

Coalition Formation under Uncertainty: Bargaining Equilibria and the Bayesian Core Stability Concept

Georgios Chalkiadakis^{*}
Dept. of Computer Science
University of Toronto
Toronto, ON, Canada
gc2@ecs.soton.ac.uk

Evangelos Markakis[†]
Center for Mathematics and
Computer Science (CWI)
Amsterdam, The Netherlands
vangelis@cwi.nl

Craig Boutilier
Dept. of Computer Science
University of Toronto
Toronto, ON, Canada
cebly@cs.toronto.edu

ABSTRACT

Coalition formation is a problem of great interest in AI, allowing groups of autonomous, rational agents to form stable teams. Furthermore, the study of coalitional stability concepts and their relation to equilibria that guide the strategic interactions of agents during bargaining has lately attracted much attention. However, research to date in both AI and economics has largely ignored the potential presence of uncertainty when studying either coalitional stability or coalitional bargaining. This paper is the first to relate a (cooperative) stability concept under uncertainty, the *Bayesian core (BC)*, with (non-cooperative) equilibrium concepts of coalitional bargaining games. We prove that if the BC of a coalitional game (and of each subgame) is non-empty, then there exists an equilibrium of the corresponding bargaining game that produces a BC element; and conversely, if there exists a coalitional bargaining equilibrium (with certain properties), then it induces a BC configuration. We thus provide a *non-cooperative justification of the BC stability concept*. As a corollary, we establish a sufficient condition for the existence of the BC. Finally, for small games, we provide an algorithm to decide whether the BC is non-empty.

Categories and Subject Descriptors

I.2 [Artificial Intelligence]: Distributed Artificial Intelligence

General Terms

Economics

Keywords

Coalition formation, multilateral bargaining

^{*}Currently at the School of Electronics and Computer Science, University of Southampton.

[†]This work was done while the second author was a postdoctoral fellow at the Department of Computer Science, University of Toronto.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee.

AAMAS'07 May 14–18 2007, Honolulu, Hawai'i, USA.
Copyright 2007 IFAAMAS.

1. INTRODUCTION

Coalition formation, widely studied in game theory and economics [11], has attracted much attention in AI as means of dynamically forming partnerships or teams of cooperating agents. Most models of coalition formation assume that the values of potential coalitions are known with certainty, implying that agents possess knowledge of the capabilities of their potential partners, or at least that this knowledge can be reached via communication (e.g., see [19, 20]). However, in many natural settings, rational agents must employ *bargaining* in order to form coalitions and divide the generated value without knowing a priori what this value may be or how suitable their potential partners are for the task at hand. For instance, an enterprise must often choose subcontractors without full certainty of their capabilities. The creation of *virtual organizations* has been anticipated as an important consequence impact of agent coalition technologies on e-commerce [12]; this cannot be achieved without addressing this type of uncertainty.

Cooperative game theory deals with the coalition formation problem, focusing on the question of stability of formed coalitions. The presence of uncertainty regarding agent abilities poses interesting theoretical questions, such as the discovery of analogs of the traditional concepts of stability. Furthermore, it suggests opportunities for agents to learn about each others' abilities through repeated interaction, refining how coalitions are formed over time.¹ As a consequence, realistic models of coalition formation must be able to deal with situations in which the presence of uncertainty regarding the effects of coalitional actions and the *types* (i.e., the capabilities) of the potential partners is translated into uncertainty about the values of various coalitions [3, 4].

To this end, we propose a model of coalition formation in which agents must derive coalitional values by reasoning about the types of other agents and the uncertainty inherent in the actions a coalition may take. The concept of the *Bayesian core (BC)* was proposed by [3] as a suitable *cooperative* solution concept in this setting. In this paper, we elaborate on the definition of the BC concept—providing three versions of it (the *weak*, the *strict* and the *strong* BC). Furthermore, for small games, we propose an algorithmic method to establish the existence of the Bayesian core (or non-existence thereof): we formulate the problem as a constraint satisfaction problem (CSP) with polynomial

¹We do not deal with any learning issues in this paper, but we refer to [3] and [4] which do describe relevant approaches.

constraints, and check whether it has a solution.

Cooperative coalition formation largely disregards the bargaining processes by which the coalitions emerge. Non-cooperative approaches, on the other hand, focus on the strategic interactions of the players, and the *equilibrium* solutions of the coalitional bargaining games that lead to the formation of coalitions. In recent years, considerable work has established connections between the outcomes arising from equilibrium play in coalitional bargaining games and the core of the underlying coalition formation problem [15, 10, 9, 8, 18, 25]. The goal of this line of research is to show that the equilibrium outcomes in particular coalitional bargaining games correspond to core allocations and more generally to prove equivalence of cooperative and non-cooperative solutions. Such results further justify the use of the core as a solution concept by contributing to the *non-cooperative justification of the core*.

The main contribution of our paper is a non-cooperative justification of the Bayesian core. We establish the first theoretical results relating a cooperative stability concept for coalition formation under *type uncertainty* with appropriate non-cooperative bargaining solutions. We prove that if the BC of a coalitional game is non-empty, then there exists an equilibrium of the corresponding bargaining game that produces a BC element, provided that the BC of each subgame is nonempty; and conversely, we show that if there exists a coalitional bargaining equilibrium with certain properties, then it leads to a BC configuration. As a corollary, we establish a sufficient condition for the non-emptiness of the BC. In the process of deriving our main results, we introduce appropriate non-cooperative games of coalitional bargaining under uncertainty, and define relevant equilibrium concepts.

2. BACKGROUND AND RELATED WORK

Cooperative game theory deals with situations where players act together in a cooperative equilibrium selection process involving some form of bargaining, negotiation, or arbitration [11]. The problem of coalition formation is the fundamental area of study within cooperative game theory.

Let $N = \{1, \dots, n\}$, $n > 2$, be a set of players (or “agents”). A subset $C \subseteq N$ is called a coalition. A *coalition structure* is a partition of the set of all agents into disjoint coalitions. Coalition formation is the process by which individual agents form such coalitions, generally to solve a problem by coordinating their efforts. The coalition formation problem can be seen as being composed of the following activities [17]: (a) the search for an optimal coalition structure; (b) the solution of a joint problem facing members of each coalition—to solve the joint problem, the agents have to agree on a *coalitional action* to perform; and (c) division of the value of the generated solution among the coalition members. These activities interact with each other, and the agents should reach agreement on all those issues through negotiations. For example, agents such as carpenters, plumbers and electricians may come together to engage in construction projects. They must form coalitions, collectively decide on a choice of coalitional action (e.g., project choice, as which style of house), and agree to a share of the payoff induced by that choice.

