

Distributed Negotiation under Uncertainty as a Foundation for Theories of Social Welfare

Cédric Dégremont

Institute for Artificial Intelligence and Cognitive Engineering

University of Groningen

Email: `cedric@ai.rug.nl`

Ulle Endriss

Institute for Logic, Language and Computation

University of Amsterdam

Email: `ulle.endriss@uva.nl`

18 May 2011

Abstract

We study a model where a group of self-interested agents negotiate over a set of resources. The agents believe that at the end of the negotiation phase a certain randomisation will take place: either the bundles of resources the agents have accumulated will get reassigned to other agents, or the valuation functions the agents use to assess the values of these bundles will get reassigned, or both. The uncertainty as to which combination of valuation function and bundle an agent will end up with will influence her negotiation strategy. For certain types of uncertainty of this kind and for certain assumptions on the attitude towards risk of the agents involved, it is possible to show that negotiation is guaranteed to converge to an allocation with certain desirable properties, such as maximising egalitarian or utilitarian social welfare. This model of distributed negotiation under uncertainty thus provides a new perspective on theories of social welfare.

1 Introduction

A central idea in the study of distributive justice is that *fair* is what rational agents would agree upon in the face of uncertainty (or ignorance) of their own identity. A first aspect of this idea is that value judgments are “nonegoistic impersonal judgments of preference” (Harsanyi, 1953), and that ignorance of your particular position in society and of your particular personal preferences “nullify the effects of specific contingencies which put men at odds and tempt

them to exploit social and natural circumstances to their advantage” (Rawls, 1999, §24). Therefore, such ignorance seems to provide sufficient conditions for the required nonegoistic impersonal judgments. The second aspect of the idea is that agreement should come as the outcome (as the equilibrium point) of a process in which “agreements [are] freely struck between willing traders” who are “rational individuals with certain ends” and who are trying to secure these ends “in view of their knowledge of the circumstances” (Rawls, 1999, §20).

Recent work in Computer Science, and more specifically in Multiagent Systems, has taken up some of these ideas and tried to devise negotiation frameworks for autonomous software agents that would allow such agents to reach socially desirable allocations of resources in an interactive and distributed manner by means of negotiation, rather than by imposing an optimal allocation computed by a benevolent dictator (Endriss et al., 2006; Chevaleyre et al., 2010). This line of work has mostly been concerned with questions of a computational nature: How can we design negotiation protocols that ensure convergence to an allocation that is optimal according to our social welfare concept of choice? What can be said about bounds on the length of negotiation processes before our objective is achieved? What is the computational complexity of the reasoning tasks required from the agents along the way?

Here, instead, we propose to use this distributed negotiation framework to provide concrete “operational models” for thought experiments proposed to analyse what decisions rational agents would make in an original position, under the veil of ignorance regarding their own identities, similar to that put forward by Rawls or Harsanyi. Our aim is not to account for some specific existing philosophical position, rather we would like to identify a simple negotiation framework with uncertainty in which agents might or might not converge to allocation of resources that can be deemed fair or efficient, according to standard welfare-theoretic notions. Having a simple, concrete, and clearly defined procedure of negotiation allows us to study the precise effect of a given set of assumptions regarding the model on the agreements the agents will reach. The basic idea underlying our approach is this: assuming that the agents believe that some randomisation of important parameters (such as the bundles of resources they receive, or the valuation functions they use to assess those bundles) will take place, we can say that some social welfare concept is grounded by a thought experiment based on some theory of decision-making and some operational model, if in this model self-interested agents acting according to the decision-making theory will converge to an allocation that maximises this social welfare concept.

There are at least four basic ways of introducing uncertainty into negotiation:

- (1) Agents are uncertain about the *bundle of resources* they will receive at the end of the negotiation process.
- (2) Agents are uncertain about the *valuation function* they will be assigned at the end

of the negotiation process, and thus how much they will be able to benefit from the resources they have accumulated.

- (3) Agents are uncertain about the pair of bundle and valuation function that will be assigned to them, but they do know that these are tied together, i.e., they will obtain the pair held by one of the agents in the system at the end of the negotiation process. In other words, agents are uncertain about their eventual *identity*.
- (4) Agents are uncertain about both the bundle and the valuation function they will obtain, and they expect these two parameters to be assigned *independently* from each other.

Conceptually speaking, these four options correspond to natural candidates for defining fairness in terms of negotiation uncertainty. The first option corresponds to the idea that fair is what agents would agree upon if they were to believe that you will randomise what they eventually obtain. The second option corresponds to the idea that fair is what agents would agree if they were uncertain about what their preferences would ultimately turn out to be. The third and the fourth option are the two most basic forms of combining these two types of uncertainty.

As we shall see, we get a tight connection between decision-making theories and theories of social welfare when we randomise identities, i.e., pairs of bundles and valuation functions tied together (the third option in our list above). In particular, we can account for the intuition that agents that are interested in maximising their payoff in the worst case would want to choose an allocation with maximal egalitarian social welfare (i.e., an allocation that maximises the well-being of the worst-off agent). In our framework, we do not only see that it will be in the interest of such agents to adopt an egalitarian solution, but we can show that any sequence of deals in a process of rational negotiation will always *converge* to such a solution. Similarly, we can account for an argument, similar to Harsanyi's, that expected-value maximisers should choose an allocation with maximal utilitarian social welfare (i.e., an allocation that maximises average utility), which in our framework also manifests itself in a convergence result. This point does shed some light on the fact that the difference is precisely located in the theory of decision-making that is being assumed.

As for the other three types of reassignment, namely to randomise only the bundles, only the valuation functions, or to randomise both parameters, albeit independently from each other, some of these scenarios turn out to lead to immediate termination of the negotiation process, while others can lead to infinite sequences of deals. Also such “negative” results are interesting. They show, for instance, that some notions that may appear to make intuitive sense can turn out not to be helpful in practice. For instance, the idea of terminal allocation (“what agents would agree upon”) might not even be well-defined, namely when we can prove that negotiation will continue indefinitely. In other cases, when we can prove that

no negotiation will take place at all, the content of the idea of terminal allocation will be well-defined but also uninformative.

