

“The maximum principle for the fulfilling of wishes guides the building up of the world by requiring that it be the best possible.”

– Kurt F. Gödel

University of Alberta

**OPTIMAL MECHANISMS FOR MACHINE LEARNING
A GAME-THEORETIC APPROACH TO DESIGNING MACHINE LEARNING
COMPETITIONS**

by

Mohammad Mahdi Ajallooeian

A thesis submitted to the Faculty of Graduate Studies and Research
in partial fulfillment of the requirements for the degree of

Master of Science

Department of Computing Science

©Mohammad Mahdi Ajallooeian
Spring 2013
Edmonton, Alberta

Permission is hereby granted to the University of Alberta Libraries to reproduce single copies of this thesis and to lend or sell such copies for private, scholarly or scientific research purposes only. Where the thesis is converted to, or otherwise made available in digital form, the University of Alberta will advise potential users of the thesis of these terms.

The author reserves all other publication and other rights in association with the copyright in the thesis, and except as herein before provided, neither the thesis nor any substantial portion thereof may be printed or otherwise reproduced in any material form whatever without the author's prior written permission.

Dedicated to Csaba.
A patient teacher and a good friend.

Abstract

In this thesis we consider problems where a self-interested entity, called the principal, has private access to some data that she wishes to use to solve a prediction problem by outsourcing the development of the predictor to some other parties. Assuming the principal, who needs the machine learning solution, and the potential providers of the solution are two independent, self-interested agents, which is the case for many real-world situations, this then becomes a game-theoretic problem. We propose mathematical models for variants of this problem by borrowing techniques from the literature of mechanism design and provide principled solutions. We consider experimental design when there are multiple self-interested agents involved in developing a solution for a machine learning problem. A first case is when there is a public competition, each agent offers a single solution and solutions are available off-the-shelf to the agents: there is no development cost included. The problem considered is to find a set of payment rules that guarantees to maximize the profit of the principal on expectation even when the developers are self-interested. The solution depends on the distribution of the skill-level of developers available, which is assumed to be known. To deal with our problem, the standard mechanism design techniques are revisited and extended in a number of ways. In particular, a general approach is given that allows the design of payment rules (more generally, mechanisms) when such payment rules must depend on some quantity that becomes known only after the mechanism is executed. This extension plays a key role in our solution to the machine learning payment-rule design problem, where data must be kept private (otherwise the developers could submit “overfitting” predictors), yet the principal’s profit (and thus the payment) should depend on the performance of the predictor chosen on the

private data. Then, we address other interesting variants of the problem and provide solutions : when a single developer can submit multiple solutions, and when the solution is to be developed in multiple stages, or when the development cost is non-zero.

Acknowledgements

My gratitude goes towards my wonderful supervisor, Csaba Szepesvári, my co-supervisor Dale Schuurmans and also András György for his big helps. I would also like to thank my family for their enduring support and love.

Contents

1	Introduction	1
1.1	Contributions	2
1.2	An Example	3
1.3	The Problem	4
1.4	Thesis Organization	6
2	Basics of Mechanism Design	7
2.1	The Formal Definition of Games	8
2.1.1	Equilibrium Concepts	11
2.2	Bayesian Games: Towards Mechanism Design	14
2.3	Mechanism Design	18
2.3.1	Bayesian Mechanism Design	20
2.3.2	Mechanism Design with Voluntary Participation	22
2.3.3	The Revelation Principle	24
2.3.4	Voluntary Participation	29
2.4	Mechanisms in the World of Money	31
2.5	Solution Methodologies	31
2.6	Summary	32
2.6.1	Known Limitations	32
2.6.2	Future Work	33
3	Machine Learning Solution Procurement	35
3.1	Problem Description	36
3.2	The Form of the Optimal Reverse Auction	38
3.3	The Proof of Optimality	40
3.4	Discussion	45
3.5	Summary	47
4	Machine Learning Competitions with Public Test Results	48
4.1	The Formal Problem Definition	49
4.1.1	Bayesian Optimal Mechanism Design with Information Leakage	49
4.1.2	Formal Problem Definition for Machine Learning Auctions with Public Testing	54
4.2	Solution	56
4.2.1	Reduction of Solving a BOMD with Ex-Ante Information Leakage to Solving Standard BOMDs	57
4.2.2	BOMD with Ex-Ante Information Leakage and Voluntary Participation	59
4.2.3	Solution to the Machine Learning Procurement Problem	62
4.3	Diversion: Ex-post Information Leakage	63
4.4	Summary	65
4.4.1	Known Limitations	66

4.4.2	Future Work	66
5	Mechanism Design with Exogenous Effects	67
5.1	Games with Exogenous Signals	67
5.2	Mechanism Design with Exogenous Signals	73
5.3	Application to Machine Learning Auctions	77
5.4	Summary	78
5.4.1	Known Limitations	79
5.4.2	Future Work	79
6	Developers have Multiple Machine Learning Packages	80
6.1	Problem Description	81
6.2	The Form of the Optimal Multi-Unit Auction	82
6.3	Proof of Theorem 6.2.1	85
6.4	Summary	90
6.4.1	Future Work	90
7	A Comprehensive Procurement Process	91
7.1	Solution to the Example of Section 1.2	91
8	Conclusions	93
8.1	Future Work	93
	Bibliography	95

List of Figures

2.1	The RP, voluntary participation and individual rationality . .	30
4.1	Protocols of interaction with information leakage	50

Chapter 1

Introduction

Machine learning competitions are becoming increasingly popular and more companies are seeking to find predictors in such a way. There are several contracts signed every day where finding predictors to be used on some data is outsourced. A characteristic property of such deals is that the party that is going to use the solution to achieve/maximize a real-world objective and the party developing the solution are different. A good example is the Netflix prize: an open competition launched in October 2006 which awarded the team with the best (collaborative filtering) algorithm predicting user film ratings based on previous ratings a prize of one million dollars.

A natural question to ask then is what is the best way of organizing a competition. Do current practices achieve the desired results? How would a competition designed to maximize profit (or social welfare) be run? Surely, an ad-hoc procedure will not yield the best results possible, but it may and even if it does, there are no guarantees. This thesis is an attempt to find principled ways to tackle the problem of outsourcing the development of a machine learning solution.

More precisely, in this thesis we focus on issues that arise when a company wants to buy a machine learning solution from several developers. How should then the price be set? What protocol yields to maximal profit for the company? This is what we call the machine learning solution procurement problem. The aim of this thesis is to initiate a rigorous way to study these problems. More precisely, the goals of the thesis are to:

1. Propose a mathematical framework that can be used to study machine learning solution procurement problems so that the alternative solutions can be analyzed and studied using the tools of mathematics in a rigorous fashion.
2. Provide answers to simple, stylized instances of the machine learning solution procurement problem.
3. Develop general tools and techniques to tackle the specific issues that come up when solving these problems.

The framework used to achieve the above goals is that of *mechanism design*, a subfield of game theory that was born in the 1960s (Arrow and Debreu, 1954) and has been the subject of extensive study since then, mainly in the field of economics (Bolton and Dewatripont, 2005; Laffont and Martimort, 2002). Within the framework of mechanism design, we study stylized versions of machine learning procurement problems. Extension to related problems, such as the problem of optimizing machine learning competitions, will be shortly discussed at the end of the thesis.

1.1 Contributions

The contributions of the thesis are the following:

1. A single unified equilibrium concept is developed that unifies several existing equilibrium concepts (Nash, Dominant Strategies and Bayesian-Nash equilibria) is offered (Definition 2.2.6). The benefit of the unified concept is that open up a new alley of thinking about equilibria in games. Whereas previous practice in game theory was to prove a statement individually for each equilibrium concept, the unified concept allows a single proof, as we shall demonstrate on several occasions. Further, the new definition allows one to discover the key distinguishing properties of equilibria that are needed in these proofs and thus it increases our depth of understanding equilibria.

2. The concept of admissible counter-strategy map schemas is introduced (Definition 2.3.8) and it is shown that admissibility of a counter-strategy map schema $\mathfrak{C}$ is a sufficient condition for the revelation principle to hold for the generalized equilibria the uses $\mathfrak{C}$ (Theorem 2.3.1).
3. A characteristic feature of the machine learning procurement problem, unpreventable information leakage, is formalized (Definition 4.1.1). Two sub-cases, the ex-ante and ex-post information leakage problem are identified and formalized (Section 4.1.1).
4. It is shown how the procurement problem of machine learning solutions can be formalized as an example of a mechanism design problem with information leakage (Section 4.1.2).
5. A reduction-based solution for mechanism design problems with ex-ante information leakage is offered (Theorem 4.2.2) and applied to solve the problem of machine learning procurement (Section 4.2.3).
6. A model of mechanism design problems with exogenous signals is introduced (Definition 5.1.1 and Definition 5.1.2)
7. Two reductions for solving mechanism design problems with exogenous signals are introduced: A general reduction (Theorem 5.2.1) and a simplified reduction that applies only in a specific, yet practical setting (Theorem 5.2.2). As an example, it is shown how a company who wants to hide its profit predictions can do so without a decrease of the profit that can be achieved in a machine learning solutions procurement problem.
8. The procurement problem when developers have multiple machine learning packages to offer is formalized (Section 6.1) and solved (Theorem 6.2.1).

1.2 An Example

One option for the company then is to hold a competition with a sizeable monetary award to be given to the competitor who submits the best solution.

However, a sizeable portion of the sales of the company happens in Nonextraenusia, a country in which advertising is prohibited due to social reasons and so if the company declared the name of the company who did the best job, this would be considered as advertisement and Aleph Corp. would lose its sales in Nonextraenusia. If Aleph Corp. still wants to run the competition, it won't be able to publicly announce the identity of the winner. As a result, no organization would participate in such a competition unless it could at least recover the cost of the submitted solution on expectation. (In economic terms, it can be said that there exist no reasons external to the trade procedure to participate in the competition, thus participation will be based on pure utilitarian reasoning.) Thus, participation may be hard to predict and the selection of the prize amount may dramatically influence the surplus of the competition. This choice is made especially difficult since the quality of the data, which is unknown at the time when the prize needs to be announced, will influence how good the recommendations will be, and thus, the profit. This creates a vicious circle, making it impossible to select the prize to optimize the expected profit.

Aleph Corp. has two alternatives to holding a competition: (i) to procure a solution in a process where the suppliers would be incentivized to participate, or (ii) to enter into a work contract with some of the suppliers selected in some principled fashion. However, Aleph Corp. prefers a one-shot deal, i.e., a procurement so as to avoid complication of contracting (verification of work, legal issues, etc.) and because it needs a solution in a short amount of time. In this thesis we will consider how Aleph Corp. should conduct its procurement process so as to maximize its expected profit.

1.3 The Problem

The goal of Aleph Corp. is to obtain the best quality solution for the least price on expectation. In this thesis we suggest a game theoretic approach this problem, and apply the methodology of *mechanism design*. To manage to get the best solution-price trade-off, there are several issues for the company to

face with. How can the subcontractors be incentivized to obtain fair prices? How should the data be used in deciding who to buy the solution from? In this work, we provide some initial answers to these questions for some simplified versions of this problem. In this problem we assume that the subcontractors have possibly multiple off-the-shelf solutions, which they agree to submit for testing to the company. The problem is a special case of buying some good from a set of sellers.

In the standard setting, the sellers are interested in maximizing their profit, which conflicts with the goal of the buyer who wants to maximize her own profit. The (standard) solution to such procurement problems (to be described in detail in Chapter 3) is to run a so-called reverse sealed-bid auction where the sellers submit their bids privately and simultaneously to the buyer. The buyer then calculates the “virtual” profitability of each of the bids with a specific formula where the actual profits are decreased by an amount that depends on the buyer’s initial uncertainty concerning the “fair price” of the individual goods offered for sale. The seller whose bid has the highest virtual value is the winner and he receives a payment, which is guaranteed to be at least as large as the price he submitted. The difference between the price submitted and the actual payment is the price that the buyer needs to pay to make up for her lack of knowledge of the “true fair prices” of the goods offered for sale.

In machine learning procurement, one difference to the standard procurement problem is that through submitting their solutions (for testing), the developers reveal information to the company. How should then the above solution be modified? Another difference is that each developer might offer multiple solutions. Finally, the selection of the winner depends on predictions of the profitability of the individual solutions. However, companies might be reluctant to reveal this information. However, by hiding this information the company cannot reveal the rules of selecting the winner of the reverse auction. Is it possible for a company to hide this information and yet achieve the best possible profit or is there some price to be paid?

1.4 Thesis Organization

This thesis is organized as follows: Chapter 2 explains the basics of game theory and mechanism design: we start by an introduction of game theory in section 2.1 and follow up with mechanism design in Section 2.3. Chapter 3 explains the problem of reverse auction as a basic model for the problem of machine learning procurement and presents the well-known optimal solution for it. The first issue specific to machine learning solution procurement problem, the unpreventable information leakage due to the developers submitting their solutions for evaluation, is formally defined in Chapter 4, where a reduction-based solution is also given. A second specific issue, when the company does not want to publicize his profit predictions, is considered in the next chapter (Chapter 5). More generally, in this chapter we introduce the concept of games with exogenous signals. Again, a reduction-based solution is presented. Chapter 6 considers the problem when the developers can submit multiple solutions. The results of Chapter 3 are extended to this case. In Chapter 7. we give a comprehensive solution to the case when the developers may submit multiple solutions which get evaluated on some test data, but the company withholds its profit predictions. We conclude the thesis in Chapter 8.

Chapter 2

Basics of Mechanism Design

Game theory is defined as “the study of mathematical models of conflict and cooperation between intelligent rational decision-makers” by Myerson (1997). In a *game* several *players* (also called agents) interact in some well-defined way, governed by the *rules* of the game: the players choose *actions*, and these choices result in some *outcomes* of the game. The players have their own assessment about the value of different outcomes; the difference in evaluating the outcomes may come from different positions of the players, or simply because the players are of different *kinds* (e.g., the value of an object in an auction is likely to be different for different participants).

In classical game theory the agents are assumed to act *rationally*, that is, in a *selfishly optimal* fashion to reach the outcome most desirable for themselves. This strategically rational behavior of the agents, which also takes into account that other agents will behave rationally, allows us to define equilibria in games: in an equilibrium, no single agent can improve the outcome of the game (for herself) by modifying her actions. A large part of the literature in game theory concerns the analysis of equilibrium in games. Throughout this thesis we will assume that the agents act rationally (an alternative to the rationality assumption is to model the agents’ behavior, leading to behavioral game theory). In the thesis we will consider a special area of game theory, called *mechanism design*, which studies *methodologies of designing games in order to achieve a specified objective*. Mechanism design suites well to the original problem we wish to consider: to design the rules of machine learning

procurement problems.

In this chapter we introduce the basic concepts of game theory and mechanism design (for more detailed description of these topics the reader is referred to, e.g., the books by Myerson, 1997 and Mas-Colell et al., 1995). Furthermore, we define and analyze a unified equilibrium concept.

2.1 The Formal Definition of Games

We define multi-agent games with a fixed number of players, also called *agents*. To keep things simple, let us assume that the number of players n is finite. In the games considered, the player at a fixed position has a fixed set of actions, or a fixed set of *messages*. The player positions are numbered from 1 to n and we let $I = \{1, \dots, n\}$ to denote the set of player positions. For $i \in I$, let Σ_i be the set of messages available to the agent playing at position i (these sets may depend on the position of the player in the game). Let O be the set of outcomes of the game. Then, the rules of the game simply specify the outcome as a function of the individual messages chosen by the individual agents playing the game. Thus, the rules of the game can be specified using a mapping $g : \Sigma \rightarrow O$, where $\Sigma = \times_{i \in I} \Sigma_i$ is the Cartesian product of the sets Σ_i , also called the *joint message space*.

In summary, we have the following definition:

Definition 2.1.1 (Game). A *game* is a tuple $G = (\Sigma, O, g)$ where given a set I indexing the agents, $\Sigma = \times_{i \in I} \Sigma_i$ is the set of *messages* agents send, O is the set of *outcomes* for the game and $g : \Sigma \rightarrow O$ is a function called the *outcome function*.

To specify the agents' preferences we will use utility functions. According to a theorem of von Neumann and Morgenstern (1947), if an agent's preference relations over experiments with two probabilistic outcomes satisfy four (reasonable) assumptions then the agent's preferences over arbitrary lotteries (experiments with probabilistic outcomes) can be described as the expectation of utilities assigned to the elementary outcomes.

Definition 2.1.2. A function that maps outcomes to real values is called a *utility function*.

By consensus, it is assumed that agents prefer outcomes with higher utilities. Formally, a utility function is thus a mapping $u : O \rightarrow \mathbb{R}$. In a game, each agent may have their own utility functions: The utility function of the agent playing at position i will be denoted by u_i (at this stage it is not yet important whether the utility functions belong to the agent or the position).

Since agents may choose their actions randomly and in general we will be interested in the resulting expected utility (from the point of view of some of the agents), we introduce some extra notation to avoid writing expectations or integrals in an explicit form. Also, this increases the burden on the reader at this stage, the notation introduced here will lead to a more comprehensible presentation of the concepts that come next.

Notation.

Distributions, densities, support:

If V is a measurable space¹, the space of probability distributions over V will be denoted by $M_1(V)$. If a distribution $P \in M_1(V)$ is absolutely continuous w.r.t. some fixed reference measure (that will always be clear from the context and is usually a Lebesgue measure when V is a subset of a Euclidean space), we denote the resulting density by the corresponding lower case letter, as in p . Lastly, we let $\text{supp}(P)$ denote the support of the distribution P .

Overloaded functions:

Let $\gamma : V \rightarrow \mathbb{R}$ be an arbitrary measurable real-valued function. We introduce a convention to overload the notation of function application for the case where a distribution $P \in M_1(V)$ is used as the argument of the function:

$$\gamma(P) = \int_V \gamma(v) dP(v) .$$

¹For readers not familiar with measure theory, it suffices to say at this stage that countable sets can be viewed as measurable space with no further difficulties. Further, Euclidean spaces and their subsets are also standardly equipped with the appropriate structure to make them “measurable” spaces where the usual calculus with probability measures works.

We can use this definition for functions of two (or more) variables v and w where both are respectively replaced by distributions P_v and P_w . In that case the order of integration will be denoted with numbers, e.g.:

$$\gamma(P_v^{\boxed{2}}, P_w^{\boxed{1}}) = \int_V \int_W \gamma(v, w) dP_w(w) dP_v(v) .$$

Note. We will drop the ordering if the ordering of integrals can be interchanged (due to an application of Fubini's theorem), as it will be common. In this case we will just write $\gamma(P_v, P_w)$, meaning that both orders are acceptable. In particular, when the function to be integrated is bounded, the result of the integral is guaranteed to be independent of the order of integration. Since we will use this notation in connection to utility functions, in what follows *we assume that all utility functions have a bounded range.*

Note. From now on, when we use probability distributions over some space V , we automatically assume that an appropriate measurability structure for the space V was fixed in advance. In particular, in what follows we assume that Σ_i ($i \in I$) and O are equipped with an appropriate measurability structure.

We need one more bit of notation before moving to the definition of “rational play”. As said before, we shall allow agents to use randomization to choose their messages. In effect, the agent playing at position $i \in I$ chooses a distribution from the $S_i = M_1(\Sigma_i)$, the space of distributions over the message space Σ_i .

Terminology 2.1.1 (Strategies). The elements of $S_i = M_1(\Sigma_i)$ will be called *mixed strategies* (or simply, *strategies*) and are usually shown by $s_i \in S_i$. We also introduce $S = \times_{i \in I} S_i$, which will be viewed as a subset of $M_1(\Sigma)$ the set of distributions over Σ with the natural embedding $\iota : S \rightarrow M_1(\Sigma)$, where $(s_1, \dots, s_{|I|}) \in S$ is viewed as a product measure over Σ . At times, we will denote S by $M_1^\times(\Sigma)$. Simple actions (degenerate, “Dirac” distributions) of the agents are also called *pure strategies*.

2.1.1 Equilibrium Concepts

Next, we formalize the notion of agents who act ‘selfishly optimal’: We start some more notation that will prove to be useful later.

Notation.

Components:

Subscripts are used to denote a value belonging to an agent, e.g., v_i for $i \in I$ denotes the v belonging to agent with index i .

Profiles:

Unsubscripted letters, most of the times, are used to denote vectors, e.g., $v = (v_i)_{i \in I}$, and are usually called profiles, e.g., if v_i s ($i \in I$) are ‘values’ for different agents, v is usually called a ‘value profile’. When the vectors belong to some linear space, they will be viewed as column vectors, unless otherwise declared.

Counter-components:

It is customary in game theory literature to denote by v_{-i} , the vector of all v_i s except for v_i :

$$v_{-i} = (v_1, v_2, v_3, \dots, v_{i-1}, v_{i+1}, \dots, v_{|I|}) = (v_{i'})_{i' \in I \setminus \{i\}} \quad .$$

Sets and ranges:

If a value is denoted by a letter, e.g., v , the set it belongs to is usually shown by the capital letter, V in this case. Some examples would be $v_i \in V_i$, $v \in V$ and $v_{-i} \in V_{-i}$. Also if $V \subset \mathbb{R}$ has a range, the infimum of V is denoted by $\underline{V}$ and the supremum of V is shown by $\overline{V}$.

Complementary pairs:

For any symbol v , for all $i \in I$, at a slight abuse of notation, we define $(v_i, v_{-i}) = v$. Likewise, $V_i \times V_{-i} = V$.

An agent playing at position i (in what follows: agent i) knowing the distributions (strategies) that its opponent play would naturally select a distribution

(strategy) that leads to the largest possible expected utility for himself. Any strategy with this property is called a best-response strategy to the opponent's joint strategy:

Definition 2.1.3 (Best-response). Fix an outcome function $g : \Sigma \rightarrow O$ and the utility functions $u = (u_i)_{i \in I}$. For $i \in I$, $s_{-i} \in S_{-i}$, introduce the set

$$B_i(g, u_i, s_{-i}) = \left\{ s_i^* \in S_i : u_i(g(s_i^*, s_{-i})) \geq \sup_{s'_i \in S_i} u_i(g(s'_i, s_{-i})) \right\},$$

the set of best responses for agent i given g, u_i and the strategies s_{-i} for the other agents.

We call an agent *rational* when for *any* given g, u_i, s_{-i} , the agent chooses a strategy from $B_i(g, u_i, s_{-i})$ (here, g essentially stands for the game).

Since agents in general cannot know how the other agents are going to act, it is not possible for the agents simply play a best-response. To find a way out of this apparent dead-end, remember that our goal is mainly to design games so that some desired outcome happens when rational agents are playing the game. Hence, our main concern is to predict properties of the joint strategy profile when rational agents play a game.

For starters, assume that there exists a joint strategy profile such that for each agent, the strategy of the agent in the profile is a best-response to *every possible counterstrategy*. It is very likely that if such a joint profile exists, a set of rational agents would adapt this profile. A joint strategy profile with this property is called a dominant strategies equilibrium:

Definition 2.1.4 (Dominant strategies equilibrium). Fix $G = (\Sigma, g, O)$ and $u = (u_i)_{i \in I}$ and let S be the joint strategy space corresponding to Σ . A strategy $s^* \in S$ is said to be a *dominant strategies equilibrium* for (G, u) if, for all $i \in I$, it holds that

$$s_i^* \in \bigcap_{s'_{-i} \in S_{-i}} B_i(g, u_i, s'_{-i}). \quad (2.1)$$

A dominant strategies equilibrium is also called *strategy-proof*. The problem with this approach is that many games lack a dominant strategies equilibrium.

Another possibility, which does not suffer from this latter problem (at least in the case of finite games), goes back to von Neumann and Morgenstern (1947) and completed by Nash (1951). Here, the assumption is only that when sufficiently intelligent agents play a game then in the strategy profile arising none of the players will have any incentive to deviate from the profile. In this case, we say that the joint strategy profile constitutes a *Nash equilibrium*. Formally, we have the following definition:

Definition 2.1.5 (Nash equilibrium). Fix $G = (\Sigma, g, O)$ and $u = (u_i)_{i \in I}$ and let S be the joint strategy space corresponding to Σ . A strategy $s^* \in S$ is said to be a *Nash equilibrium* for (G, u) if, for all $i \in I$, it holds that

$$s_i^* \in B_i(g, u_i, s_{-i}^*) . \quad (2.2)$$

Both of these concepts distinguish a certain subset of S . More generally, *game solution concepts* define special subsets of S with various desirable properties (Fudenberg and Tirole, 1991; Leyton-Brown, 2008; Mertens, 1989; Shoham and Leyton-Brown, 2008). Since in many cases the various statements that we will be interested in remain valid for at least the two above concepts (and some generalizations of them considered later), we introduce a common generalization of these two equilibrium concepts:

Definition 2.1.6 (Generalized equilibria). Fix $G = (\Sigma, g, O)$ and $u = (u_i)_{i \in I}$ and let S be the joint strategy space corresponding to Σ . For each $i \in I$, fix a mapping $+_{-i}: S \rightarrow 2^{S_{-i}}$, which we shall call the *counter-strategy mapping* and let $C = (C_{-i})_{i \in I}$ the collection of these mappings. We say that $s^* \in S$ is a *generalized equilibrium* (GE) strategy of (G, u, C) if for all $i \in I$,

$$s_i^* \in \bigcap_{s'_{-i} \in C_{-i}(s^*)} B_i(g, u_i, s'_{-i}) .$$

The set of all generalized equilibrium strategies shall be denoted by $\mathcal{E}(G, u, C)$.

In words, the idea is that agent *i*'s expectation is to be able to play the best response to all the counterstrategies in $C_{-i}(s^*)$.

The following result is an immediate corollary of the definitions:

Proposition 2.1.1. *The generalized equilibrium defined above encompasses both Nash and dominant strategies equilibria. In particular, to get the Nash equilibria choose $C_{-i}(s) = C_{-i}^N(s) = \{s_{-i}\}$ ($s \in S$), while to get the dominant strategies equilibria choose $C_{-i}(s) = C_{-i}^{DS}(s) = S_{-i}$ ($s \in S$).*

Proof. The proof follows immediately from the definitions. □

2.2 Bayesian Games: Towards Mechanism Design

In the problems we are interested in, the utility functions of the agents are typically unknown. More precisely, each of the agents are assumed to know their own utility functions, but neither the designer of the game, nor the other agents know the utility function of a specific agent. How can one then design games for the agents so that if the agents are rational the outcome of the game will be as desired when the agents play some type of equilibrium? In general, this is only possible when the requirements are relaxed.

