

The Clarke Tax as a Consensus Mechanism Among Automated Agents

Eithan Ephrati

Jeffrey S. Rosenschein

Computer Science Department, Hebrew University

Givat Ram, Jerusalem, Israel

tantush@cs.huji.ac.il, jeff@cs.huji.ac.il

Abstract

When autonomous agents attempt to coordinate action, it is often necessary that they reach some kind of consensus. Reaching such a consensus has traditionally been dealt with in the Distributed Artificial Intelligence literature via the mechanism of negotiation. Another alternative is to have agents bypass negotiation by using a voting mechanism; each agent expresses its preferences, and a group choice mechanism is used to select the result. Some choice mechanisms are better than others, and ideally we would like one that cannot be manipulated by an untruthful agent.

One such non-manipulable choice mechanism is the Clarke tax [Clarke, 1971]. Though theoretically attractive, the Clarke tax presents a number of difficulties when one attempts to use it in a practical implementation. This paper examines how the Clarke tax could be used as an effective "preference revealer" in the domain of automated agents, reducing the need for explicit negotiation.

Background and Motivation

When autonomous agents attempt to coordinate action, it is often necessary that they reach some kind of consensus. Multi-agent activity is obviously facilitated by, and sometimes requires, agreement by the agents as to how they will act in the world. Reaching such a consensus has traditionally been dealt with in the Distributed Artificial Intelligence literature via the mechanism of negotiation [Rosenstein and Genssereth, 1985; Durfee, 1988; Sycara, 1988; Kuwabara and Lesser, 1989; Conry *et al.*, 1988; Kreifelts and von Martial, 1990; Kraus and Wilkenfeld, 1990; Laasri *et al.*, 1990].

One scenario [Zlotkin and Rosenschein, 1990b] that has been addressed in the research on negotiation involves a group of agents and a negotiation set. The role of negotiation is to reach consensus by allowing the agents to choose one element of this set. The main concern of a negotiation protocol is usually that the agreed-upon decision will be optimal in some sense.

A basic assumption of the negotiation process is that each of the participating agents has a private preference relation over the set of alternatives. Optimality is measured with respect to these preferences. Given the agents' preferences and the optimality criterion, determining the optimal choice is a matter of direct computation. Thus, the substantive role of the negotiation process is to reveal preferences. If there existed another method of revealing the true preferences of agents, the need for negotiation would be lessened.¹

There have been several attempts, both inside of Artificial Intelligence (AI) and outside, to consider market mechanisms as a way of revealing agents' true preferences (and thus efficiently allocate resources). Notable among the AI work is that of Smith's Contract Net [Smith, 1978], Malone's Enterprise system [Malone *et al.*, 1988], and the work of Miller and Drexler on Agoric Open Systems [Miller and Drexler, 1988]. Following in this line of work, we present an alternative method for revealing agents' preferences, the Clarke tax, and consider how it could be used among automated agents.

The General Framework

Assume a group of N agents $\mathcal{A}$ operating in a world currently in the state s_0 , facing the decision of what to do next. One way of formulating this problem is to consider that the agents are trying to agree into which member of the set $\mathcal{S}$ of m possible states the current world should be moved. Each agent in $\mathcal{A}$ has a *worth*, or utility, that he associates with each state; that worth gives rise to a preference relation over states. Agent i 's true worth for state k will be denoted by $W(i, k)$. However, the preferences *declared* by an agent might differ from his true preferences. The decision procedure that chooses one state from $\mathcal{S}$ is a function from the agents' declared preferences to a member of the set $\{1, \dots, m\}$. It maps the agents' declared preferences into a group decision as to how the world will be transformed.

¹This assumes the agents' preference relations are static during the negotiation process. Otherwise, the negotiation itself could cause the agents to acquire new information and alter their preferences, thus remaining useful.

There are many decision procedures that reach a pareto optimal decision, but they suffer from two major drawbacks. First, they are *manipulable*, which means that an agent can benefit by declaring a preference other than his true preference.² Second, they only take into consideration the *ordinal preferences* of the agents, i.e., the order in which an agent ranks choices, without assigning relative weights.