Some scenarios, finally, do provide guarantees for convergence to allocations with attractive properties, but only under somewhat restrictive assumptions, concerning either the number of agents or the class of valuation functions that agents may hold. Results of this kind suggest that a discussion of fairness in terms of uncertainty must pay close attention to the precise parameters of the domain in which we are operating.

Our model of distributed negotiation under uncertainty is chiefly intended as a simple “implementation” of the kinds of thought experiments in the spirit of those originally put forward by Harsanyi and Rawls. Rather than showing that a certain theory of social welfare would suit agents best under certain assumptions, we show how negotiation amongst rational and self-interested agents does, under similar such assumptions, lead to an adoption of the theory of social welfare in question. Beyond its interest as a pure thought experiment, the specific case where only the bundles of resources get randomised also has some practical interest, given that this kind of randomisation can actually be implemented in practice (while randomising the agents’ valuation functions is hardly possible in the real world).

Paper overview. The remainder of the paper is organised as follows. Section 2 defines the model we shall be working with. In particular, this involves defining the notion of rational negotiation for different types of decision-makers. Section 3 recalls a number of results for the perfect information case. We then present our results for a number of different settings. Section 4 considers the case where only the resource bundles get randomised, but each agent will keep her own valuation function. This kind of randomisation (unlike randomisation involving valuations) can also be implemented in practice. Section 5 then briefly discusses the (somewhat less interesting) case where only valuations get reassigned. Section 6 analyses the case where agents are uncertain about what their final identity will be, but they do know that bundles and valuation functions will stay paired the way they are, and that they will each receive one of the identities (bundle/valuation-pair) present in society at the end of the negotiation process. This turns out to be the scenario leading to the clearest alignment between attitudes towards risk of individual decision makers and theories of social welfare. Before concluding, Section 7 discusses the remaining case, where we announce to the agents that we will reassign both valuations and bundles, but that we will perform each assignment independently from the other.

The proofs of all the formal results stated in the text may be found in the appendix.

2 The Model

In this section we define the model we shall be working with. It is based on a negotiation framework that has been used to study the dynamics of resource allocation in multiagent systems in the Artificial Intelligence literature (Endriss et al., 2006; Chevaleyre et al., 2010); we enrich it here by introducing a notion of uncertainty over possible reassignments of resource bundles and valuation functions.

2.1 Multiagent Resource Allocation

Let $\mathcal{N}$ be a finite set of *agents* and let $\mathcal{R}$ be a finite set of (indivisible) *resources*. An *allocation* is a function $\alpha : \mathcal{R} \rightarrow \mathcal{N}$ mapping each resource to an agent. The *bundle* (set of resources) held by agent $i \in \mathcal{N}$ under allocation α is $\alpha^{-1}(i) = \{x \in \mathcal{R} \mid \alpha(x) = i\}$.

From an initial allocation, our agents can negotiate a sequence of deals regarding the exchange of some of the resources in their position. A single deal may involve any number of agents and any number of resources; formally, a *deal* $\delta = (\alpha, \alpha')$ is simply a pair of (distinct) allocations, describing the situation before and after the deal has taken place. We denote the set of agents *involved* in the deal $\delta = (\alpha, \alpha')$, i.e., the set of agents whose bundle changes when we move from α to α' , as $\mathcal{N}^\delta = \{i \in \mathcal{N} \mid \alpha^{-1}(i) \neq \alpha'^{-1}(i)\}$.

2.2 Valuation Functions

Each agent $i \in \mathcal{N}$ is equipped with a *valuation function* $v_i : 2^{\mathcal{R}} \rightarrow \mathbb{R}$ mapping the bundles she might receive to the reals. The preferences of agents will be based on these valuation functions. Of special interest is the simple case in which valuation functions are *additive*, i.e., in which there are no synergies between resources and agents can compute the value of a bundle of resources simply by adding up the values of the individual items in that bundle.

Definition 1 (Additive valuations). *A valuation function $v : 2^{\mathcal{R}} \rightarrow \mathbb{R}$ is said to be additive if $v(B) = \sum_{x \in B} v(\{x\})$ for all bundles $B \subseteq \mathcal{R}$.*

Valuations declared over bundles of resources naturally extend to valuations over allocations: we write $v_i(\alpha) := v_i(\alpha^{-1}(i))$ for the value agent i attaches to the bundle she would receive under allocation α . That is, we make the assumption that there are *no externalities*; the value an agent assigns to an allocation only depends on the set of resources she receives under that allocation.

2.3 Social Welfare Criteria

To measure the quality of an allocation, we can employ a number of criteria from the welfare economics and distributive justice literature. We focus on two criteria that can be formulated

in terms of collective utility functions, egalitarianism and classical utilitarianism (Moulin, 1988; Sen, 1970).

Definition 2 (Utilitarianism). *The utilitarian social welfare of an allocation α is defined as:*

$$sw_u(\alpha) = \sum_{i \in \mathcal{N}} v_i(\alpha)$$

Definition 3 (Egalitarianism). *The egalitarian social welfare of an allocation α is defined as:*

$$sw_e(\alpha) = \min_{i \in \mathcal{N}} v_i(\alpha)$$

Furthermore, we shall make use of the well-known criterion of Pareto efficiency.

Definition 4 (Pareto efficiency). *An allocation α is Pareto efficient if there exists no other allocation α' such that $v_i(\alpha') \geq v_i(\alpha)$ for all $i \in \mathcal{N}$ and $v_i(\alpha') > v_i(\alpha)$ for some $i \in \mathcal{N}$.*

2.4 Types of Decision-Makers

In this paper, we consider situations where both the bundles agents hold and the valuation functions they use to rate these bundles may be randomised at the end of negotiation. That is, each agent will receive a valuation function and a bundle from a finite set $X \subseteq (2^{\mathcal{R}} \rightarrow \mathbb{R}) \times 2^{\mathcal{R}}$ of “identities”, according to a *lottery* that is dependent on the final allocation of resources. For example, in case we uniformly randomise bundles (but leave valuation functions intact), the possible outcomes X for agent i will be the set of all pairs, where the first element is v_i and the second element is a bundle allocated to one of the agents in the current allocation.