One (standard) way of relaxing the requirements is as follows: In the new approach, it is assumed that for each position $i \in I$, the agent playing at that position is chosen from a known, fixed set of agents Θ_i such that for each $\theta_i \in \Theta_i$, the utility function for position i and agent θ_i is known. Furthermore, it is assumed that the distribution from which the agent for position i is chosen is known both to the game designer and all the agents who will play the game (i.e., everyone knows all the $|I|$ distributions). Since the agents do not know the utility functions of the other agents playing the game, but they know the distribution over these, the equilibrium concepts are adapted so that the agents play a best-response “in expectation” with respect to these distributions. Since the agents know their own identity (θ_i is known to agent playing at position i), the equilibrium strategy in general must depend on $\theta = (\theta_i)_{i \in I}$.

For the formal definitions, let us start with the generalization of best-responses. However, first let us fix some extra notation. In what follows, an element of Θ_i shall be called a possible *type* for the agents playing at position i . The utility functions will depend on the position i and also on which agent (what agent with what type) is playing at the given position. Thus, from now on, the utility functions will take the form

$$u_i : O \times \Theta_i \rightarrow \mathbb{R}, \quad i \in I.$$

For $\theta_i \in \Theta_i$ we will also use the notation $u_{i,\theta}$ to denote the function $o \mapsto u_i(o, \theta_i)$. We also call u_i a typed utility function when we want to emphasize that the utilities also depend on the agent types.

With this, we arrive at the definition of Bayesian games:

Definition 2.2.1 (Bayesian Game). A *Bayesian game* is defined by a tuple (G, Θ, u, P_θ) , where $G = (\Sigma, O, g)$ is a game, $\Theta = \times_{i \in I} \Theta_i$ is the set of agent type profiles, $u = (u_i)_{i \in I}$ are typed utility functions with $u_i : O \times \Theta_i \rightarrow \mathbb{R}$ and $P_\theta = (P_{\theta_1}, \dots, P_{\theta_{|I|}}) \in \times_{i \in I} M_1(\Theta_i)$ is a product distribution over the space of types.

The counterpart of best-response for “simple games” assumes that the agent knows what strategy its opponent would play. Since the opponent’s strategy is a function of their own types, this means that the best-response is defined for a $s_{-i} : \Theta_{-i} \rightarrow S_{-i}$ type-to-strategy mapping. Even though s_{-i} is given, the agent does not have access to the opponent’s types θ_{-i} . An agent playing a Bayesian best-response addresses this by playing a response which is best on expectation when the uncertainty over θ_{-i} is integrated out using $P_{\theta_{-i}}$. Finally, we need a best response for each type θ_i (the agent knows his own type!), leading to the concept of Bayesian best-response maps:

Definition 2.2.2 (Bayesian best-response maps). Let $\mathfrak{B} = (G, \Theta, u, P_\theta)$ be a Bayesian game and let $s_{-i} : \Theta_{-i} \rightarrow S_{-i}$ be a function mapping (incomplete)

type-profiles to (incomplete) strategy profiles. Given $i \in I$ and s_{-i} define

$$B_i(g, u_i, s_{-i}, P_{\theta_{-i}}) = \left\{ s_i^* : \Theta_i \rightarrow S_i : \begin{aligned} & u_{i, \theta_i}(g(s_i^*(\theta_i), s_{-i}(P_{\theta_{-i}}))) \geq \sup_{s'_i \in S_i} u_{i, \theta_i}(g(s'_i, s_{-i}(P_{\theta_{-i}}))), \\ & \theta_i \in \text{supp}(P_{\theta_i}) \end{aligned} \right\},$$

the set of *Bayesian best-response* maps for agent i .

As before, an agent playing in a Bayesian game $\mathfrak{B}$ is called rational if knowing that his opponents will play according to s_{-i} would play $s_i^*(\theta_i)$ for some Bayesian best-response map $s_i^* \in B_i(g, u_i, s_{-i}, P_{\theta_{-i}})$, where θ_i is his own type.

With this, we can generalize previous equilibrium concepts. We start with the the generalization of dominant strategies equilibria. However, first we need the notation of separable maps:

Definition 2.2.3 (Separable map). Given the product spaces $A = \times_{i \in I} A_i$, $B = \times_{i \in I} B_i$, a function $s : A \rightarrow B$ is called separable if for some functions $s_i : A_i \rightarrow B_i$, $s = (s_1, \dots, s_{|I|})$ in the sense that for any $a \in A$ it holds that $s(a) = (s_1(a_1), \dots, s_{|I|}(a_{|I|}))$. We let $\mathfrak{S}(A, B)$ denote the set of separable $s : A \rightarrow B$ maps. Further, for a separable map s and $i \in I$ we can define $s_{-i} : A_{-i} \rightarrow B_{-i}$ by $s_{-i} = (s_1, \dots, s_{i-1}, s_{i+1}, \dots, s_{|I|})$, i.e., $s_{-i}(a_{-i}) = (s_1(a_1), \dots, s_{i-1}(a_{i-1}), s_{i+1}(a_{i+1}), \dots, s_{|I|}(a_{|I|}))$, $a_{-i} \in A_{-i}$.

We can now give the definition of Bayesian dominant strategies equilibria.

Definition 2.2.4 (Bayesian dominant strategies equilibria). Let $\mathfrak{B} = (G, \Theta, u, P_\theta)$ be a Bayesian game. A strategy map $s^* : \Theta \rightarrow S$ is said to be a *Bayesian dominant strategies equilibria* map if s^* is separable and for all $i \in I$ it holds that²

$$s_i^* \in \bigcap_{s_{-i} \in (S_{-i})^{\Theta_{-i}}} B_i(g, u_i, s_{-i}, P_{\theta_{-i}}). \quad (2.3)$$

² If Θ_{-i} has a non-trivial measurability structure, instead of $(S_{-i})^{\Theta_{-i}}$ the set of all measurable mappings should be used. For simplicity, we disregard measurability issues. The reader who is worried about this may assume that Θ_{-i} is countable.

We note in passing that this definition is somewhat vacuous: It is not hard to see that an equilibrium of a Bayesian game (G, Θ, u, P_θ) is a Bayesian dominant strategies equilibrium if and only if it is a dominant strategies equilibrium of the pair (G, u) . Nevertheless, for the sake of completeness we keep this definition.

The generalization of Nash equilibrium leads to the following definition (Harsanyi, 1968a):

Definition 2.2.5 (Bayesian Nash Equilibria). Let $\mathfrak{B} = (G, \Theta, u, P_\theta)$ be a Bayesian game. A strategy map $s^* : \Theta \rightarrow S$ is said to be a *Bayesian Nash Equilibrium* map if s^* is separable and for all $i \in I$ it holds that

$$s_i^* \in B_i(g, u_i, s_{-i}^*, P_{\theta_{-i}}). \quad (2.4)$$

Thanks to Nash (1951), it is well known that the set of Nash equilibria is always non-empty. It is also known that the set of Bayesian Nash equilibria is non-empty (Shoham and Leyton-Brown, 2008, Section 6.1).

Finally, the notion of generalized equilibria can also be extended:

Definition 2.2.6 (Generalized Bayesian Equilibria (GBE)). Let $\mathfrak{B} = (G, \Theta, u, P_\theta)$ be a Bayesian game. For each $i \in I$, fix a mapping $C_{-i} : \mathfrak{S}(\Theta, S) \rightarrow 2^{S_{-i}^{\Theta_{-i}}}$, which we shall call the *counter-strategy mapping* and let $C = (C_{-i})_{i \in I}$ the collection of these mappings. We say that $s^* \in \mathfrak{S}(\Theta, S)$ is a *generalized Bayesian equilibrium* (GBE) strategy map of $(\mathfrak{B}, C)$ if for all $i \in I$,

$$s_i^* \in \bigcap_{s'_{-i} \in C_{-i}(s^*)} B_i(g, u_i, s'_{-i}, P_{\theta_{-i}}).$$

Given $\mathfrak{B}, C$, the set of all generalized Bayesian equilibrium strategy maps shall be denoted by $\mathcal{E}(\mathfrak{B}, C)$.

The following result is an immediate corollary of the definition:

Proposition 2.2.1. *The generalized Bayesian equilibrium defined above encompasses both Bayesian dominant strategy equilibrium maps and Bayesian Nash equilibrium maps. In particular, Bayesian dominant strategy equilibria maps can be obtained by choosing $C_{-i}^{\text{BDS}}(s) = (S_{-i})^{\Theta_{-i}}$, while Bayesian*

Nash equilibrium maps can be obtained by choosing $C_{-i}^{\text{BN}}(s) = \{s_{-i}\}$. Further, $\mathcal{E}(\mathfrak{B}, C^{\text{DS}}) \subset \mathcal{E}(\mathfrak{B}, C^{\text{BN}})$.

Proof. The statement follows immediately from the definitions. $\square$

It is also easy to see that the Bayesian equilibrium concepts generalize the previous concepts:

Proposition 2.2.2. *The generalized Bayesian equilibrium concept encompasses the generalized equilibrium concept.*

Proof. Simply define Θ_i to be a singleton for each i . $\square$

The significance of these results is that they show that if the generalized Bayesian equilibrium maps are shown to satisfy some desired property then it follows that all the other (specialized) equilibria concepts satisfy this property, too. Thus, it suffices to prove such statements for the generalized Bayesian equilibrium maps.

Building on top of game theory a new field of mechanism design was established with a paper by Vickrey (1961). Mechanism design is an attempt to find games such that an objective is met or some quantitative criterion is maximized. A short introduction on mechanism design will be given next. A good reference for further study would be (Myerson, 2008) and (Fudenberg and Tirole, 1991, Chapter 7).

2.3 Mechanism Design

Having laid out the basic concepts of game theory, we are now ready to address the problem of mechanism design.

To design a mechanism means to find or design a game such that a certain objective is achieved. Without loss of generality, we will assume that there is a designer of the game, whom, following the economics literature we shall call the *principal*. In a nutshell, the principal chooses a game and announces the rules publicly. The rules are binding for both the principal and the agents.³

³In the real-world this could be difficult to arrange because of privacy and security concerns. However, in theory secure cryptographic protocols can often be used to guarantee no cheating without revealing valuable extra information.

The game specification, as usual, involves the messages that the agents may send, together with the rules that determine the outcome given the messages. (The game to be designed may have external constraints on it. For example, oftentimes the agents can decide on their own whether they want to participate in the game. This can be captured, for example, by allowing the agents to send the message “I opt out”, which puts a constraint on the games.) The agents, after some deliberation, simultaneously send their messages to the principal, who, after receiving all the messages announces the outcome by following the rules of the game. The outcome bears various consequences to the principal, as well as to the agents (which are known beforehand to all parties involved): For example, as a “tangible” consequence of the outcome some agent (or agents) may do some work for the principal, and/or the principal may pay some money to one (or more) agents, or the principal may obtain some information that she did not possess beforehand. Just like the agents preferences, the preferences of the principal will also be captured using a utility function, $u_0 : O \rightarrow \mathbb{R}$, assuming that the outcome space is O .

To make inferences about the outcome of the designed game, the principal typically assumes that the players are rational, which, in accordance with the previous section, is captured using a notion of the equilibrium. When the utility functions of the agents are unknown, the principal’s objective is to maximize her utility no matter what agents “show up” (of course, the maximum utility that can be achieved will be limited by the “types” of agents that show up). When agents are chosen at random from some distributions, and the distributions for each player position are common knowledge, one possibility, which we shall follow, is to adopt the generalized Bayesian equilibrium concept to capture the notion of rational agents.

This leads to the formal definition of Bayesian mechanism design problems.

Note. One thing that should be minded for selecting equilibrium concepts is that while a Nash equilibrium will definitely exist and may be easier to find, it is too demanding on the agents and may give rise to problems of coordination (it is not a robust equilibrium concept). Dominant strategies, on

the other hand, are not very demanding from the agents and are very robust but require more work on the part of mechanism designer. Moreover, they might not always exist. Bayesian equilibria stand somewhat in the middle ground between these two ends, but they require the distribution of types to be known both to the principal and the agents which may prove to be a big assumption. It is also notable that in some cases, the mechanism designed will be very sensitive to these distributions and any inaccuracy in the distribution may take a big toll from the principal or the agents.

2.3.1 Bayesian Mechanism Design

The purpose of this section is to formally introduce the concepts and key techniques of (Bayesian) mechanism design. Let us start with the definition of a mechanism. In what follows, we shall assume that an outcome space O and a player position set I are given.

Definition 2.3.1 (Mechanism). A *mechanism* for an outcome space O and player position set I is a pair $M = (\Sigma, g)$, where $\Sigma = \times_{i \in I} \Sigma_i$ is a message space and $g : \Sigma \rightarrow O$ is an outcome function.

As suggested beforehand, the outcomes are rated from the perspective of the principal using the principal's utility function,

$$u_0 : O \rightarrow \mathbb{R}.$$

Assume for now that the agents utility function participating in the game are unknown, but that the agents' types $(\theta_i)_{i \in I} \in \times_{i \in I} \Theta_i = \Theta$ for each position $i \in I$ are selected at random from some distribution P_{θ_i} , independently from each other. It is assumed that the distributions, $P_{\theta} = (P_{\theta_i})_{i \in I}$ are common knowledge. Furthermore, the utility functions $u_i : O \times \Theta_i \rightarrow \mathbb{R}$ are also common knowledge.

The principal adopts a generalized Bayesian equilibrium concept for modeling rational play. For this, the principal needs to select for each Σ its desired counter-strategy map collection $C_{\Sigma} : \mathfrak{S}(\Theta, S) \rightarrow 2^{S^{\Theta_{-i}}}$.⁴ For a game

⁴For example, if the principal wants to adopt Bayesian Nash equilibria, she can use

$G = (\Sigma, O, g)$, it will be assumed that the rational players whose randomly selected, unknown to the principal, joint type is θ will play a joint $s^*(\theta)$ strategy, where s^* is an equilibrium map, $s^* \in \mathcal{E}(\mathfrak{B}, C_\Sigma)$, where $\mathfrak{B} = (G, \Theta, u, P_\theta)$ (cf. Definition 2.2.6). When it is important to emphasize that $\mathfrak{B}$ depends on $M = (\Sigma, g)$, we will use $\mathfrak{B}_M$.

With this, we can capture the notion of the expected utility that the principal can achieve with some fixed mechanism M :

$$u_0^*(M) = \sup_{s^* \in \mathcal{E}(\mathfrak{B}_M, C_\Sigma)} \int u_0(g(s^*(\theta))) dP_\theta.$$

By taking an optimistic viewpoint, the optimum is defined by taking the best of the possible equilibria. The main reason is technical, though in general it is commendable to avoid mechanisms with multiple equilibria because then coordination problems can also arise.⁵⁶ Finally, note that for some equilibrium concepts (i.e., counter-strategy maps) and some mechanisms M , the set $\mathcal{E}(\mathfrak{B}_M, C_\Sigma)$ may be empty. As usual, we use the convention $\sup \emptyset = -\infty$, thus effectively restricting the set of mechanism considering in the optimization problem to those for which $\mathcal{E}(\mathfrak{B}_M, C_\Sigma)$ is non-empty. We also note in passing that further restrictions on the considered equilibria are also possible, but we consider these out of the scope for the time being.

Let the space of all mechanisms with index set I and outcome space O be $\mathcal{M}(I, O)$:

$$\mathcal{M}(I, O) = \left\{ (\Sigma, g) : \Sigma = \bigtimes_{i \in I} \Sigma_i, g : \Sigma \rightarrow O \right\}.$$

With this, we can define the Bayesian optimal mechanism design problem:

$\Sigma \mapsto C_\Sigma$, where $C_\Sigma = (C_{-i})_{i \in I}$ (for brevity Σ is dropped from C_{-i}) with $C_{-i} : \mathfrak{S}(\Theta, S) \rightarrow 2^{S_{-i}}$, $C_{-i}(s) = \{s_{-i}\}$ ($s \in \mathfrak{S}(\Theta, S)$), where $S = \bigtimes_{i \in I} M_1(\Sigma_i)$.

⁵⁶The coordination problem is what the agents face when they play a game with multiple equilibria. When the principal, the designer of the game, also has the ability to communicate the agents which equilibrium they all should play, the rational agents, given this information and trusting that all other agents will act rationally and that the principal communicated in a consistent fashion, would have no incentive to deviate from their respective parts of the equilibrium strategy. Hence, when such pre-game communication is possible, the presence of multiple equilibria should not be of no major concern.

⁶*Implementation theory* is the subfield of game theory where the focus is on tracking all equilibria, as opposed in mechanism design, where the problem studied is to find a game that has an equilibrium with some desired property (Jackson, 2001; Maskin and Sjöström, 2002).

Definition 2.3.2 (Bayesian optimal mechanism design problem (BOMD)). Given the tuple $(I, O, \Theta, u, P_\theta, u_0)$ and a counter-strategy map schema $\mathfrak{C} : \Sigma \mapsto C_\Sigma$, the *Bayesian optimal mechanism design problem* is to find $M^* = (\Sigma^*, g^*)$ for which

$$u_0^*(M^*) = \sup_{M \in \mathcal{M}(I, O)} u_0^*(M).$$

We call M^* the optimal mechanism for $(I, O, \Theta, u, P_\theta, u_0)$ and $\mathfrak{C}$.

Note that since we took an optimistic viewpoint in the definition of $u_0^*(M)$, if a strong guarantee is required then a separate argument will be needed e.g. to show that there exists mechanism which achieves the optimal value even when its worst equilibrium is chosen.

A large body of works on optimal mechanism design is available in the literature. From this, we mostly build on the results of Myerson (1981) on optimal auction design. Laffont, Maskin, and Rochet (1987) address the problem of optimal nonlinear pricing or single-agent principal-agent mechanism design problem with two-dimensional type spaces. Dasgupta and Spulber (1990) design the optimal auction mechanism for single sourcing and multiple sourcing a contract but the private information is one-dimensional. Chen (2007) develops a two-stage implementation of (Dasgupta and Spulber, 1990). Che (1993) designs 2-dimensional optimal reverse auctions where the sellers bid price and quality but quality preferences are common knowledge and prices are the only private information. Also, Rochet and Stole (2003) wrote a good survey on multi-dimensional screening.

2.3.2 Mechanism Design with Voluntary Participation

As was mentioned before, oftentimes participation is voluntary, i.e., the agents can opt out from the game. In those cases, the agents can be incentivized to participate in the game. Assume that an agent of type θ_i playing at position i by opting out from the game would achieve a utility of $\underline{u}_i(\theta_i)$. Here, $\underline{u}_i : \Theta_i \rightarrow \mathbb{R}$ is called the *reservation utility* function for position i . In these cases one option

is to consider the constrained problem:

$$\begin{aligned}
u_0^*(M) &\rightarrow \max && \text{s.t.} \\
M &= (\Sigma, g) \in \mathcal{M}(I, O), \\
\exists s &\in \mathcal{E}(\mathfrak{B}_M, C_\Sigma) \text{ s.t.} && (2.5a)
\end{aligned}$$

$$u_0^*(M) = u_0(g(s(P_\theta))) \text{ and} \quad (2.5b)$$

$$u_i(g(s(\theta_i, P_{\theta_{-i}})), \theta_i) \geq \underline{u}_i(\theta_i), \quad \text{for all } \theta_i \in \Theta_i \text{ and } i \in I. \quad (2.5c)$$

The new condition (2.5c) constrains the mechanism such that at equilibrium, any agent's so-called *expected interim utility*, $u_i(g(s(\theta_i, P_{\theta_{-i}})), \theta_i)$, is at least as large as his reservation value. Since the condition requires that the interim utility, i.e., the utility a rational agent who knows his type⁷ can expect to achieve when participating, should exceed the value incurred when he decided to opt out, a (Bayesian) rational agent when presented with the option of participating in the game with the given rules should never prefer to opt out.

Naturally, adding the extra constraints (2.5a)–(2.5c) to the optimization problem reduces the optimal value that the principal can achieve (as compared to the problems where these constraints are not present). The constraints (2.5a)–(2.5c) are called the *individual rationality (IR)* or *participation* constraints and thus we call the above problem a *BOMD problem with IR constraints*.

The IR constraints make sure that all rational agents will participate irrespective of their “qualities”. It is, however, possible that in certain cases agents with “poor qualities” will not contribute in any useful way to the final outcome. However, taking this into account seems to lead a challenging optimization problem with a combinatorial structure and hence, following the mechanism design literature, we will keep things simple (and feasible) by adding the above constraint when necessary.

Note that if we allow $\underline{u}_i(\theta_i)$ to take on the value $-\infty$, then the class of

⁷ Note that in the economics literature, in connection to how utilities are calculated, “ex-ante” (from neo-Latin, meaning “before the event”) means the time period when none of the agents know their types, “interim” means the period when each agent knows their own types, while “ex-post” (meaning, “after the event”) means the period when all agents know all the types.

BOMD problems with voluntary participation will subsume the class of BOMD problems.

2.3.3 The Revelation Principle

For the remainder of this section we fix the counter-strategy map schema $\mathfrak{C} : \Sigma \mapsto C_\Sigma$ thus fixing the generalized Bayes equilibrium (GBE) concept. Fix a mechanism $M = (\Sigma, g)$. The mechanism gives rise to the set-valued map

$$f_M : \Theta \rightarrow 2^O, \quad f_M(\theta) = \{g(s^*(\theta)) : s^* \in \mathcal{E}(\mathfrak{B}_M, C_\Sigma)\}.$$

Definition 2.3.3 (Implementation of social choice functions). A mechanism M is said to implement a function $f : \Theta \rightarrow O$ in $\text{GBE}(\mathfrak{C})$ if $f(\theta) \in f_M(\theta)$ holds for each $\theta \in \Theta$.

In line with the mechanism design literature, we call functions of the form $f : \Theta \rightarrow O$ *social choice functions* (SCFs). Note that a game having multiple equilibria and hence a mechanism may implement many SCFs. Extra effort is needed usually to ensure the uniqueness of equilibria (when this is desired). A natural question to ask is whether a given SCF can be implemented using a mechanism.

A key idea in mechanism design is that, without any loss of generality, the search space for the mechanisms can be restricted to the so-called truthful direct-revelation mechanisms where the actions of the agents are type declarations and rational agents truthfully declare their types. This is known as the *revelation principle* and we will give a short proof of it later in this section. The idea of the principle is that if there exists a mechanism that implements some SCF then this mechanism can be used to construct the rules in a new mechanism where the agents have to report their types. The new mechanism requests the agents' types. The outcome is determined by the outcome that results from playing the equilibrium strategy corresponding to the reported type in the old game (more specifically, the equilibrium that implemented the chosen SCF should be used). Then, there is no reason for an agent to misreport his type, since this would only result in the new mechanism playing the part of that agent suboptimally in the old game.

Definition 2.3.4 (Direct-revelation mechanisms). A mechanism $M = (\Sigma, g)$ is called a *direct revelation mechanism* when $\Sigma_i = \Theta_i$ for each $i \in I$.

We will use id_Θ to denote the identity map over Θ : $\text{id}_\Theta : \Theta \rightarrow \Theta$ and $\text{id}_\Theta(\theta) = \theta$ ($\theta \in \Theta$).

Definition 2.3.5 (Incentive compatibility). We say that the direct-revelation mechanism $M = (\Theta, g)$ is *GBE($\mathfrak{C}$) incentive compatible* if $\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_M, C_\Theta)$.

In words, in an incentive compatible direct-revelation mechanism M telling the truth (truthfulness) is always a rational strategy. A related concept is as follows:

Definition 2.3.6 (Truthful mechanism). Given a direct-revelation mechanism $M = (\Theta, g)$ that implements an SCF $f : \Theta \rightarrow O$ in $\text{GBE}(\mathfrak{C})$, it is said to be a *truthful* mechanism if the identity map id_Θ is an equilibrium strategy map for M : $\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_M, C_\Theta)$ and $g(\theta) (= g(\text{id}_\Theta(\theta))) = f(\theta)$. The identity equilibrium strategy map is called the truthful map.

Note that the definition does not rule out that there are equilibrium strategy maps other than the “truthful map” id_Θ .

Definition 2.3.7 (Payoff equivalence). Given two games $M = (\Sigma, g)$, $M' = (\Sigma', g')$, two respective strategy profile maps $s : \Theta \rightarrow S$, $s' : \Theta \rightarrow S'$ are called *payoff equivalent* if for any $\theta \in \Theta$ and $i \in I$,

$$u_i(g(s(\theta)), \theta_i) = u_i(g'(s'(\theta)), \theta_i).$$

We prove the revelation principle for counter-strategy map-schemas which are admissible:

Definition 2.3.8 (Admissible counter-strategy maps). The counter-strategy map schema $\mathfrak{C} = (C_\Sigma)_\Sigma$ is called *admissible* if for any Σ , $s \in \mathfrak{S}(\Theta, S)$, $i \in I$, it holds that

$$\left\{ s_{-i} \circ \vartheta_{-i} : \vartheta_{-i} \in C_{-i}^{(\Theta)}(\text{id}_\Theta) \right\} \subset C_{-i}^{(\Sigma)}(s), \quad (2.6)$$

where we use $C_{-i}^{(\Sigma)}$ to denote the i th component of C_Σ .

Note that the maps $C_{-i}^{(\Sigma)}(s) = C_{-i}^N(s) = \{s_{-i}\}$ (defining Nash equilibria) and $C_{-i}^{(\Sigma)}(s) = C_{-i}^{DS}(s) = (S_{-i})^{\Theta_{-i}}$ (defining dominant strategy equilibria) are trivially admissible.

With this, we are ready to state the revelation principle:

Theorem 2.3.1 (The revelation principle (RP)). *Fix an admissible counter-strategy map schema $\mathfrak{C} = (C_\Sigma)_\Sigma$ and assume that some mechanism $M = (\Sigma, g)$ implements a social choice function f in $GBE(\mathfrak{C})$. Then, there exists a truthful direct-revelation mechanism $M' = (\Theta, g')$ which implements f . Further, the truthful map in M' and the equilibrium map of M that implements f are payoff equivalent.*

Proof. Let $s^* \in \mathcal{E}(\mathfrak{B}_M, C_\Sigma)$ be the equilibrium map for which it holds that $f(\theta) = g(s^*(\theta))$ and choose

$$g'(\theta) = g(s^*(\theta)).$$

We claim that it suffices to show that truthfulness is an equilibrium strategy of M' , i.e., that $\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_{M'}, C_\Theta)$. Indeed, we have $g'(\text{id}_\Theta(\theta)) = g'(\theta) = g(s^*(\theta)) = f(\theta)$. Thus, if $\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_{M'}, C_\Theta)$ then in M' implements f and is truthful. Payoff equivalence then follows from $g'(\text{id}_\Theta(\theta)) = g(s^*(\theta))$. Thus, it remains to show that $\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_{M'}, C_\Theta)$.

