ■

Following the discussion of the previous Section, we also have to show that the set of equilibria is not enlarged by restricting to direct mechanisms.

**Theorem 4** *Consider a message space  $M = \times_{i \in N} M_i$  such that  $\Theta \subseteq M_i$  for all  $i \in N$ . If  $\sigma^*$  is a PBE in the common agency game  $\Gamma_\Theta^e$ , then  $\sigma^*$  is a PBE in the common agency game associated with message space  $M$ , that is  $\Gamma_M^e$ , then every outcome that can be supported as a PBE of the direct mechanism game  $\Gamma_\Theta^e$  can be as well implemented in the game  $\Gamma_M^e$*

Proof. See the Proof of Theorem 2. ■

Finally, it is important to remark that all the results of Sections 2 and 3 can be extended to allow for mixed strategies in a natural way; we just need to consider the separability requirement for lotteries over deterministic contracts. Obviously, if the separability requirement holds for lotteries, it holds as well for deterministic contracts. If we look at the mixed extension of the common agency game, then any agent's optimal strategy  $\sigma_0^*$  is an element of the set of probability measures over  $M$  times the set of probability measures over  $E$ :

$$\sigma_0^* \in \Delta(M = \times_{i \in N} M_i) \times \Delta E$$

If the strengthened separability condition holds, then  $\sigma_0^* \in \times_{i \in N} \Delta(M_i) \times \Delta E$ . Observe that

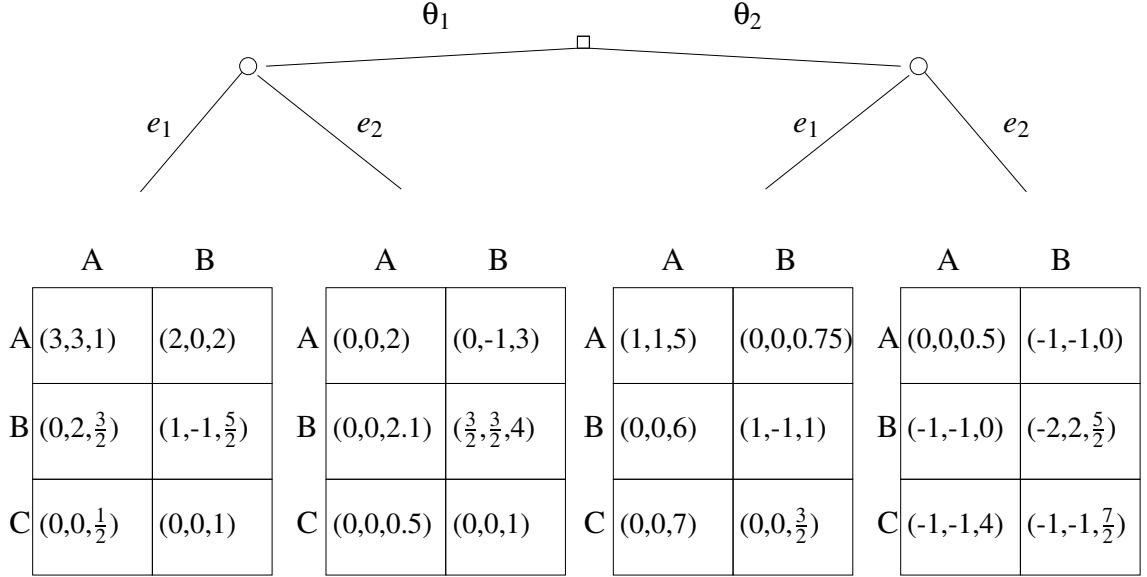
$$\times_{i \in N} \Delta(M_i) \times \Delta E \subseteq \Delta(\times_{i \in N} M_i) \times \Delta E$$

What the strong separability gives us is the following: this condition holding, the agent can look for the optimal strategy in the smaller set. However, by restricting his attention to the smaller set the agent ceases to remain a correlating device for the principals. Hence, principals can restrict attention to direct mechanisms without loss of generality.

## 4.2 Contractible actions

Allowing for the action  $e$  to be contractible introduces a new source of externality among principals. Our Strong Separability condition does not suffice to re-establish the Revelation Principle in this setting, precisely because the contractible effort choice introduces an essential problem of non-separability. This links the principals' choices and gives them an incentive to design more complex mechanisms.

We provide here a counter-example showing that when  $e$  is contractible, then an outcome implemented by standard indirect mechanisms (menus) *cannot* be supported by looking at direct mechanisms only, even if the strong separability condition is fulfilled.