Coalition formation can be abstracted into a fairly simple model assuming *transferable utility*, or the existence of a (divisible) commodity (e.g., money) that players can freely transfer among themselves. The *characteristic function* of a coalitional game with transferable utility (TU-game) [11],

$v : 2^N \Rightarrow \mathbb{R}$, defines the *value* $v(C)$ of each coalition C [24]. Intuitively, $v(C)$ represents the maximal payoff the members of C can jointly receive by cooperating effectively. An *allocation* is a vector of payoffs $\vec{x} = (x_1, \dots, x_n)$ assigning some payoff to each $i \in N$. An allocation is *feasible* with respect to coalition structure CS if $\sum_{i \in C} x_i \leq v(C)$ for each $C \in CS$, and is *efficient* if this holds with equality. The *reservation value* rv_i of an agent i is the amount it can attain by acting alone (in a *singleton* coalition): $rv_i = v(\{i\})$.

A characteristic function is called *superadditive* if any pair of disjoint coalitions C and T can improve its payoff by merging into one coalition: $v(C \cup T) \geq v(C) + v(T)$ [11]. In a superadditive environment, an optimal solution can be determined by creating of the *grand coalition* containing all N players (with some division of payoff among its members). We make no superadditivity assumptions in our work.

When rational agents seek to maximize their individual payoffs, the *stability* of the underlying coalition structure becomes critical. A structure would be stable only if the outcomes attained by the coalitions and the payoff combinations agreed by the agents are such that both individual and group rationality are satisfied in some way. Research in coalition formation has developed several notions of stability, among the strongest being the *core* [11]. The core of a characteristic function game is a set of *payoff configurations* $\langle CS, \vec{x} \rangle$, where each $\vec{x}$ is a vector of payoffs to the agents in coalition structure CS , which are such that no subgroup of agents is motivated to depart from CS :

DEFINITION 1. Let CS be some coalition structure, and let $\vec{x} \in \mathbb{R}^n$ be some allocation of payoffs to the agents. The core is the set of payoff configurations

$$\{ \langle CS, \vec{x} \rangle \mid \forall C \subseteq N, \sum_{i \in C} x_i \geq v(C) \text{ and } \sum_{i \in N} x_i = \sum_{C \in CS} v(C) \}$$

A core allocation $\langle CS, \vec{x} \rangle$ is both feasible and efficient, and no subgroup of players can guarantee all of its members a higher payoff. As such, no coalition would ever “block” the proposal for a core allocation. Unfortunately, the core is exponentially hard to compute [16]. Moreover, in many cases the core is empty, as there exist games for which it is impossible to divide the utility in such a way that the coalition structure becomes stable. If the core exists, however, it is desirable to establish processes that lead to it.

Diekmann and Schwalbe [7] provide a dynamic formation process (a bargaining process that induces an underlying Markov process) that can be shown to converge to the core, if it exists. They do not, however, deal with the problem of the agents agreeing on actions for the coalitions to take—nor do they deal with the issue of establishing coalitional bargaining equilibria solutions.

Suijs *et al.* [23, 22] do not address formation processes, but do describe a notion of the core under coalition value uncertainty. Payoffs in their model are stochastic, and depend on the coalition action taken. To deal with payoff stochasticity, they use *relative* shares for the allocation of the residual of the stochastic coalitional values. It is assumed that agents have *common expectations* regarding expected coalitional values. Blankenburg *et al.* [2] propose a *kernel* stability concept under coalitional value uncertainty. More recently, Yokoo *et al.* [26] have coined two core concepts for contexts where the agents’ skills are private information reported by the agents to a special “mechanism designer” agent; but these concepts do not take into account the beliefs of the agents regarding others’ types.

As for non-cooperative approaches, Okada [13] suggests a form of coalitional bargaining where agreement can be reached in one bargaining round if the proposer is chosen *randomly*, while Chatterjee *et al.* [5] present a bargaining model with a fixed proposer order, both focusing on subgame-perfect equilibria (SPE) [14] in superadditive environments. Neither model deals with uncertainty or the selection of coalitional actions.

The work that is most relevant to ours is that of Moldovanu and Winter [10]. They show that if a strategy profile is an *order independent equilibrium (OIE)* (an SPE that remains an equilibrium and leads to the same payoff allocation for any choice of proposer in a sequential coalitional bargaining game), then the resulting payoffs must be in the core—and conversely, if the coalition formation game has subgames with nonempty cores, then for each payoff vector there exists an OIE with the same payoff. The model of [10] differs from ours: it is deterministic, does not assume random proposers, and assumes superadditive, non-transferable utility. However, we use a form of OIE in our work, generalized to incorporate an agent’s (uncertain) beliefs about the abilities (or types) of potential partners and lack of superadditivity.

Other recent work [15, 9, 8, 18, 25] has also established connections between the equilibrium outcomes of coalitional bargaining and the core in deterministic environments. In establishing similar types of results, not only do we deal with uncertainty related to agent types (or abilities) and coalitional action effects, we do so without making additivity assumptions w.r.t. coalition value. Our model is thus richer and more realistic than existing ones.

3. A BAYESIAN COALITION FORMATION MODEL

The need to address type uncertainty, reflecting an agent’s uncertainty about the abilities of potential partners, is critical to the modeling of realistic coalition formation problems. For instance, if a carpenter wants to find a plumber and electrician with whom to build a house, his decision to propose (or join) such a partnership, to engage in a specific type of project, and to accept a specific share of the surplus generated should all depend on his (probabilistic) assessment of their abilities. To capture this, we start by introducing the problem of *Bayesian coalition formation* under type uncertainty. We then show how this type uncertainty can be translated into coalitional value uncertainty, adopting the model introduced by [3]. We then elaborate on the BC concept, a version of which first appeared in [3].

DEFINITION 2. A *Bayesian coalition formation problem (BCFP)* is a coalition formation problem that is characterized by a set of agents, N ; a set of types per agent, T_i ; a set A_C of coalitional actions per coalition C ; a set $\mathcal{O}$ of stochastic outcomes (or states); a reward function $R : \mathcal{O} \rightarrow \mathbb{R}$; and agent beliefs over types of potential partners.

We now describe each of the BCFP components in turn: We assume a set of agents $N = \{1, \dots, n\}$, and for each agent i a finite set of possible types T_i . Each agent i has a specific type $t \in T_i$, which intuitively captures i ’s “abilities”. We let $T = \times_{i \in N} T_i$ denote the set of type profiles. For any coalition $C \subseteq N$, $T_C = \times_{i \in C} T_i$, and for any $i \in N$, $T_{-i} = \times_{j \neq i} T_j$. Each i knows its own type t_i , but not those of other agents. Agent i ’s beliefs B_i comprise a joint distribution

over T_{-i} , where $B_i(\vec{t}_{-i})$ is the probability i assigns to other agents having type profile $\vec{t}_{-i}$. We use $B_i(\vec{t}_C)$ to denote the marginal of B_i over any subset C of agents.