By definition 2.2.6, for this we need to show that for any $i \in I$,

$$\text{id}_{\Theta, i} \in \bigcap_{\vartheta_{-i} \in C_{-i}^{(\Theta)}(\text{id}_\Theta)} B_i(g', u_i, \vartheta_{-i}, P_{\theta_{-i}}).$$

By definition 2.2.2, this is equivalent to requiring that for all $i \in I$, $\vartheta_{-i} \in C_{-i}^{(\Theta)}(\text{id}_\Theta)$ and $\theta'_i \in \Theta_i$,

$$u_{i, \theta_i}(g'(\text{id}_{\Theta, i}(\theta_i), \vartheta_{-i}(P_{\theta_{-i}}))) \geq u_{i, \theta_i}(g'(\theta'_i, \vartheta_{-i}(P_{\theta_{-i}}))). \quad (2.7)$$

From now on fix $i \in I$, $\vartheta_{-i} \in C_{-i}^{(\Theta)}(\text{id}_\Theta)$ and $\theta'_i \in \Theta_i$. By the definition of id_Θ , the above inequality holds if and only if

$$u_{i, \theta_i}(g'(\theta_i, \vartheta_{-i}(P_{\theta_{-i}}))) \geq u_{i, \theta_i}(g'(\theta'_i, \vartheta_{-i}(P_{\theta_{-i}}))).$$

Now, by the definition of g' , and using that s^* is separable,

$$u_{i,\theta_i}(g'(\theta_i, \vartheta_{-i}(P_{\theta_{-i}}))) = u_{i,\theta_i}(g(s_i^*(\theta_i), s_{-i}^*(\vartheta_{-i}(P_{\theta_{-i}})))).$$

Similarly,

$$u_{i,\theta_i}(g'(\theta'_i, \vartheta_{-i}(P_{\theta_{-i}}))) = u_{i,\theta_i}(g(s_i^*(\theta'_i), s_{-i}^*(\vartheta_{-i}(P_{\theta_{-i}})))).$$

Introduce $s'_{-i} = s_{-i}^* \circ \vartheta_{-i}$ and $s''_i = s_i^*(\theta'_i)$. Thus, (2.7) is equivalent to

$$u_{i,\theta_i}(g(s_i^*(\theta_i), s'_{-i}(P_{\theta_{-i}})))) \geq u_{i,\theta_i}(g(s''_i, s'_{-i}(P_{\theta_{-i}}))))). \quad (2.8)$$

Now, since $s^* \in \mathcal{E}(\mathfrak{B}_M, C_\Sigma)$, we have

$$s_i^* \in \bigcap_{\hat{s}'_{-i} \in C_{-i}^{(\Sigma)}(\text{id}_\Theta)} B_i(g, u_i, \hat{s}'_{-i}, P_{\theta_{-i}}).$$

Equivalently, for any $\hat{s}'_{-i} \in C_{-i}^{(\Sigma)}(s^*)$ and $\hat{s}''_i \in S_i$,

$$u_{i,\theta_i}(g(s_i^*(\theta_i), \hat{s}'_{-i}(P_{\theta_{-i}})))) \geq u_{i,\theta_i}(g(\hat{s}''_i, \hat{s}'_{-i}(P_{\theta_{-i}}))))).$$

Since $\mathfrak{C}$ is admissible, $s'_{-i} \in C_{-i}^{(\Sigma)}(s^*)$ and thus by choosing $\hat{s}'_{-i} = s'_{-i}$ and $\hat{s}''_i = s''_i$, we see that (2.8) indeed holds. $\square$

Background: The revelation principle was introduced by Gibbard (1973) for dominant strategies. Later on, the revelation principle was proved for Bayesian Nash equilibria by Holmström (1977), Dasgupta et al. (1979) and Myerson (1979). Technically, the proof given here is new: We prove a general version of the revelation principle for the generalized Bayes equilibrium concept that we introduced. Nevertheless, unsurprisingly, the proof follows closely previous proofs. Although the proof is slightly more complicated, it has the benefit of showing that a single proof suffices, while pointing out the key property (admissibility of the counter-strategy maps) of equilibrium concepts that is sufficient for guaranteeing that the revelation principle remains true.

As noted above, the key consequence of the revelation principle is that it allows the narrowing of the search for the mechanisms. However, the revelation principle on its own does not guarantee that a truthful direct-revelation

mechanism would exist that implements only the chosen SCF. Hence, other type of mechanisms can still be interesting if one is interested in such strong implementations (or extra effort is needed to make sure that this property holds). Very often, however, the revelation principle is used to prove negative results: If some SCF cannot be implemented using a direct mechanism, then it cannot be implemented by any kind of mechanism. On the positive side, by narrowing down the search space, the revelation principle can (sometimes) be used to find a useful mechanism. In many cases, this is done by transforming the search problem into an optimization problem.

Proposition 2.3.2. *Fix $D = (I, O, \Theta, u, P_\theta, u_0, \mathfrak{C})$, where $\mathfrak{C}$ is admissible. Let $g^* : \Theta \rightarrow O$ be the solution to the optimization problem:*

$$\int u_0(g(\theta)) dP_\theta \rightarrow \max \quad s.t. \quad (2.9a)$$

$$g : \Theta \rightarrow O,$$

$$\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_{(\Theta, g)}, C_\Theta). \quad (2.9b)$$

Then, (Θ, g^) is Bayesian optimal mechanism for D .*

Proof. The result follows from the definition and the revelation principle. $\square$

The constraint (2.9b) is known as the *incentive compatibility (IC) constraint* as this constraint guarantees that truthfulness is rational.

Remembering that in the definition of the value of a mechanism an optimistic viewpoint was taken, if a stronger guarantee is required then one should try to achieve that id_Θ is the only equilibrium of (Θ, g^*) (or that any other equilibrium's expected payoff is at least as large as that of id_Θ).

For later reference, it is worthwhile to expand the abstract IC constraint of (2.9b). Investigating the definitions, we see that this constraint is equivalent to the following one:

$$\begin{aligned} u_{i, \theta_i}(g(\theta_i, s_{-i}(P_{\theta_{-i}}))) &\geq u_{i, \theta'_i}(g(\theta'_i, s_{-i}(P_{\theta_{-i}}))) \\ \text{for all } i \in I, \theta_i, \theta'_i \in \Theta_i, \text{ and } s_{-i} \in C_{-i}^{(\Theta)}(\text{id}_\Theta). \end{aligned} \quad (2.10)$$

When considering implementations in Bayesian Nash equilibria, $C^{(\Theta)}(s) = C^{(\Theta), BN}(s) = \{s_{-i}\}$, the constraint further simplifies to

$$u_{i, \theta_i}(g(\theta_i, P_{\theta_{-i}})) \geq u_{i, \theta'_i}(g(\theta'_i, P_{\theta_{-i}})) \quad \text{for all } i \in I \text{ and } \theta_i, \theta'_i \in \Theta_i. \quad (2.11)$$

When considering implementation in dominant strategies equilibria, $C^{(\Theta)}(s) = C^{(\Theta),BDS}(s) = (\Theta_{-i})^{\Theta_{-i}}$, the constraint becomes

$$u_{i,\theta_i}(g(\theta_i, \theta_{-i})) \geq u_{i,\theta_i}(g(\theta'_i, \theta_{-i})) \quad \text{for all } i \in I, \theta_{-i} \in \Theta_{-i}, \text{ and } \theta_i, \theta'_i \in \Theta_i \quad (2.12)$$

(as it was noted earlier, the distributions P_θ play no role when implementation in dominant strategies is chosen).

2.3.4 Voluntary Participation

Recall the BOMD problem with voluntary participation from Section 2.3.2. The counterpart of Proposition 2.3.2 for BOMD problems with voluntary participation is as follows:

Proposition 2.3.3. *Fix $D = (I, O, \Theta, u, \underline{u}, P_\theta, u_0, \mathfrak{C})$, where $\mathfrak{C}$ is admissible and $\underline{u} = (\underline{u}_i)_{i \in I}$, $\underline{u}_i : \Theta_i \rightarrow \mathbb{R}$ (for $i \in I$, $\underline{u}_i(\theta_i)$ is the reservation utility of agent i with type $\theta_i \in \Theta_i$). Let $g^* : \Theta \rightarrow O$ be the solution to the optimization problem:*

$$\int u_0(g(\theta)) dP_\theta \rightarrow \max \quad s.t. \quad (2.13a)$$

$$g : \Theta \rightarrow O,$$

$$\text{id}_\Theta \in \mathcal{E}(\mathfrak{B}_{(\Theta, g)}, C_\Theta),$$

$$u_i(g(\theta_i, P_{\theta_{-i}}), \theta_i) \geq \underline{u}_i(\theta_i) \text{ for all } i \in I \text{ and } \theta_i \in \Theta_i. \quad (2.13b)$$

Then, (Θ, g^) is the solution to the BOMD problem with voluntary participation.*

Proof. The result follows from the definitions and the revelation principle noting that the equilibrium in the IR constraint (2.5a)–(2.5c) is chosen to be id_Θ . $\square$

Note. There is an issue with Bayesian optimal mechanisms with voluntary participation that should be noted: By definition Bayesian optimal mechanisms with voluntary participation provide agents with an “opt-out” message. However, direct revelation mechanisms require the space of messages to be exactly equal to agent’s type space. This is a reason why we cannot search for a

Bayesian optimal mechanism with voluntary participation using the revelation principle. Thus, we are forced to solve a surrogate problem of Bayesian optimal mechanisms with individual rationality constraint. However, this might cause the resulting mechanism to be sub-optimal as there is no guarantee that the a Bayesian optimal mechanisms with IR constraint would achieve the same value as its counterpart Bayesian optimal mechanisms with voluntary participation (see Figure 2.1).

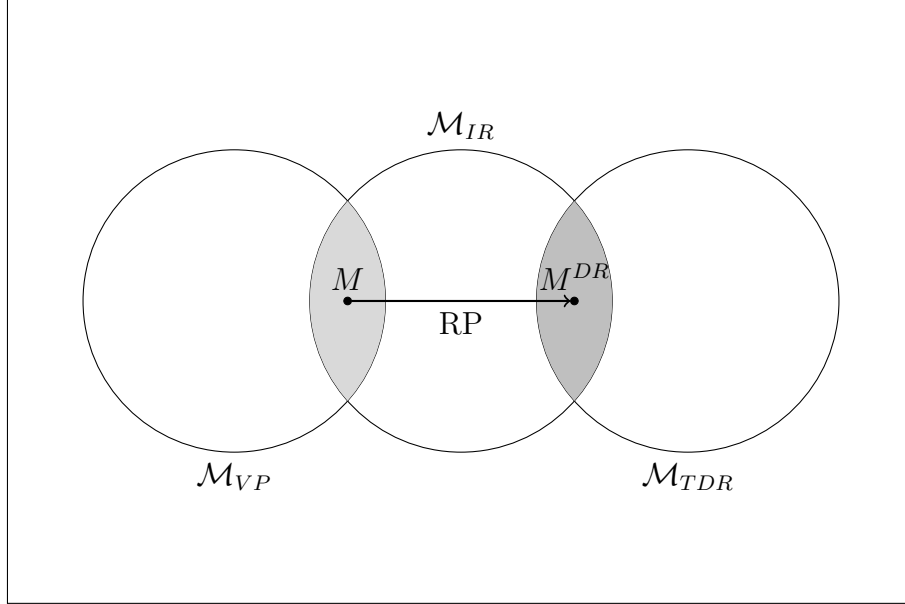


Figure 2.1: The problem with applying the revelation principle to voluntary mechanisms: Bayesian optimal mechanisms with voluntary participation are defined as to offer the agents the option to *opt out*: such actions must be propagated by the rules to the outcome space and taken into account appropriately in the principal's utility function. The set of these mechanisms is labeled as $\mathcal{M}_{VP}$ on the figure. On the other hand, (truthful) direct revelation mechanisms accept no message other than agent's type (the set of these is labeled by $\mathcal{M}_{TDR}$ on the figure). Thus, the intersection of these sets is empty and thus the revelation principle fails. A Bayesian optimal mechanisms with voluntary participation *may* also be a Bayesian optimal mechanisms with an IR constraint, but in the process of applying the RP, it will nonetheless lose the voluntary participation feature.

2.4 Mechanisms in the World of Money

In a world with money, the outcome space has the form $O = E \times \times_{i \in I} T_i$ where $T_i \subset \mathbb{R}$ and the agents' utility functions $(u_i)_{i \in I}$ are assumed to take the *quasi-linear* form ($i \in I$):

$$u_i(o, \theta_i) = v_i(e, \theta_i) + t_i, \quad o = (e, t_1, \dots, t_{|I|}) \in O, \theta_i \in T_i,$$

for some functions $v_i : E \times \Theta_i \rightarrow \mathbb{R}$. In particular, in a world with money, an outcome has the form $(e, t_1, \dots, t_{|I|})$, where t_i is the amount of money transferred to agent i and the agent's utility functions are linear in the “money component”. The utility function is called *linear* when θ_i is an element of a linear space and $v_i(e, \cdot)$ is also (affine) linear. In later chapters, we will often make the natural assumption that the world has money and agents have quasi-linear, or linear utility functions.

2.5 Solution Methodologies

How to solve a BOMD problem? From the point of view of mathematics, once the problem is formally specified, it can be considered solved! However, of course, in practice we want solutions that are actionable. Traditionally, mechanism design problems were studied by economists, where the actual form of the mechanisms and that of the equilibria is also of major interest (cf. Gibbard 1973; Hurwicz 1969, 1973; Myerson and Satterthwaite 1983; Satterthwaite 1975). However, closed-form solutions are only possible (or known) when some rather specific assumptions are made on the problem (such as assumptions on the distributions involved, or the form of the utility functions). Although we will argue that in many cases these assumptions are reasonable, another possibility is to use some numerical technique to solve the optimization problem (after applying the revelation principle, or simplifying, re-representing the problem in some other ways perhaps). In particular, note that after the revelation principle is applied the problem takes the form of a one-stage reinforcement learning problem with constrained policies: in this viewpoint the mapping g becomes a “policy”, with the goal being, using this

new terminology, to find a policy that maximizes the principal’s expected utility while the policy should satisfy some stochastic constraints (the constraints are stochastic in the sense that they involve integrals). This is also an instance of stochastic programming. With the language of stochastic programming, the constraints would be described as “chance constraints” and the problem would be described as an instance of expected value maximization subject to chance constraints. When the outcome space is a convex subset of a vector space and the agents’ utility functions are linear in the outcome (e.g., see the example in the next chapter), both the IC and IR constraints are linear in g , assuming that u_0 is concave the problem becomes an instance of convex optimization.

In this thesis, for the sake of simplicity, we decided not deal with such computational approaches and thus will follow the economics literature in specifying closed-form solutions at the price of making extra assumptions. However, it remains an interesting and exciting avenue for future research to the computational issues that arise from mechanism design.

2.6 Summary

In this chapter, the basic concepts of game theory and mechanism design were explained to be used in next chapters. We started with basic definitions of game theory, introduced a new construct for generalizing equilibrium concepts and then we went on to describe mechanism design basics. After that, we defined the problem of Bayesian optimal mechanism design (BOMD) and elaborated about voluntary participation. Finally, we introduced the revelation principle (RP) and how it translates to limiting our search for mechanisms to the space of incentive-compatible mechanisms.

2.6.1 Known Limitations

- In Bayesian Nash equilibria, it is assumed that agent types are chosen independently. *Collusion* is when this is not the case. The framework is by definition unable to handle this.
- We assume the agents to be perfectly rational. This is most probably not

the case in the real world. What would be more accurate is to assume *bounded rationality*, which would make the problem of mechanism design challenging or provide a relaxation of the concept of bounded rationality. For example, a very good and realistic assumption would be to assume that the agents are not necessarily able to compute the equilibria of a game but are able to verify equilibria when presented to them. This would be a less challenging extension, but is still challenging.

- We assume priors and just do an expectation analysis. Many questions remain, like: how big the variance is, for example, or what is the worst-case? These questions are not addressed in this framework.
- Algorithmic game theory (Nisan et al., 2007) considerations are left unanswered (competitive mechanism).

2.6.2 Future Work

- Here we are addressing only normal-form games. This is not necessarily adequate. We may have to extend this equilibrium concept in order to be able to deal with extensive-form games with incomplete information. A good follow-up would be to integrate the equilibrium concepts of extensive-form games, e.g., sequential equilibrium concept by Kreps and Wilson (1982), into this framework and:
 - Provide GBE for extensive-form games;
 - Study the effect on the revelation principle (Myerson, 1986).
 - In case of extensive-form games or multi-stage mechanisms, the problem will become very similar to reinforcement learning. A study of dynamic programming-like algorithms for solving this would be well justified.
- Sometimes, information flow in is not under control by an agent (or the principal) in the game. This is specially true for the case of machine learning auctions. These unpreventable information flows should be accounted for. This is the main focus of Chapter 4 and Chapter 5.

- A characterization of equilibrium concepts under which the revelation principle holds. The new equilibrium concept opens up ways to study such problems, but a further study is needed. For example and question would be on the need for admissibility for RP to hold.
- Randomized mechanisms have not been studied but are a very rational step.
- Does voluntary participation conceals a back-door for an ad-hoc solution? The problem is that the revelation principle by definition does not allow for opt-out actions by the agents. There are a number of questions such as:
 - You could hurt the agents with the individual rationality. But with IR, maybe the search space is too constrained. Is this suboptimal?
 - The relations between individual rationality and voluntary participation and multi-stage mechanisms need further study.
 - Can we formulate a version of the revelation principle with extra actions like opting out?
 - Ex-post voluntary participation (agents realizing that participation was indeed voluntary after everything is said and done) is another interesting subject to study.
- Existence of an equilibrium by itself is not a guarantee that the mechanism acts as intended. Indeed this is the subject of study of *implementation theory* and *problems of coordination*. For that we need to make sure that the set of equilibria is a singleton and the equilibrium is unique. A (naive) question is how to achieve that the equilibrium is unique so that problems of coordination are avoided?
- Computational approaches to mechanism design is yet another interesting subject that has not been addressed in this thesis and merits to be studied.

Chapter 3

Machine Learning Solution Procurement

Assume that a company (the principal) owns some data and needs an algorithm that solves a prediction problem defined based on the data. The company decides to buy an off-the-shelf machine learning solution from some developers that the company is in contact with instead of investing into in-house development. Which solution should be bought? How much should be paid for the solution? In this chapter we will study a simplified model where we assume that each solution has already been evaluated on the data, and the evaluations results have been transformed into how profitable each solution might be (throughout the thesis we assume that the principal, and also the developers are risk-neutral). The company then may decide for a number of procurement formats (i.e., how the developers should submit their prices, what information should flow between the participants, etc.). In any case, at the end the company will need to make a decision of which solution to buy and how much to pay. The decision should maximize the company's profit. To model this as an optimization problem, we will use the last chapter's framework and we require that the implementation should be in Bayesian Nash equilibria.

An idea we will pursue in this thesis is casting this procurement of machine learning solutions into the form of an auction. An *auction* is a form of trade in which different parties bid for the acquirement of a good. A machine learning problem is more properly cast into the form of a reverse auction. A *reverse auction* is an auction where a principal wants to procure an item from a set

of agents offering different items. Each item has a different value for the principal and this is really what sets reverse auction apart from a normal ‘forward’ auction.

We are interested in this problem because it can model a very basic scenario when we want to acquire a machine learning solution: A principal wants to find the most profitable solution to her among what’s offered. Different agents are offering different off-the-shelf solutions (no tailoring is made) and each solution has a different quality for the principal and comes with a different price.

This problem has been addressed many times in the literature of economics. We adopt the solution by Myerson (1981) to the case of reverse auctions, which is minor change to his results. Solutions to reverse auctions in our setting where each item sold has the same value to the principal are described, e.g., in Narahari (2012).

3.1 Problem Description

An outline of the problem we solve is as follows: We assume that the principal knows the profitability ρ_i of the individual solutions offered by the developers (the developers are indexed by $i \in I$). Let θ_i denote the price that agent i values his solution at. The outcome of the mechanism is who wins and how much each agent gets paid. The only information concerning agent i that enters the agent’s utility function is θ_i : Thus we will take θ_i to be the type of agent i . We assume that θ_i is drawn at random from P_{θ_i} , where P_{θ_i} (for all $i \in I$) is known for all participants.

The flow of information in the mechanism is as follows:

1. Agent i ’s price θ_i is drawn from P_{θ_i} , $i \in I$;
2. The principal declares the rules: who will win and how much is each agent paid – as a function of ρ_i and the submitted bids
3. Agents offer their solutions with their bids $\hat{\theta}_i$ ($i \in I$);

4. Based on the previously announced rules, the winner is determined, the right to freely use his solution is transferred to the principal and the payments are made.

Let us now define the underlying BOMD problem formally. Instead of simply determining a winner whose solution is taken in a deterministic fashion, we consider a randomized choice (randomization can actually help in increasing the value that the principal can achieve and is by no means a restriction and will help with the mathematical developments): Thus, upon receiving all the information (for example bids and the test results, or any other messages the actual mechanism requires), the principal determines a subprobability distribution $\pi \in M_{\leq 1}(I)$ over I and a vector $t \in \mathbb{R}^{|I|}$ of payments whose i th component determines the money to be transferred to agent i (here, $M_{\leq 1}(I) = \{\pi \in [0, 1]^I : \sum_{i \in I} \pi_i \leq 1\}$ is the space of subprobability distributions over set I). Once, the outcome (π, t) is determined, the following events take place: With probability $1 - \sum_{i \in I} \pi_i$, the principal rejects all the offers, while with probability π_i the offer of agent i is accepted (that is, zero or one offer is accepted). As said before, if the offer on agent i is accepted, the principal gains the right to use the solution q_i of agent i . In any case, all agents i are transferred the respective amounts t_i . Thus, the outcome space is $O = M_{\leq 1}(I) \times T$, where $T = \times_{i \in I} T_i$ with $T_i \subset \mathbb{R}$ being the set of possible payment values for agent i .

The utility function of the principal is defined by

$$u_0(\pi, t) = \left(1 - \sum_{i \in I} \pi_i\right) \theta_0 + \sum_{i \in I} \pi_i \rho_i - \sum_{i \in I} t_i, \quad (3.1)$$

where $\theta_0 \in \mathbb{R}$ is the profit (possibly negative) that the principal makes when no solution is accepted, and ρ_i is the profitability of solution offered by agent i . Note that as usual with randomized mechanisms, the utility function determines the expected utility where the expectation is over the randomization of the mechanism.

As hinted on before, we represent the type θ_i of developer (or agent) i by the price $\lambda_i \in \mathbb{R}$ the agent would be “happy with”. With this, the utility of

agent with type θ_i , given the outcome $(\pi, t) \in O$ is

$$u_{i,\theta_i}(\pi, t) = -\pi_i \theta_i + t_i, \quad (3.2)$$

i.e., the agent is giving up the value θ_i when he is winning the competition, while he receives the amount t_i independently of whether he won the competition. Again, the utility function determines the expected utility for the agent where the expectation is over the randomization of the mechanism.

To fully specify the optimal mechanism design problem it remains to choose some counter-strategy map schema, $\mathfrak{C}$. In this chapter, we consider the case when $\mathfrak{C}$ is the Nash-choice: $C_{-i}(s) = C_{-i}^N(s) = \{s_{-i}\}$, giving rise to implementations in Bayesian Nash equilibria.

Problem 3.1.1. For $i \in I$, let $\Theta_i = [\underline{\theta}_i, \bar{\theta}_i] \subset \mathbb{R}$, $\underline{\theta}_i < \bar{\theta}_i$, P_{θ_i} be a distribution supported on Θ_i . Consider a BOMD problem where the outcome space is $O = M_{\leq 1}(I) \times T$, $T = \mathbb{R}^{|I|}$. Let the utility of agent i with type $\theta_i \in \Theta_i$, given $(\pi, t) \in O$, be

$$u_{i,\theta_i}(\pi, t) = -\pi_i \theta_i + t_i,$$

while the utility of the principal be

$$u_0(\pi, t) = \left(1 - \sum_{i \in I} \pi_i\right) \theta_0 + \sum_{i \in I} \pi_i \rho_i - \sum_{i \in I} t_i,$$

where θ_0 and ρ_i ($i \in I$) are given fixed real numbers. Let $\mathfrak{C}$ be defined by $C_\Sigma(s) = C_\Sigma^N(s) = \{s_{-i}\}$ (implementation in Bayes-Nash equilibria). The problem is to find a solution to the BOMD problem described by $(I, O, \Theta, P_\theta, u, u_0, \mathfrak{C})$.

Under some extra assumptions, this reverse auction problem has a well-known, closed-form solution Myerson (1981), which is described in the next section.

3.2 The Form of the Optimal Reverse Auction

In what follows we will reuse the symbol P_{θ_i} to denote the cumulative distribution function $P_{\theta_i} : \mathbb{R} \rightarrow [0, 1]$ corresponding to the distribution P_{θ_i} : $P_{\theta_i}(x) = \int_{-\underline{\theta}_i}^x dP_{\theta_i}$, $x \in \mathbb{R}$ (the meaning of P_{θ_i} should remain clear from the context). We make the following assumption concerning these distributions:

Assumption 3.2.1. For each $i \in I$, the distribution P_{θ_i} has a density with respect to the Lebesgue measure on Θ_i . Further, the density is bounded away from zero on Θ_i .¹

We will denote the resulting density by p_{θ_i} . WLOG (without loss of generality) we seek the optimal mechanism amongst the set of truthful direct-revelation mechanisms (that this can be done holds because of the revelation principle, cf. Theorem 2.3.1). Thus, $\Sigma = \Theta$, i.e., in our case the agents' messages will be prices. For $i \in I$, introduce the function $V_i : \Theta_i \rightarrow \mathbb{R}$ defined by

$$V_i(x) = \rho_i - x - \frac{P_{\theta_i}(x)}{p_{\theta_i}(x)}. \quad (3.3)$$

The function V_i assigns a “virtual value” to a price x submitted by agent i : The agent's offer is compared to the profit to be made and is adjusted by $\frac{P_{\theta_i}(x)}{p_{\theta_i}(x)}$ that reflects the uncertainty regarding the type of agent i (this term is also known in the literature as the *information rent* (Myerson, 1981) as it decreased the profitability by an amount that reflects the uncertainty of θ_i).

These functions form the basis of the solution to the reverse auction problem. In particular, the solution will take the form $M = (\Theta, g^*)$, $g^* = (\pi^*, t^*)$ with $\pi^* : \Theta \rightarrow M_{\leq 1}(I)$, $t^* : \Theta \rightarrow \mathbb{R}$ are specified as follows: For every vector of submitted prices θ , the mechanism will select a winner amongst the agents, with the possibility that no agent is selected as a winner. We let $w^* : \Theta \rightarrow I \cup \{0\}$ denote the function that determines the winner: the value $0 \notin I$ is used to allow the mechanism to reject all offers. To make the definition of w^* more concise (and maybe easier to comprehend) define

$$V_0(x) = x$$

Then, for $\theta \in \Theta$,²

$$w^*(\theta) = \arg \max_{i \in I \cup \{0\}} V_i(\theta_i), \quad (3.4)$$

where ties should be broken in an arbitrary, but systematic fashion independently of x (i.e., by ordering $I \cup \{0\}$ in some way and in the case of ties choosing

¹This is the reason Θ_i has to be a *bounded* interval.