Definition 5 (Lotteries and allocation-dependent lotteries). *Let $X \subseteq (2^{\mathcal{R}} \rightarrow \mathbb{R}) \times 2^{\mathcal{R}}$ be a finite set of pairs of valuation functions and bundles. An $\langle \mathcal{N}, X \rangle$ -lottery $\ell \in \Delta(X^{\mathcal{N}})$ is a probability distributions over the set $X^{\mathcal{N}}$ of all functions from the set of agents $\mathcal{N}$ to the set of outcomes X . An allocation-dependent $\langle \mathcal{N}, X \rangle$ -lottery $L : \mathcal{N}^{\mathcal{R}} \rightarrow \Delta(X^{\mathcal{N}})$ is a mapping from allocations to $\langle \mathcal{N}, X \rangle$ -lotteries.*

That is, an allocation-dependent lottery L together with an allocation α determines a lottery (probability distribution) over $X^{\mathcal{N}}$. We write L_α as a shorthand for $L(\alpha)$.

Given an agent’s valuation function, we cannot infer what that agent will prefer when it comes to choosing between lotteries. Depending on their attitude towards risk, two agents with the same valuation function may not pick the same lottery. In this paper, we will always assume that agents are aware of the current allocation and thus of the probability distribution that would be used if randomisation were to take place now. This assumption will simplify the following definitions.

Given an $\langle \mathcal{N}, X \rangle$ -lottery ℓ , we write $\ell(i, v, B) := \sum_{f \in X^{\mathcal{N}} : f(i) = (v, B)} \ell(f)$ for the probability that agent i will be assigned valuation v and bundle B under ℓ . We also write $\text{supp}_i(\ell) := \{(v, B) \in X \mid \ell(i, v, B) > 0\}$ for the set of support of ℓ from the point of view of agent i .

Definition 6 (Expected valuation). *Let ℓ be an $\langle \mathcal{N}, X \rangle$ -lottery. Then the expected valuation of ℓ for agent $i \in \mathcal{N}$ is defined as follows:*

$$\bar{v}_i(\ell) = \sum_{(v,B) \in X} \ell(i, v, B) \cdot v(B)$$

Definition 7 (Safety level). *Let ℓ be an $\langle \mathcal{N}, X \rangle$ -lottery. Then the safety level of ℓ for agent $i \in \mathcal{N}$ is defined as follows:*

$$SL_i(\ell) = \min_{(v,B) \in \text{supp}_i(\ell)} v(B)$$

We are now ready to define the types of decision makers we shall consider. Fixing the type of an agent amounts to fixing her preferences over alternative lotteries.

Definition 8 (Types of decision makers). *The type of an agent determines her preferences over alternative $\langle \mathcal{N}, X \rangle$ -lotteries:*

- (i) *An agent i is called a minmaximiser (an MM-agent) if she weakly prefers lottery $\ell \in \Delta(X^{\mathcal{N}})$ to lottery $\ell' \in \Delta(X^{\mathcal{N}})$ if and only if $SL_i(\ell) \geq SL_i(\ell')$.*
- (ii) *An agent i is called a expected-valuation maximiser (an EVM-agent) if she weakly prefers lottery $\ell \in \Delta(X^{\mathcal{N}})$ to lottery $\ell' \in \Delta(X^{\mathcal{N}})$ if and only if $\bar{v}_i(\ell) \geq \bar{v}_i(\ell')$.*

We furthermore say that an agent strictly prefers ℓ to ℓ' if she weakly prefers ℓ to ℓ' , and not *vice versa*.

2.5 Rational Negotiation

Negotiation takes place against the backdrop of an allocation-dependent lottery L . Whether or not an agent is willing to propose a particular deal $\delta = (\alpha, \alpha')$ and whether or not other agents are willing to accept that deal depends on their preferences over L_α and $L_{\alpha'}$, the lotteries induced by L and the allocations before and after δ . We say that such a deal is *rational* if all of the agents involved weakly prefer the new lottery $L_{\alpha'}$ over the old lottery L_α and if at least one of them (say, the proposer) prefers it strictly.

Definition 9 (Rational deals). *Given an allocation-dependent $\langle \mathcal{N}, X \rangle$ -lottery L , a deal $\delta = (\alpha, \alpha')$ is called rational if all involved agents $i \in \mathcal{N}^\delta$ weakly prefer $L_{\alpha'}$ over L_α and some involved agent $j \in \mathcal{N}^\delta$ strictly prefers $L_{\alpha'}$ over L_α .*

We assume that agents will only agree on deals that are rational in this sense. Therefore, which deals are possible depends on L and on the types of the agents involved. It is important to note that we are here assuming that agents are *not* strategising. The same kind of analysis could in theory be carried out for agents that do strategise, but this would call for a

drastically different (and more complex) game-theoretical framework, that would lead beyond our original goal of devising a simple and intuitive implementation of thought experiments that have been used to back up concepts of social welfare.

3 No Uncertainty: Negotiation in Multiagent Systems

We first consider the perfect information case, where no randomisation takes place. Here, L_α assigns to each agent i her valuation function v_i and her bundle $\alpha^{-1}(i)$ with certainty. This is the case that has been treated in previous work on negotiation in multiagent systems (Endriss et al., 2006; Chevaleyre et al., 2010). Note that in this case, the safety level and the expected valuation of a lottery will always coincide, so MM-agents and EVM-agents will behave in the same manner, and we do not need to make this distinction here.

Much of the work in the multiagent systems literature has addressed scenarios where agents can use monetary side payments when making deals. One such result shows that rational negotiation with side payments will always converge in an allocation with maximal utilitarian social welfare when agents have quasi-linear utility functions, composed of a valuation function as in our model here and a linear component representing money (Endriss et al., 2006). Other results show that when the range of valuation function is restricted (e.g., to additive functions), then structurally simple deals can be sufficient to obtain the same kind of convergence result (Chevaleyre et al., 2010).