A coalition C has available to it a finite set of *coalitional actions* A_C . When an action is taken, it results in some outcome or *state* $s \in \mathcal{O}$. The odds with which an outcome is realized depends on the types of the coalition members (e.g., the outcome of building a house will depend on the capabilities of the team members). We let $\Pr(s|\alpha, \vec{t}_C)$ denote the probability of outcome s given that coalition C takes action $\alpha \in A_C$ and member types are given by $\vec{t}_C \in T_C$. Finally, we assume that each stochastic state s results in some *reward* $R(s)$, divisible/transferable among the members.

The *value* of coalition C with members of type $\vec{t}_C$ is:

$$V(C|\vec{t}_C) = \max_{\alpha \in A_C} \sum_s \Pr(s|\alpha, \vec{t}_C) R(s) = \max_{\alpha \in A_C} Q(C, \alpha|\vec{t}_C)$$

Unfortunately, this coalition value cannot be used in the coalition formation process if the agents are uncertain about the types of their potential partners. However, each i has beliefs about the value of any coalition based on its expectation of this value w.r.t. other agents’s types:

$$V_i(C) = \max_{\alpha \in A_C} \sum_{\vec{t}_C \in T_C} B_i(\vec{t}_C) Q(C, \alpha|\vec{t}_C) = \max_{\alpha \in A_C} Q_i(C, \alpha)$$

Note that $V_i(C)$ is not simply the expectation of $V(C)$ w.r.t. i ’s belief about types. The expectation Q_i of action values (i.e., Q -values) cannot be moved outside the max operator: a single action must be chosen which is useful *given* i ’s uncertainty. Of course, i ’s estimate of the value of a coalition, or any coalitional action, may not conform with those of other agents. This leads to additional complexity when defining suitable stability concepts. (We turn to this issue later in this section.) However, i is certain of its *reservation value*, the amount it can attain by acting alone: $rv_i = V_i(\{i\}) = \max_{\alpha \in A_{\{i\}}} \sum_s \Pr(s|\alpha, t_i) R(s)$.

We define a BCFP *subgame* as follows:

DEFINITION 3. Let N be a set of agents, and $S \subseteq N$. An S -agents subgame of a BCFP game with N agents, is the BCFP with S agents whose sets of types, beliefs, coalitional actions, outcomes and reward function are the restriction of their corresponding elements in the N -agents problem.

Thus, for example, an agent i in the S -agent subgame has the same beliefs regarding potential partners in S as it has in the N -agent game.

We define an analog of the traditional core concept for a BCFP. The notion of stability is made somewhat more difficult by the uncertainty associated with actions: since the payoffs associated with coalitional actions are stochastic, allocations must reflect this [23, 22]. Stability is rendered more complex still by the fact that different agents have potentially different beliefs about the types of other agents.

Because of the stochastic nature of payoffs, we assume that players join a coalition with certain *relative payoff demands* [23, 22]. Intuitively, since the agents cannot expect to have an accurate estimate of the coalition payoffs (and, consequently, the payoff shares of coalition members), it is more natural for them to take into consideration relative demands; these correspond to the perceived “power structure” within the coalition and can be used for the allocation of stochastic gains or losses.

Let $\mathbf{x}$ represent the payoff demand vector $\langle x_1, \dots, x_n \rangle$, and $\mathbf{x}_C$ the demands of those players in coalition C , assuming that these demands are observable by all agents. For any agent $i \in C$ we define the relative demand of agent to be $d_i = \frac{x_i}{\sum_{j \in C} x_j}$. If reward R is received by coalition C as a result of its choice of action, each i receives payoff $d_i R$. This means that the gains or losses due to reward stochasticity are allocated to the agents in proportion to their agreed upon demands. As such, each agent has beliefs about any other agent's *expected* payoff given a coalition structure and demand vector. Specifically, agent i 's beliefs about the *expected stochastic payoff* of some agent $j \in C$ is denoted

$$\bar{p}_j^i = d_j V_i(C)$$

with d_j being the relative demand of agent j given the stated demands of the agents in C , and $V_i(C)$ the value that i expects C to have. (Henceforth, whenever we write “demand”, we imply a “relative demand” unless stated otherwise.) Similarly, if $i \in C$, i believes its *own* expected payoff to be $\bar{p}_i^i = d_i V_i(C)$.

A difficulty with using $V_i(C)$ in the above definition of expected stochastic payoff is that i 's assessment of the best (expected reward-maximizing) action for C is not necessarily shared by the rest of the agents. Therefore, we suppose instead that coalitions are formed using a process by which some coalitional action α is agreed upon, in the same way that the demand vector is agreed upon. In this case, i 's beliefs about j 's expected payoff is $\bar{p}_j^i(\alpha, C) = d_j Q_i(C, \alpha)$. Finally, we let $\bar{p}_j^i(C, \mathbf{d}_C, \alpha)$ denote i 's belief about j 's expected payoff if j is a member of $C \subseteq N$ with relative demand vector $\mathbf{d}_C$ taking action α :

$$\bar{p}_j^i(C, \mathbf{d}_C, \alpha) = d_j Q_i(C, \alpha)$$

Intuitively, if a coalition structure and payoff allocation are stable, we would expect: (a) no agent believes it will receive a payoff (in expectation) that is less than its reservation value; and (b) based on its beliefs, no agent will have an incentive to suggest that the coalitional agreement changed—specifically, there is no alternative coalition it could reasonably expect to join that offers it a better payoff than it expects to receive given the action choice and allocation agreed upon by the coalition to which it belongs.

We define the BC as follows:

DEFINITION 4 (WEAK BAYESIAN CORE). Let $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ be a coalition structure-demand vector-action vector triplet, with C_i denoting the $C \in CS$ of which i is a member. Then $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ is in the weak Bayesian core of a BCFP iff there is no coalition $S \subseteq N$, demand vector $\mathbf{d}_S$ and action $\beta \in A_S$ s.t. $\bar{p}_i^i(S, \mathbf{d}_S, \beta) > \bar{p}_i^i(C_i, \mathbf{d}_{C_i}, \alpha_{C_i})$, $\forall i \in S$, where $\mathbf{d}_{C_i}, \alpha_{C_i}$ is the projection of $\mathbf{d}, \mathbf{a}$ on the coalition C_i .