²Note that θ does *not* include θ_0

the index that precedes all the other tied indices in the chosen ordering). According to (3.4), the winner is selected as the agent whose virtual valuation at the price he submitted is the largest. Now, define π^* by

$$\pi_i^*(\theta) = \mathbb{I}_{\{w^*(\theta)=i\}}, \quad i \in I. \quad (3.5)$$

To define the payment function t^* , first define the functions $z_i^* : \Theta_{-i} \rightarrow \mathbb{R}$, $i \in I$:

$$z_i^*(\theta_{-i}) = \sup \{ \theta_i \in \Theta_i : w^*(\theta) = i \}.$$

That is, $z_i^*(\theta_{-i})$ specifies the largest price agent i can submit and still win given that the other agents submit the prices θ_{-i} . With this, define

$$t_i^*(\theta) = \begin{cases} z_i^*(\theta_{-i}), & \text{if } w^*(\theta) = i; \\ 0, & \text{otherwise.} \end{cases} \quad (3.6)$$

Note that agent i gets paid if and only if he wins. When the agent wins, he gets paid $z_i^*(\theta_{-i})$, which is guaranteed to be more than θ_i , otherwise he would not have won.

Let us introduce one more technical assumption:

Assumption 3.2.2. The virtual valuation functions, V_i are strictly decreasing.

Note that for some common probability distributions, such as the uniform or exponential distributions, the virtual valuation functions are indeed strictly decreasing. Myerson (1981) described a technique (“ironing”) that allows one to avoid this assumption.

We can now state the main result of this section:

Theorem 3.2.1. *Let Assumptions 3.2.1 and 3.2.2 hold and let $g^* = (\pi^*, t^*)$, where the functions (π^*, t^*) are defined above. Then, the mechanism (Θ, g^*) is a solution to Problem 3.1.1.*

3.3 The Proof of Optimality

In this section we give a proof of Theorem 3.2.1. For a given pair of functions $\pi : \Theta \rightarrow M_{\leq 1}(I)$, $t : \Theta \rightarrow \mathbb{R}^{|I|}$ let $g_{\pi,t} : \Theta \rightarrow O$ be defined by $g_{\pi,t}(\theta) =$

$(\pi(\theta), t(\theta))$. Further, for $i \in I$, $\theta_i, \theta'_i \in \Theta_i$, let

$$\begin{aligned} u_i(\pi, t, \theta'_i, \theta_i) &= u_{i, \theta_i}(g_{\pi, t}(\theta'_i, P_{\theta_{-i}})), \\ U_i(\pi, t, \theta_i) &= u_i(\pi, t, \theta_i, \theta_i). \end{aligned}$$

Thus, $u_i(\pi, t, \theta'_i, \theta_i)$ is the interim expected utility of agent i when he chooses to send θ'_i while his type is θ_i and $U_i(\pi, t, \theta_i)$ is his interim expected utility when he chooses to be truthful.

By Proposition 2.3.3 and also using that by our choice of Bayes-Nash implementation the IC constraint (2.9b) is equivalent to (2.11), Problem 3.1.1 is equivalent to the following functional optimization problem:

$$\begin{aligned} \int u_0(g_{\pi, t}(\theta)) dP_\theta &\rightarrow \max \text{ s.t.} & (\text{OPT-1}) \\ \pi : \Theta &\rightarrow M_{\leq 1}(I), t : \Theta \rightarrow \mathbb{R}^{|I|}, \\ U_i(\pi, t, \theta_i) &\geq u_i(\pi, t, \theta'_i, \theta_i) \text{ for all } i \in I \text{ and } \theta_i, \theta'_i \in \Theta_i & (\text{IC-1}) \\ U_i(\pi, t, \theta_i) &\geq 0 \text{ for all } i \in I \text{ and } \theta_i \in \Theta_i. & (\text{IR-1}) \\ & & (3.7a) \end{aligned}$$

For $i \in I$, $\pi : \Theta \rightarrow M_{\leq 1}(I)$, $\theta_i \in \Theta_i$ define

$$E_i(\pi, \theta_i) = \pi_i(\theta_i, P_{\theta_{-i}}).$$

Note that with this definition we can write $u_i(\pi, t, \theta'_i, \theta_i) = t(\theta'_i, P_{\theta_{-i}}) - E_i(\pi, \theta'_i)\theta_i$. We claim that π, t satisfies (IC-1), (IR-1) if and only if it satisfies the following constraints:

$$\begin{aligned} E_i(\pi, \cdot) &\text{ is decreasing for all } i \in I, & (\text{DEC-2}) \\ U_i(\pi, t, \theta_i) &= U_i(\pi, t, \bar{\theta}_i) + \int_{\theta_i}^{\bar{\theta}_i} E_i(\pi, \hat{\theta}_i) d\hat{\theta}_i \text{ for all } i \in I, \theta_i \in \Theta_i, & (\text{INT-2}) \\ U_i(\pi, t, \bar{\theta}_i) &\geq 0 \text{ for all } i \in I. & (\text{IR-2}) \end{aligned}$$

The proof of this equivalence actually holds for each index $i \in I$, separately and follows immediately from the following analysis lemma:

Lemma 3.3.1. *Let $X = [a, b]$ be a closed subinterval of the real line, $t : X \rightarrow \mathbb{R}$ and let $e : X \rightarrow [0, \infty)$ be a function that is also integrable. For $x, y \in X$,*

define $u(x, y) = t(x) - e(x)y$, $U(x) = u(x, x)$. Then the inequalities

$$U(y) \geq u(x, y) \text{ for all } x, y \in X, \quad (\text{IC-ENV})$$

$$U(x) \geq 0 \text{ for all } x \in X \quad (\text{IR-ENV})$$

are satisfied if and only if the constraints

$$e \text{ is decreasing,} \quad (\text{DEC-ENV})$$

$$U(x) = U(b) + \int_x^b e(x)dx \text{ for all } x \in X, \quad (\text{INT-ENV})$$

$$U(b) \geq 0 \quad (\text{IR2-ENV})$$

are satisfied.

Note that the significance of the transformation of the constraints is that in (INT-ENV) the effect of function t only shows up in $U(b)$, i.e., though an additive term, which does not depend on x .

We use the following theorem in the proof of Lemma 3.3.1:

Theorem 3.3.2 (Envelope Theorem, Theorem 2 of (Milgrom and Segal, 2002)). *Let $f : X \times [0, 1] \rightarrow \mathbb{R}$ be a function, $V(t) = \sup_{x \in X} f(x, t)$, $X^*(t) = \{x \in X : f(x, t) = V(t)\}$. Let f_t denote the partial derivative of f w.r.t. t . Suppose that the following hold:*

- a) $f(x, \cdot)$ is absolutely continuous for all $x \in X$;
- b) $f(x, \cdot)$ is differentiable for all $x \in X$;
- c) $|f_t(x, t)| \leq b(t)$ for all $x \in X$ and almost all $t \in [0, 1]$, where $b : [0, 1] \rightarrow \mathbb{R}$ is integrable;
- d) $X^*(t) \neq \emptyset$ almost everywhere on $[0, 1]$.

Then, for any measurable selection $x^*(t) \in X^*(t)$,

$$V(t) = V(0) + \int_0^t f_t(x^*(s), s) ds. \quad (3.11)$$

Lemma 3.3.1. We first show that (IC-ENV) is equivalent to

$$U(y) - U(x) \geq (x - y)e(x) \text{ for all } x, y \in X. \quad (\text{IC2-ENV})$$

Indeed, using the definition of u , we see that

$$U(y) \geq u(x, y) = t(x) - e(x)y = U(x) + e(x)(x - y),$$

and reordering the terms gives the required equivalence. Also, note that (IC2-ENV) clearly implies that U is decreasing thanks to $e \geq 0$.

$\Rightarrow$: Clearly, (IR2-ENV) is implied by (IR-ENV). Swapping x and y in (IC2-ENV) gives $U(x) - U(y) \geq (y - x)e(y)$. Combining this with (IC2-ENV) we get

$$(x - y)e(x) \leq U(y) - U(x) \leq (x - y)e(y) \text{ for all } x, y \in X. \quad (3.12)$$

This implies that e is decreasing (i.e., (DEC-ENV)). The plan now is to apply Theorem 3.3.2 to $f(x, t) = u(x, \tau(t))$, where $\tau : [0, 1] \rightarrow X$ is an affine linear function such that $\tau(0) = x_0$ with some $x_0 \in X$, $\tau(1) = b$ to show (INT-ENV). With the notation of the theorem, $V(t) = \max_{x \in X} f(x, t) = u(\tau(t), \tau(t))$. Choose $x^*(t) = \tau(t)$. The conditions of the theorem can be readily verified. Further, $f_t(x, t) = \frac{\partial}{\partial t}(t(x) - e(x)\tau(t)) = -e(x)\tau'(t)$. Hence, $U(b) = V(1) = V(0) + \int_0^1 f_t(x^*(t), t)dt = U(x_0) - \int_0^1 e(\tau(t))\tau'(t)dt$ and thus by substituting $s = \tau(t)$ ($ds = \tau'(t)dt$) we get $U(x_0) = U(b) + \int_{\tau(0)}^{\tau(1)} e(s)ds = U(b) + \int_{x_0}^b e(s)dds$. Since x_0 was arbitrary, we get (INT-ENV). This finishes the direction $\Rightarrow$.

$\Leftarrow$: Since e is decreasing, starting from (INT-ENV) we get that

$$U(x) = U(y) - \int_y^x e(z)dz \geq U(y) - \int_y^x e(y)dz = U(y) - (x - y)e(y)$$

holds for any $x, y \in X$. This implies (IC2-ENV), which was seen to be equivalent to (IC-ENV) and to imply that U is decreasing. Since U is decreasing, (IR2-ENV) implies (IR-ENV). $\square$

Let us now return to the optimization problem. Using the function U_i , we can rewrite the objective function as

$$u_0(g_{\pi, t}(P_\theta)) = \lambda_0 + \sum_{i \in I} (\rho_i - \theta_i - \lambda_0) \pi_i(P_\theta) - \sum_{i \in I} U_i(\pi, t, P_{\theta_i}). \quad (3.13)$$

Let us now write $U_i(\pi, t, P_\theta)$ in a form that allows the separation of the terms that involve t . Take any π, t satisfying the constraints (DEC-2), (INT-2), (IR-2). Due to (IR-2) and the definition of E_i and Assumption 3.2.1,

$$\begin{aligned}
U_i(\pi, t, P_{\theta_i}) &= \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(U_i(\pi, t, \bar{\theta}_i) + \int_{\theta_i}^{\bar{\theta}_i} E_i(\pi, \theta'_i) d\theta'_i \right) dP_{\theta_i}(\theta_i) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \int_{\theta_i}^{\bar{\theta}_i} E_i(\pi, \theta'_i) d\theta'_i dP_{\theta_i}(\theta_i) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(\int_{\underline{\theta}_i}^{\theta'_i} dP_{\theta_i}(\theta_i) \right) E_i(\pi, \theta'_i) d\theta'_i \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(P_{\theta_i}(\theta'_i) \int_{\Theta_{-i}} \pi_i(\theta'_i, \theta_{-i}) dP_{\theta_{-i}}(\theta_{-i}) \right) d\theta'_i \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\Theta} P_{\theta_i}(\theta_i) \pi_i(\theta) dP_{\theta_{-i}}(\theta_{-i}) \frac{P_{\theta_i}(\theta_i)}{p_{\theta_i}(\theta_i)} d\theta_i \quad (*) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\Theta} \frac{P_{\theta_i}(\theta_i)}{p_{\theta_i}(\theta_i)} \pi_i(\theta) dP_{\theta}(\theta).
\end{aligned}$$

Note that we have indeed separated the term that includes t . The equation where we used the positivity of p_{θ_i} over its domain is denoted by (*). Plugging the expression obtained for $U_i(\pi, t, P_{\theta_i})$ into (3.13) and using the functions

$$\hat{V}_i(x) = \rho_i - x - \lambda_0 - \frac{P_{\theta_i}(x)}{p_{\theta_i}(x)} \quad (x \in \Theta_i),$$

we get

$$u_0(g_{\pi, t}(P_\theta)) = \lambda_0 - \sum_{i \in I} U_i(\pi, t, \bar{\theta}_i) + \int \sum_{i \in I} \hat{V}_i(\theta_i) \pi_i(\theta) dP_\theta(\theta). \quad (3.14)$$

For π fixed, let us maximize this in t subject to the constraints (DEC-2), (INT-2), (IR-2). Since only the second term depends on t and in fact this term has a negative sign, we maximize the objective if we minimize $\sum_{i \in I} U_i(\pi, t, \bar{\theta}_i)$. Let us consider the i th term of this sum for some i fixed. By (INT-2) and plugging in the definitions of U_i and E_i , for any $\theta_i \in \Theta_i$ we get

$$\begin{aligned}
U_i(\pi, t, \bar{\theta}_i) &= U_i(\pi, t, \theta_i) - \int_{\theta_i}^{\bar{\theta}_i} E_i(\pi, \theta'_i) d\theta'_i \\
&= t_i(\theta_i, P_{\theta_{-i}}) - \pi_i(\theta_i, P_{\theta_{-i}}) \theta_i - \int_{\theta_i}^{\bar{\theta}_i} \pi_i(\theta'_i, P_{\theta_{-i}}) d\theta'_i \\
&= \left[t_i(\theta) - \pi_i(\theta) \theta_i - \int_{\theta_i}^{\bar{\theta}_i} \pi_i(\theta'_i, \theta_{-i}) d\theta'_i \right]_{\theta_{-i} \leftarrow P_{\theta_{-i}}},
\end{aligned}$$

where $[\cdot]_{\theta_{-i} \leftarrow P_{\theta_{-i}}}$ is used to denote the substitution of θ_{-i} by $P_{\theta_{-i}}$ (and hence, taking the integral of the expression). Thus, $U_i(\pi, t, \bar{\theta}_i)$ depends on t only through t_i . By (IR-2), all feasible pairs (π, t) must satisfy $U_i(\pi, t, \bar{\theta}_i) \geq 0$. Hence, the minimum of $U_i(\pi, t, \bar{\theta}_i)$ is zero. This minimum is achieved if we choose

$$t_i^{(\pi)}(\theta) = \pi_i(\theta)\theta_i + \int_{\theta_i}^{\bar{\theta}_i} \pi_i(\theta'_i, \theta_{-i}) d\theta'_i$$

and by choosing t this way (as a function of π), (INT-2), (IR-2) are satisfied for any π . Thus, it remains to choose π .

When we choose $t = t^{(\pi)}$, we see that the only term that still depends on π in (3.14) is the last term. Call this term: $\Upsilon(\pi, P_\theta) = \int \sum_{i \in I} \hat{V}_i(\theta_i) \pi_i(\theta) dP_\theta(\theta)$. Taking into account that $\pi(\theta)$ is a subprobability distribution, we see that for any feasible π

$$\Upsilon(\pi, P_\theta) \leq \left[\max(0, \max_{i \in I} \hat{V}_i(\theta_i)) \right]_{\theta \leftarrow P_\theta}$$

(if $\max_{i \in I} \hat{V}_i(\theta_i) < 0$ $\pi_i(\theta) = 0$, $i \in I$ achieves zero inside the integral at θ). Further, the upper bound on $\Upsilon(\pi, P_\theta)$ can be achieved by any π when $\pi(\theta)$ assigns zero to all indices $i \in I$ such that $\hat{V}_i(\theta_i) < 0$ and assigns nonnegative values to indices in $W(\theta) = \left\{ i \in I : \hat{V}_i(\theta_i) \geq 0, \hat{V}_i(\theta_i) = \max_{j \in I} \hat{V}_j(\theta_j) \right\}$. Now, if Assumption 3.2.2 is satisfied then it can be shown that by choosing a single nonzero entry from $W(\theta)$ will result in π that satisfies (DEC-2). Denoting the resulting choice π^* and letting $t^* = t^{(\pi^*)}$, after elementary transformation we arrive at the desired statement, thus finishing the proof of Theorem 3.2.1.

3.4 Discussion

Let us return to the discussion of the optimal mechanism. An important feature of the optimal mechanism is that the price submitted by an agent influences his payment only through whether he wins or not: If the agent does not win, he receives no payment, while if the agent wins, his payment is determined solely by the prices submitted by the *other* agents. Furthermore, since the virtual valuation function is a monotone decreasing function of the price submitted, the higher the price that the agent submits, the smaller is

the chance he is going to win the competition. Hence, to maximize his utility, the agent should declare the minimum price that he still finds acceptable. In fact, under these rules, even if the agent was given the prices submitted of the other agents, the agent could not increase his profit by manipulating the price he submits: The strategy of honest declaration of everyone's prices is an *ex-post Nash* equilibrium. This simple reasoning can be used by the principal to convince each agent that they are best off by honestly submitting their "reservation prices". Thus, the agents need not do any calculations to verify that honest declaration is in their best interests. An important consequence of this is that honest declaration is actually a Bayes-Nash equilibrium even if the agents beliefs about the prior distribution of prices differ from each other, or the prior used by the principal. Thus, if the principal chooses a prior which is different from the probability distribution the agents' types are drawn from, this "incorrect choice" will only influence the principal's expected profit, but not how the agent's act as long as they are rational.

Let us now turn to the discussion of some limitations of the model used. One such limitation is the assumption that both the agents and the principal are risk-neutral (in games with money, a risk-neutral entity is one whose utility is linear in money, a risk-seeking entity is one whose utility is superlinear in money, and a risk-averse entity is an entity with a utility sublinear in money). The assumption of risk-neutrality may very well be unrealistic as agents "in the real world" may not risk-neutral: big companies are often risk-seeking and individuals with low wealth are often risk-averse (Markowitz, 1959).

Another issue worth noting is the importance of effective communication of the mechanisms' incentivizing structure to the agents. The principal has to argue informally, and also formally if need be, and convince the agents that honesty is in equilibrium due to the structure of the mechanism. Even though we have argued above that the argument that the principal has to make is not complicated, some agents may still not "accept" such an argument.

3.5 Summary

When the agents have different distributions, the agent with the largest virtual valuation is not necessarily the agent with the lowest bid. Thus, the optimal auction need to be what is called allocatively efficient,³ and therefore, need not be ex-post efficient.⁴

In the case when no information is available about the participants of the machine learning competition, perhaps it is more reasonable to assume complete symmetry (i.e., $\Theta_i = \Theta_j$ and $P_{\theta_i} = P_{\theta_j}$ for all $i, j \in I$). In this case, the winner is the agent with the lowest bid. Further, the payment coincides with the payment rules in second-price reverse auctions. In other words, the second price reverse auction is an optimal auction when the agents are homogeneous. As a result, oftentimes, the optimal auction is also known as a modified Vickrey auction (Vickrey, 1961)⁵.

Here, we assume that agent types are drawn independently. When this assumption is not true, we say that we have a case of “collusion”.

³ A mechanism is allocatively efficient when it maximizes the sum of interim expected utilities of the agents, also known as the total social surplus. The reverse auction is called allocatively efficient when the winner is the agent with the lowest price.

⁴This means, that an omnipotent agent who has access to all the private information can choose outcomes such that all agents are not worse off than with the current set of rules and some of them are at least sometimes strictly better off.

⁵A Vickrey auction, or more correctly, a Vickrey–Clarke–Groves (VCG) auction, also called a second-price, sealed-bid auction (SPA) is an auction where the aim is to maximize a social welfare function defined as the sum of the utilities of all agents. See (Clarke, 1971; Groves, 1973; Vickrey, 1961).

Chapter 4

Machine Learning Competitions with Public Test Results

In this chapter we look into the first version of our problem. In this simplified version it is assumed that the agents participating in the competition will not do any tailoring to their methods, but will submit off-the-shelf solutions. Once the solutions are obtained, the principal tests the solutions on his data and makes the results public so that the agents can adjust their prices accordingly. More precisely, the participants interact as follows:

1. Agents are randomly drawn;
2. The principal declares the rules of determining the winner and that of the payments;
3. The agents submit their solutions for evaluation to the principal;
4. The principal, upon evaluating the solutions, makes the results public;
5. The agents submit their offers to the principal;
6. The winner and payments are determined according to the rules declared at the beginning based on the information that was made public earlier and the offers of the agents;
7. Payments are made to the agents.

In addition to the rules of the game, this protocol is also public knowledge and it is assumed that the participants will not deviate from the rules and

they trust each other. Again, trust, privacy and security can be guaranteed by appropriate cryptographic methods.

In the next chapter we will consider the case when the developers still submit their solutions, but they do not receive information about how well their solutions are doing on the data. The rest of this chapter is organized as follows: in Section 4.1, we introduce the problem of machine learning auctions: in Section 4.1.1 we model the situation of this problem with the concept of ‘information leakage’. Next in Section 4.1.2 we show how this model fits our problem of machine learning competitions. We proceed to solve this problem in Section 4.2. A summary of the chapter is given in Section 4.4. Finally conclusions for this chapter are stated in Section 4.4.

4.1 The Formal Problem Definition

The purpose of this section is formalize the problem in a rigorous fashion. Because information is leaked about the agent’s private information (i.e., they need to submit their solutions whose results will be made public), the framework of Chapter 2 has to be appropriately extended.

4.1.1 Bayesian Optimal Mechanism Design with Information Leakage

Consider a mechanism design problem with the following structure: Let I be the finite set of player positions, O an outcome space, $\Theta = \times_{i \in I} \Theta_i$ a type-space, $P_\theta \in M_1^\times(\Theta)$ (a product distribution over Θ), $u = (u_i)_{i \in I}$ be the utility functions of the agents, $\mathfrak{C}$ a counter-strategy map schema. Further, assume that a separable function $\varepsilon^{\text{fn}} : \Theta \rightarrow \mathcal{E}$, $\mathcal{E} = \times_{i \in I} \mathcal{E}_i$, is given and the standard protocol of interaction is modified such that the value of $\varepsilon^{\text{fn}}(\theta)$ gets revealed once the types θ are assigned to the agents (i.e., agent i cannot help the leakage of $\varepsilon_i^{\text{fn}}(\theta_i)$ about his type). Here, we consider two protocols: In one case, the information is leaked after the agents learn their types and to all the agents, while in the second case the information leaked is not available to the agents at the time they have to decide about their messages. We call the first

Ex-ante information leakage

1. A type profile is drawn:
 $\theta \sim P_\theta$;
2. Agents learn the rules of the game;
3. The information $\varepsilon^{\text{fn}}(\theta)$ is leaked to everyone;
4. Agents submit their messages;
5. The outcome is determined.

Ex-post information leakage

1. A type profile is drawn:
 $\theta \sim P_\theta$;
2. Agents learn the rules of the game;
3. Agents submit their messages;
4. The information $\varepsilon^{\text{fn}}(\theta)$ is leaked to everyone;
5. The outcome is determined.

Figure 4.1: Protocols of interaction with information leakage. The two protocols differ in that when is the information is leaked. Note that in the ex-post case, from the point of the strategies of the agent it does not actually matter they also learn the information $\varepsilon^{\text{fn}}(\theta)$ as they cannot use it in making their decisions.

case a game with *ex-ante* information leakage, while the second a game with *ex-post* information leakage.¹ Figure 4.1.1 shows the two protocols. In both type of games the rules of the game, g , depend on the revealed information:

$$g : \Sigma \times \mathcal{E} \rightarrow O.$$

When agent i having type θ_i decides to send message s_i , the outcome of the game is $g(s, \varepsilon^{\text{fn}}(\theta))$ with $s = (s_i)_{i \in I}$. The utility functions of agent i is allowed to depend on the revealed information:

$$u_i : O \times \Theta_i \times \mathcal{E} \rightarrow \mathbb{R}.$$

When considering optimal mechanism design, the principal's utility is also allowed to depend on $\varepsilon^{\text{fn}}(\theta)$:

$$u_0 : O \times \mathcal{E} \rightarrow \mathbb{R}.$$

¹ The modifier ex-ante and ex-post are concerned with the event when the agents commit to a message. This use of these modifiers differs from when these modifiers refer to utilities. However, since the context makes it clear what these modify (viz. information leakage, or utility), this should not lead to any confusion.

Formally, a game with information leakage is thus defined as follows:

Definition 4.1.1 (Bayesian games with information leakage). Let Θ be a space of type profiles, P_θ the corresponding type distribution, Σ a compatible message space, O a space of outcomes, $\varepsilon^{\text{fn}} \in \mathfrak{S}(\Theta, \mathcal{E})$ an information leakage function. Further, let $g : \Sigma \times \mathcal{E} \rightarrow O$ be the mapping determining the rules of the game $G = (\Sigma, O, g)$. Then, $(G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ is called a *Bayesian game with information leakage*, where $u = (u_i)_{i \in I}$ and $u_i : O \times \Theta_i \times \mathcal{E} \rightarrow \mathbb{R}$ ($i \in I$).

Since each agent knows that information about their types is involuntarily revealed and also that this information is used to determine the output, definition of best-response maps and equilibrium concepts need to be adjusted.

In this section we focus on games with ex-ante information leakage. In this case, even the definition of strategy profile maps need to be modified as the agents can use the information leaked to adjust their strategies. In this case a strategy profile map will be a function of type

$$s : \Theta \times \mathcal{E} \rightarrow S,$$

such that for any fixed $\varepsilon \in \mathcal{E}$, $\theta \mapsto s(\theta, \varepsilon)$ is a separable function. The set of these maps will be denoted by $\mathfrak{S}_{\mathcal{E}}(\Theta, S)$.

Let us now define the best-response maps. For the definition, we introduce $P_{\theta_{-i}|\varepsilon_{-i}}^{(\varepsilon'_{-i})}(\text{d}\theta_{-i}) = P_{\theta_{-i}|\varepsilon_{-i}}(\text{d}\theta_{-i}|\varepsilon'_{-i})$, i.e., the conditional distribution of θ_{-i} given the condition that $\varepsilon_{-i}^{\text{fn}}(\theta_{-i}) = \varepsilon'_{-i}$, provided that $\theta_{-i} \sim P_{\theta_{-i}}$. Further, define $P_{\theta_i|\varepsilon_i}^{(\varepsilon'_i)}(\text{d}\theta_i) = P_{\theta_i|\varepsilon_i}(\text{d}\theta_i|\varepsilon'_i)$ and $P_{\theta|\varepsilon}^{(\varepsilon')}(\text{d}\theta) = P_{\theta|\varepsilon}(\text{d}\theta|\varepsilon')$. We have the following simple observation:

Proposition 4.1.1. $P_{\theta|\varepsilon}^{(\varepsilon')}$ is the product of its marginals $(P_{\theta_i|\varepsilon_i}^{(\varepsilon'_i)})_{i \in I}$: $P_{\theta|\varepsilon}^{(\varepsilon')}(\text{d}\theta) = \prod_{i \in I} P_{\theta_i|\varepsilon_i}^{(\varepsilon'_i)}(\text{d}\theta_i)$.