Attempts to overcome this latter drawback motivated the development of voting procedures based on *cardinal orderings* over alternatives (that is, allowing agents to weight their choices, including negative weights). The most straightforward procedure, “sealed bidding,” allows each voter to specify an amount of money (positive or negative) for each alternative. The alternative that has the maximal sum is chosen. Positive bids are then collected, and some of this money is then handed over to those agents (if any) who gave negative bids with respect to the chosen alternative.

Although a voter can guarantee his max-min value [Dubins, 1977], he does have an incentive to underbid—if he assumes other agents will cause some alternative to win even without the full strength of his vote, he can underbid, get what he wants, and pay less. However, since the agent might be mistaken as to how others will vote, a sub-optimal alternative might be chosen. In the literature of Economics, this problem is known as the *free rider problem*; for many years it was believed to be unsolvable.

A solution to the problem was presented by E. H. Clarke in 1971 [Clarke, 1971; Clarke, 1972; Straffin, 1980]. In the following sections, we present Clarke’s scheme and analyze ways in which it can be used by communities of automated agents.

The Clarke Tax

The basic idea of Clarke’s solution is to make sure that each voter has only one dominant strategy, telling the truth. This phenomenon is established by slightly changing the sealed-bid mechanism: instead of simply collecting the bids, each agent is fined with a tax. The tax equals the portion of his bid that made a difference to the outcome. The example in Figure 1 shows how to

²Unfortunately, a theorem due to Gibbard [Gibbard, 1973] and Satterthwaite [Satterthwaite, 1975] states that any non-manipulable choice function that ranges over more than two alternatives is dictatorial. This means that there is no choice function (other than one corresponding strictly to one of the agents’ preferences), that motivates all participating agents to reveal their true desires. This is related to Arrow’s famous “impossibility theorem” [Arrow, 1963], which showed how a group of reasonable criteria could not be simultaneously met by any social decision function (a function that produces a complete social preference ordering over alternatives, not just a single “best” choice). The technique presented in this paper, the Clarke tax, is distinguished in the literature as a “voting procedure” rather than as a pure decision function; it includes a kind of “incentive mechanism.”

calculate this tax. Each row of the table shows several pieces of information regarding an agent. First, his preferences for each state are listed. Then, the total score that each state would have gotten, had the agent *not* voted, are listed. An asterisk marks the winning choice in each situation.

	True worth of each state			Sum for each state without <i>i</i>			Tax for <i>i</i>
	<i>s</i> ₁	<i>s</i> ₂	<i>s</i> ₃	<i>s</i> ₁	<i>s</i> ₂	<i>s</i> ₃	
<i>a</i> ₁	27	−33	6	−46	*23	*23	0
<i>a</i> ₂	−36	12	24	*17	−22	5	12
<i>a</i> ₃	−9	24	−15	−10	−34	*44	0
<i>a</i> ₄	−18	−15	33	−1	*5	−4	9
<i>a</i> ₅	17	2	−19	−36	−12	*48	0
Sum	−19	−10	*29				

Figure 1: Calculating the Clarke Tax

For example, when all the agents voted, state *s*₃ was chosen. If *a*₂ had not voted, *s*₁ would have been chosen. The score in this situation would have been (17, −22, 5), and *s*₁ would have beaten *s*₃ by 12. Thus, agent *a*₂ has affected the outcome by his vote, and he has affected it by a “magnitude” of 12; he is therefore fined 12. Agents *a*₁, *a*₃, and *a*₅ are not fined because even if they had not voted, *s*₃ would still have been chosen.

Given this scheme, revealing true preferences is the dominant strategy. An agent that overbids (so that some given state will win) risks having to pay a tax larger than his true preferences warrant. Similarly, the only way to pay less tax is to actually change the outcome—and any agent that underbids (to change the outcome and save himself some tax) will always come out behind; the saved tax will never compensate him for his lost utility. For a proof that revealing true preferences is the dominant strategy, see [Ephrati and Rosenschein, 1991].

Using the Clarke tax in communities of automated agents brings into focus new problems that did not arise when the system was first developed. In the following sections we examine how the Clarke tax might be used as an effective “preference revealer” in the domain of intelligent agents, reducing the need for explicit negotiation.

Calculation of States and Preferences

We now specify our model and show how the Clarke tax can be used.