In words, there exists no coalition *all* of whose members each believe that they (personally) can be better off in it (in terms of expected payoffs, given some choice of action) than they currently are (within the current weak Bayesian core configuration). The agents' beliefs, in every $C \in CS$, “coincide” in the weak sense that there is a payoff allocation $\mathbf{d}_C$ and some coalitional action α_C that is commonly believed to ensure a better payoff. This doesn't mean that $\mathbf{d}_C$ and α_C is what each agent believes to be best. But an agreement on these is enough to keep any other coalition S from forming (even if one proposed its formation, others would disagree, not expecting to become strictly better off themselves).

We can define a stronger version of the Bayesian core, by demanding that there is *no* agent who believes that there exists a coalitional agreement that can make it strictly better off while not hurting the other members of the coalition, according to their own beliefs:

DEFINITION 5 (STRICT BAYESIAN CORE). Let $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ be a coalition structure-demand vector-action vector triplet, with C_i denoting the $C \in CS$ of which i is a member. Then $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ is in the strict Bayesian core of a BCFP iff there is no coalition $S \subseteq N$, demand vector $\mathbf{d}_S$ and action $\beta \in A_S$ s.t., for some $i \in S$, $\bar{p}_i^i(S, \mathbf{d}_S, \beta) > \bar{p}_i^i(C_i, \mathbf{d}_{C_i}, \alpha_{C_i})$ and $\bar{p}_j^j(S, \mathbf{d}_S, \beta) \geq \bar{p}_j^j(C_j, \mathbf{d}_{C_j}, \alpha_{C_j}) \forall j \in S, j \neq i$, where $\mathbf{d}_{C_i}, \alpha_{C_i}$ is the projection of $\mathbf{d}, \mathbf{a}$ on the coalition C_i .

The following is an obvious fact:

OBSERVATION 1. The strict Bayesian core is a subset of the weak Bayesian core.

In a similar manner we can define a different stability concept, which we call the *strong BC*. The strong BC requires that there is no agent who believes there is an agreement that can make it better off and that it expects all partners to accept based on (its subjective view of) their expected payoff. This differs from the strict (and the weak) BC in that the agent assesses its own beliefs about the value of an agreement to its partners.

DEFINITION 6 (STRONG BAYESIAN CORE). $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ is in the strong Bayesian core iff there is no coalition $S \subseteq N$, demand vector $\mathbf{d}_S$ and action $\beta \in A_S$ s.t. for some $i \in S$ $\bar{p}_i^i(S, \mathbf{d}_S, \beta) > \bar{p}_i^i(C_i, \mathbf{d}_{C_i}, \alpha_{C_i})$ and $\bar{p}_j^j(S, \mathbf{d}_S, \beta) \geq \bar{p}_j^j(C_j, \mathbf{d}_{C_j}, \alpha_{C_j}) \forall j \in S, j \neq i$.

The strong BC describes a notion of stability that is more tightly linked to the agents' subjective views on the potential acceptability of their proposals and is thus more “endogenous” in nature. By comparison, stability in the strict BC concept (and weak BC) is somewhat distinct. In an element of the strict BC, there may be an agent i who believes he would be strictly better off in some other coalition, and who believes all of his proposed partners would be better off as well; but the coalition may be considered unacceptable to some proposed partner j , since his beliefs about the value of the coalition are different than those of i .

The traditional notion of the *core* in deterministic settings is a special case of the weak BC when all agents know the actual types of other agents, in which case, the strict and the strong BC coincide and they are a subset of the weak BC. Since the core does not always exist, the three versions of the BC do not always exist either. We now make the following observation, which we will use later in the paper:

OBSERVATION 2. Let $\langle CS_N, \mathbf{d}, \mathbf{a} \rangle$ be an element of the BC of a BCFP with agents N . If $S \in CS_N$, $L = N \setminus S$, $CS_L = CS_N \setminus S$ and $\mathbf{d}_L, \alpha_L$ is the restriction of $\mathbf{d}, \mathbf{a}$ to the agents in L , the tuple $\langle CS_L, \mathbf{d}_L, \alpha_L \rangle$, which is contained in the $\langle CS_N, \mathbf{d}, \mathbf{a} \rangle$ configuration, is an element of the BC of the corresponding BCFP subgame with L agents.

Obviously, if an agent i in the $L = N \setminus S$ subset believed that he could do strictly better in some coalition other than his current C_i , he would have believed this when the S coalition was present, as well; in that case, $\langle CS_N, \mathbf{d}, \mathbf{a} \rangle$ couldn't have been a (strong) BC element. Notice that Observation 2 holds for all versions of the BC.

4. A NON-COOPERATIVE JUSTIFICATION OF THE BAYESIAN CORE

In this section we present the main results of our paper. We first define a coalitional bargaining game that deals with the non-cooperative aspects of a BCFP, and then prove our main propositions, which relate the BC with appropriate non-cooperative solution concepts for coalition formation under uncertainty.

4.1 Bayesian Coalitional Bargaining

While coalition structures and allocations can sometimes be computed centrally, in many situations they emerge as the result of some bargaining process among the agents, who propose, accept and reject partnership agreements. We now define a (*Bayesian*) coalitional bargaining game (BCBG) for the Bayesian model above, as a *Bayesian extensive game with observable actions* [14], adopting the approach of [4].

The game proceeds in stages, with a randomly chosen agent proposing a coalition, a coalitional action and an allocation of payments to partners, who then accept or reject the proposal. A finite set of *bargaining actions* is available to the agents. A bargaining action corresponds to either: (a) some *proposal* $\pi = \langle C, P_C \rangle$ to form a coalition C with a specific payoff configuration P_C (specifying payoff shares d_i to each $i \in C$ and a suggested coalitional action α_C for C to perform); or (b) the acceptance or rejection of such a proposal. The game proceeds in stages, and initially all agents are active. At the beginning of stage t , one of the (say n) active agents i is chosen randomly with probability $\gamma = \frac{1}{n}$ to make a proposal $\langle C, P_C \rangle$ (with $i \in C$). Each other $j \in C$ either accepts or rejects this proposal. If all $j \in C$ accept, the agents in C are made inactive and removed from the game. Value $V_i(t_C) = \delta^{t-1}Q(C, \alpha_C | t_C)$ is realized by C at s , and split according to P_C , where $\delta \in (0, 1]$ is the discount factor. If any $j \in C$ rejects the proposal, the agents remain active (no coalition is formed). At the end of a stage, the responses are observed by all participants, and the agents can update their beliefs regarding others using Bayes rule (as described in [4]). If the game is finite-horizon, at the end of the final stage F , any i not in any coalition receives its discounted reservation value $\delta^{F-1}V(i)$. In this paper, however, we will be assuming an infinite horizon (and thus, due to discounting, agents' payoffs are zero in the long run, if they reach no agreement with others in the meantime).