Proof. This follows from the fact that P_θ is the product of its marginals and that ε is separable. $\square$

With this, we are ready to define best-responses:

Definition 4.1.2 (Bayesian best-response maps with ex-ante information leakage). Let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage and let $s_{-i} : \Theta_{-i} \times \mathcal{E} \rightarrow S_{-i}$ be a function mapping (incomplete) type-profiles to (incomplete) strategy profiles. Define $u_{i,\theta}(o, \varepsilon) = u_i(o, \theta, \varepsilon)$, $(o, \varepsilon) \in O \times \mathcal{E}$. Given $i \in I$, s_{-i} , define

$$\begin{aligned} B_i^<(g, u_i, s_{-i}, P_{\theta_{-i}}, \varepsilon) = \\ \left\{ s_i^* : \Theta_i \rightarrow S_i : \right. \\ \left. u_{i,\theta_i} \left(g(s_i^*(\theta_i), s_{-i}(P_{\theta_{-i}|\varepsilon_{-i}}^{(\varepsilon_{-i})}), \varepsilon), \varepsilon \right) \geq \sup_{s'_i \in S_i} u_{i,\theta_i} \left(g(s'_i, s_{-i}(P_{\theta_{-i}|\varepsilon_{-i}}^{(\varepsilon_{-i})}), \varepsilon), \varepsilon \right), \right. \\ \left. \theta_i \in \text{supp}(P_{\theta_i}) \right\}, \end{aligned}$$

the set of *Bayesian best-response* maps for agent i under ex-ante information leakage.

The reader may be surprised that, despite that it was emphasized that best-responses should depend on the information leaked, the elements of $B_i^<(\dots)$ do not depend on ε . However, this is done only for convenience with later proofs: in the definition the dependence of the maps on the revealed information still exist: it is implicit in that we let the $B_i^<(\dots)$ depend on ε . Intuitively, agent i learns ε and given this information and some assumed counter-strategy map s_{-i} calculates the best responses.

The following simple observation that characterizes these best-response will prove to be useful later:

Proposition 4.1.2. *Let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage, where $G = (\Sigma, g)$, $g : S \times \mathcal{E} \rightarrow O$. For each value of $\varepsilon \in \mathcal{E}$, define the Bayesian game $\mathfrak{B}_\varepsilon = ((\Sigma, g_\varepsilon), u^{(\varepsilon)}, P_{\theta|\varepsilon}^{(\varepsilon)})$, where $g_\varepsilon : \Sigma \rightarrow O$ is defined as $g_\varepsilon(\cdot) = g(\cdot, \varepsilon)$ and $u^{(\varepsilon)} = (u_i^{(\varepsilon)})_{i \in I}$, $u_i^{(\varepsilon)}(o, \theta) = u_i(o, \theta, \varepsilon)$, $(o, \theta) \in O \times \Theta$. Fix $i \in I$ and $s_{-i} : \Theta_{-i} \rightarrow S_{-i}$, $\varepsilon \in \mathcal{E}$. Then*

$$B_i^<(g, u_i, s_{-i}, P_{\theta_{-i}}, \varepsilon) = B_i(g_\varepsilon, u_i^{(\varepsilon)}, s_{-i}, P_{\theta_{-i}|\varepsilon}^{(\varepsilon)}). \quad (4.1)$$

Note that $\mathfrak{B}_\varepsilon$ is well-defined thanks to our previous observation that $P_{\theta|\varepsilon}^{(\varepsilon)}$ is the product of its marginals (cf. Proposition 4.1.1).

Proof. The statement follows immediately from the definitions. $\square$

Definition 4.1.3. The collection $(\mathfrak{B}_\varepsilon)_{\varepsilon \in \mathcal{E}}$ is called the *ex-ante decomposition* of the Bayesian game $\mathfrak{B}$.

The definition of generalized Bayesian equilibria under ex-ante information leakage is as follows:

Definition 4.1.4 (Generalized Bayesian Equilibria (GBE) with ex-ante information leakage). Let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage and let C be a collection of counter-strategy mappings. Define

$$\mathcal{E}^<(\mathfrak{B}, C) = \left\{ s^* \in \mathfrak{S}_\mathcal{E}(\Theta, S) : s_i^*(\cdot, \varepsilon) \in \bigcap_{s'_{-i} \in C_{-i}(s^*(\cdot, \varepsilon))} B_i^<(g, u_i, s'_{-i}, P_{\theta_{-i}}, \varepsilon), i \in I, \varepsilon \in \mathcal{E} \right\}.$$

This set may be empty for some choices of C . The following result, however, gives an explicit construction that shows exactly when this equilibrium set is empty. The results follows easily from Proposition 4.1.2:

Theorem 4.1.3. Let C be a collection of counter-strategy mappings and let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage, where $G = (\Sigma, g)$, $g : S \times \mathcal{E} \rightarrow O$. Consider the ex-ante decomposition $(\mathfrak{B}_\varepsilon)_{\varepsilon \in \mathcal{E}}$ of $\mathfrak{B}$. Then

$$\mathcal{E}^<(\mathfrak{B}, C) = \{s \in \mathfrak{S}_\mathcal{E}(\Theta, S) : s(\theta, \varepsilon) = s_\varepsilon(\theta), s_\varepsilon \in \mathcal{E}(\mathfrak{B}_\varepsilon, C), \varepsilon \in \mathcal{E}, \theta \in \Theta\}.$$

Proof. Proposition 4.1.2 and the definition $\mathcal{E}^<(\mathfrak{B}, C)$ combined give

$$\mathcal{E}^<(\mathfrak{B}, C) = \left\{ s \in \mathfrak{S}_\mathcal{E}(\Theta, S) : s_i(\cdot, \varepsilon) \in \bigcap_{s'_{-i} \in C_{-i}(s(\cdot, \varepsilon))} B_i(g_\varepsilon, u_i^{(\varepsilon)}, s_{-i}, P_{\theta_{-i}|\varepsilon_{-i}}^{(\varepsilon_{-i})}), i \in I, \varepsilon \in \mathcal{E} \right\},$$

from which the result follows immediately. $\square$

Let $\mathcal{M}(I, O, \mathcal{E}) = \{(\Sigma, g) : g : \Sigma \times \mathcal{E} \rightarrow O\}$ be the space of possible mechanism for Bayesian mechanism design with information leakage. The *Bayesian*

Optimal Mechanism Design problem with ex-ante information leakage is defined as the problem of finding a mechanism $M = (\Sigma, g) \in \mathcal{M}(I, O, \mathcal{E})$ such that

$$\sup_{s \in \mathcal{E}^{<}(\mathfrak{B}_M, C_\Sigma)} \int u_0\left(g(s(\theta, \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta)\right) dP_\theta$$

is maximal over $\mathcal{M}(I, O, \mathcal{E})$, where $\mathfrak{B}_M = ((\Sigma, O, g), u, P_\theta, \varepsilon)$.

4.1.2 Formal Problem Definition for Machine Learning Auctions with Public Testing

In this section we formalize the design problem of machine learning auctions as described in the introduction in the framework of the previous section.

Let the finite set of player positions be I and for $i \in I$ let $Q_i \subset Q$ be the set of solutions that i may possess, where Q is the set of possible solutions. Further, let $\varepsilon^{\text{fn}} : Q \rightarrow \mathcal{E}$ be a mapping that is used by the principal to evaluate solutions. Thus, ε^{fn} depends on the data. The evaluation result can be a single numerical quantity (e.g., accuracy of predictions), or it can be a more detailed set of values (precision, recall, a ROC curve, etc.). The key property of ε^{fn} is that the principal, who is running the auction, should be able to ascertain the profitability of a particular solution $q \in Q$ based on the results of the $\varepsilon^{\text{fn}}(q)$. That is, it is assumed that the profitability ρ_i of q_i (which will enter the utility function u_0 of the principal) depends on q_i only through $\varepsilon_i = \varepsilon^{\text{fn}}(q_i)$, i.e., there exists a function $\rho_{\mathcal{E}} : \mathcal{E} \rightarrow \mathbb{R}$ such that $\rho_i = \rho_{\mathcal{E}}(\varepsilon_i)$.

As earlier in Chapter 3, instead of simply determining a winner whose solution is taken in a deterministic fashion, we consider a randomized choice (randomization can actually help in increasing the value that the principal can achieve and is by no means a restriction): Upon receiving all the information (bids and the test results), the principal determines a subprobability distribution $\pi \in M_{\leq 1}(I)$ and a vector $t \in \mathbb{R}^{|I|}$ of payments whose i th component determines the money to be transferred to agent i . Once, the outcome (π, t) is determined, the following events take place: With probability $1 - \sum_{i \in I} \pi_i$, the principal rejects all the offers, while with probability π_i the offer of agent i is accepted (that is, zero or one offer is accepted). If the offer on agent i is

accepted, the principal gains the right to use the solution q_i of agent i . In any case, all agents i are transferred the respective amounts t_i . Thus, the outcome space is $O = M_{\leq 1}(I) \times T$, where $T = \times_{i \in I} T_i$ with $T_i \subset \mathbb{R}$ being the set of possible payment values for agent i .

The utility function of the principal is defined by

$$u_0(\pi, t, \varepsilon) = \left(1 - \sum_{i \in I} \pi_i\right) \lambda_0 + \sum_{i \in I} \pi_i \rho_i - \sum_{i \in I} t_i, \quad (4.2)$$

where $\lambda_0 \in \mathbb{R}$ is the profit (possibly negative) that the principal makes when no solution is accepted, and $\rho_i = \rho_{\mathcal{E}}(\varepsilon_i)$. (Thus, the difference to (3.1) is that the utility function now depends on ε .) Note that as usual with randomized mechanism, the utility function determines the expected utility where the expectation is over the randomization of the mechanism.

We will represent the type θ_i of developer (or agent) i by a pair $(q_i, \lambda_i^{\text{fn}})$, where $q_i \in Q_i$ and $\lambda_i^{\text{fn}} : \mathcal{E}^{|I|} \rightarrow \mathbb{R}$ is a pricing function. We let the set of possible pricing functions for agent i be Λ_i (i.e., $\Theta_i = Q_i \times \Lambda_i$). Here, $\lambda_i^{\text{fn}}(\varepsilon)$ determines the minimum price developer i is willing to accept given the knowledge of the result of evaluating everyone's solution.² With this, the utility of agent with type $\theta_i = (q_i, \lambda_i^{\text{fn}})$, given the outcome $(\pi, t) \in O$ is

$$u_{i, \theta_i}(\pi, t, \varepsilon) = -\pi_i \lambda_i^{\text{fn}}(\varepsilon) + t_i, \quad (4.3)$$

i.e., the agent is giving up the value λ_i^{fn} ‘when he is winning the auction, while he receives the amount t_i independently of whether he won the auction. (Thus, the difference to (3.2) is that the utility function now depends on ε .) Again, the utility function determines the expected utility for the agent where the expectation is over the randomization of the mechanism.

With this, the problem is put into the framework of the previous section: Obviously, the information leakage function is $\varepsilon_i^{\text{fn}} : \Theta_i \rightarrow \mathcal{E}$, $\varepsilon_i^{\text{fn}}(q_i, \lambda_i^{\text{fn}}) = \varepsilon^{\text{fn}}(q)$. At this point modeling P_θ may look very demanding. As we will see soon, however, the optimal solution depends only on lower-dimensional distributions, which may be easier to model.

²In a special case, λ_i^{fn} is the constant function, i.e., when the developer is “happy” with a fixed price independently of how well everyone's algorithms are doing.

To put the problem into the Bayesian framework, one needs the distribution over the agent types. Modeling the whole joint distribution over Θ_i of “random developers” might look quite demanding. However, luckily, as we will see soon, this is not necessary as the solution will only depend on certain marginals of the full distribution and in fact these marginals might be much easier to estimate from past data.

To fully specify the optimal mechanism design problem it remains to choose some counter-strategy map schema, $\mathfrak{C}$. In this chapter, we consider the case when $\mathfrak{C}$ is the Nash-choice: $C_{-i}(s) = C_{-i}^N(s) = \{s_{-i}\}$, giving rise to implementations in Bayesian Nash equilibria.

4.2 Solution

In this section we present the solution to the Bayesian optimal mechanism design (BOMD) problem of Section 4.1.2. The solution is based on the observation that in this case the problem is closely related to what is studied under the name *procurement problems* in the economics literature (see, e.g., Laffont and Tirole 1993 and (Bajari and Tadelis, 2001)). In a standard procurement problem the principal is buying an “item” from the suppliers (agents). The principal knows his valuation of the item, but by an appropriate mechanism he wants to maximize his profit by keeping the cost of buying the item as low as possible. Each agent may have a different expectation for the lowest price they are willing to accept for the item. The standard solution to this problem is a reverse auction where the agents privately submit their offers to the principal who then selects the winner (the rules would then be chosen to maximize the principal’s profit, while still incentivizing the agents to participate). The principal then buys the item from the winner, i.e., the winner gives up the right to the item and the principal makes the payments. In a slightly more complicated version of the problem each player position is associated with a different item that the principal assigns different values to. The key to solving problems like this is the realization that WLOG the agents’ types can be chosen to be their private prices (no other qualities of the agent influence

the solution to the BOMD). Thus, each agent's private information is a single one-dimensional quantity, which greatly simplifies the analysis. The resulting class of problems, known as *single-parameter optimal mechanism design problems*, is extensively studied in the economics literature. The book by Laffont and Martimort (2002) provides a good starting point.

The difference of the problem of the previous section to the standard procurement problem is that in our case the “items offered for sale” come with the agents, i.e., our problem is a BOMD with information leakage. However, as we will next show, this problem can be reduced to solving a collection of standard procurement problems with no information leakage. This is the subject of the next section. Next, the solution to the standard procurement (or reverse auction) problem is given.

With this, the problem is put into the framework of the previous section: Obviously, the information leakage function is $\varepsilon_i^{\text{fn}} : \Theta_i \rightarrow \mathcal{E}$, $\varepsilon_i^{\text{fn}}(q_i, \lambda_i^{\text{fn}}) = \varepsilon^{\text{fn}}(q)$. At this point modeling P_θ may look very demanding. As we will see soon, however, the optimal solution depends only on lower-dimensional distributions, which may be easier to model.

Thus, the subproblem (4.8) to be solved, after simplifying the notation, is equivalent to the following BOMD problem:³

4.2.1 Reduction of Solving a BOMD with Ex-Ante Information Leakage to Solving Standard BOMDs

The purpose of this section is to show how the BOMD problem with ex-ante information leakage can be reduced to solving a collection of standard BOMD problems. In particular, the reduction will use the solutions to the BOMD problems defined using $D_\varepsilon = (I, O, \Theta, P_{\theta|\varepsilon}, u^{(\varepsilon)}, u_0^{(\varepsilon)}(\cdot), \mathfrak{C})$ for each different value of $\varepsilon \in \mathcal{E}$. Here, $P_{\theta|\varepsilon}$ is the product distribution of the distributions $P_{\theta_i|\varepsilon_i}$, where $P_{\theta_i|\varepsilon_i}$ is the conditional of P_{θ_i} given $\varepsilon_i^{\text{fn}}(\theta_i)$. Further, $u_0^{(\varepsilon)} : O \rightarrow \mathbb{R}$ is defined by

$$u_i^{(\varepsilon)}(o) = u_0(o, \varepsilon)$$

³Note that we are redefining the meaning of the quantities involved in the subproblem.

and $u_i^{(\varepsilon)}(o, \theta_i) = u_i(o, \theta_i, \varepsilon)$, $i \in I$ and $u^{(\varepsilon)} = (u_i^{(\varepsilon)})_{i \in I}$.

Fix $\varepsilon \in \mathcal{E}$ and let $M_\varepsilon^* = (\Sigma_\varepsilon, g_\varepsilon^*)$ be the mechanism that solves the BOMD problem D_ε . Assume that $\mathfrak{C}$ is admissible. Then, thanks to the revelation principle (Theorem 2.3.1), WLOG we can assume that $\Sigma_\varepsilon = \Sigma$, i.e., that the action spaces do not depend on ε . Consider the mechanism M^* defined by $M^* = (\Sigma, g^*) \in \mathcal{M}(I, O, \mathcal{E})$, where

$$g^*(\cdot, \varepsilon) = g_\varepsilon^*(\cdot). \quad (4.4)$$

The following theorem is the main result of this section.

Theorem 4.2.1 (Reduction of BOMD with ex-ante information leakage). *Assume that $\mathfrak{C}$ is admissible. Then M^* , as defined in the previous paragraph, solves the BOMD problem under ex-ante information leakage specified by*

$$(I, O, \Theta, P_\theta, u, u_0, \varepsilon, \mathfrak{C}).$$

Proof. We need to show that

$$M \mapsto \sup_{s \in \mathcal{E}^<(\mathfrak{B}_M, C_{\Sigma'})} \int u_0(g(s(\theta, \varepsilon^{\text{fn}}(\theta))), \varepsilon^{\text{fn}}(\theta)) dP_\theta$$

is maximized by M^* .

To prove this, first notice that by Theorem 4.1.3,

$$\mathcal{E}^<(\mathfrak{B}_{M^*}, C_\Sigma) = \{s \in \mathfrak{S}_\mathcal{E}(\Theta, S) : s(\theta, \varepsilon) = s_\varepsilon(\theta), s_\varepsilon \in \mathcal{E}(\mathfrak{B}_{M_\varepsilon^*}, C), \varepsilon \in \mathcal{E}, \theta \in \Theta\}. \quad (4.5)$$

To show that M^* is the solution of the BOMD problem under ex-ante information leakage, pick *some* mechanism $\hat{M} = (\hat{\Sigma}, \hat{g}) \in \mathcal{M}(I, O, \mathcal{E})$ and some $\hat{s} \in \mathcal{E}^<(\mathfrak{B}_{\hat{M}}, C_{\hat{\Sigma}})$. Write

$$\begin{aligned} & \int u_0(g(\hat{s}(\theta, \varepsilon^{\text{fn}}(\theta))), \varepsilon^{\text{fn}}(\theta)) dP_\theta = \\ & \int \left\{ \int u_0(g(\hat{s}(\theta, \varepsilon^{\text{fn}}(\theta))), \varepsilon^{\text{fn}}(\theta)) P_{\theta|\varepsilon}(d\theta|\varepsilon) \right\} P_\varepsilon(d\varepsilon). \end{aligned}$$

Now, fix some $\varepsilon \in \mathcal{E}$. We claim that

$$\begin{aligned} & \int u_0(g(\hat{s}(\theta, \varepsilon), \varepsilon)) P_{\theta|\varepsilon}(d\theta|\varepsilon) \leq \\ & \sup_{M' = (\Sigma', g') \in \mathcal{M}(I, O)} \sup_{s \in \mathcal{E}^<(\mathfrak{B}_{M'}, C_{\Sigma'})} \int u_0^{(\varepsilon)}(g'(s(\theta))) P_{\theta|\varepsilon}(d\theta|\varepsilon). \end{aligned}$$

Indeed, take the ex-ante decomposition $(\hat{\mathfrak{B}}_\varepsilon)_{\varepsilon \in \mathcal{E}}$ of $\mathfrak{B}_{\hat{M}}$. According to Theorem 4.1.3, $\hat{s}(\cdot, \varepsilon) \in \mathcal{E}(\hat{\mathfrak{B}}_\varepsilon, C_{\hat{\Sigma}})$. Since $\hat{\mathfrak{B}}_\varepsilon = \mathfrak{B}_{M'}$ for some $M' = (\hat{\Sigma}, g') \in \mathcal{M}(I, O)$, the above inequality indeed holds. Now, by definition the right-hand side of the above inequality equals

$$\sup_{s \in \mathcal{E}(\mathfrak{B}_{M_\varepsilon^*}, C_\Sigma)} \int u_0^{(\varepsilon)}(g_\varepsilon^*(s(\theta))) P_{\theta|\varepsilon}(d\theta|\varepsilon).$$

Asume for simplicity that in this expression the supremum is taken at some equilibrium map, say $s_\varepsilon^* \in \mathcal{E}(\mathfrak{B}_{M_\varepsilon^*}, C_\Sigma)$ (if the optimum is not taken, the result can still be obtained by taking limits). Let s^* be defined by

$$s^*(\theta, \varepsilon) = s_\varepsilon^*(\theta), \quad \theta \in \Theta, \varepsilon \in \mathcal{E}.$$

By Theorem 4.1.3, $s^* \in \mathcal{E}^<(\mathfrak{B}_{M^*}, C_\Sigma)$, with $M^* = (\Sigma, g^*)$, g^* defined by (4.4).

Putting things together, we obtain

$$\begin{aligned} \int u_0(g(\hat{s}(\theta), \varepsilon^{\text{fn}}(\theta))) dP_\theta &\leq \int \left\{ \int u_0^{(\varepsilon)}(g_\varepsilon^*(s_\varepsilon^*(\theta))) P_{\theta|\varepsilon}(d\theta|\varepsilon) \right\} P_\varepsilon(d\varepsilon) \\ &= \int \left\{ \int u_0(g^*(s_\varepsilon^*(\theta), \varepsilon)) P_{\theta|\varepsilon}(d\theta|\varepsilon) \right\} P_\varepsilon(d\varepsilon) \\ &= \int u_0(g^*(s_{\varepsilon^{\text{fn}}(\theta)}^*(\theta), \varepsilon^{\text{fn}}(\theta))) P_\theta(d\theta) \\ &= \int u_0(g^*(s^*(\theta, \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta))) P_\theta(d\theta) \\ &\leq \sup_{s \in \mathcal{E}^<(\mathfrak{B}_{M^*}, C_\Sigma)} \int u_0(g^*(s(\theta, \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta))) P_\theta(d\theta). \end{aligned}$$

Since $\hat{M}$ and $\hat{s}$ were arbitrary, the optimality of M^* is proven. $\square$

Note. We could prove this even without relying on the revelation principle if we define $\Sigma = (\times_{\varepsilon \in \mathcal{E}} \Sigma_\varepsilon)$ and g to be $g(\sigma, \varepsilon) = g_\varepsilon(\sigma_\varepsilon)$, noting that $\sigma_\varepsilon \in \Sigma_\varepsilon$ holds for any $\varepsilon \in \mathcal{E}$. The benefit is that we would not need to assume that $\mathfrak{C}$ is admissible.

4.2.2 BOMD with Ex-Ante Information Leakage and Voluntary Participation

When agent participation is voluntary (cf. Section 2.3.2, the principal needs to incentivize the agents to participate. In the standard BOMD case, this is

done by means of constraining the space of acceptable mechanism to the ones where in equilibrium the so-called IR constraints are satisfied. The effect of these constraints is that all feasible mechanisms will have the property that all rational agents will be better off by participating than by opting out of the game (i.e., the game becomes a “free ride”).

In this section we consider BOMD problems with ex-ante information leakage when agent participation is voluntary. In this case there are two points in time when the agents may decide to quit the game: Before the game starts, or after they receive the extra information. A principal who cannot enforce the agents to stay in the game has thus to incentivize the agents to stay in the game at both decision points. When the agents incur no cost by the time they receive the extra information, there is no reason they should be incentivized at the beginning if the principal commits to incentivize them given the extra information they have received. This is because any mechanism with this property will result in game that is guaranteed to be a free ride for all agents who act rationally.

Formally, this means that the principal needs to consider the following optimization problem: Remember that $\mathcal{M}(I, O, \mathcal{E}) = \{(\Sigma, g) : g : \Sigma \times \mathcal{E} \rightarrow O\}$ is the space of possible mechanism for Bayesian mechanism design with information leakage. For $M = (\Sigma, g) \in \mathcal{M}(I, O, \mathcal{E})$ and $s \in \mathcal{E}^{<}(\mathfrak{B}_M, C_\Sigma)$ let

$$u_0(g, s) = \int u_0\left(g(s(\theta, \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta)\right) dP_\theta$$

and let

$$u_*(M) = \sup_{s \in \mathcal{E}^{<}(\mathfrak{B}_M, C_\Sigma)} u_0(g, s).$$

As before let $u_{i,\theta}(o, \varepsilon) = u_i(o, \theta, \varepsilon)$, $(o, \varepsilon) \in O \times \mathcal{E}$, $i \in I$ and consider the problem

$$u_0^*(M) \rightarrow \max \quad \text{s.t. } M = (\Sigma, g) \in \mathcal{M}(I, O, \mathcal{E}) \text{ and} \quad (4.6a)$$

$$\exists s \in \mathcal{E}^{<}(\mathfrak{B}_M, C_\Sigma) \text{ s.t.} \quad (4.6a)$$

$$u_0^*(M) = u_0(g, s) \text{ and} \quad (4.6b)$$

$$u_{i,\theta_i}\left(g(s_i^*(\theta_i), s_{-i}(P_{\theta_{-i}|\varepsilon_{-i}}^{(\varepsilon_{-i})}), \varepsilon), \varepsilon\right) \geq \underline{u}_i(\theta_i), \quad (4.6c)$$

$$\text{for all } \varepsilon \in \mathcal{E}, \theta_i \in \Theta_i \text{ and } i \in I.$$

Condition (4.6c), a version of the individual rationality (RO) constraint prevents the selection of mechanism when a rational agent's expected utility given the received information would be below his reservation utility. We call this problem the *Bayesian Optimal Mechanism Design (BOMD) problem with ex-ante information leakage under voluntary participation*.