For scenarios without money, i.e., the case we are interested in here, the known results are weaker. One basic observation is the fact that any sequence of rational deals will always converge to a Pareto efficient allocation. To see this, note that in the perfect information case allocation-dependent lotteries reduce to plain allocations, and in that case Definition 9 simply expresses that a deal $\delta = (\alpha, \alpha')$ is rational if and only if α' constitutes a Pareto-improvement over α . Thus, if we let agents negotiate rational deals until no more such deals are possible, they must have necessarily arrived at a Pareto efficient allocation.

The only known result regarding converge to an allocation that maximises egalitarian social welfare (Endriss et al., 2006), relies on a different (and arguably much less convincing) rationality criterion than the one given in Definition 9. That is, to obtain strong convergence results to allocations that are optimal according to criteria that are more demanding than Pareto efficiency we either need to introduce a monetary component into our model or we need to design new rationality criteria, which may not always be fully satisfactory. In this paper, instead, we will show that introducing a level of uncertainty into the model can also help us achieve such convergence results.

4 Reassigning Bundles

In the first scenario involving uncertainty that we shall consider the case where agents are told that the bundles they have assembled will get reassigned at the end of the negotiation process (but each agent keeps her own valuation function). Note that the composition of the bundles will remain intact; only the ownership of those bundles will change (or rather, the agents are made to believe that the ownership of bundles will change, as dictated by the allocation-dependent lottery in place). We then assume that the agents will evaluate the attractiveness of deals, and behave accordingly, believing that negotiation could stop and randomisation take place at the very next moment.

Besides the case in which no randomisation at all is taking place, this is the only case that can be implemented with real agents. It is therefore not only of theoretical, but also of some practical interest.

It turns out that for this scenario we cannot usually expect that agents will negotiate a socially optimal allocation, except in some very specific cases. In a nutshell, the reason is that, as agents do not know which bundle they will end up owning, they have no incentive to try and direct goods to those agents that will benefit the most from obtaining them. Rather, their interest is to achieve an optimal “spread” of the goods amongst the agents, making the worst bundle (in the case of MM-agents) or the average bundle (in the case of EVM-agents) attractive to themselves.

4.1 MM-Agents

We first consider MM-agents. If the resources have *the same value* for all agents, then as long as they all believe that they could receive any of the current bundles with non-zero probability the negotiation process will converge to an allocation maximising egalitarian social welfare.

Proposition 1. *Negotiation between MM-agents with the same valuation function who believe that they could receive any bundle currently allocated with non-zero probability will always converge to an allocation that maximises egalitarian social welfare.*

Now, if we drop this very strong assumption and allow agents to have different valuation functions, then we will not be able to show convergence to an egalitarian allocation anymore. Even for the case of negotiation between just *two agents* we are only able to establish *termination* of the process, but the terminal allocation need not have maximal egalitarian social welfare, nor does it need to be Pareto efficient.

Proposition 2. *Negotiation between two MM-agents believing that they could receive any bundle currently allocated with non-zero probability will always terminate, but the final allocation need not be Pareto efficient or have maximal egalitarian social welfare.*

If there are *three or more agents* involved, then we cannot even guarantee termination any longer. Instead, the negotiation process might loop, i.e., there might be an infinite sequence of deals that are rational for MM-agents.

Proposition 3. *Negotiation between three or more MM-agents believing that they could receive any bundle currently allocated with non-zero probability need not terminate.*

4.2 EVM-Agents

Next, we consider the scenario in which EVM-agents negotiate and each of them believes she may receive any of the currently allocated bundles with the same probability. First, observe that, in analogy to Proposition 1, it is not hard to see that in case all agents share the same valuation function, then we will observe convergence to an allocation with maximal utilitarian social welfare. We do not spell out the details here. Instead, consider the case of agents with *additive* valuation functions. Interestingly, in this case no rational negotiation is possible, in the sense that negotiation will terminate immediately. The reason is that, whatever the initial allocation, no deal will be rational for our agents under such circumstances.

Proposition 4. *Negotiation between EVM-agents with additive valuation functions believing that they could receive any bundle currently allocated with equal probability will terminate immediately.*

5 Reassigning Valuations

The next option to consider would be to randomise the valuation function at the end of the negotiation process, but to let the agents keep the resource bundles they have accumulated. Like the remaining scenarios we shall consider later on, this scenario, clearly, is not one that could actually be implemented in practice. While still potentially interesting as a thought experiment from a conceptual point of view, it turns out that at the technical level not much can be said about this particular scenario. Convergence to an allocation with attractive properties will typically not be possible.

We only include one simple technical observation here. For EVM-agents with *additive* valuations, we obtain a similar result as for the case in which we have only randomised bundles (Proposition 4). In this case, the expected value of a resource is the same for each agent (it is the average of the actual values each agent assigns to it), and thus, due to additivity, all agents assign the same expected value to any given bundle. Therefore, negotiation must terminate immediately, as for any deal, if one agent is strictly better off, at least one other agent needs to be strictly worse off.

Proposition 5. *Negotiation between EVM-agents believing that they could receive any of the valuation functions of the agents in society with equal probability will terminate immediately, if those valuation functions are additive.*

6 Reassigning Identities: Bundles and Valuations Together

In the setting we analyse next, agents are assumed to believe that their *identities* will be randomised at the end of the negotiation process. By that we mean that they believe that the current pairing of valuation functions and bundles will stay the same, but that they might receive a different pair than the one they are currently holding. In short, bundles and valuation functions are tied together. Naturally, this, again, has to be taken as a thought experiment. One can randomise bundles, but the agents have to believe that their preferences can also be randomised. If we do so, the connection works out in a mathematically neat way.