This bargaining game focuses on the strategic interactions of the rational players, and, thus, provides a *non-cooperative* view of the BCFP. The suitable solution concept for the game above is, as proposed by [4], a *perfect Bayesian equilibrium* (PBE) [14]. In this paper, we are interested in establishing connections between non-cooperative, equilibrium solutions of coalitional bargaining games under uncertainty, and the cooperative solution concepts (i.e., the BC) of the underlying BCFP. However, the BC, being a stability concept, implicitly assumes that agents' beliefs are settled to specific values before it can be defined. Thus, we make the simplifying assumption that agents' beliefs remain *fixed* throughout bargaining, and we define an appropriate sequential equilibrium concept for this game. Certainly, the fact that an agent i makes a specific proposal to a set of other agents (or accepts or rejects a particular proposal) can influence j 's beliefs about i 's type, and thus j 's behavior at future rounds of the BCBG. However, this form of belief dynamics over multiple rounds is not (and most probably

cannot be) reflected in a static cooperative solution concept such as the BC. Hence the motivation for this restriction.

The appropriate solution concept for a BCBG game under fixed beliefs, is a *sequential equilibrium under fixed beliefs*:

DEFINITION 7. *Let G be a BCBG throughout the course of which the agents' beliefs are assumed to remain fixed. Then, a profile of (possibly mixed) strategies, one for each player in N , is a sequential equilibrium under fixed beliefs (SEFB) for G , if, for each $i \in N$ and each history h , i 's strategy continuation after h is optimal, given the strategies of other players and i 's fixed beliefs.*

The SEFB is therefore defined as an extension of SPE and a restriction of PBE equilibria. This is appropriate for our fixed-beliefs bargaining game, which is an extensive form game that *does* incorporate beliefs (and the fact that different agents can have widely varying beliefs about the value of any coalition) unlike the SPE; the beliefs are merely held fixed throughout the bargaining process (unlike the PBE).

4.2 Equivalence of the Cooperative and Non-Cooperative Solutions

We now show that the existence of stable coalition structures in a coalition formation problem under uncertainty implies the existence of an equilibrium bargaining profile that leads to their formation; and also that "optimal" coalitional bargaining under uncertainty leads to stable coalitions.

We first define a subclass of bargaining games that we will be interested in.

DEFINITION 8. *Let $\mathcal{C}$ be the class of N -player BCFP with the following properties:*

1. *All subgames have a nonempty strict BC.*
2. *For every member $B_N = \langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$ of the strict BC, where CS_N is of the form $\{S_1, \dots, S_k\}$, every subgame with set of players $T \subseteq N$ has an element in its strict BC in which the coalition structure is of the form $\{T \cap S_1, \dots, T \cap S_k\}$ and also the demand vector is the projection of $\mathbf{d}_N$ to the corresponding coalition.*

In the above definition we ignore the empty sets that may arise if T does not intersect any of the S_j 's. Note that by Observation 2, the properties of Definition 8 are already satisfied by subgames in which the set of players is a union of some of the coalitions of CS_N . Hence the definition simply imposes the same properties for other subsets of N as well.

We are now ready to prove our first relevant Proposition:

PROPOSITION 1. *Let $\mathcal{P} \in \mathcal{C}$ be an N -player BCFP. Then, for every member $B_N = \langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$ of the strict BC of $\mathcal{P}$, there exists an SEFB equilibrium $\sigma^* = \sigma^*(B_N)$ in pure strategies, of the corresponding BCBG G (under fixed beliefs) with N players and random proposers, such that, the coalition structure induced by σ^* is exactly $\langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$.*

PROOF. Let $B_N = \langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$ be an arbitrary element of the BC of $\mathcal{P}$ and let BC_S represent the strict BC of the subgame where the set of agents is S . Let $CS_N = \{S_1, S_2, \dots, S_k\}$ for some k , where $\cup S_i = N$ and $S_i \cap S_j = \emptyset$ for every i, j . For each S we choose an element $B_S = \langle CS_S, \mathbf{d}_S, \mathbf{a}_S \rangle \in BC_S$. In particular, we choose such an element according to Definition 8. For example if S is the union of some of the S_i 's, i.e., if it consists precisely of some

of the coalitions of CS_N , then we let B_S be the restriction of $\langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$ to S , which by Observation 2 lies in BC_S . Our way of choosing these core elements for each subgame ensures the following, easy to verify, fact:

FACT 1. *Let $T \subseteq N$ and suppose the coalition structure in B_T is $\{T_1, \dots, T_l\}$. Then for the subgame where the set of players is $S = T \setminus T_1$, the structure in the corresponding core element B_S is $\{T_2, \dots, T_l\}$ and the demand and action vectors are the same as in T .*

Given a triplet $\langle C, \mathbf{d}_C, \alpha_C \rangle$, with $i \in C$, we will denote by $\bar{p}_i^i(C, \mathbf{d}_C, \alpha_C)$ the expected payoff of i from the formation of coalition C , according to i 's beliefs. We will also use $\bar{p}_i^i(B_S)$ to denote $\bar{p}_i^i(C_i, \mathbf{d}_{C_i}, \alpha_{C_i})$, where C_i is the coalition in the coalition structure of B_S that i belongs to.

Consider now the following strategy σ_i^* for a player i :

- (i) If i is the proposer in some round of the game, and the set of agents still present is S , let C denote i 's coalition in the coalition structure of B_S , and let $\mathbf{d}_C, \alpha_C$ be its corresponding demand vector and action, i.e., we look at the projection of B_S to coalition C that contains i . Then, i proposes $\langle C, \mathbf{d}_C, \alpha_C \rangle$ to the rest of the agents in C^2 .
- (ii) If i is a responder, the subset of agents still present is S , and the standing proposal is $\langle T, \mathbf{d}_T, \alpha_T \rangle$ (with $i \in T$), then i accepts iff $\bar{p}_i^i(\langle T, \mathbf{d}_T, \alpha_T \rangle) \geq \bar{p}_i^i(B_S)$.

Let σ^* be the profile of the strategies of all players. It is clear that, no matter what the nature's choice of proposers in the game is, if σ^* is played then the outcome of the game is exactly B_N . To see this, suppose that the first random proposer at round 1, say i , belongs to S_1 (recall $B_N = \langle CS_N, \mathbf{d}_N, \mathbf{a}_N \rangle$ and $CS_N = \{S_1, S_2, \dots, S_k\}$). Then i will propose $\langle S_1, \mathbf{d}_{S_1}, \alpha_{S_1} \rangle$, with $\mathbf{d}_{S_1}, \alpha_{S_1}$ being the projection of $\mathbf{d}_N, \mathbf{a}_N$ to S_1 . By the definition of σ^* all other members of S_1 will accept and the game will go to round 2 with the remaining players $L = N \setminus S_1$. Suppose agent j is now the proposer and $j \in S_2$. By the way we defined B_L , we know that $CS_L = \{S_2, \dots, S_k\}$, (because L consists of a collection of coalitions of CS_N) therefore j will propose $\langle S_2, \mathbf{d}_{S_2}, \alpha_{S_2} \rangle$ with $\mathbf{d}_{S_2}, \alpha_{S_2}$ being the projection of $\mathbf{d}_N, \mathbf{a}_N$ on S_2 . The other members of S_2 will accept the proposal and the game will continue in the same manner. Hence after k rounds the game will end and the outcome will be B_N .