In the example above the agent's type space is  $\Theta = \{\theta_1, \theta_2\}$  and she is allowed to take the action  $e \in E = \{e_1, e_2\}$ , which is assumed to be contractible. There are two principals: principal 1 has the three available actions  $\{A, B, C\}$ , while principal 2 can choose between  $A$  and  $B$  only. The outcome  $\{(\frac{3}{2}, \frac{3}{2}, 4)$  if  $\theta = \theta_1$ ,  $(1, 1, 5)$  if  $\theta = \theta_2\}$  cannot be sustained as an equilibrium in the direct mechanism game. If it were an equilibrium, then principals' strategies should involve the following behaviors: whenever the message  $\theta_1$  is sent and  $e_2$  is performed, then both principals play  $B$ ; whenever the message  $\theta_2$  is sent and  $e_1$  is performed, they both play  $A$ .

Principal 2 has a profitable deviation in playing  $B$  irrespectively of both the received message and the performed action. If he takes such a strategy, then the agent has an incentive to misreport her type and choose  $e_2$  when she is  $\theta_2$ . In such a case principal 1 is playing  $B$  and principal 2 earns a payoff of 4 instead of 3.

We stress here that the same outcome can indeed be supported by allowing principals to offer menus. Let us consider the following strategy profiles: both principals play  $A$  if  $e_1$  is chosen; when the agent plays  $e_2$ , then principal 2 plays  $B$  and principal 1 offers the menu  $\{B, C\}$ . If principal 2 deviates to a strategy that prescribes to play  $B$  irrespectively of  $e$ , then the  $\theta_2$  agent will play  $e_2$  and ask for  $C$  from principal 1, that makes the deviation unprofitable.

Even with separability, therefore, there is a potential role for the agent as a correlating device among principals. Indirect mechanisms and, in particular, menus may help to exploit this, while direct mechanisms cannot.

## 5 Discussion

This work deals with a general class of games. In contrast with the most traditional way of representing common agency interactions, we explicitly allowed for *direct* externalities among principals. That is, principal  $i$ 's payoff directly depends on principal  $j$ 's strategy. The most popular works in common agency (e.g. Bernheim and Whinston, 1986) are focused on *indirect* externalities: the payoff of every single principal is supposed to depend on other principals' strategies only through the effect of those strategies on agent's actions. That is, whenever principal  $i$  is deviating, he will not be punished by changes in principal's  $j$  strategies unless the agent modifies her choices.

A recent paper by Peters (2003) is also investigating the relationship between players' preferences and the set of sustainable equilibria in a common agency game where the agent is taking actions (either contractible or non-contractible). His main concern is to identify conditions that enable to implement the equilibrium outcomes of a general game where principals compete by offering arbitrarily complex menus by simple take it or leave it offers. The relevant condition is referred to as *no-externality* and it assumes principal  $i$ 's payoff to depend only on his own strategies and on the agent's action. Furthermore, the agent is supposed to have a weak preference ordering over principal  $i$ 's offers, conditional on her effort choice, that should be *independent* of her type and of principal  $j$ 's offers, for every  $j \neq i$ . This last condition holding, there is a clear correspondence between take it or leave it offers and direct mechanisms.

Peters' analysis differs from ours in several important respects. These differences can be more clearly understood if we consider independently the incomplete information and the moral hazard with incomplete information frameworks.

When applied to a pure incomplete information problem, Peters' condition is much more stringent than ours. First, we do not need to assume any restriction on principals' preferences; moreover, Peters requires that the agent's ranking over alternatives has to be independent of her type. This way, no real revelation occurs through the contract, as it is a *take it or leave it offer*, unconditional on agent's messages. Notice that we have just assumed that the agent has a weak preference ordering over each principal's action; this ordering needs not to be independent of the agent's type. Under our weak separability assumption, the Revelation Principle can therefore be reestablished in a large class of incomplete information games.

When moral hazard issues are embedded in the analysis, our results and those of Peters can be seen as complements. Our condition is less restrictive in two dimensions. First, it allows for direct externalities among principals. Second, it does not require the agent's preferences ordering over principals' actions to be independent of her type. Peters *no-externality* condition, on the contrary, is less demanding in terms of the agent's actions, as the agent's ranking over the principals' action is conditional on the agent's choice of effort. Our *strong-separability* condition forces this ranking to be independent of effort.

To summarize: whenever incomplete information is an important issue for the analysis, then our *separability* conditions provide a way to overcome the conflict between externality and revelation and enables to restrict the analysis to direct mechanisms without loss of generality. A clear drawback of our analysis is that it is only suitable for interactions where the agent takes uncontractible actions. If actions are assumed to be contractible, then a new source of externalities among principals is introduced. Our condition does not suffice in this setting, precisely because of the non-separability that the (contractible) action choice introduces in the problem.<sup>7</sup> Contractible actions seem therefore to be incompatible with revelation of the agent's information through a direct mechanism. Hence, if we focus on a general analysis of incentive provision, Peters *no-externality* condition turns out to be the relevant reference. Notice that Peters condition is able to encompass the case of contractible actions precisely because his requirement eliminates any possible revelation of the agent's type through the contract.