6.1 MM-Agents

Our first result for this setting shows that, if agents are in some sense maximally risk-averse (in the sense that they care only about the worst possible outcome when evaluating a lottery) and if they believe that their identities will be randomised (according to a lottery that is positive everywhere), then the negotiation process will always end up with an allocation that maximises egalitarian social welfare.

Proposition 6. *Negotiation between MM-agents believing that they could receive any current identity with non-zero probability will always converge to an allocation that maximises egalitarian social welfare.*

Proposition 6 establishes, for our specific mathematical framework, the kind of result we would expect to see in view of Rawls' *veil-of-ignorance* argument (Rawls, 1999). What our result can offer on top of the well-understood link between risk averseness and egalitarianism is a concrete implementation of this link, by providing a negotiation protocol that agents can use to agree on a final allocation of resources in a sequence of small steps, each of which could be initiated by anyone of them.

6.2 EVM-Agents

If agents that are only concerned with the expected value of the lottery engage in negotiation, then they will eventually converge to an allocation that maximises utilitarian social welfare, provided they think that the final randomisation will take place according to a *uniform* probability distribution.

Proposition 7. *Negotiation between EVM-agents believing that they could receive any current identity with equal probability will always converge to an allocation that maximises utilitarian social welfare.*

In the light of the work of Harsanyi (1953), this is the kind of link between expected-utility maximisation and classical utilitarianism does not come as a surprise. What is interesting about Proposition 7 is, again, that it offers a concrete implementation of this link in terms of a simple negotiation framework.

Observe that if all valuations are *additive*, then Proposition 7 can be strengthened to show that any sequence of rational deals involving only a single resource at a time will converge to an allocation with maximal utilitarian social welfare. We omit a formal statement (and the proof) of this simple result; technically it is closely related to analogous results on negotiation in multiagent systems without uncertainty (Chevaleyre et al., 2010).

7 Reassigning Bundles and Valuations Independently

The last remaining case we shall discuss is the scenario where we randomise *both* valuations and bundles, and where these two parameters are reassigned to the agents *independently* from each other. This introduces so much uncertainty into the system that we cannot expect agents to negotiate very attractive allocations. Indeed, for this scenario we have not been able to obtain any results that would establish convergence to a socially optimal allocation, even under strong restrictions on the valuation functions. However, it *is* possible to show, at least, that the outcome of negotiation will not be worse than a randomly chosen allocation. This is true for the utilitarian perspective and for EVM-agents.

Proposition 8. *Negotiation between EVM-agents believing that bundles and valuation functions will be reassigned independently according to uniform probability distributions will always converge to an allocation that has a utilitarian social welfare that is at least as high as the expected utilitarian social welfare we obtain if we choose an allocation at random using a uniform probability distribution over all possible allocations.*

While this form of convergence is very weak, it still does show that negotiation will, at least in expectation, have a positive effect.

8 Conclusion

We have argued that distributed negotiation under uncertainty can provide an attractive model for analysing the connections between theories of individual decision making, on the one hand, and theories of social welfare, on the other. We have borrowed a simple model of

negotiation from the literature on multiagent systems in Computer Science, in which myopic agents implement a series of exchanges of resources and accept an individual deal if and only if it does not diminish their immediate payoff. The core of our proposal is to enrich this model of negotiation with an element of uncertainty: the negotiating agents are made to believe that negotiation could stop at any point in time, after which some relevant parameters will get randomised, in a manner that is dependent on the allocation they have agreed upon. The two parameters we have considered here are the bundle of resources held by each agent and the valuation function used by each agent to evaluate those bundles. For each type of lottery, which determines how the relevant parameters may change, we can analyse the incentives of the agents in negotiation. If a particular set of assumptions on individual agent behaviour (e.g., risk averseness) and on the nature of the lottery will cause agents to always negotiate an allocation that is optimal according to a particular theory of social welfare, then, we argue, this can serve as a justification for that theory of social welfare.

The insight that we can link theories of individual decision making and theories of social welfare by introducing a certain degree of uncertainty is not new and goes back to the seminal work of Harsanyi and Rawls. Our interest here has been in investigating to what extent we can devise a concrete and simple implementation of this insight. Our model of negotiation provides this implementation. Our model also allows us to precisely pinpoint which types of uncertainty do and do not allow us to make the well-known connections between individual decision making and social welfare.

We have considered four different scenarios, each making a different combination of parameters subject to uncertainty. In broad terms, our findings are as follows. If only bundles get randomised, then negotiation will not usually converge to a socially attractive allocation; the only exception is the case where risk-averse agents with identical valuation functions engage in negotiation. If only valuation functions get randomised, then the situation is even more bleak; we have not been able to identify positive results for this scenario. If, on the other hand, bundles and valuations get reassigned in pairs, i.e., if we randomise identities, then risk averse agents will negotiate allocations that are optimal in view of egalitarian social welfare, while agents maximising their expected payoff will negotiate allocations that are optimal in view of utilitarian social welfare. Finally, if bundles and valuations get randomised independently, then we can at least establish a weak form of convergence showing that under these conditions, from an utilitarian point of view, negotiation between agents maximising their expected payoff is superior to selecting an allocation at random.

References

- Y. Chevaleyre, U. Endriss, and N. Maudet. Simple negotiation schemes for agents with simple preferences: Sufficiency, necessity and maximality. *Journal of Autonomous Agents and Multiagent Systems*, 20(2):234–259, 2010.
- U. Endriss, N. Maudet, F. Sadri, and F. Toni. Negotiating socially optimal allocations of resources. *Journal of Artificial Intelligence Research*, 25:315–348, 2006.
- J. C. Harsanyi. Cardinal utility in welfare economics and in the theory of risk-taking. *Journal of Political Economy*, 61:434–435, 1953.
- H. Moulin. *Axioms of Cooperative Decision Making*. Cambridge University Press, 1988.
- J. Rawls. *A Theory of Justice*. Harvard University Press, 2nd edition, 1999.
- A. K. Sen. *Collective Choice and Social Welfare*. Holden Day, 1970.