It remains to show that σ^* is an *SEFB* equilibrium.

Assume i is a proposer at some round of the game, the set of active players is S and all players apart from i will play according to σ^* from this point of the game onwards. Let $\langle C, \mathbf{d}_C, \alpha_C \rangle$ be the triplet of B_S with $i \in C$. We want to show that i cannot gain by deviating from σ^* . Suppose i deviates by proposing $\langle T, \mathbf{d}_T, \beta_T \rangle$, different from $\langle C, \mathbf{d}_C, \alpha_C \rangle$, where C is the coalition of B_S that i belongs to. Consider first the case that $\bar{p}_i^i(T, \mathbf{d}_T, \beta_T) > \bar{p}_i^i(B_S)$. Note that in this case T cannot be a singleton since that would contradict the fact that $B_S \in BC_S$. Hence $|T| \geq 2$. Then the proposal is accepted *only if* for all agents $j \in T, j \neq i$, it is the case that $\bar{p}_j^j(T, \mathbf{d}_T, \beta_T) \geq \bar{p}_j^j(B_S)$ (since they follow σ^*). However, if this is the case, then we have found a coalition, namely T , along with a demand vector and an action such that agent i believes he is strictly better off and no other agent believes that he is worse off. This contradicts the fact that $B_S \in BC_S$, hence the proposal is never accepted and i cannot gain from such a deviation.

²We must assume the non-emptiness of BC_S in order to define this strategy at any bargaining round.

Consider now the case that $\bar{p}_i^i(T, \mathbf{d}_T, \beta_T) \leq \bar{p}_i^i(B_S)$. Agent i cannot gain from such a deviation either. If the proposal is accepted he does not gain more than his payoff under σ_i^* . If the proposal is rejected, then the game moves to the next round without any coalition forming. In the next round, if the proposer is some other member of C , then the proposal for i will be $\langle C, \mathbf{d}_C, \alpha_C \rangle$, which does not give him a better payoff than $\bar{p}_i^i(B_S)$. If i is chosen to be a proposer again, then we already know that he cannot propose a coalition that gives him better payoff. Now suppose the chosen proposer, say j , does not belong to C and let $C_j \subseteq S$ be the coalition of B_S that j belongs to. Since every player apart from i follows σ^* , j will propose C_j which will be accepted and the game will move to the next round where the set of players is $S \setminus C_j$. By Fact 1 the core element $B_{S \setminus C_j}$ for this subgame still contains $\langle C, \mathbf{d}_C, \alpha_C \rangle$. Hence by repeating the above arguments, i cannot gain more than $\bar{p}_i^i(B_S)$. Therefore, whenever nature i becomes a proposer, i cannot gain a better payoff than the payoff he obtains if he follows σ^* .

Assume now that i is a *responder* to some offer $\langle T, \mathbf{d}_T, \beta_T \rangle$ and the current set of active players is S . Suppose that all the agents in T that responded before i have already accepted the proposal and let U be the set of agents who are to decide after agent i .

Case 1: $\bar{p}_i^i(T, \mathbf{d}_T, \beta_T) \geq \bar{p}_i^i(B_S)$. Then according to σ^* agent i should accept the proposal. If i deviates from σ^* and rejects the proposal then there are two subcases to consider. If all agents in U are going to accept the proposal then i would receive a payoff of at least $\bar{p}_i^i(B_S)$ had he followed σ^* . Since he rejected the proposal, no coalition forms and the game goes to the next round. In the next round either he is the proposer, in which case we know by the previous arguments that he cannot gain more than $\bar{p}_i^i(B_S)$ or someone else is the proposer in which case again, using Fact 1, he cannot gain more than $\bar{p}_i^i(B_S)$ because all other agents follow σ^* . If on the other hand some agent in U will reject the proposal then it does not matter whether i accepts or rejects. The game moves to the next round and agent i cannot obtain a payoff better than the payoff under σ^* .

Case 2: $\bar{p}_i^i(T, \mathbf{d}_T, \beta_T) < \bar{p}_i^i(B_S)$. Then according to σ^* agent i should reject the proposal. If i deviates from σ^* and accepts the proposal then if all agents in U also accept, the coalition T forms and agent i receives a payoff which is less than $\bar{p}_i^i(B_S)$. However, if he had followed σ^* , the proposal would have been rejected and in the future he would have obtained $\bar{p}_i^i(B_S)$.³ If some agent in U will reject the proposal then as in Case 1, i cannot profit by the deviation from σ^* .

Overall, agent i cannot benefit by any deviation from σ_i^* and, thus, σ^* is a *SEFB* of the corresponding game G . $\square$

Proposition 1 remains true if we use the weak BC instead of the strict one. In that case we only have to modify the strategy σ^* accordingly. Due to lack of space we omit the details here. For the reverse direction (can an *SEFB* give rise to a configuration that belongs to the core?), we cannot hope to always have a positive answer since the BC does not always exist. However we can provide a positive answer if the bargaining game possesses equilibria whose outcomes do not depend on the random choice of the proposers. The following definition is a generalization of the one given in [10].

DEFINITION 9. *An SEFB equilibrium in pure strategies is order independent if, whenever it is played, it leads to*

³Notice that this argument holds only for $\delta = 1$.

the same outcome $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ regardless of the choice of proposers.

Note that the equilibrium defined in the proof of Theorem 1 is also order independent. In the following result, we show that order independent equilibria lead to outcomes that belong to the weak BC. We are not aware at the moment if this result is true for the strict or the strong BC.

PROPOSITION 2. *Let σ^* be an order independent SEFB equilibrium strategy profile in pure strategies for a BCBG G with random proposers and discount factor δ . Then, the outcome of σ^* , $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ must be in the weak BC of the corresponding BCFP.*

PROOF. Let $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ be the outcome of the game if the equilibrium σ^* is played. Assume, contrary to the proposition, that $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ is not in the weak BC. Let $\bar{p}_i^j(\sigma^*; t = 1)$ denote i 's expected payoff under σ^* (i.e., if everybody follows σ^* right from the first round). Since $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ derived by σ^* is not in the weak BC, there exists a coalition $S \subseteq N$, a demand vector $\mathbf{d}_S$ and an action α_S such that:

$$\bar{p}_i^j(S, \mathbf{d}_S, \alpha_S) > \bar{p}_i^j(\sigma^*; t = 1) \quad \forall j \in S \quad (1)$$

Consider now an agent $i \in S$, and consider the following strategy for i :

- (i) If i is chosen by nature to be the first proposer, then i proposes $\langle S, \mathbf{d}_S, \alpha_S \rangle$.
- (ii) In all other cases, i follows σ_i^* .