Finally, an important property of our approach is that it can deal with mixed strategies in a natural way. That is, if principals are allowed to randomize over mechanisms, then strengthening the notion of separability for agent's preferences still gives a sufficient condition for the validity of the Revelation Principle. The same property does not hold in Peters (2001) framework: if he considered principals randomizing over mechanisms, then the correspondence between equilibria in menu games and in *take it or leave it offer* games cannot be maintained.<sup>8</sup>

## 6 Concluding Remarks

We provided a sufficient condition that determines a class of *common agency* games in which the Revelation Principle retains all its power. Every allocation supported as an equilibrium inside an arbitrarily huge class of indirect mechanisms can therefore be implemented through simple direct mechanisms. We focused on games where the utility of the agent satisfies a *separability property* with respect to the actions taken by all the principals and by herself.

In an incomplete information framework, the requirement is that of *weak separability*. If this property holds, then messages sent by the agent to every principal will, at equilibrium, depend only on her type. Thus, a direct mechanism is a sufficiently powerful tool, and nothing is gained from allowing the principals to design more complex indirect mechanisms.

In a more general setting in which the agent is allowed to take also some, non-contractible action, we have shown that the same result holds, provided the condition of separability is appropriately defined. In this general setting, agent's preferences are required to satisfy a *strong separability condition*. The result could

---

<sup>7</sup>Our results can be extended to the contractible effort case only in a limited sense. If the agent takes independent effort choices for each principal, then the separability property is preserved. As a consequence our main theorems can still be applied.

<sup>8</sup>This is still a consequence of the allowance for contractible actions for the agent (see Theorem 4 in Peters, 2003).

therefore potentially apply to auctions where the agent has an unknown reservation value for the object he earns or produces, and the principals make offers for the object (or quantity) that the agent accepts or rejects. Applying our result to this framework ensures us that the Revelation Principle holds, if there is no limitation in the amount of units of the object that can be produced. If only one unit is produced, then changing the offer of one principal might break the agent's ranking over the others and the separability requirement violated.

The class of games we have identified has some appealing properties. First, no condition is imposed on the principals' preferences, which can depend on any other principal's action. Second, the condition we impose on the agent's preferences is not a pure additive separability condition, but a milder separability requirement.

We believe, therefore, that the results in this paper can be a useful tool for researchers to check whether the Revelation Principle can be applied in their models under study.

## References

- BERNHEIM, B. D., AND M. D. WHINSTON (1986): "Common Agency," *Econometrica*, 54(4), 923–942.
- BIAIS, B., D. MARTIMORT, AND J.-C. ROCHET (2000): "Competing mechanisms in a common value environment," *Econometrica*, 78(4), 799–837.
- DIXIT, A., G. M. GROSSMAN, AND E. HELPMAN (1997): "Common agency and coordination: General theory and application to government policy making," *Journal of Political Economy*, 105(4), 752–769.
- EPSTEIN, L. G., AND M. PETERS (1999): "A revelation principle for competing mechanisms," *Journal of Economic Theory*, 88(1), 119–160.
- KAHN, C. M., AND D. MOOKHERJEE (1998): "Competition and Incentives with Nonexclusive Contracts," *RAND Journal of Economics*, 29(3), 443–465.
- MARTIMORT, D. (1992): "Multi-principaux avec anti-selection," *Annales d'Economie et de Statistique*, 0(28), 1–37.
- MARTIMORT, D., AND L. A. STOLE (2002): "The revelation and delegation principles in common agency games," *Econometrica*, 70(4), 1659–1673.
- PARLOUR, C. A., AND U. RAJAN (2001): "Competition in Loan Contracts," *American Economic Review*, 91(5), 1311–1328.
- PAVAN, A., AND G. CALZOLARI (2003): "A markovian revelation principle for common agency games," mimeo Northwestern University.
- PECK, J. (1997): "A note on competing mechanisms and the revelation principle," mimeo Ohio State University.

- PETERS, M. (2001): “Common Agency and the Revelation Principle,” *Econometrica*, 69(5), 1349–1372.
- (2003): “Negociation and Take-it-or-leave-it in common agency,” *Journal of Economic Theory*, 111(1), 88–109.
- ROCHET, J.-C. (1986): “Le contrôle des équations aux dérivées partielles issues de la théorie des incitations,” Ph.D. thesis, Université Parix IX.
- SEGAL, I., AND M. D. WHINSTON (2003): “Robust predictions for bilateral contracting with externalities,” *Econometrica*, 71(3), 757–791.