Appendix: Proofs

Proof of Proposition 1

We start by proving a lemma relating rationality and increases in social welfare:

Lemma 1. *Suppose all agents have the same valuation function v^* and let L be an allocation-dependent $\langle \mathcal{N}, X \rangle$ -lottery with the property that, for any allocation $\alpha \in \mathcal{N}^{\mathcal{R}}$, $L_\alpha(i, v, B) > 0$ if and only if $v = v^*$ and there exists an agent $j \in \mathcal{N}$ such that $B = \alpha^{-1}(j)$. Then the following two statements are equivalent:*

- (i) δ is a deal that is rational for MM-agents.
- (ii) δ is a deal that increases egalitarian social welfare.

Proof. Under the assumptions stated in the lemma, the set of support of a lottery L_α from the perspective of any agent i is $\text{supp}_i(L_\alpha) = \{(v^*, \alpha^{-1}(j)) \mid j \in \mathcal{N}\}$.

Now consider any rational deal $\delta = (\alpha, \alpha')$. As all valuation functions are identical, this is equivalent to saying that all agents *strictly* prefer the lottery $L_{\alpha'}$ to the lottery L_α . This, in turn is equivalent to $\min_{(v,B) \in \text{supp}_i(L_\alpha)} v(B) < \min_{(v,B) \in \text{supp}_i(L_{\alpha'})} v(B)$ for all $i \in \mathcal{N}$. But given our earlier remark about the set of support under the assumptions made here this is the same as $sw_e(\alpha) < sw_e(\alpha')$, so we are done. $\square$

Proposition 1 now follows easily:

Proof. By Lemma 1, every deal results in an increase in egalitarian social welfare. Thus, after a finite number of rational deals, negotiation is bound to terminate.

Now, for the sake of contradiction, assume that the terminal allocation α does *not* maximise egalitarian social welfare. Then there is an allocation α' such that $sw_e(\alpha') > sw_e(\alpha)$. But then Lemma 1 is again applicable and entails that the deal $\delta = (\alpha, \alpha')$ must be rational for MM-agents and thus possible, contradicting the assumption that α is terminal. $\square$

Proof of Proposition 2

Proof. We start by proving termination. For the sake of contradiction, assume the negotiation process does not terminate. Since there are only finitely many allocations, this means that there is a sequence of allocations $\alpha_1 \dots \alpha_k \dots \alpha_1$. But since there are only two agents, they have to be involved in every deal and thus every new allocation has to be weakly preferable to both of them and strictly preferable to one of them. This clearly is impossible.

A very simple example shows that the resulting allocation need not be either Pareto efficient or have maximal egalitarian social welfare. Suppose agent 1 holds item a and agent 2 holds item b , agent 1 prefers b to a and agent 2 prefers a to b (and further suppose both agents place no value on the empty bundle, so any allocation giving all items to just one agent is clearly unattractive). Then the agents do not have an incentive to swap the items they hold (since they believe that bundles will be randomised, they are indifferent between the two allocations), but the actual allocation is neither Pareto efficient nor does it maximise egalitarian social welfare. $\square$

Proof of Proposition 3

Proof. We give an example for a negotiation scenario with three MM-agents and four resources $\{a, b, c, d\}$ that permits an infinite sequence of rational deals. For simplicity we write a for $\{a\}$ and ab for $\{a, b\}$ etc. Suppose that all agents have strictly monotonic valuation functions (getting more resources is better) and that, as far as singletons are concerned, their valuation functions satisfy the following constraints:

- $v_1(c) < v_1(b) < v_1(d)$
- $v_2(d) < v_2(c) < v_2(b)$
- $v_3(b) < v_3(d) < v_3(c)$

Then the following sequence of deals will be rational:

	α_1	α_2	α_3	$\alpha_4 = \alpha_1$	$\dots$
Agent 1:	ab	b	b	ab	$\dots$
Agent 2:	c	ac	c	c	$\dots$
Agent 3:	d	d	ad	d	$\dots$

To see this, recall that only the (two) agents involved in a deal need to agree with that deal and that for MM-agents who believe that bundles will get randomised only the value of the least desirable bundle held by any agent is relevant. So, for instance, only agents 1 and 2 are involved in the deal from allocation α_1 to allocation α_2 , agent 2 is indifferent between the two allocations (the worst bundle, d , does not change), and agent 1 believes that b is more valuable than c , i.e., the worst bundle is improved by the deal. $\square$

Proof of Proposition 5

The idea of the proof is that, since agents are all evaluating according to the same “meta-valuation” and since there are no complementarity effects, there can be no win-win trade.

Proof. Given our assumptions, the expected value an agent i places on an allocation α is $\frac{1}{|\mathcal{N}|} \sum_{j \in \mathcal{N}} v_j(\alpha^{-1}(i))$. This is the average value (over the valuation functions of all the agents in the system) given to the bundle she obtains in allocation α .

Now, let α be the initial allocation. For the sake of contradiction, assume there exists another allocation β such that the deal $\delta = (\alpha, \beta)$ is rational. That is, one agent strictly prefers β over α and all others weakly prefer it. Thus:

$$\sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \sum_{j \in \mathcal{N}} v_j(\alpha^{-1}(i)) < \sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \sum_{j \in \mathcal{N}} v_j(\beta^{-1}(i))$$

For a bundle of resources B and a resource r , let $B(r) = 1$ if $r \in B$ and $B(r) = 0$ otherwise. Now, by virtue of the additivity of the valuation functions, we can rewrite above inequality as follows:

$$\sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \sum_{j \in \mathcal{N}} \sum_{r \in \mathcal{R}} \alpha^{-1}(i)(r) \cdot v_j(\{r\}) < \sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \sum_{j \in \mathcal{N}} \sum_{r \in \mathcal{R}} \beta^{-1}(i)(r) \cdot v_j(\{r\})$$