Let us now state the analogue of Theorem 4.2.2 for this case. For this, define for each value value of $\varepsilon \in \mathcal{E}$ the BOMD problems with voluntary participation: $D_\varepsilon = (I, O, \Theta, P_{\theta|\varepsilon}, u^{(\varepsilon)}, \underline{u}, u_0^{(\varepsilon)}(\cdot), \mathfrak{C})$.⁴ Fix $\varepsilon \in \mathcal{E}$ and let $M_\varepsilon^* = (\Sigma_\varepsilon, g_\varepsilon^*)$ be the mechanism that solves the problem D_ε . Assume that $\mathfrak{C}$ is admissible. Then, thanks to the revelation principle (Theorem 2.3.1), WLOG we can assume that $\Sigma_\varepsilon = \Sigma$, i.e., that the action spaces do not depend on ε . Consider the mechanism M^* defined by $M^* = (\Sigma, g^*) \in \mathcal{M}(I, O, \mathcal{E})$, where

$$g^*(\cdot, \varepsilon) = g_\varepsilon^*(\cdot). \quad (4.7)$$

With this, we have the following counterpart of Theorem 4.2.1:

Theorem 4.2.2 (Reduction of BOMD with ex-ante information leakage and voluntary participation). *Assume that $\mathfrak{C}$ is admissible. Then M^* , as defined in the previous paragraph, solves the BOMD problem under ex-ante information leakage and information leakage specified by*

$$(I, O, \Theta, P_\theta, u, \underline{u}, u_0, \varepsilon, \mathfrak{C}).$$

Proof. The proof follows that of Theorem 4.2.1 with the modifications that now one needs to also argue that (i) M^* satisfies the IR constraints (4.6c) and (ii) the decomposition of a feasible mechanism $\hat{M}$ gives rise feasible mechanisms for the BOMDP problems D_ε with voluntary participation constraint. However, these are obvious from the definitions. Then, the argument of the proof of Theorem 4.2.1 goes through. $\square$

⁴Mathematically, we would be able to deal with the case when the reservation utilities, $\underline{u}$, are allowed to depend on the revealed information (by incorporating this into (4.6c)). However, this would not add much and would just complicate the presentation further, and hence we decided not to pursue this.

4.2.3 Solution to the Machine Learning Procurement Problem

To finish the solution to the machine learning auctions problem defined in Section 4.1.2, it remains to consider the problems

$$D_\varepsilon = (I, O, \Theta, P_{\theta|\varepsilon}, u^{(\varepsilon)}, u_0^{(\varepsilon)}(\cdot), \mathfrak{C}). \quad (4.8)$$

In this section, we will first make the form of these problems explicit and then, by introducing some extra assumptions, we will present closed-form solutions for the resulting problems.

Remember that Q_i is the solution space for agent i , q_i is the solution submitted by agent i , $\varepsilon^{\text{fn}} : Q \rightarrow \mathcal{E}$ is the function that provides the evaluation results for the submitted solutions, $\rho_\mathcal{E} : \mathcal{E} \rightarrow \mathbb{R}$ is the function that maps evaluation results into profit predictions used by the principal, $\rho_i = \rho_\mathcal{E}(\varepsilon^{\text{fn}}(q_i))$ is the profitability of the i th solution ($i \in I$). Further, the outcome space is $O = M_{\leq 1}(I) \times T$, where $T = \times_{i \in I} T_i$ is the space of payments for the agents, and $M_{\leq 1}(I)$ is the space of subprobability distributions over I . An individual outcome (π, t) is implemented by choosing to refuse all the solutions of the agents with probability $1 - \sum_{i \in I} \pi_i$, and choosing the solution of agent i with probability π_i . In any case, agent i receives t_i as his payment.

Given an outcome and the results $\varepsilon_i = \varepsilon^{\text{fn}}(q_i)$, $\varepsilon = (\varepsilon_i)_{i \in I}$ of the evaluation of the submitted solutions, the principal's utility is

$$u_0(\pi, t, \varepsilon) = \left(1 - \sum_{i \in I} \pi_i\right) \lambda_0 + \sum_{i \in I} \pi_i \rho_i - \sum_{i \in I} t_i,$$

(cf. (4.2)), while the utility of agent i is

$$u_{i, \theta_i}(\pi, t, \varepsilon) = -\pi_i \lambda_i^{\text{fn}}(\varepsilon) + t_i$$

(cf. (4.3)). For technical reasons that will be explained later, we will assume that

$$\lambda_i^{\text{fn}}(\varepsilon) = \lambda_i^{\text{fn}}(\varepsilon_i)$$

where we overloaded λ_i^{fn} .⁵ The meaning of this assumption is that the developers decide about the prices they are willing to accept as a function of how

⁵That is, λ_i^{fn} denotes two functions, one mapping $\mathcal{E}^I$ to the reals and another one mapping $\mathcal{E}$ to the reals, which are distinguished based on their arguments.

well their own solutions was doing on the tests, and independently of how well the other agents' solutions performed. However, we should remind the reader here that the developers can still take ε into account in their strategic decisions when they are submitting an “offer” (which, a rational agent would indeed do by adjusting their models of the type distributions of the other agents). The revealed information in our setting is ε . Notice that given ε , the solutions q themselves do not play any role in the optimal mechanism: Given ε , the utility functions depend on q only through $\rho = \rho_{\mathcal{E}}(\varepsilon)$ and $\lambda_i = \lambda_i^{\text{fn}}(\varepsilon)$. Here, the only remaining random quantity is λ_i . Hence, in the subproblems (4.8), we may replace the conditional distribution of type i of agent given the revealed information with the conditional distribution of λ_i given ε_i , i.e., by $P_{\lambda_i|\varepsilon_i}$.

Now, this (Theorem 4.2.2) makes the individual subproblems instances of the problem studied in Chapter 3. In short, the optimal mechanism for the machine learning problem works as follows:

1. The principal announces the rules (which are as follows):
2. Agents submit their solutions (q_i for all $i \in I$);
3. The principal evaluates the agent's solutions (principal gets ε_i for all $i \in I$);
4. Each agent learns how well the solutions did (ε_i for all $i \in I$ is publicly declared);
5. The agents submit their bids (λ_i^{fn} for all $i \in I$);
6. The principal computes the virtual valuations of the agent's bids with (3.3). The agent with the highest virtual valuation wins (3.5). The payment to the agent equals to the highest bid that would have still allowed the agent to win the auctions (3.6).

4.3 Diversion: Ex-post Information Leakage

In this section, for completeness, we define the equilibria in Bayesian games with ex-post information leakage. When the leaked information remains hid-

den to the agents by the time they have to submit their messages, the agents can only reason using the “expected” information about $\varepsilon_{-i}^{\text{fn}}(\theta_{-i})$ (they may reason about this information as they know that this information got revealed and will be used in determining the outcome of the game). Before jumping into the definitions, note that since $\varepsilon_{-i}^{\text{fn}}(\theta_{-i})$ is not available to the agents, the agents’ utility functions do not depend on it, i.e., the utility function of agent i is of the type $u_i : O \times \Theta_i \rightarrow \mathbb{R}$, as opposed to the case of ex-ante information leakage. The best-response definition now becomes:

Definition 4.3.1 (Bayesian best-response maps with ex-post information leakage). Let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage and let $s_{-i} : \Theta_{-i} \rightarrow S_{-i}$ be a function mapping (incomplete) type-profiles to (incomplete) strategy profiles. Given $i \in I$, s_{-i} and g define $u_{i,s_{-i},g,\varepsilon^{\text{fn}}} : S_{-i} \times \Theta \rightarrow \mathbb{R}$ by

$$u_{i,s_{-i},g,\varepsilon^{\text{fn}}}(s_i, \theta) = u_{i,\theta_i} \left(g(s_i, s_{-i}(\theta_{-i}), \varepsilon^{\text{fn}}(\theta)) \right),$$

the utility of agent i when the agent plays the strategy s_i and the type profile of the agents is θ given s_{-i} . The set

$$\begin{aligned} B_i^>(g, u_i, s_{-i}, P_{\theta_{-i}}, \varepsilon^{\text{fn}}) = \\ \left\{ s_i^* : \Theta_i \rightarrow S_i : \quad u_{i,s_{-i},g,\varepsilon^{\text{fn}}}(s_i^*(\theta_i), \theta_i, P_{\theta_{-i}}) \geq \sup_{s'_i \in S_i} u_{i,s_{-i},g,\varepsilon^{\text{fn}}}(s'_i, \theta_i, P_{\theta_{-i}}) \right. \\ \left. \theta_i \in \text{supp}(P_{\theta_i}) \right\}, \end{aligned}$$

is called the set of *Bayesian best-response* maps for the game $\mathfrak{B}$ under ex-post information leakage.

Note that in the ex-post information leakage scenario agents know less at the time they have to figure out their messages. In an optimal mechanism design scenario, this is expected to drive up the cost for the principal. This is what is known as the *linkage principle* in auction theory.⁶

⁶ In auction theory, the linkage principle asserts that when the principal possesses private information that is a function of the types of the agents, the principal’s utility is enhanced (on average over the principal’s information) when the principal commits to a policy of always revealing her private information. Milgrom and Weber (1982) were the first to

Generalized Bayesian equilibria with ex-post information leakage is defined next:

Definition 4.3.2 (Generalized Bayesian Equilibria (GBE) with information leakage). Let $\mathfrak{B} = (G, \Theta, u, P_\theta, \varepsilon^{\text{fn}})$ be a Bayesian game with information leakage and let C be a collection of counter-strategy mappings. Define

$$\mathcal{E}^>(\mathfrak{B}, C) = \left\{ s^* \in \mathfrak{S}(\Theta, S) : s_i^* \in \bigcap_{s'_{-i} \in C_{-i}(s^*)} B_i^>(g, u_i, s'_{-i}, P_{\theta_{-i}}, \varepsilon^{\text{fn}}), i \in I \right\},$$

the set of generalized Bayesian equilibrium strategy maps with ex-post information leakage.

Let $\mathcal{M}(I, O, \mathcal{E}) = \{(\Sigma, g) : g : \Sigma \times \mathcal{E} \rightarrow O\}$ be the space of possible mechanism for Bayesian mechanism design with information leakage. The *Bayesian Optimal Mechanism Design problem with ex-post information leakage* is defined as the problem of finding a mechanism $M = (\Sigma, g)$ such that

$$\sup_{s \in \mathcal{E}^>(\mathfrak{B}_M, C_\Sigma)} \int u_0(g(s(\theta), \varepsilon^{\text{fn}}(\theta)), \varepsilon^{\text{fn}}(\theta)) dP_\theta$$

is maximized, where $\mathfrak{B}_M = ((\Sigma, O, g), u, P_\theta, \rho)$.

4.4 Summary

In this chapter we tried to reduce the problem of procurement of machine learning solutions to a reverse auction.

Our contribution in this chapter was in the introduction of the concept of information leakage and presenting a way to handle ex-ante information leakage. We followed that by reducing the problem of procurement of machine learning solution to a reverse auction using that concept.

Much remains to be done. This was only a very simple first model. Here, principal had to indirectly reveal her utility by declaring the rules of the game. In the next chapters, we extend the results to the case where this is not required and when multiple solutions are offered by each agent.

uncover the linkage principle and proved that the principle holds when a single indivisible unit is sold through either first-price (sealed-bid) auction or a Vickrey auction. Although for some time the linkage principle was widely considered one of the fundamental lessons provided by auction theory, later work by Perry and Reny (1999) has shown that it fails even if in two-unit auction, i.e., if one deviates only slightly from the original setting.

4.4.1 Known Limitations

We rely on the revelation principle in the reduction (Theorem 4.2.1). The theorem continues to hold even when the revelation principle is not used (this is important when the RP cannot be used).

4.4.2 Future Work

- A reduction for mechanism design problems with ex-post information leakage is left.
- An interesting question that follows is a comparison between mechanisms with ex-post information leakage versus mechanisms with ex-ante information leakage: Does any of these two have an upper-hand in providing more utility to the principal? Our conjecture is that ex-ante mechanisms might be preferable for the principal because it might provide more profit due to similar problems discussed under the term *linkage principle* in economics (Milgrom and Weber, 1982).
- Next chapter considers a similar, yet different problem when the principal does not want to reveal her profit information to the agents.⁷

⁷ In the problem considered by Myerson (1983), the principal has also a type which determines her utility. Further, the principal also takes part in the mechanism by sending some message which also influences the outcome, the principal also has a type and agent participation is voluntary. Hence, the agents' expected utilities (conditioned on what they know when they decide about their messages) will depend on what they know about the principal's type. The principal, by announcing the mechanism may reveal information about her type (which he knows). Myerson, considering incentive compatible direct-revelation mechanisms, then concludes the need for additional IC constraints that take into account the information revealed by the announcement of the mechanism. Thus, the problem is similar to ex-ante information leakage. Myerson argues that the principal should select with no loss of generality an incentive compatible direct-revelation mechanism that reveals no information about the principal's type (i.e., the principal should be "inscrutable"). He then puts the optimal mechanism design problem into a two-stage sequential framework, where he considers implementation in sequential (Bayesian Nash) equilibria. He shows the existence of solutions and considers various refined concepts.

Chapter 5

Mechanism Design with Exogenous Effects

In the last chapter we tried to model the problem of procurement of machine learning solution to a reverse auction. However, by declaring the rules of the game, the principal is her profit predictions to the agents. This may be due to different reasons, for example, revealing utility information might have negative side-effects. A similar case is when the principal does not possess the data yet but wants to procure a solution nonetheless, or the agents might not have yet developed a solution, or tuned their solutions to the specific machine learning problem. In fact, by running a “blind-auction” the company can possibly eliminate bidders whose bids are dominated before the bidders even would have a solution that could be submitted for evaluation. This is advantageous from the point of the eliminated bidders (they do not need to work) and the company as well (as fewer solutions need to be evaluated). This chapter addresses this problem.

5.1 Games with Exogenous Signals

The concept of random exogenous signals that could affect the desirability of outcomes is something that comes up as a part of designing mechanism for machine learning solutions. We extend mechanism design with such a concept.

In this chapter we introduce the concept of mechanism design for problems with exogenous signals. In these problems, an exogenous signal (selected by

“Nature”) enters both the utility of the principal and the outcome of the game. The exogenous signal is not under the control of the agents, nor is it under the control of the principal. None of the parties know about the value this exogenous signal takes by the time they take their actions (in the case of the principal, this means the time when she designs the auction).

To develop the theory for this case, we will first consider games with exogenous signals and introduce an equilibrium concept. In a game with exogenous signals, the outcome of the game is introduced by the exogenous signal. Again, the agents have no knowledge of the value of this signal, nor can they control the value it takes. In the equilibrium concept we propose the agents follow strategies such that none of them would have an incentive to deviate in an unilateral fashion *no matter what the value the exogenous signal takes*. Although this may look restrictive (the set of such “strong” equilibria may be empty for many games), we show that this concept actually gives rise to a strong, pointwise, solution to the optimal mechanism design problem.

From now on, we use I to denote the index set of player positions, while we will use R to denote the set of values that the exogenous signal may take.

Definition 5.1.1 (Game with exogenous signals). A 4-tuple $G_R = (\Sigma, g, O, R)$ is called a *game with exogenous signals* R if $\Sigma = \times_{i \in I} \Sigma_i$ and $g: \Sigma \times R \rightarrow O$.

The information flow in the game is as follows:

1. The types of the agents are drawn from P_θ ;
2. The agents choose their actions simultaneously, leading to the joint action $\sigma \in \Sigma$;
3. Nature picks the exogenous signal $\rho \in R$;
4. The outcome of the game, $g(\sigma, \rho)$ is announced to all the players.

By saying that it is Nature who is picking the exogenous signal, our goal is to make it clear that the signal is picked by a disinterested party.

As before, we add types (Θ, P_θ) and type-dependent utility functions $u = (u_i)_{i \in I}$, $u_i: O \times \Theta_i \rightarrow \mathbb{R}$ to arrive at Bayesian-games:

Definition 5.1.2 (Bayesian game with exogenous signals). Let $\Theta = \times_{i \in I} \Theta_i$ be a type-space and P_θ a distribution over the types that has a product form: $P_\theta(d\theta) = \prod_{i \in I} P_{\theta_i}(d\theta_i)$. Further, let $u = (u_i)$ be a collection of typed-utility functions. Let G_R be a game with exogenous signals. Then, we call the 4-tuple $\mathfrak{B} = (G_R, \Theta, u, P_\theta)$ is a *Bayesian game with exogenous signals*.

As before with Bayesian games, a Bayesian game starts with every agent learning their own type θ_i , which are drawn independently from each other from the respective distributions, P_{θ_i} . Then, the game continues as before.

Next, we define our desired equilibrium concept. Our aim is to define a strong equilibrium concept since, as we will see soon, from an implementation point of view even this strong concept will be tractable. For the definition, let us first define the decomposition of game with exogenous signals:

Definition 5.1.3 (Decomposition of game with exogenous signals). Let $\mathfrak{B} = (G_R, \Theta, u, P_\theta)$ be a Bayesian game with exogenous signals, where $G_R = (\Sigma, O, g, R)$. For $\rho \in R$, define $g_\rho(\cdot) = g(\cdot, \rho)$. We call the collection of Bayesian games $(\mathfrak{B}_\rho)_{\rho \in R}$ the *decomposition of $\mathfrak{B}$* if $\mathfrak{B}_\rho = (G_\rho, \Theta, u, P_\theta)$ where $G_\rho = (\Sigma, O, g_\rho)$.

Note that when R is a singleton, we get back to the usual definition that does not concern exogenous signals. This property will remain true throughout this section, i.e., the theory developed here for the case of exogenous signals is a strict generalization of the previous theory.

With this, we are ready to introduce our equilibrium concept:

Definition 5.1.4 (Strong equilibria for games with exogenous signals). Let $\mathfrak{B}$ be a Bayesian game with exogenous signals R , action set Σ and outcome set O and let $(\mathfrak{B}_\rho)_{\rho \in R}$ be the decomposition of $\mathfrak{B}$. Let C be a counter-strategy map over the message set Σ . Then, $\sigma \in S$ is an equilibrium action profile for $\mathfrak{B}$ w.r.t. C if it is an equilibrium action profile for $\mathfrak{B}_\rho$ and C for every $\rho \in R$. The set of these equilibria is denoted by $\mathcal{E}(\mathfrak{B}, C)$. Thus,

$$\mathcal{E}(\mathfrak{B}, C) = \bigcap_{\rho \in R} \mathcal{E}(\mathfrak{B}_\rho, C). \quad (5.1)$$

Note that $\sigma \in \mathcal{E}(\mathfrak{B}, C)$ means that σ_i is simultaneously a best-response for all the different games that arise from the exogenous signal. Hence, for many games $\mathcal{E}(\mathfrak{B}, C)$ could well be empty. The essence of the definition is for the agents who have no knowledge or control of the exogenous signal it is rational to use any strategy map $\sigma \in \mathcal{E}(\mathfrak{B}, C)$ assuming that such a map exists.

We now define the *product* of a collection of Bayesian games, a construction that will allow us to solve optimal mechanism design problems in a strong sense (to be discussed in the next section). For the construction, let $(\mathfrak{B}_\rho)_{\rho \in R}$ be a collection of Bayesian games that share the same action set, outcomes, utility function and priors: $\mathfrak{B}_\rho = (G_\rho, \Theta, u, P_\theta)$ with $G_\rho = (\Sigma, O, g_\rho)$.¹ Now, consider the Bayesian game with exogenous signals, $\mathfrak{B} = (G, \Theta, u, P_\theta)$, which is obtained as follows: $G = (\Sigma', O, g, R)$, where $\Sigma' = \times_{i \in I} (\Sigma_i)^R$ and where $g: \Sigma' \times R \rightarrow O$ is defined follows: Pick $\sigma' = (\sigma'_i)_{i \in I} \in \Sigma'$ and $\rho \in R$. Note that $\sigma'_i: R \rightarrow \Sigma_i$. Define $\sigma'(\rho) = (\sigma'_i(\rho))_{i \in I}$. Then, we define²

$$g(\sigma', \rho) = g_\rho(\sigma'(\rho)), \quad \sigma' \in \Sigma', \rho \in R. \quad (5.2)$$

Definition 5.1.5. We say that G as defined in the previous paragraph is the *product* of the games $(G_\rho)_{\rho \in R}$. Similarly, $\mathfrak{B}$ is said to be the *product* of the games $(\mathfrak{B}_\rho)_{\rho \in R}$.

It will prove to be useful to define the *projection maps* $\mathfrak{P}_\rho: \mathfrak{S}(\Theta, S') \rightarrow \mathfrak{S}(\Theta, \Sigma)$ that take a strategy map $s' \in \mathfrak{S}(\Theta, S')$ and project it by fixing the exogenous signal to ρ : $(\mathfrak{P}_\rho s')(\theta) = s'(\theta)(\rho)$ (note that $s'(\theta)$ itself is a $R \rightarrow S$ function). We will also need the analogously defined maps for the i th component of strategy maps: $(\mathfrak{P}_\rho s'_i)(\theta) = s'_i(\theta)(\rho)$, as well as for counter-strategy maps: $(\mathfrak{P}_\rho s'_{-i})(\theta) = s'_{-i}(\theta)(\rho)$. We further introduce the convention that for any function $G: A \rightarrow B$ and any set $U \subset A$, $G(U)$ will denote the set

¹ Notice that we assume that all the games share the same action set. This is not a crucial assumption and is adopted only for the sake of simplifying the notation. The definitions and results below can be easily modified to accommodate the general case when each game G_ρ has its own action set Σ_ρ .

² When the action set of G_ρ is Σ_ρ , Σ_i^R in the definition of Σ' must be replaced by the set $\{f: R \rightarrow \cup_{\rho \in R} \Sigma_\rho : f(\rho) \in \Sigma_\rho\}$, where as assumed WLOG that the sets $(\Sigma_\rho)_{\rho \in R}$ are disjoint. Then, the rest goes through with some obvious modifications.

$\{Gx : x \in A\}$. When no ambiguity arises, we will often write GU instead of $G(U)$.

We have the following result:

Theorem 5.1.1 (Equilibria of Product Games). *Let $\mathfrak{B}$ be the product of the games $(\mathfrak{B}_\rho)$ and let $\mathfrak{C}$ be a counter-strategy map schema that commutes with $\mathfrak{P}$. in the sense that*

$$\mathfrak{P}_\rho C_{-i}^{(\Sigma')}(s') = C_{-i}^{(\Sigma)}(\mathfrak{P}_\rho s') \text{ for all } s' \in \Sigma'. \quad (5.3)$$

Then,

$$\mathcal{E}(\mathfrak{B}, C_{\Sigma'}) = \{s' \in \mathfrak{S}(\Theta, S') : \mathfrak{P}_\rho s' \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma), \quad \rho \in R\}. \quad (5.4)$$

In particular, s' is an equilibrium of $\mathcal{E}(\mathfrak{B}, C_{\Sigma'})$ if and only if it can be written in the form

$$s'(\theta)(\rho) = s_\rho(\theta), \quad \rho \in R, \theta \in \Theta \quad (5.5)$$

with some $s_\rho \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma)$. Equivalently,

$$\mathfrak{P}_\rho \mathcal{E}(\mathfrak{B}, C_{\Sigma'}) = \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma), \quad \text{for all } \rho \in R. \quad (5.6)$$

From the second part of the theorem we obtain an explicit way of constructing the equilibria of the product game $\mathfrak{B}$ given the equilibria of its constituent games by “patching them up” using (5.5). Note that the counter-strategy map schemas defining dominant equilibria and Bayesian Nash equilibria both commute with $\mathfrak{P}$.

Proof. The second part (i.e., the result about “patching”) immediately follows from (5.4). It is also clear that (5.6) is equivalent to (5.4). Hence, it remains to prove (5.4).

For this, let $(\mathfrak{B}_\rho^{(d)})_{\rho \in R}$ be the decomposition of $\mathfrak{B}$ and let $g_\rho^{(d)} : \Sigma' \rightarrow O$ be the outcome map of $\mathfrak{B}_\rho^{(d)}$. Then, thanks to the construction of $\mathfrak{B}$,

$$g_\rho^{(d)}(\sigma') = g(\sigma(\rho), \rho). \quad (5.7)$$

Fix $\rho \in R$, $i \in I$, $s'_{-i} \in \mathfrak{S}(\Theta_{-i}, S'_{-i})$. By the definitions and (5.7), we find that

$$\begin{aligned}
& B_i(g_\rho^{(d)}, u_i, s'_{-i}, P_{\theta_{-i}}) \\
&= \bigcap_{\theta_i \in \text{supp}(P_{\theta_i})} \left\{ s'_i : \Theta_i \rightarrow S'_i : \right. \\
&\quad \left. u_{i, \theta_i}(g_\rho^{(d)}(s'_i(\theta_i), s'_{-i}(P_{\theta_{-i}}))) \geq \sup_{s''_i \in S'_i} u_{i, \theta_i}(g_\rho^{(d)}(s''_i, s'_{-i}(P_{\theta_{-i}}))) \right\} \\
&= \bigcap_{\theta_i \in \text{supp}(P_{\theta_i})} \left\{ s'_i : \Theta_i \rightarrow S'_i : \right. \\
&\quad \left. u_{i, \theta_i}(g(s'_i(\theta_i)(\rho), s'_{-i}(P_{\theta_{-i}})(\rho), \rho)) \geq \sup_{s''_i \in S'_i} u_{i, \theta_i}(g(s''_i, s'_{-i}(P_{\theta_{-i}})(\rho), \rho)) \right\} \\
&= \{ s'_i : \Theta_i \rightarrow S'_i : \mathfrak{P}_\rho s'_i \in B_i(g_\rho, u_i, \mathfrak{P}_\rho s'_{-i}, P_{\theta_{-i}}) \}. \tag{5.8}
\end{aligned}$$

Now, let us expand the definition of $\mathcal{E}(\mathfrak{B}, C_{\Sigma'})$. We have

$$\mathcal{E}(\mathfrak{B}, C_{\Sigma'}) = \bigcap_{\rho \in R} \mathcal{E}(\mathfrak{B}_\rho^{(d)}, C_{\Sigma'}).$$

Hence, $s' \in \mathcal{E}(\mathfrak{B}, C_{\Sigma'})$ if and only if $s' \in \mathcal{E}(\mathfrak{B}_\rho^{(d)}, C_{\Sigma'})$ for all $\rho \in R$. Fix some $\rho \in R$ and $s' \in \mathfrak{S}(\Theta, S')$.

We have have that $s' \in \mathcal{E}(\mathfrak{B}_\rho^{(d)}, C_{\Sigma'})$ holds if and only if

$$s'_i \in B_i(g_\rho^{(d)}, u_i, s'_{-i}, P_{\theta_{-i}}) \quad \text{for all } i \in I, s'_{-i} \in C_{-i}^{(\Sigma')}(s'). \tag{5.9}$$

By (5.8), $s'_i \in B_i(g_\rho^{(d)}, u_i, s'_{-i}, P_{\theta_{-i}})$ holds if and only if $\mathfrak{P}_\rho s'_i \in B_i(g_\rho, u_i, \mathfrak{P}_\rho s'_{-i}, P_{\theta_{-i}})$ holds. Hence, by the condition that $\mathfrak{P}_\rho$ commutes with $\mathfrak{C}$ (cf. (5.3)), (5.9) is equivalent to

$$\mathfrak{P}_\rho s'_i \in B_i(g_\rho, u_i, s_{-i}, P_{\theta_{-i}}) \quad \text{for all } i \in I, s_{-i} \in C_{-i}^{(\Sigma)}(\mathfrak{P}_\rho s'),$$

which in turn is equivalent to

$$\mathfrak{P}_\rho s' \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma),$$

where we used that $(\mathfrak{P}_\rho s')_i = \mathfrak{P}_\rho s'_i$ holds for any $s' \in \mathfrak{S}(\Theta, S)$, $\rho \in R$ and $i \in I$.