We will show that this deviation from σ^* benefits agent $i \in S$ in expectation when the other agents play according to σ^* , and therefore σ^* cannot be an SEFB equilibrium.

Assume that i was chosen by nature to be the first proposer. Then, i proposes $\langle S, \mathbf{d}_S, \alpha_S \rangle$ with the above property. Note that $|S| \geq 2$, otherwise we would already have a contradiction to σ^* being an equilibrium.

All other agents $j \in S$ follow their σ^* -equilibrium strategies. Consider a responder $j \in S$ and consider the subgame that starts at the node where j is to decide whether to accept or reject the proposal and assume every other agent in $S \setminus \{i, j\}$ has accepted. Note that from this point onwards every agent (including i) plays according to σ^* , which is an equilibrium for this subgame (since σ^* is a sequential equilibrium). We will show that i is optimal for j to accept.

If j rejects the proposal, then the game moves to round 2 where all agents are present and from then on they all play σ^* . Since σ^* is order independent, the configuration $\langle CS, \mathbf{d}, \mathbf{a} \rangle$ will form and therefore agent j can get a payoff of at most $\bar{p}_j^j(\sigma^*; t = 1)$ (possibly discounted). On the other hand if j accepts the proposal he obtains a better payoff by (1). Hence rejecting the proposal cannot be optimal for j . By backward induction, the proposal of agent i must be accepted by all agents of S , therefore the coalition S will form and i will obtain a better payoff. This implies that σ^* is not an equilibrium, a contradiction. $\square$

REMARK 1. *Note that in Proposition 2 we allow the bargaining game to have an arbitrary discount factor $\delta \leq 1$, whereas in the game defined in the proof of Proposition 1 we did not allow any discounting ($\delta = 1$).*

A consequence of Proposition 2 is the following corollary, which provides a condition for the existence of the weak BC.

COROLLARY 1. *If an order independent SEFB equilibrium strategy profile exists, then the weak BC cannot be empty.*

5. TESTING FOR THE EXISTENCE OF THE BAYESIAN CORE

Even in a deterministic setting where there is no uncertainty, testing for the nonemptiness of the core is an intractable problem. In fact, even in superadditive games, where the coalition structure is simply the grand coalition N and we only need to find an allocation of payoffs to the agents, the problem is NP-hard (see, e.g., [6], where a concise game representation is used and superadditivity is assumed). In the absence of superadditivity, as is our case there are even worse lower bounds on the complexity of the problem. Sandholm *et al.* [16] show that in order to find an approximately optimal coalition structure (i.e., in order to be able to establish a bound from the optimum), exponentially many coalition structures have to be searched (at least 2^{n-1} , if n agents). All these suggest that it is unlikely to find an efficient algorithm for verification of the BC's existence.

Here we show that for a relatively small number of agents it is feasible to check for the nonemptiness of the BC without employing a brute-force approach that simply searches over all coalition structures. We formulate the problem as a constraint satisfaction problem (CSP), where the constraints are polynomial equalities or inequalities. We can then use existing algorithms that solve such CSP's (e.g. see [1, 21]).

Before we present the polynomial program that solves our problem, we make some simplifying assumptions. First we assume that for each coalition there is a finite (and not very large) number of possible demand vectors that one could propose (i.e., there is a finite number of possible ways in which the agents will split the payoff of the coalition). W.l.o.g., assume that each coalition has k actions available.

The CSP that we present below tests the nonemptiness of the weak BC and has 4 types of variables. Similar CSPs can be written for the strict and the strong BC too. For each coalition S , there is an indicator variable X_S which indicates whether coalition S will form in the coalition structure that we are looking for. We also have a variable r_i for each agent i that indicates the share that i will have in the coalition that he belongs. Furthermore, let $Q_i(S, j)$ denote the payoff that coalition S gets if the j th action is taken (recall there are k available actions). Then for each coalition S and action j , $j = 1, \dots, k$ we have an indicator variable α_j^S that indicates whether action j is taken or not (if coalition S forms). Finally, for each possible deviation from the core, say $\langle T, \mathbf{d}, \beta \rangle$, where $\mathbf{d}$ is a $|T|$ -dimensional demand vector ($d_1, \dots, d_{|T|}$) and β is one of the k available actions to coalition T we have an auxiliary variable $Z_{T, \mathbf{d}, \beta}$ whose role is to ensure that it cannot be the case that all agents gain more expected payoff if they deviate to T .

$$X_S(1 - X_S) = 0 \quad \forall S \subseteq N \quad (2)$$

$$X_S(1 - \sum_{i \in S} r_i) = 0 \quad \forall S \subseteq N \quad (3)$$

$$r_i \geq 0 \quad \forall i \in N \quad (4)$$

$$\sum_{S: i \in S} X_S = 1 \quad \forall i \in N \quad (5)$$

$$\alpha_j^S(1 - \alpha_j^S) = 0 \quad \forall S \subseteq N, j \in \{1, \dots, k\} \quad (6)$$

$$\prod_{i \in T} [Z_{T, \mathbf{d}, \beta} - (d_i Q_i(T, \beta) - r_i \sum_{S: i \in S} X_S \sum_{j=1}^k \alpha_j^S Q_i(S, j))] = 0 \quad \forall T, \mathbf{d}, \beta \in \{1, \dots, k\} \quad (7)$$

$$Z_{T,\mathbf{d},\beta} \leq d_i Q_i(T, \beta) - r_i \sum_{S:i \in S} X_S \sum_{j=1}^k \alpha_j^S Q_i(S, j) \quad \forall T, \mathbf{d}, \beta, i \in T \quad (8)$$

$$Z_{T,\mathbf{d},\beta} \leq 0 \quad \forall T, \mathbf{d}, \beta \in \{1, \dots, k\} \quad (9)$$

PROPOSITION 3. *The above problem is feasible iff the weak BC of the corresponding game is nonempty.*