This can be simplified to yield:

$$\begin{aligned} \sum_{i \in \mathcal{N}} \sum_{j \in \mathcal{N}} \sum_{r \in \mathcal{R}} [\alpha^{-1}(i)(r) - \beta^{-1}(i)(r)] \cdot v_j(\{r\}) &< 0 \\ \sum_{j \in \mathcal{N}} \sum_{r \in \mathcal{R}} \left(\sum_{i \in \mathcal{N}} [\alpha^{-1}(i)(r) - \beta^{-1}(i)(r)] \right) \cdot v_j(\{r\}) &< 0 \end{aligned}$$

But the term in the large pair of parentheses must be equal to 0, because any resource r will occur in exactly one bundle belonging to any given allocation α or β . Thus, we have obtained a contradiction. $\square$

Proof of Proposition 6

We start by proving the following lemma:

Lemma 2. *Suppose L is an allocation-dependent $\langle \mathcal{N}, X \rangle$ -lottery with the property that, for any allocation $\alpha \in \mathcal{N}^{\mathcal{R}}$, $L_{\alpha}(i, v, B) > 0$ if and only if there exists an agent $j \in \mathcal{N}$ such that $v = v_j$ and $B = \alpha^{-1}(j)$. Then the following two statements are equivalent:*

- (i) δ is a deal that is rational for MM-agents.
- (ii) δ is a deal that increases egalitarian social welfare.

Proof. We first show that (i) implies (ii). Assume that $\delta = (\alpha, \alpha')$ is rational for MM-agents. Then, by Definitions 8 and 9, there exists an agent $j \in \mathcal{N}^{\delta}$ such that the following holds:

$$SL_j(L_{\alpha'}) > SL_j(L_{\alpha})$$

By Definition 7, this entails:

$$\min_{(v,B) \in \text{supp}_j(L_{\alpha'})} v(B) > \min_{(v,B) \in \text{supp}_j(L_{\alpha})} v(B)$$

By definition of set of support for a lottery, together with our assumption that precisely the identities (v, B) that can be found amongst the agents under the current allocation are those that have non-zero probability, this in turn entails:

$$\min_{i \in \mathcal{N}} v_i(\alpha') > \min_{i \in \mathcal{N}} v_i(\alpha)$$

Note that, due to the randomisation of identities, we were able to draw this conclusion concerning *all* agents $i \in \mathcal{N}$ from the first inequality above concerning only *one* agent j . By Definition 3, this last inequality is equivalent to $sw_e(\alpha') > sw_e(\alpha)$, so we are done.

For the other direction, observe that all of the above steps can be reversed. That is, from $sw_e(\alpha') > sw_e(\alpha)$, we can infer $SL_j(L_{\alpha'}) > SL_j(L_{\alpha})$ for *all* agents j (including all involved agents). Thus, $\delta = (\alpha, \alpha')$ must be rational whenever α' is an improvement over α in terms of egalitarian social welfare. $\square$

We are now ready to prove Proposition 6:

Proof. The proof parallels that of Proposition 1, using Lemma 2 in place of Lemma 1. $\square$

Proof of Proposition 7

We again start by proving a lemma relating rationality of deals to increases in social welfare:

Lemma 3. *Suppose L is an allocation-dependent $\langle \mathcal{N}, X \rangle$ -lottery with the property that, for any allocation $\alpha \in \mathcal{N}^{\mathcal{R}}$, $L_\alpha(i, v, B) = \frac{|\{j \in \mathcal{N} \mid v=v_j \text{ and } B=\alpha^{-1}(j)\}|}{|\mathcal{N}|}$. Then the following two statements are equivalent:*

- (i) δ is a deal that is rational for EVM-agents.
- (ii) δ is a deal that increases utilitarian social welfare.

Proof. We first show that (i) implies (ii). The proof is similar to the proof of Lemma 2. Assume that $\delta = (\alpha, \alpha')$ is rational for EVM-agents. Then there exists an agent $j \in \mathcal{N}^\delta$ such that the following is true:

$$\begin{aligned}
\bar{v}_j(L_{\alpha'}) &> \bar{v}_j(L_\alpha) \\
\sum_{(v,B) \in X} L_{\alpha'}(j, v, B) \cdot v(B) &> \sum_{(v,B) \in X} L_\alpha(j, v, B) \cdot v(B) \\
\sum_{(v,B) \in X} \frac{|\{i \in \mathcal{N} \mid v=v_i \text{ and } B=\alpha'^{-1}(i)\}|}{|\mathcal{N}|} \cdot v(B) &> \sum_{(v,B) \in X} \frac{|\{i \in \mathcal{N} \mid v=v_i \text{ and } B=\alpha^{-1}(i)\}|}{|\mathcal{N}|} \cdot v(B) \\
\sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \cdot v_i(\alpha') &> \sum_{i \in \mathcal{N}} \frac{1}{|\mathcal{N}|} \cdot v_i(\alpha) \\
\sum_{i \in \mathcal{N}} v_i(\alpha') &> \sum_{i \in \mathcal{N}} v_i(\alpha) \\
sw_u(\alpha') &> sw_u(\alpha)
\end{aligned}$$

Note, again, that we are moving from a statement about an individual agent to a statement about society as a whole. This is possible, because all agents evaluate the decision in the exact same way. Note also how we have made use of the assumption that the probability distribution over identities is uniform (leading to the term $\frac{1}{|\mathcal{N}|}$ on both sides of the inequality, which we were then able to eliminate).

The other direction is similar: above transformation can be applied from the bottom to the top; so from $sw_u(\alpha') > sw_u(\alpha)$ we can infer that $\bar{v}_j(L_{\alpha'}) > \bar{v}_j(L_\alpha)$ for *all* agents, i.e., $\delta = (\alpha, \alpha')$ will certainly be rational for a society of EVM-agents. $\square$

Proposition 7 now follows easily:

Proof. The proof parallels that of Proposition 1, using Lemma 3 in place of Lemma 1. (Note that requiring $L_\alpha(i, v, B) = \frac{|\{j \in \mathcal{N} \mid v=v_j \text{ and } B=\alpha^{-1}(j)\}|}{|\mathcal{N}|}$ is just another way of saying that all agents assign equal probability to being assigned any of the current identities.) $\square$

Proof of Proposition 8

Proof. We start by observing that negotiation will always terminate when agents believe that bundles and valuation functions will be independently randomised using a uniform probability distribution. The argument is as follows. As valuations are assumed to be randomised, each agent will evaluate any situation according to the same expected valuation function. And as bundles are assumed to be randomised, each agent will be only be interested in the partitioning of the goods into subsets, but not in who receives which bundle. Thus, with each deal the partitioning must change and improve strictly according to all agents, which is a process that is bound to terminate.