- Agents are capable of performing actions that transform the world from one state into another. Each action has an associated cost.
- The symbol *s* will stand for a set of predicates that demarcates a group of fully specified states. For simplicity, we will refer to *s* itself as a “state.”

- Each agent a_i has its own goal g_i , which is a set of predicates.
- $C(a, s \rightsquigarrow g)$ denotes the minimal cost that it would take for agent a , in state s , to bring about any state that satisfies g . $C(s_0 \rightsquigarrow s_1)$ is the minimal cost needed for moving the world from s_0 into s_1 , using any combination of agents' actions.
- $V(a, g)$ is the value that agent a assigns to goal g .

Since telling the truth is the dominant strategy when the Clarke tax is being used, it is in each agent's interest to compute the true worth he associates with each of the potential alternative states.

As an example, consider a simple scenario in the blocks world as described in Figure 2. There are three agents, with the following goals: $g_1 = \{At(G, 3), At(W, 2)\}$, $g_2 = \{On(W, G), On(R, W)\}$, $g_3 = \{On(B, W), At(W, 3)\}$. Assume that each *Move* action costs 1, and that $V(a_i, g_i) = C(a_i, s_0 \rightsquigarrow g_i)$. Thus, $V(a_1, g_1) = 2$, $V(a_2, g_2) = 3$, and $V(a_3, g_3) = 4$. As shown in Figure 2, the agents in state s_0 are faced with choosing among six alternative future states (we will later discuss how alternatives are to be generated).

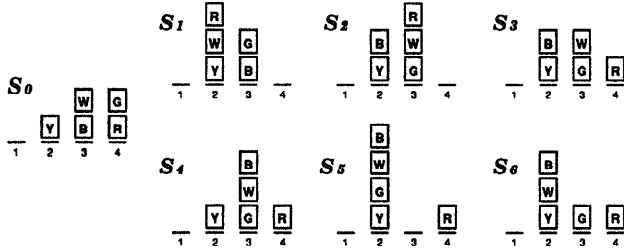


Figure 2: A Blocks World Example

Assessment of Worth

We suggest three conceptually different approaches for an agent to determine the worth of a given state.

According to the “all-or-nothing” approach, the agent assigns the full value of his goal to any state that satisfies it, and zero otherwise. In the example above, s_4 would be chosen, causing a_3 to pay a tax of 3. In the general case, the state that satisfies the single most valuable private goal will be chosen, unless there is a state that fully satisfies more than one goal. This approach suffers from the fact that an agent cannot assign relative weights to the alternatives, and no mutual compromise can be achieved.

A more flexible approach (“partial satisfaction”) is for the agent to give each state a worth that represents the portion of the agent's goal that the state satisfies, i.e., which predicates in the agent's composite goal are satisfied in the state. Assume that each of the agents' goal predicates contributes equally to the worth associated with a state. In the example, s_4 is again chosen, but a_3 pays a tax of only 1.5. This approach is superior in the sense that compromise can be achieved

via a state that partially satisfies a group of different goals. But in addition to preventing the agent from ranking bad alternatives (since there are no negative valuations), the method can be misleading. Consider, for example, a_2 . His evaluation of s_1 (1.5) is based on the fact that s_1 satisfies $On(R, W)$, while any attempt to achieve his *other* subgoal, $On(W, G)$, will require the violation of this predicate.

Yet a third approach (“future cost”) is to evaluate a state by taking into consideration the cost of the agent's eventually achieving his full goal, given that state: $W(i, k) = V(a_i, g_i) - C(a_i, s_k \rightsquigarrow g_i)$. Consider a_3 calculating the worth of s_1 . Given s_0 , he could achieve his goal using four *Move* operations; our assumption is thus that his goal's value is 4. Given s_1 , however, he would need five *Move* operations, *Move(R, 4)*, *Move(G, R)*, *Move(B, G)*, *Move(W, 3)* and *Move(B, W)*. He is therefore “worse off” by 1, and gives s_1 a worth of -1 . In the example above this yields the following true worths for each agent: $\langle 2, 0, 1, 0, -2, 2 \rangle$, $\langle 0, 3, 2, 1, 1, 0 \rangle$, $\langle -1, 2, 3, 4, 1, 1 \rangle$.

s_3 (which is only one *Move* operation distant from *all* the agents' goals) is chosen, and no tax is collected.