PROOF. Suppose that the program is feasible and consider a solution. Then constraints (2) and (6) ensure that the variables X_S and α_j^S are integer 0/1 variables. Hence we can see them as indicator variables, indicating which coalitions were chosen by the solution and which action was taken (if $X_S = 1$ we consider that S forms). Constraints (5) ensure that the coalitions that form make up a coalition structure; each agent belongs to exactly one of them. Constraints (3) and (4) ensure that for any coalition that forms the demands r_i for $i \in S$ form a valid demand vector. The rest of the constraints ensure that there is no coalition T , demand vector $\mathbf{d}$ and action β that would make all agents of T better off. For a coalition T and an agent $i \in T$, let ϵ_i be the amount by which i 's payoff changes if he deviates from the solution to the program to $(T, \mathbf{d}, \beta)$. Constraints (7) and (8) make sure that the variable $Z_{T,\mathbf{d},\beta}$ is equal to $\min_i \epsilon_i$ because the expression $r_i \sum_{S:i \in S} X_S \sum_{j=1}^k \alpha_j^S Q_i(S, j)$ is equal to the expected payoff of agent i under the feasible solution of the program (recall only one of the variables X_S with $i \in S$ is set to 1 and the rest are 0). Finally the last constraint ensures that for any $(T, \mathbf{d}, \beta)$, $\min_i \epsilon_i \leq 0$, which means that there is no coalitional agreement that can make all agents strictly better off. The reverse direction is straightforward. $\square$

The number of variables in the program above is $O(kd2^n)$, where k is the number of actions, d the number of demand vectors, and n the number of agents. Although the worst case running time guarantees for such a program can be prohibitively high in most realistic settings [1], it can be solved heuristically for small problem sizes [21] and thus is a better tool than a brute-force approach. Moreover, this program offers a concise way to describe existence conditions, since there is no obvious way to define *balancedness* [11] in an uncertain, non-superadditive environment.

6. CONCLUSIONS

We have further developed the foundations of coalition formation under a realistic model of uncertainty by proposing new stability concepts for BCFPs and assessing the related bargaining processes. We have established strong connections between a cooperative coalitional stability concept under type uncertainty and non-cooperative equilibrium solutions of the corresponding bargaining games. We proved that if the BC of a coalitional game is non-empty, then there exists an equilibrium of the corresponding bargaining game that induces an element of the BC; and we showed that if an order independent coalitional bargaining equilibrium exists, then it leads to a BC configuration. This provides a non-cooperative justification of the BC. We also established a sufficient condition for the existence of the weak BC and—for small games—an algorithm to decide whether the BC is non-empty. Current and future work includes investigating these issues for the stronger stability concepts (e.g., strong BC) and other bargaining models.

7. REFERENCES

- [1] S. Basu, R. Pollack, and M.-F. Roy. *Algorithms in Real Algebraic Geometry*. Springer-Verlag, 2003.
- [2] B. Blankenburg, M. Klusch, and O. Shehory. Fuzzy Kernel-Stable Coalitions Between Rational Agents. In *Proc. of AAMAS'03*, 2003.
- [3] G. Chalkiadakis and C. Boutilier. Bayesian Reinforcement Learning for Coalition Formation Under Uncertainty. In *Proc. of AAMAS'04*, 2004.
- [4] G. Chalkiadakis and C. Boutilier. Coalitional Bargaining with Agent Type Uncertainty. In *Proc. of IJCAI-07*, 2007.
- [5] K. Chatterjee, B. Dutta, and K. Sengupta. A Noncooperative Theory of Coalitional Bargaining. *Review of Economic Studies*, 60:463–477, 1993.
- [6] V. Conitzer and T. Sandholm. Complexity on constructing solutions in the core based on synergies among coalitions. *Artificial Intelligence*, 170:607–619, 2006.
- [7] T. Dieckmann and U. Schwalbe. Dynamic Coalition Formation and the Core, 1998. Economics Department Working Paper Series, Department of Economics, National University of Ireland - Maynooth.
- [8] R. Evans. Coalitional Bargaining with Competition to Make Offers. *Games and Econ. Behavior*, 19:211–220, 1997.
- [9] S. Hart and A. Mas-Colell. Bargaining and Value. *Econometrica*, 64(2):357–380, Mar. 1996.
- [10] B. Moldovanu and E. Winter. Order Independent Equilibria. *Games and Econ. Behavior*, 9:21–34, 1995.
- [11] R. Myerson. *Game Theory: Analysis of Conflict*. Harvard University Press, 1991.
- [12] T. Norman, A. Preece, S. Chalmers, N. Jennings, M. Luck, V. Dang, T. Nguyen, V. Deora, J. Shao, A. Gray, and N. Fiddian. Agent-based formation of virtual organisations. *Knowledge Based Systems*, 17:103–111, 2004.
- [13] A. Okada. A Noncooperative Coalitional Bargaining Game With Random Proposers. *Games and Economic Behavior*, 16:97–108, 1996.
- [14] M. Osborne and A. Rubinstein. *A course in game theory*. MIT Press, 1994.
- [15] M. Perry and P. Reny. A Noncooperative View of Coalition Formation and the Core. *Econometrica*, 62(4):795–817, July 1994.
- [16] T. Sandholm, K. Larson, M. Andersson, O. Shehory, and F. Tohme. Coalition Structure Generation with Worst Case Guarantees. *Artificial Intelligence*, 111(1–2):209–238, 1999.
- [17] T. Sandholm and V. Lesser. Coalitions Among Computationally Bounded Agents. *Artificial Intelligence*, 94(1):99 – 137, 1997.
- [18] R. Serrano and R. Vohra. Non-cooperative implementation of the core. *Social Choice and Welfare*, 14:513–525, 1997.
- [19] O. Shehory and S. Kraus. Methods for Task Allocation via Agent Coalition Formation. *Artificial Intelligence*, 101(1–2):165–200, 1998.
- [20] O. Shehory and S. Kraus. Feasible Formation of Coalitions among Autonomous Agents in Nonsuperadditive Environments. *Comput. Intelligence*, 15:218–251, 1999.
- [21] B. Sturmfels. *Solving Systems of Polynomial Equations*. 2002.
- [22] J. Suijs and P. Borm. Stochastic cooperative games: superadditivity, convexity and certainty equivalents. *Journal of Games and Econ. Behavior*, 27:331–345, 1999.
- [23] J. Suijs, P. Borm, A. D. Wagenaere, and S. Tijs. Cooperative games with stochastic payoffs. *European Journal of Operational Research*, 113:193–205, 1999.
- [24] J. von Neumann and O. Morgenstern. *Theory of Games and Economic Behavior*. Princeton University Press, 1944.
- [25] H. Yan. Noncooperative selection of the core. *International Journal of Game Theory*, 31(4):527–540, Sept. 2003.
- [26] M. Yokoo, V. Conitzer, T. Sandholm, N. Ohta, and A. Iwasaki. Coalitional games in open anonymous environments. In *Proc. of AAAI-05*, 2005.