This shows that

$$\mathcal{E}(\mathfrak{B}_\rho^{(d)}, C_{\Sigma'}) = \{ s' \in \mathfrak{S}(\Theta, S') : \mathfrak{P}_\rho s' \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma) \}.$$

Hence,

$$\begin{aligned}
\mathcal{E}(\mathfrak{B}, C_{\Sigma'}) &= \bigcap_{\rho \in R} \mathcal{E}(\mathfrak{B}_\rho^{(d)}, C_{\Sigma'}) \\
&= \bigcap_{\rho \in R} \{s' \in \mathfrak{S}(\Theta, S') : \mathfrak{P}_\rho s' \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma)\} \\
&= \{s' \in \mathfrak{S}(\Theta, S') : \mathfrak{P}_\rho s' \in \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma), \rho \in R\} ,
\end{aligned}$$

finishing the proof. $\square$

An easy corollary of the previous result is as follows:

Corollary 1. *If all the games $\mathfrak{B}_\rho$ have a unique equilibrium, then the product game $\mathfrak{B}$ has also a unique equilibrium.*

5.2 Mechanism Design with Exogenous Signals

Let us now consider the problem of optimal mechanism design for problems when an exogenous signal that is neither known, nor controlled by either the agents or the principal is *bound to influence the utility of the principal*. What can the principal achieve in this case?

Below we pick up a rather bold goal: we search for a mechanism that achieves the best possible ex-ante expected utility for the principal *for all values of the exogenous signal*, i.e., to maximize the principal's utility in a *pointwise fashion*. We formulate this optimization problem next.

Let the principal's utility function be $u_0 : O \times R \rightarrow \mathbb{R}$ and fix some counter-strategy map schema $\mathfrak{C}$. Let us denote by $\mathcal{M}_R(I, O)$ the space of mechanisms of the form $M = (\Sigma, g)$, where $g : O \times R \rightarrow \mathbb{R}$. Given $M = (\Sigma, g) \in \mathcal{M}_R(I, O)$ let $\mathfrak{B}_M = (G_R, \Theta, u, P_\theta)$, where $G_R = (\Sigma, O, g_M, R)$. Define the principal's expected utility under M , for a given value of $\rho \in R$ and for a given equilibrium strategy map $s \in \mathcal{E}(\mathfrak{B}_M, C_\Sigma)$ by

$$u_0(M, s, \rho) = \int u_0(g_M(s^\theta), \rho), \rho) dP_\theta .$$

For $\rho \in R$, define

$$u_0^*(\rho) = \sup_{M=(\Sigma, g) \in \mathcal{M}_R(I, O)} \sup_{s \in \mathcal{E}(M, C_\Sigma)} u_0(M, s, \rho) ,$$

the largest utility that the principal can achieve given the exogenous signal ρ . We define the Bayesian optimal mechanism design problem for games with exogenous signals as follows:

Definition 5.2.1 (Bayesian optimal mechanism design problem for games with exogenous signals (BOMDX)). Given the set R of exogenous signals, the tuple $(I, O, \Theta, u, P_\theta, u_0)$ ($u_0 : O \times R \rightarrow \mathbb{R}$) and a counter-strategy map schema $\mathfrak{C} : \Sigma \mapsto C_\Sigma$, we say that $M^* = (\Sigma^*, g^*) \in \mathcal{M}_R(I, O)$ is a solution to the *Bayesian optimal mechanism design problem* with exogenous signal from R if there exists an equilibrium strategy map $s^* \in \mathcal{E}(M^*, C_{\Sigma^*})$ such that

$$u_0^*(M^*, s^*, \rho) = u_0^*(\rho), \quad \text{for all } \rho \in R.$$

We call M^* an optimal mechanism for the corresponding BOMDX problem.

The extension of BOMDX to problems with voluntary participation holds no surprises: When considering voluntary participation, we need to constrain the optimization problem to those mechanisms that have an optimizing equilibria that yields expected utilities that exceed the agents' respective reservations utilities. However, in this case these “individual rationality” constraints must be met for all values of the exogenous signal. Since the extension is trivial, we omit the details (Theorem 5.2.1 given below can be extended to this case without any problems).

The next result shows that the goal of finding pointwise maximizing mechanism is not that bold. Intuitively, the principal can ask the agents to send messages where they describe how they will behave for each value of the exogenous signal and then find the best mechanism for each value of the exogenous signal. Note that for this mechanism, the principal does not need to know the value of the exogenous signal. The same idea also works when the goal is to hide some details of the principal's utility function, which may be interesting if the utility function could reveal sensitive information.³

³ Myerson (1983) essentially describes the same idea for the case when R and Θ are finite and for a slightly different setting when the mechanism is symmetric in that the principal also sends a message. He calls this the case of “informed principal” and a principal “inscrutable” if the principal's mechanism does not reveal information about her type.

An explicit solution to the BOMDX problem can be obtained as follows:
For $\rho \in R$, let $u_0^{(\rho)} : O \rightarrow \mathbb{R}$ be given by

$$u_0^{(\rho)}(\cdot) = u_0(\cdot, \rho)$$

and let $M_\rho^* = (\Sigma_\rho^*, g_\rho^*)$ be the solution to the BOMD problem specified by the tuple $(I, O, \Theta, u, P_\theta, u_0^{(\rho)})$ and $\mathfrak{C}$. We assume that $\mathfrak{C}$ is admissible, hence, by Theorem 2.3.1, WLOG we can assume that $\Sigma_\rho^* = \Sigma^*$ (i.e., the action sets do not depend on ρ).⁴ Let $G_\rho^* = (\Sigma^*, O, g_\rho^*)$ and let $G^* = (\Sigma_\times^*, O, g_\times^*)$ be the product of $(G_\rho^*)_{\rho \in R}$. Let $M^* = (\Sigma_\times^*, g_\times^*)$ be the mechanism underlying G . We have the following result:

Theorem 5.2.1. *Let B denote the BOMDX problem specified by the set R , the tuple $(I, O, \Theta, u, P_\theta, u_0)$ and a counter-strategy map schema $\mathfrak{C}$, where $u_0 : O \times R \rightarrow \mathbb{R}$ is the principal's utility. Assume that $\mathfrak{C}$ is admissible and commutes with $\mathfrak{P}$ in the sense of (5.3). Further, assume that for any $\rho \in R$, M_ρ^* is well-defined and assume that there exists $s_\rho^* \in \mathcal{E}(M_\rho^*, C_\Sigma^*)$ such that*

$$\int u_0^{(\rho)}(g_\rho^*(s_\rho^*(\theta))) dP_\theta = \sup_{s \in \mathcal{E}(M_\rho^*, C_\Sigma^*)} \int u_0^{(\rho)}(g_\rho^*(s(\theta))) dP_\theta. \quad (5.10)$$

Then M^ as defined in the preceding paragraph is an optimal solution to the BOMDX problem B .*

Proof. Let $M = (\Sigma, g_M) \in \mathcal{M}_R(I, O)$ be any mechanism for the BOMDX problem and let $s \in \mathcal{E}(B_M, C_\Sigma)$ be any strong equilibrium of the resulting Bayesian game. Define $(M_\rho)_{\rho \in \mathcal{E}}$ to be the decomposition of M : $M_\rho = (\Sigma, g_M^{(\rho)}) \in \mathcal{M}(I, O)$, where $g_M^{(\rho)} : \Sigma \rightarrow O$ is given by $g_M^{(\rho)}(\sigma) = g_M(\sigma, \rho)$, $\sigma \in \Sigma$, $\rho \in R$.

By (5.1) (i.e., the definition of strong equilibria), $s \in \mathcal{E}(M_\rho, C_\Sigma)$. Define $s' \in \mathcal{E}(M^*, C_{\Sigma_\times^*})$ by $\mathfrak{P}_\rho s' = s_\rho^*$. Note that s' is well-defined by Theorem 5.1.1.

⁴ In fact, this assumption is not necessary as it was hinted upon earlier.

We have

$$\begin{aligned}
u_0(M, s, \rho) &= \int u_0(g_M(s(\theta), \rho), \rho) dP_\theta \\
&= \int u_0^{(\rho)}(g_M(s(\theta), \rho)) dP_\theta \\
&= \int u_0^{(\rho)}(g_M^{(\rho)}(s(\theta))) dP_\theta \\
&\leq \sup_{s \in \mathcal{E}(M_\rho^*, C_\Sigma^*)} \int u_0^{(\rho)}(g_\rho^*(s(\theta))) dP_\theta \\
&= \int u_0^{(\rho)}(g_\rho^*(s_\rho^*(\theta))) dP_\theta \\
&= \int u_0(g_{M^*}(s'(\theta), \rho), \rho) dP_\theta,
\end{aligned}$$

where the inequality follows since M_ρ^* is an optimal solution to the BOMD problem specified by $(I, O, \Theta, u, P_\theta, u_0^{(\rho)})$ and $\mathfrak{C}$, while the last equality uses the definition of M^* and s' . Taking the supremum of both sides w.r.t. M and $s \in \mathcal{E}(B_M, C_\Sigma)$, we get that

$$u_0^*(\rho) \leq \int u_0(g_{M^*}(s'(\theta), \rho), \rho) dP_\theta = u_0(M^*, s', \rho),$$

which proves the statement. $\square$

By the revelation principle, a mechanism is often implemented by truthful declarations. If, such a mechanism is used for all problem M_ρ $\rho \in R$, then the message space of the agents can be significantly simplified. Intuitively, if independently of the value ρ , the same equilibrium strategy “works” then the principal can explain to the agents that although she could require them to send a vector of actions with one component for each possible value of ρ , since they would send the same action each time (more precisely, same strategy), they may just compress this and send the action only once. This is made formal by the following result:

Theorem 5.2.2. *Assume that $\mathfrak{C}$ is admissible and commutes with $\mathfrak{P}$. in the sense of (5.3). Further, assume that for any $\rho \in R$, $M_\rho^* = (\Sigma, g_\rho^*) \in \mathcal{M}(I, O)$ is well-defined and the optimizing equilibrium maps, s_ρ^* defined by (5.10) can be all taken to be equal:*

$$s_\rho^* = s_{\rho'}^*, \quad \rho, \rho' \in R. \quad (5.11)$$

Define $M^* = (\Sigma, g^*) \in \mathcal{M}_R(I, O)$ by $g^* : \Sigma \times R \rightarrow O$, $g^*(\cdot, \rho) = g_{\rho^*}(\cdot)$. Then M^* is an optimal solution to the BOMDX problem B of Theorem 5.2.1.

Proof. Let M^* be as in the statement. Let $\mathfrak{B} = \mathfrak{B}_{M^*}$ be the corresponding Bayesian game and let $(\mathfrak{B}_\rho)_{\rho \in R}$ be its decomposition. Then, for any $\rho \in R$,

$$\mathcal{E}(\mathfrak{B}_\rho, C_\Sigma) = \mathcal{E}(\mathfrak{B}_{M_\rho^*}, C_\Sigma),$$

where $\mathfrak{B}_{M_\rho^*}$ is the Bayesian game corresponding to M_ρ^* . and hence

$$\mathcal{E}(\mathfrak{B}, C_\Sigma) = \bigcap_{\rho \in R} \mathcal{E}(\mathfrak{B}_\rho, C_\Sigma) = \bigcap_{\rho \in R} \mathcal{E}(\mathfrak{B}_{M_\rho^*}, C_\Sigma).$$

Let s^* be the common value of the optimizing equilibrium maps in (5.11).

Then, $s^* \in \mathcal{E}(\mathfrak{B}_{M_\rho^*}, C_\Sigma)$ for each $\rho \in R$ and hence $s^* \in \mathcal{E}(\mathfrak{B}, C_\Sigma)$.

Now, proceeding as in the proof of Theorem 5.2.1, let $M = (\Sigma, g_M) \in \mathcal{M}_R(I, O)$ be any mechanism for the BOMDX problem and let $s \in \mathcal{E}(B_M, C_\Sigma)$ be any strong equilibrium of the resulting Bayesian game. Then,

$$\begin{aligned} u_0(M, s, \rho) &\leq \int u_0^{(\rho)}(g_\rho^*(s_\rho^*(\theta))) dP_\theta \\ &= \int u_0^{(\rho)}(g_\rho^*(s^*(\theta))) dP_\theta \\ &= \int u_0(g^*(s^*(\theta), \rho), \rho) dP_\theta. \end{aligned}$$

Taking the supremum of both sides with respect to M and $s \in \mathcal{E}(M, C_\Sigma)$, we get

$$u_0^*(\rho) \leq \int u_0(g^*(s^*(\theta), \rho), \rho) dP_\theta.$$

Since $s^* \in \mathcal{E}(\mathfrak{B}, C_\Sigma)$, it follows that M^* is indeed an optimal mechanism. $\square$

5.3 Application to Machine Learning Auctions

Let us now apply the results of this chapter to machine learning procurement problems. For example, take the situation that was addressed in Chapter 3. In addition, let $\rho_i \in R$ for all $i \in I$ denote possible profit values of solution of agent i which are now unknown to the agents. This too can be a single numerical quantity, or it can be a more detailed set of values. The utility function are exactly as defined in (3.1) and (3.2).

We can solve this problem using Theorem 5.2.1, but in fact, even the solution, which is more convenient for the parties involved, of Theorem 5.2.2 is applicable. The reason is because the mechanism is such that for *any value* of $(\rho_i)_{i \in I}$, at equilibrium, the agents should declare the “true” price of their solutions honestly. Hence, Theorem 5.2.2 applies.

The solution then will be as follows:

1. The principal announces the rules (which are as follows):
2. Agents submit their bids;
3. The principal evaluates the agent’s solutions (principal gets ρ_i for all $i \in I$, but he keeps these values secretly);
4. The principal computes the virtual valuations of the agent’s bids with (3.3). The agent with the highest virtual valuation wins (3.5). The payment to the agent equals to the highest bid that would have still allowed the agent to win the auctions (3.6).

Now, the principal also explains to the agents that the mechanism is designed in such a way that at equilibrium they should submit the fair price of their solutions, or they are risking losing utility.

5.4 Summary

What we did in this chapter is similar to what was done in Chapter 4. There are differences though: BOMD with ex-ante information leakage reveals the principal’s utility information but has a short and concise form and is guaranteed optimal. There, agents change their expectations (ex-ante expected utilities change) according to distributional information about the signal that will be used at the end in the calculation of the outcome. BOMD with exogenous signals, on the other hand, is used for the cases when the principal wants to hide her utility information and may not be optimal in the sense of BOMD with ex-ante information leakage. In this case, there is no such distributional information available to the agents and rationality is defined as

playing best response for all possible values of the exogenous signal. One could use the BOMDX framework to address the ex-ante problem, in theory, but the solution may not be optimal. This is analogous to the certainty equivalence principle in control theory in the fact that the agents do not have the “true” model, but only an estimate, and yet they react as if they had the model. This is similar to using BOMDX and asking the agents to plan for all possible values of the signal (but clearly, the solution arising may not be optimal in expectation).

5.4.1 Known Limitations

A known limitation is the use of the revelation principle. This has been discussed in previous chapters why revelation principle might provide a back-door and the solution may be sub-optimal when faced with voluntary participation.

5.4.2 Future Work

- Because we used a reduction, we can combine the results of this chapter with the results from the last chapter. For example, principal wants to hide profit figures but prefers ex-ante evaluation of the solutions of the developers. The result is the following mechanism: She would use BOMDX as an outer shell and inside, she would employ the results from Chapter 4 for each case.
- Another things that would be interesting to see is a common generalization of the mechanism design with exogenous information to extensive form-games and multi-stage auctions: A sequential framework, for example.

Chapter 6

Developers have Multiple Machine Learning Packages

Let us consider again a machine learning solution procurement problem. Imagine now, however, that each of the developers have multiple off-the-shelf solutions that worth different amounts (e.g., because some of them cost less to develop than others). In the previous chapters we assumed that each developer submits a single solution. How should a developer decide which solution to submit? If the cost of evaluating solutions is negligible then instead of requiring each developer to submit a single solution, we may ask them to submit all of their solutions, leaving it to the principal to choose the (single) solution that represents the best trade-off between cost and profit. How should then the principal change the allocation and payment rules of the auction to maximize her expected profit? This is the question that is answered in this chapter. If we abstract away the details of this problem, the resulting auctions should perhaps be called *multi-item reverse auctions* (see for the formal definition below). Note that multi-item (reverse) auctions should not be confused with the commonly studied multi-unit auctions where multiple copies of *identical* items are offered for sale (respectively, are procured). Although multi-unit auctions are well researched, we were not able to identify any relevant papers for the multi-item problem studied here. Again, this is only probably because of our unfamiliarity with the literature, although maybe the problem described is specific enough for machine learning and it does not come up that often in the economics literature and thus it remained unstudied.

Before diving into the details, let us first briefly consider whether the solution of Chapter 3 would work for this setting. In this solution the principal decides about which solution to buy by first computing the virtual values for each of the solutions and then, assuming that there exists a solution whose virtual value is above the reservation utility, buys the solution whose virtual value is the highest. The price that the principal pays for the solution is the price that would make the same item still win. At this stage we may suspect that this payment rule will not work as the price could be manipulated by a developer by manipulating the prices of the solutions that he submits. However, as we will see the actual optimal solution to the multi-item optimal reverse auction will be very close to this solution, but the amount to be paid will be the maximum amount under which the current winner will stay remain a winner. The proof is a fairly straightforward generalization of the proof presented in Chapter 3 (with some minimal technical difficulties) and is presented mainly for the sake of completeness.

Note that the solution of multi-item reversed auctions presented in the next section can, naturally and in an obvious fashion, be used in the more complicated settings of the previous two chapters. In fact, this is a nice illustration of the strength of the reduction approach followed in the thesis. The obvious details of these combinations are left for the reader.

6.1 Problem Description

Fix I , the set of developer positions. For $i \in I$, let $J_i = \{1, \dots, |J_i|\}$ be indices used to index the items offered by developer i for sale (i.e., the machine learning solutions). Let

$$K = \{(i, j) : i \in I, j \in J_i\}$$

be the index set for the solutions (items). The outcome space is $O = M_{\leq 1}(K) \times \mathbb{R}^{|I|}$: For some $(\pi, t) \in O$ where $\pi = (\pi_{i,j})_{i,j \in K}$ and $t = (t_i)_{i \in I}$ the meaning of π and t are as follows: The item $(i, j) \in K$ is bought with probability $\pi_{i,j}$ and the payment to developer i is t_i . (With probability $1 - \sum_{(i,j) \in K} \pi_{i,j}$ the principal buys none of the items.) We will also use $\pi_i = (\pi_{i,j})_{j \in J_i}$ and we will

treat this as a column vector of $\mathbb{R}^{|J_i|}$. The inner product of two vectors, a, b , of identical dimensions will be denoted by $\langle a, b \rangle$. Further, we extend The vector whose components are one is denoted by $\mathbf{1}$ (the same symbol is used independently of the dimension of the vector – the dimension should be clear from the context).

The utility function of the principal is given by

$$u_0(\pi, t) = \left(1 - \sum_{i \in I} \langle \pi_i, \mathbf{1} \rangle\right) \lambda_0 + \sum_{(i,j) \in K} \pi_{i,j} \rho_{i,j} - \sum_{i \in I} t_i, \quad (6.1)$$

where $\lambda_0 \in \mathbb{R}$ is the profit that the principal makes when no solution is accepted, and $\rho_{i,j}$ is the profit made when item (i, j) offered by agent i is acquired. Note that as usual with randomized mechanisms, the utility function determines the expected utility where the expectation is over the randomization of the mechanism.

We will represent the type of developer (or agent) i by a vector $\theta_i = (\theta_{i,j})_{j \in J_i} \in \Theta_i = \times_{i \in I} \Theta_{i,j}$, where $\Theta_{i,j} = [\underline{\theta}_{i,j}, \bar{\theta}_{i,j}] \subset \mathbb{R}$, $\underline{\theta}_{i,j} < \bar{\theta}_{i,j}$. Thus, Θ_i is a box in $\mathbb{R}^{|J_i|}$. The utility function of developer i will be

$$u_{i,\theta_i}(\pi, t) = -\langle \pi_i, \theta_i \rangle + t_i, \quad (\pi, t) \in O \quad (6.2)$$

i.e., the agent is giving up the value $\theta_{i,i}$ when the principal buys item $(i, j) \in K$, while he receives the amount τ_i independently of whether he won the auction. Again, the utility function determines the expected utility for the agent where the expectation is over the randomization of the mechanism.

To fully specify the optimal mechanism design problem it remains to choose some counter-strategy map schema, $\mathfrak{C}$. In this chapter, we consider the case when $\mathfrak{C}$ is the Nash-choice: $C_{-i}(s) = C_{-i}^N(s) = \{s_{-i}\}$, giving rise to implementations in Bayes-Nash equilibria.

6.2 The Form of the Optimal Multi-Unit Auction

In what follows we will reuse the symbol $P_{\theta_{i,j}}$ to denote the cumulative distribution function $P_{\theta_{i,j}} : \mathbb{R} \rightarrow [0, 1]$ corresponding to the distribution $P_{\theta_{i,j}}$:

$P_{\theta_{i,j}}(x) = \int_{-\theta_{i,j}}^x dP_{\theta_{i,j}}$, $x \in \mathbb{R}$ (the meaning of $P_{\theta_{i,j}}$ should remain clear from the context). We make the following assumption concerning these distributions:

Assumption 6.2.1. For each $(i, j) \in K$, the distribution $P_{\theta_{i,j}}$ has a density with respect to the Lebesgue measure on $\Theta_{i,j}$. Further, the density is bounded away from zero on $\Theta_{i,j}$.¹

We will denote the resulting density by $p_{\theta_{i,j}}$. We also introduce $P_{\theta_i}(\theta_i) = \prod_{j \in J_i} P_{\theta_{i,j}}(\theta_{i,j})$ and $p_{\theta_i}(\theta_i) = \prod_{j \in J_i} p_{\theta_{i,j}}(\theta_{i,j})$.

WLOG we seek the optimal mechanism amongst the set of direct-revelation mechanisms that are truthful (that this can be done holds because of the revelation principle, cf. Theorem 2.3.1). Thus, $\Sigma = \Theta$, i.e., in our case the agents' messages will be prices. For $i \in I$, introduce the function $V_{i,j} : \Theta_i \rightarrow \mathbb{R}$ defined by

$$V_{i,j}(x) = \rho_{i,j} - x - \frac{P_{\theta_i}(x)}{p_{\theta_i}(x)}. \quad (6.3)$$

The function $V_{i,j}$ assigns a “virtual value” to a price-vector x submitted by agent i for solution $j \in J_i$: The function compares the agent's offer to the profit to be made if the item is accepted and is adjusted by $\frac{P_{\theta_i}(x)}{p_{\theta_i}(x)}$ that reflects the uncertainty regarding the type θ_i of agent i , i.e., the information rent that decreases the profit that can be extracted from the given agent.

These functions form the basis of the solution to the reverse auction problem. In particular, the solution will take the form $M = (\Theta, g^*)$, $g^* = (\pi^*, t^*)$ with $\pi^* : \Theta \rightarrow M_{\leq 1}(K)$, $t^* : \Theta \rightarrow \mathbb{R}$ are specified as follows: For every “table” of submitted prices θ , the mechanism will select the item to be bought, with the possibility that no item will be selected. We let $w^* : \Theta \rightarrow K \cup \{0\}$ denote the function that determines the item bought: the value $0 \notin K$ is used to allow the mechanism to reject all offers. To make the definition of w^* more concise define

$$V_{0,0}(x) = x$$

and define $\theta_0 = \lambda_0$ so that $V_{0,0}(\theta_0)$ becomes the profit of the principal when

¹This is the reason $\Theta_{i,j}$ has to be a *bounded* interval.

she decides not to buy the item. Then, for $\theta \in \Theta$,²

$$w^*(\theta) = \arg \max_{(i,j) \in K \cup \{(0,0)\}} V_{i,j}(\theta_i), \quad (6.4)$$

where ties should be broken in an arbitrary, but systematic fashion independently of x (i.e., by ordering $I \cup \{0\}$ in some way and in the case of ties choosing the index that precedes all the other tied indices in the chosen ordering). According to (6.4), the item to be bought is selected as the one whose virtual valuation at the submitted price is the largest (with item $(0,0)$ representing no trade). Now, define π^* by

$$\pi_i^*(\theta) = \mathbb{I}_{\{w_1^*(\theta)=i\}}, \quad i \in I, \quad (6.5)$$

where $w_1^*(\theta)$ denotes the first component of $w^*(\theta)$. To define the payment function t^* , first define the functions $z_i^* : \Theta_{-i} \rightarrow \mathbb{R}$, $i \in I$:

$$z_i^*(\theta_{-i}) = \sup \left\{ \theta_{w_2^*(\theta)} : \theta_i \in \Theta_i, w_1^*(\theta) = i \right\}.$$

That is, $z_i^*(\theta_{-i})$ specifies the largest price agent i can submit and still win given that the other agents submit the prices θ_{-i} (and regardless of which item of him makes him win). With this, define

$$t_i^*(\theta) = \begin{cases} z_i^*(\theta_{-i}), & \text{if } w_1^*(\theta) = i; \\ 0, & \text{otherwise.} \end{cases} \quad (6.6)$$

Note that agent i gets paid if and only if he wins. When the agent wins, he gets paid $z_i^*(\theta_{-i})$, which is guaranteed to be more than $\theta_{i,w_2^*(\theta)}$, the price of the item bought, otherwise he would not have won.

Let us introduce one more technical assumption:

Assumption 6.2.2. The virtual valuation functions, $V_{i,j}$ are strictly decreasing.

Note that for some common probability distributions, such as the uniform or exponential distributions, the virtual valuation functions are indeed strictly decreasing. Again, if the assumption is not met, the “ironing” technique of Myerson (1981) can be used.

We can now state the main result of this section:

²Note that θ does *not* include θ_0

Theorem 6.2.1. *Let Assumptions 6.2.1 and 6.2.2 hold and let $g^* = (\pi^*, t^*)$, where the functions (π^*, t^*) are defined above. Then, the mechanism (Θ, g^*) is a solution to the BOMD problem of Section 3.1.*

The proof is presented in the next section.