In some sense, this last method guarantees a “fair” consensus (where all agents are approximately equally distant from their ultimate goals). If it is important that some agent's goal be fully satisfied, a coin can be tossed to determine which of the agents will continue to fulfill his complete goal. Given a distribution of labor, the utility of an agent using this scheme may be greater than it would be if we had a lottery to select one agent, then let that agent bring about his own goal.³

The Generation of Alternatives

The selection of the candidate states (among which the agents will vote) plays a crucial role in the voting process. Given a group of agents with fixed goals, choosing different candidates can result in wildly different outcomes. The question thus arises of how these candidate states are to be generated. It is desirable that this generation be based upon each of the agents' goals. The generation of candidate states should aspire to choosing states maximal with respect to the satisfaction of agents' goals. Let $\mathcal{P}_A^U = \cup_{a_i \in A} (g_i)$ be the set of all the predicates appearing in all the agents' goals. Usually this does not specify a real-world state, since in the general case there are contradicting predicates among different agents' goals (otherwise, this state is guaranteed to be chosen).

We want it to be the case that each s_k in the set of candidate states satisfies the following definition:

$$s_k = \{p \mid p \in \mathcal{P}_A^U \text{ and } p \text{ is consistent with } s_k\}.$$

³See [Zlotkin and Rosenschein, 1990a] for an example of a similar scenario. Two agents agree to cooperate to an intermediate state that satisfies neither, then flip a coin to see who, alone, continues to his own goal. There, the cooperation is brought about by negotiation instead of voting.

Thus, each s_k is a maximal feasible subset of $\mathcal{P}_A^U$, a fixed point with respect to the predicates' consistency.

In order to check consistency, we assume a set of axioms over the domain predicates by which inconsistency can be discovered. In the example above we might have

$On(Obj, t) \Rightarrow At(Obj, t)$
 $[At(Obj_1, t) \wedge On(Obj_2, Obj_1)] \Rightarrow At(Obj_2, t)$
 $[At(Obj, t_1) \wedge At(Obj, t_2) \wedge (t_1 \neq t_2)] \Rightarrow False$

to establish the inconsistency of a set such as $\{At(W, 2), On(W, R), At(R, 3)\}$.

Note that this generation has several features. First, this procedure guarantees for each i the existence of at least one s_k such that $g_i \subseteq s_k$. Second, each agent is motivated to hand the generator his true goal. Declaring $\tilde{g}_i \supset g_i$ might prevent the generation of compromise states that benefit a_i , or cause the generation of states preferable to other agents (resulting in the selection of a worse alternative than otherwise would have been chosen). Declaring $\{\tilde{g}_i \mid (\tilde{g}_i \cap g_i) \subset g_i \text{ or } (\tilde{g}_i \cap g_i) = \emptyset\}$ may prevent the generation of any s_k that satisfies g_i , as well as preventing the generation of other states preferred by a_i which otherwise could have been chosen. In either case, a_i cannot hope to improve on his utility.

Note that the phase of candidate generation is completely distinct from the voting phase that follows it. An agent could declare goals that are used in generating candidates, and then vote in ways that contradict its declared desires. Note also that the technique above assumes the collection of information regarding agents' goals in a central location. This, of course, may be undesirable in a distributed system because of bottlenecks and communication overhead. [Ephrati and Rosenschein, 1991] develops several techniques for distributing the generation of alternatives among agents.

Additional Criteria

Candidate state generation can be refined by taking into consideration several additional criteria that avoid dominated states. These additional criteria sometimes depend upon the approach agents will be using to evaluate the worth of candidate states.

First, the generator can exclude states $\tilde{s}_k$ such that $\exists s_k [\forall i (W(i, k) \geq W(i, \tilde{k})) \wedge (C(s_0 \rightsquigarrow s_k) < C(s_0 \rightsquigarrow \tilde{s}_k))]$. The generator thus excludes a candidate state if there is another of equivalent value that is easier to reach. In the example, this test causes the elimination of the state $\{At(B, 3), At(G, 2), On(W, G), On(R, W)\}$ in favor of s_2 .