Hence, there exists a (at least one) terminal allocation. We want to prove that the following inequality holds for *any* terminal allocation α^* :

$$\frac{1}{|\mathcal{N}|^{|\mathcal{R}|}} \cdot \sum_{\alpha \in \mathcal{N}^{\mathcal{R}}} \sum_{i \in \mathcal{N}} v_i(\alpha^{-1}(i)) \leq \sum_{i \in \mathcal{N}} v_i(\alpha^{*-1}(i))$$

The righthand side of this inequality is the utilitarian social welfare of allocation α^* . The lefthand side is the average utilitarian social welfare over *all* possible allocations α . in other words, if we pick an allocation at random, using a uniform probability distribution, then this will be the expected utilitarian social welfare.

Let $\text{Perm}(\mathcal{N})$ be the set of all permutations $\sigma : \mathcal{N} \rightarrow \mathcal{N}$.

We shall derive above inequality starting from the following fact, which holds for any terminal allocation α^* and any arbitrary allocation α :

$$\frac{1}{|\mathcal{N}|!} \cdot \sum_{\sigma \in \text{Perm}(\mathcal{N})} \frac{1}{|\mathcal{N}|} \cdot \sum_{i \in \mathcal{N}} v_i(\alpha^{-1}(\sigma(i))) \leq \frac{1}{|\mathcal{N}|!} \cdot \sum_{\sigma \in \text{Perm}(\mathcal{N})} \frac{1}{|\mathcal{N}|} \cdot \sum_{i \in \mathcal{N}} v_i(\alpha^{*-1}(\sigma(i)))$$

To see that this is true, observe that the lefthand side of the inequality represents the expected valuation for any one of the agents (recall that they all have the same expected valuation function) for allocation α , given that they believe that any pair of valuation function and bundle is equally likely to be matched up (averaging over $|\mathcal{N}|!$ possible matchings) and that they are equally likely to receive any of the resulting identities (averaging over $|\mathcal{N}|$ identities). The righthand side is the corresponding value for α^* . By virtue of α^* being terminal, the agents do not want to move from α^* to α , i.e., the righthand value must be at least as great as the value on the left.

As above inequality holds for *any* allocation α , we can add up the corresponding $|\mathcal{N}|^{|\mathcal{R}|}$ inequalities to obtain the following new inequality (after first having multiplied both sides with $|\mathcal{N}|$ to eliminate the factor $\frac{1}{|\mathcal{N}|}$):

$$\sum_{\alpha \in \mathcal{N}^{\mathcal{R}}} \frac{1}{|\mathcal{N}|!} \cdot \sum_{\sigma \in \text{Perm}(\mathcal{N})} \sum_{i \in \mathcal{N}} v_i(\alpha^{-1}(\sigma(i))) \leq \frac{|\mathcal{N}|^{|\mathcal{R}|}}{|\mathcal{N}|!} \cdot \sum_{\sigma \in \text{Perm}(\mathcal{N})} \sum_{i \in \mathcal{N}} v_i(\alpha^{*-1}(\sigma(i)))$$

Next, observe that if we compose a permutation $\sigma \in \text{Perm}(\mathcal{N})$ and an allocation α^{-1} , then we obtain another allocation $\beta^{-1} = \sigma \circ \alpha^{-1}$. Furthermore, given an allocation β^{-1} , for every permutation σ there is a different allocation α^{-1} such that $\beta^{-1} = \sigma \circ \alpha^{-1}$. Thus:

$$\#\{(\sigma, \alpha) \in \text{Perm}(\mathcal{N}) \times \mathcal{N}^{\mathcal{R}} \mid \sigma \circ \alpha^{-1} = \beta^{-1}\} = |\mathcal{N}|!$$

That is, when in above inequality we are ranging over all allocations of the form $\sigma \circ \alpha^{-1}$ for all allocations α^{-1} and all permutations σ , we are in fact ranging $|\mathcal{N}|!$ times over each and every possible allocation. This insight allows us to eliminate σ from the lefthand side of above inequality and to rewrite it as follows:

$$\sum_{\alpha \in \mathcal{N}^{\mathcal{R}}} \sum_{i \in \mathcal{N}} v_i(\alpha^{-1}(i)) \leq \frac{|\mathcal{N}|^{|\mathcal{R}|}}{|\mathcal{N}|!} \cdot \sum_{\sigma \in \text{Perm}(\mathcal{N})} \sum_{i \in \mathcal{N}} v_i(\alpha^{\star-1}(\sigma(i)))$$

Now observe that the set of terminal allocations is closed under permutations. This is, again, because agents are only interested in the partitioning of the goods, not how the resulting bundles are assigned to agents. Therefore, relying on the fact that above inequality holds for *any* terminal allocation α^* , we can make the same kind of simplification also on the righthand side, and we obtain the desired inequality (again for any terminal allocation α^*):

$$\frac{1}{|\mathcal{N}|^{|\mathcal{R}|}} \cdot \sum_{\alpha \in \mathcal{N}^{\mathcal{R}}} \sum_{i \in \mathcal{N}} v_i(\alpha^{-1}(i)) \leq \sum_{i \in \mathcal{N}} v_i(\alpha^{\star-1}(i))$$

This concludes the proof. We have shown that the expected utilitarian social welfare of an allocation chosen at random (using a uniform probability distribution over all allocations) cannot be higher than the utilitarian social welfare of any allocation we might reach by allowing agents to negotiate. $\square$