6.3 Proof of Theorem 6.2.1

The proof follows closely that of Theorem 3.2.1. For a given pair of functions $\pi : \Theta \rightarrow M_{\leq 1}(K)$, $t : \Theta \rightarrow \mathbb{R}^{|I|}$ let $g_{\pi,t} : \Theta \rightarrow O$ be defined by $g_{\pi,t}(\theta) = (\pi(\theta), t(\theta))$. Further, for $i \in I$, $\theta_i, \theta'_i \in \Theta_i$, let

$$\begin{aligned} u_i(\pi, t, \theta'_i, \theta_i) &= u_{i,\theta_i}(g_{\pi,t}(\theta'_i, P_{\theta_{-i}})), \\ U_i(\pi, t, \theta_i) &= u_i(\pi, t, \theta_i, \theta_i). \end{aligned}$$

Thus, $u_i(\pi, t, \theta'_i, \theta_i)$ is the interim expected utility of agent i when he chooses to send θ'_i while his type is θ_i and $U_i(\pi, t, \theta_i)$ is his interim expected utility when he chooses to be truthful.

By Proposition 2.3.3 and also using that by our choice of Bayes-Nash implementation the IC constraint (2.9b) is equivalent to (2.11), our problem is equivalent to the following functional optimization problem:

$$\int u_0(g_{\pi,t}(\theta)) dP_\theta \rightarrow \max \text{ s.t.} \tag{OPT-M1}$$

$$\pi : \Theta \rightarrow M_{\leq 1}(I), t : \Theta \rightarrow \mathbb{R}^{|I|},$$

$$U_i(\pi, t, \theta_i) \geq u_i(\pi, t, \theta'_i, \theta_i) \text{ for all } i \in I \text{ and } \theta_i, \theta'_i \in \Theta_i \tag{IC-M1}$$

$$U_i(\pi, t, \theta_i) \geq 0 \text{ for all } i \in I \text{ and } \theta_i \in \Theta_i. \tag{IR-M1}$$

For $i \in I$, $\pi : \Theta \rightarrow M_{\leq 1}(K)$, $\theta_i \in \Theta_i$ define

$$E_i(\pi, \theta_i) = \pi_i(\theta_i, P_{\theta_{-i}}).$$

Note that π_i and E_i are both vector-valued functions. With the above definition we can write

$$u_i(\pi, t, \theta'_i, \theta_i) = t(\theta'_i, P_{\theta_{-i}}) - \langle E_i(\pi, \theta'_i), \theta_i \rangle.$$

In this section we call a vector valued function $v : \mathbb{R}^k \rightarrow \mathbb{R}^p$ decreasing if $v(x) \leq v(y)$ holds whenever $x \geq y$. Here, and in what follows, when using $\leq$ (or $<$, $\geq$, or $>$) we will mean the operator that compares the vectors in a componentwise manner. For example, for $x, y \in \mathbb{R}^d$, $x \leq y$ if $x_p \leq y_p$ holds for all $1 \leq p \leq d$.

We claim that π, t satisfies (IC-M1),(IR-M1) if and only if it satisfies the following constraints:

$$E_i(\pi, \cdot) \text{ is decreasing for all } i \in I, \quad (\text{DEC-M2})$$

$$U_i(\pi, t, \theta_i) = U_i(\pi, t, \bar{\theta}_i) + \int_{\bar{\theta}_i}^{\theta_i} \langle E_i(\pi, \hat{\theta}_i), d\hat{\theta}_i \rangle \text{ for all } i \in I, \theta_i \in \Theta_i, \quad (\text{INT-M2})$$

$$U_i(\pi, t, \bar{\theta}_i) \geq 0 \text{ for all } i \in I. \quad (\text{IR-M2})$$

The proof of this equivalence actually holds for each index $i \in I$, separately and follows immediately from the following analysis lemma:

Lemma 6.3.1 (Vector Envelope Theorem). *Let $X = \times_{i=1}^d [a_i, b_i]$ be a closed box in $\mathbb{R}^d$, $t : X \rightarrow \mathbb{R}$, $e : X \rightarrow [0, \infty)^d$ be integrable. For $x, y \in X$, define $u(x, y) = t(x) - \langle e(x), y \rangle$, $U(x) = u(x, x)$. Then the inequalities*

$$U(y) \geq u(x, y) \text{ for all } x, y \in X, \quad (\text{IC-VE})$$

$$U(x) \geq 0 \text{ for all } x \in X \quad (\text{IR-VE})$$

are satisfied if and only if the constraints

$$e \text{ is decreasing}, \quad (\text{DEC-VE})$$

$$U(x) = U(b) + \int_x^b \langle e(x), dx \rangle \text{ for all } x \in X, \quad (\text{INT-VENV})$$

$$U(b) \geq 0 \quad (\text{IR2-VE})$$

are satisfied.

Proof. We first show that (IC-VE) is equivalent to

$$U(y) - U(x) \geq \langle x - y, e(x) \rangle \text{ for all } x, y \in X. \quad (\text{IC2-VE})$$

Indeed, using the definition of u , we see that

$$U(y) \geq u(x, y) = t(x) - \langle e(x), y \rangle = U(x) + \langle e(x), x - y \rangle,$$

and reordering the terms gives the required equivalence. Also, note that (IC2-VE) clearly implies that U is decreasing thanks to $e \geq 0$.

$\Rightarrow$: Clearly, (IR2-VE) is implied by (IR-VE). Swapping x and y in (IC2-VE) gives $U(x) - U(y) \geq \langle y - x, e(y) \rangle$. Combining this with (IC2-VE) we get

$$\langle x - y, e(x) \rangle \leq U(y) - U(x) \leq \langle x - y, e(y) \rangle \text{ for all } x, y \in X. \quad (6.11)$$

This implies that e is decreasing (i.e., (DEC-VE)). The plan now is to apply Theorem 3.3.2 to $f(x, t) = u(x, \tau(t))$, where $\tau : [0, 1] \rightarrow X$ is a smooth function such that $\tau(0) = x_0$ with some $x_0 \in X$, $\tau(1) = b$ to show (INT-ENV). With the notation of the theorem, $V(t) = \max_{x \in X} f(x, t) = u(\tau(t), \tau(t))$. Choose $x^*(t) = \tau(t)$. The conditions of the theorem can be readily verified. Further, $f_t(x, t) = \frac{\partial}{\partial t}(t(x) - \langle e(x), \tau(t) \rangle) = -\langle e(x), \tau'(t) \rangle$. Hence, $U(b) = V(1) = V(0) + \int_0^1 f_t(x^*(t), t) dt = U(x_0) - \int_0^1 \langle e(\tau(t)), \tau'(t) \rangle dt$. Here, $\int_0^1 \langle e(\tau(t)), \tau'(t) \rangle dt$ is the path integral of $e(x)$ from $\tau(0) = x_0$ to $\tau(1) = b$. Hence, $U(x_0) = U(b) + \int_{x_0}^b \langle e(s), ds \rangle$. Since x_0 was arbitrary, we get (INT-ENV). This finishes the direction $\Rightarrow$.

$\Leftarrow$: Since e is decreasing, starting from (INT-ENV) we get that

$$U(x) = U(y) - \int_y^x \langle e(z), dz \rangle \geq U(y) - \int_y^x \langle e(y), dz \rangle = U(y) - \langle x - y, e(y) \rangle$$

holds for any $x, y \in X$. This implies (IC2-VE), which was seen to be equivalent to (IC-VE) and to imply that U is decreasing. Since U is decreasing, (IR2-VE) implies (IR-VE). $\square$

Let us now return to the optimization problem. Using the function U_i , we can rewrite the objective function as

$$u_0(g_{\pi, t}(P_\theta)) = \lambda_0 + \sum_{i \in I} \langle \rho_i - \theta_i - \lambda_0 \mathbf{1}, \pi_i(P_\theta) \rangle - \sum_{i \in I} U_i(\pi, t, P_{\theta_i}), \quad (6.12)$$

where we use $\rho_i = (\rho_{ij})_{j \in J_i}$. Let us now write $U_i(\pi, t, P_\theta)$ in a form that allows the separation of the terms that involve t . Take any π, t satisfying the constraints (DEC-M2), (INT-M2), (IR-M2). Due to (IR-M2) and the

definition of E_i and Assumption 6.2.1,

$$\begin{aligned}
U_i(\pi, t, P_{\theta_i}) &= \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \langle E_i(\pi, \theta'_i), d\theta'_i \rangle \right) dP_{\theta_i}(\theta_i) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \int_{\underline{\theta}_i}^{\bar{\theta}_i} \langle E_i(\pi, \theta'_i), d\theta'_i \rangle dP_{\theta_i}(\theta_i) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(\int_{\underline{\theta}_i}^{\theta'_i} dP_{\theta_i}(\theta_i) \right) \langle E_i(\pi, \theta'_i), d\theta'_i \rangle \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\underline{\theta}_i}^{\bar{\theta}_i} \left(P_{\theta_i}(\theta'_i) \int_{\Theta_{-i}} \langle \pi_i(\theta'_i, \theta_{-i}), \mathbf{1} \rangle dP_{\theta_{-i}}(\theta_{-i}) \right) d\theta'_i \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\Theta} P_{\theta_i}(\theta_i) \langle \pi_i(\theta), \mathbf{1} \rangle dP_{\theta_{-i}}(\theta_{-i}) \frac{p_{\theta_i}(\theta_i)}{p_{\theta_i}(\theta_i)} d\theta_i \quad (*) \\
&= U_i(\pi, t, \bar{\theta}_i) + \int_{\Theta} \frac{P_{\theta_i}(\theta_i)}{p_{\theta_i}(\theta_i)} \langle \pi_i(\theta), dP_{\theta}(\theta) \rangle.
\end{aligned}$$

Note that we have indeed separated the term that includes t . The equation where we used the positivity of p_{θ_i} over its domain is denoted by (*). Plugging the expression obtained for $U_i(\pi, t, P_{\theta_i})$ into (6.12) and using the functions

$$\hat{V}_{i,j}(x) = \rho_{i,j} - x_j - \lambda_0 - \frac{P_{\theta_i}(x)}{p_{\theta_i}(x)} \quad (x \in \Theta_i)$$

and for i fixed collecting these functions into the function $\hat{V}_i : \Theta_i \rightarrow \mathbb{R}^{|J_i|}$ we get

$$u_0(g_{\pi,t}(P_{\theta})) = \lambda_0 - \sum_{i \in I} U_i(\pi, t, \bar{\theta}_i) + \int \sum_{i \in I} \langle \hat{V}_i(\theta_i), \pi_i(\theta) \rangle dP_{\theta}(\theta). \quad (6.13)$$

For π fixed, let us maximize this in t subject to the constraints (DEC-M2), (INT-M2), (IR-M2). Since only the second term depends on t and in fact this term has a negative sign, we maximize the objective if we minimize $\sum_{i \in I} U_i(\pi, t, \bar{\theta}_i)$. Let us consider the i th term of this for some fixed index i . By (INT-M2) and plugging in the definitions of U_i and E_i , for any $\theta_i \in \Theta_i$ we get

$$\begin{aligned}
U_i(\pi, t, \bar{\theta}_i) &= U_i(\pi, t, \theta_i) - \int_{\theta_i}^{\bar{\theta}_i} \langle E_i(\pi, \theta'_i), d\theta'_i \rangle \\
&= t_i(\theta_i, P_{\theta_{-i}}) - \langle \pi_i(\theta_i, P_{\theta_{-i}}), \theta_i \rangle - \int_{\theta_i}^{\bar{\theta}_i} \langle \pi_i(\theta'_i, P_{\theta_{-i}}), d\theta'_i \rangle \\
&= \left[t_i(\theta) - \langle \pi_i(\theta), \theta_i \rangle - \int_{\theta_i}^{\bar{\theta}_i} \langle \pi_i(\theta'_i, \theta_{-i}), d\theta'_i \rangle \right]_{\theta_{-i} \leftarrow P_{\theta_{-i}}},
\end{aligned}$$

where $[\cdot]_{\theta_{-i} \leftarrow P_{\theta_{-i}}}$ is used to denote the substitution of θ_{-i} by $P_{\theta_{-i}}$ (and hence, taking the integral of the expression). Thus, $U_i(\pi, t, \bar{\theta}_i)$ depends on t only through t_i . By (IR-M2), all feasible pairs (π, t) must satisfy $U_i(\pi, t, \bar{\theta}_i) \geq 0$. Hence, the minimum of $U_i(\pi, t, \bar{\theta}_i)$ is zero. This minimum is achieved if we choose

$$t_i^{(\pi)}(\theta) = \langle \pi_i(\theta), \theta_i \rangle + \int_{\theta_i}^{\bar{\theta}_i} \langle \pi_i(\theta'_i, \theta_{-i}), d\theta'_i \rangle$$

and by choosing t this way (as a function of π), (INT-M2), (IR-M2) are satisfied for any π . Thus, it remains to choose π .

When we choose $t = t^{(\pi)}$, we see that the only term that still depends on π in (6.13) is the last term. Call this term $\Upsilon(\pi, P_\theta) = \int \sum_{i \in I} \langle \hat{V}_i(\theta_i), \pi_i(\theta) \rangle dP_\theta(\theta)$. Taking into account that $\pi(\theta)$ is a subprobability distribution, we see that for any feasible π

$$\Upsilon(\pi, P_\theta) \leq \left[\max(0, \max_{(i,j) \in K} \hat{V}_{i,j}(\theta_i)) \right]_{\theta \leftarrow P_\theta}$$

(if $\max_{(i,j) \in K} \hat{V}_{i,j}(\theta_i) < 0$ $\pi_i(\theta) = 0$, $i \in I$ achieves zero inside the integral at θ).

Further, the upper bound on $\Upsilon(\pi, P_\theta)$ can be achieved by any π when $\pi(\theta)$ assigns zero to all indices $(i, j) \in K$ such that $\hat{V}_{i,j}(\theta_i) < 0$ and assigns nonnegative values to indices in $W(\theta) = \left\{ (i, j) \in K : \hat{V}_{i,j}(\theta_i) \geq 0, \hat{V}_{i,j}(\theta_i) = \max_{(i',j') \in K} \hat{V}_{i',j'}(\theta_{j'}) \right\}$.

Now, if Assumption 6.2.2 is satisfied then it can be shown that by choosing a single nonzero entry from $W(\theta)$ will result in π that satisfies (DEC-M2). Denoting the resulting choice π^* and letting $t^* = t^{(\pi^*)}$, after elementary transformation we arrive at the desired statement, thus finishing the proof of Theorem 3.2.1.

Note. Although the optimality of mechanism is still intact, as the number of agents vary, the ex-post expected utility of the principal will change. For instance, as the number of participating agents tend to infinity, there will be an agent of the highest type with high probability and also an agent of the second highest type, thus the ex-post expected utility of the principal will be at the maximum possible.

6.4 Summary

In this chapter, we extended the results of Chapter 3 and Chapter 4 to the case when agents have multiple machine learning packages to offer. This is a useful result that captures the real world situation more closely. The conclusions here are similar to Chapter 3 and Chapter 4 but with this added power.

6.4.1 Future Work

In reality, the cost of evaluating solutions is *not* negligible, but we assumed so here to be able to continue the analysis. Now, if we consider a price for evaluating solutions, which may be due to both practical reasons and learning-theoretic considerations (the more solutions we try, the closer we get to having an overfit solution), perhaps by addition of a penalty for number of solutions, the impact on voluntary participation will be significant and we will face a very different problem that needs to be solved for itself.

Chapter 7

A Comprehensive Procurement Process

In this chapter we bring an example of a sample comprehensive procurement process to put the results of the thesis in context.

We used reduction in proving the main results of Chapter 4 and Chapter 5. This proves to be a powerful tool because now, we can combine the results throughout this thesis to introduce a comprehensive solution to the problem of machine learning procurement.

7.1 Solution to the Example of Section 1.2

Consider the problem faced by Aleph Corp. introduced in Section 1.2. To recap Aleph Corp. supplies books and it has gathered a database of different users' interest in books over the years. Aleph Corp. wants to obtain a book recommendation system so to increase its sales and profit. Because of the reasons mentioned in Section 1.2, Aleph Corp. decides to go for a procurement mechanism: Aleph Corp. calls for solutions, while it does not want to publicly announce its profit information. Moreover, Aleph Corp. does not want to limit the number of packages each developer has to offer to one because it might have a negative impact on its profit and will publicly declare how good each agent did, but not her profit information.

Now, Aleph Corp. can combine these solution to find an optimal procurement procedure. Then, by results of Chapter 4, Chapter 5 and Chapter 6, the

timeline for an optimal procurement procedure would be as follows:

1. Aleph Corp. announces the rules, which are as follows:
2. Developers submit all their solutions. They may offer multiple solutions;
3. Aleph Corp. evaluates the offered solutions on the dataset;
4. Aleph Corp. publicly announces the results of this evaluation;
5. Developers submit their bids for each of their solutions (they will not be different for different cases of solution performance because the optimal action for them is honesty);
6. Aleph Corp. computes the profitability of each solution; and
7. Aleph Corp. computes the virtual valuations of all solutions with (6.3). The solution with the highest virtual valuation wins (6.5). The developer that offered that solution will be compensated an amount equal to the highest bid that would have still allowed it to win according to (6.6).

In this way, Aleph Corp. has a solution that is maximizing the profits. One thing that needs to be mentioned is that Aleph Corp. needs to effectively communicate the inherent incentive-compatible structure of the mechanism to the agents. In words the Aleph Corp. will say that if the developers underbid, they may increase their chances of winning but in doing so, they will lose in expectation. Also, if they overbid, they will decrease their chances of winning and will not increase their payments in case they win, so they will lose in expectation. Thus, Aleph Corp. will show that honesty is going to be the most profitable course of action by the developers. Aleph Corp. will use the results of this thesis to communicate this in a more rigorous manner, if need be.

Chapter 8

Conclusions

In this thesis, we introduced additions to the framework of mechanism design in order to make it ready to be appropriate for finding principled ways to do procure machine learning solutions in an optimal fashion. We used reverse auctions as a base model to deal with this situation. Much work needs to be done and in fact the results in this thesis are preliminary. However, the hope is that this thesis calls the game theoretical aspects of machine learning problems to the attention of machine learning researchers. One immediate next step for future research is to address the problem of how the data held by the principal should be split to avoid overfitting. The issue here is that the principal wants to release data to give the developers a chance to build good solutions. An obvious approach is to split the data equally. However, it is far from clear whether this is the best the principal can do to achieve a good utility with high probability (or on average).

8.1 Future Work

In addition to all the issues stated in ‘Future Work’ sections of different chapters which were not addressed in later chapters, these problems seem to be good next steps:

- A logical next step to this thesis is to find and formalize a way to accept and combine multiple solutions. An important consideration, then, would be considerations of aggregation versus model selection. This will also correlate to the problems of crowdsourcing for the development of

machine learning solutions. Recently, Abernethy and Frongillo (2011) proposed a mechanism based on collaboration for solving such problems.

- Multi-stage auction also seem a logical next step. In fact, we did try to extend the framework to a multi-stage auction, but the results were too immature to be presented.
- Another problem we face in procurement of machine learning packages is that the principal may be faced by an unknown number of agents. We did some work on formalizing this problem but unfortunately, these results were not fit to be presented in this thesis in their current form.

Bibliography

- Jacob Abernethy and Rafael M. Frongillo. A Collaborative Mechanism for Crowdsourcing Prediction Problems. In *Advances in Neural Information Processing Systems 24*, pages 1–9, 2011.
- Kenneth J Arrow and Gérard Debreu. Existence of a Competitive Equilibrium for a Competitive Economy. *Econometrica*, 22(3):265–290, 1954.
- Patrick Bajari and Steven Tadelis. Incentives versus Transaction Costs: A Theory of Procurement Contracts. *The RAND Journal of Economics*, 32(3):387, January 2001. ISSN 07416261. doi: 10.2307/2696361.
- Patrick Bolton and Mathias Dewatripont. *Contract Theory*. The MIT Press, 2005.
- Yeon-Koo Che. Design Competition Through Multidimensional Auctions. *The RAND Journal of Economics*, 24(4):668–680, January 1993. ISSN 07416261. doi: 10.2307/2555752.
- Fangruo Chen. Auctioning Supply Contracts. *Management Science*, 53(10):1562–1576, October 2007. ISSN 0025-1909. doi: 10.1287/mnsc.1070.0716.
- Edward H. Clarke. Multipart Pricing of Public Goods. *Public Choice*, 11(1):17–33, 1971.
- Partha S. Dasgupta, Peter J. Hammond, and Eric S. Maskin. The Implementation of Social Choice on Results Rules : Some General Results on Incentive Compatibility. *Review of Economic Studies*, 46(2):185–216, 1979.
- Sudipto Dasgupta and Daniel F. Spulber. Managing Procurement Auctions. *Information Economics and Policy*, 4(1989190):5–29, 1990.
- Drew Fudenberg and Jean Tirole. *Game Theory*. The MIT Press, 1991. ISBN 0262061414.
- Allan Gibbard. Manipulation of Voting Schemes: A General Result. *Econometrica*, 41(4):587–601, 1973.
- Theodore Groves. Incentives in Teams. *Econometrica*, 41(4):617–631, 1973.
- John Charles Harsanyi. Games with Incomplete Information Played by “Bayesian” Players, I–III: Part I. The Basic Model. *Management science*, 14(3):159–183, 1967.
- John Charles Harsanyi. Games with Incomplete Information Played by “Bayesian” Players, I–III: Part II. Bayesian Equilibrium Points. *Management Science*, 14(5):320–334, 1968a.

- John Charles Harsanyi. Games with Incomplete Information Played by “Bayesian” Players, I-III: Part III. The basic Probability Distribution of the Game. *Management Science*, 14(7):486–502, 1968b.
- Bengt R. Holmström. *On Incentives and Control in Organizations*. PhD thesis, Stanford University, 1977.
- Leonid Hurwicz. On the Concept and Possibility of Informational Decentralization. *The American Economic Review*, 59(2):513–524, 1969.
- Leonid Hurwicz. The Design of Mechanisms for Resource Allocation. *The American Economic Review*, 63(2):1–30, 1973.
- Matthew O. Jackson. A Crash Course in Implementation Theory. *Social Choice and Welfare*, 18(4):655–708, October 2001. ISSN 0176-1714. doi: 10.1007/s003550100152.
- David M. Kreps and Robert Wilson. Sequential Equilibria. *Econometrica*, 50(4):863–894, 1982.
- Jean-Jacques Laffont and David Martimort. *The Theory of Incentives: The Principal-Agent Model*. Princeton University Press, 2002.
- Jean-Jacques Laffont and Jean Tirole. *A Theory of Incentives in Procurement and Regulation*. The MIT Press, Cambridge, MA, 1993.
- Jean-Jacques Laffont, Eric S. Maskin, and Jean-Charles Rochet. Optimal Non-linear Pricing with Two-Dimensional Characteristics. In Theodore Groves, Roy Radner, and Stanley Reiter, editors, *Information, Incentives and Economic Mechanisms; Essays in Honor of Leonid Hurwicz*, chapter 8, pages 256–266. University of Minnesota Press, Minneapolis, 1987.
- Kevin Leyton-Brown. *Essentials of Game Theory: A Concise, Multidisciplinary Introduction*. Morgan and Claypool Publishers, 2008. ISBN 1598295934.
- Harry M. Markowitz. *Portfolio Selection: Efficient Diversification of Investments*. John Wiley & Sons, New York, 1959.
- Andreu Mas-Colell, Michael D. Whinston, and Jerry R. Green. *Microeconomic Theory*. Oxford University Press, 1995.
- Eric S. Maskin and Tomas Sjöström. Implementation Theory. In *Handbook of Social Choice and Welfare*, volume 1, pages 237–288. Elsevier B.V., 2002.
- Jean-François Mertens. Stable Equilibria: A Reformulation Part I. Definition and Basic Properties. *Mathematics of Operations Research*, 14(4):575–625, 1989.
- Paul R. Milgrom and Ilya Segal. Envelope Theorems for Arbitrary Choice Sets. *Econometrica*, 70(2):583–601, 2002.
- Paul R. Milgrom and Robert J. Weber. A Theory of Auctions and Competitive Bidding. *Econometrica*, 50(5):1089–1122, 1982.
- Roger B. Myerson. Incentive-Compatibility and the Bargaining Problem. *Econometrica*, 47(1):61–73, 1979.

- Roger B. Myerson. Optimal Auction Design. *Mathematics of Operations Research*, 5(February):58–73, 1981.
- Roger B. Myerson. Mechanism Design by an Informed Principal Author. *Econometrica*, 51(6):1767–1797, 1983.
- Roger B. Myerson. Axiomatic Foundations of Bayesian Decision Theory. 1986. URL <http://ideas.repec.org/p/nwu/cmsems/671.html>.
- Roger B. Myerson. *Game Theory: Analysis of Conflict*. Harvard University Press, 1997. ISBN 0674341163.
- Roger B. Myerson. Mechanism Design. In Steven N. Durlauf and Lawrence E. Blume, editors, *The New Palgrave Dictionary of Economics*. Palgrave Macmillan, second edition, 2008.
- Roger B Myerson and Mark Allen Satterthwaite. Efficient Mechanisms for Bilateral Trading. *Journal of Economic Theory*, 29(2):265–281, April 1983. ISSN 00220531. doi: 10.1016/0022-0531(83)90048-0.
- Y. Narahari. Game Theory: Lecture Notes. 2012. URL <http://lcm.csa.iisc.ernet.in/gametheory/ln/web-md10-myerson.pdf>.
- John Nash. Non-Cooperative Games. *Annals of Mathematics*, 54(2):286–295, 1951.
- Noam Nisan, Tim Roughgarden, Eva Tardos, and Vijay V. Vazirani. *Algorithmic Game Theory*. Cambridge University Press, 2007. ISBN 0521872820.
- Motty Perry and Philip J. Reny. On The Failure of The Linkage Principle in Multi-Unit Auctions. *Econometrica*, 67(4):895–900, 1999.
- Jean-Charles Rochet and Lars A. Stole. The Economics of Multidimensional Screening. In Mathias Dewatripont, Lars Peter Hansen, and Stephen J. Turnovsky, editors, *Advances in Economics and Econometrics: Theory and Applications - Eight World Congress*, chapter 5, pages 150–197. Cambridge University Press, 2003.
- Mark Allen Satterthwaite. Strategy-Proofness and Arrow’s Conditions: Existence and Correspondence Theorems for Voting Procedures and Social Welfare Functions. *Journal of Economic Theory*, 10:187–217, 1975.
- Yoav Shoham and Kevin Leyton-Brown. *Multiagent Systems: Algorithmic, Game-Theoretic, and Logical Foundations*. Cambridge University Press, 2008. ISBN 0521899435.
- William Vickrey. Counterspeculation, Auctions, and Competitive Sealed Tenders. *Journal of Finance*, 16(1):8–37, 1961.
- John von Neumann and Oskar Morgenstern. *Theory of Games and Economic Behavior*. Princeton University Press, 1947. ISBN 0691130612.