If the agents are going to evaluate candidate states using the "partial satisfaction" criterion, the generator can exclude $\tilde{s}_k$ such that $\exists s_k [(\tilde{s}_k \cap \mathcal{P}_A^U) \subset (s_k \cap \mathcal{P}_A^U)]$. The generator will exclude a candidate that specifies states that are a superset of another candidate's states. In the example, this would exclude s_3 in favor of s_2 .

If the agents are going to evaluate candidate states using the "future cost" criterion, the generator can eliminate states $\tilde{s}_k$ such that $\exists s_k [\forall i (C(a_i, s_k \rightsquigarrow g_i) \leq$

$C(a_i, \tilde{s}_k \rightsquigarrow g_i)) \wedge \exists i (C(a_i, s_k \rightsquigarrow g_i) < C(a_i, \tilde{s}_k \rightsquigarrow g_i))]$. The generator thus excludes a candidate that, for all agents, is "more expensive" than another candidate.⁴ In the example, such a test would eliminate the state s_5 in favor of s_4 .

One might suppose that if it is known ahead of time how candidate states are going to be evaluated, actually voting becomes redundant. By extension of elimination procedures such as those above, the generator could just compute the optimal state. For instance, using the "future cost" criterion, it might directly generate the s_k that minimizes $\sum_i^n C(a_i, s_k \rightsquigarrow g_i)$, and using the "partial satisfaction" criterion, it might directly choose the s_k that is the maximal (with respect to number of predicates) consistent subset of $\mathcal{P}_A^U$.

However, such extensions to the generation method are not always desirable. If the state generator uses them, the agents will sometimes be motivated to declare false goals. For example, if a_1 declares his goal to be $\{At(G, 3), At(W, 2), On(G, B), On(R, W)\}$ (whose predicates are a superset of his original goal), s_1 becomes dominant over all the other states if the generator uses either of the two global extensions considered above. Thus s_1 would automatically be chosen, and a_1 achieves a higher utility by lying.

Handling the Tax Waste

The Clarke tax itself must be wasted [Clarke, 1971; Clarke, 1972]; it cannot be used to benefit those whose voting resulted in its being assessed. To understand why, consider again the vote established by the "future cost" approach to the problem in Figure 2. As shown, a_1 's worth for the chosen state (s_3) is 1. However, knowing that he'll get a portion of the collected tax (as a compensation to the losers, or as an equally distributed share), a_1 would be motivated to understate his relative worth for s_3 , thus raising the total amount of tax—and his share of that tax. For example, his declaring $\langle 5, 0, -2, 0, -2, 2 \rangle$ would yield a total tax of 4 (s_2 is chosen causing a_2 and a_3 to pay 2 each).

Actually, the fact that the tax must flow away from the group prevents the decision from being pareto optimal—any taxpayer could improve its own utility by the amount of the tax without changing the other voters' utilities. Note, however, that in general the total tax will decrease as the number of agents increases. When there are more agents, the chance of any single one actually changing the decision approaches zero.

Our solution to this problem of tax waste is to use the tax for the benefit of agents *outside* the voting group. For that purpose, in parallel to goal generation, the entire society should be partitioned into disjoint voting groups (including one group of indifferent agents, A_0). When a voting group A_v completes the decision process, each taxed agent $a_i \in A_v$ has to dis-

⁴Actually, more or equally expensive for all, and more expensive for at least one.

tribute his tax equally among $a_j \in \mathcal{A} - A_v$. For convenience, we might make each a_i pay to a randomly chosen $a_j \in \mathcal{A} - A_v$.

As an example, consider again the blocks world scenario of Figure 2. Assume that s_0 is the same, but there are six agents operating in the world having the following private goals: $g_1 = \{At(B, 2)\}$, $g_2 = \{At(G, 3), On(G, W)\}$, $g_3 = \{On(G, W)\}$, $g_4 = \{At(R, 1)\}$, $g_5 = \{At(W, 4)\}$, $g_6 = \{At(W, 1)\}$.

The formation of voting groups can be based upon the agents' goals. The basic idea is to group agents who share common or conflicting interests in the same group. One obvious way to do this is to consider the resources over which the agents are competing. In our simple model, these resources can be the objects specified in the agents' goals. We can thus build each voting group to be the maximal set of agents whose goals share objects in common. Denoting the set of objects that are referred to in a set of predicates P by $O(P)$, we get: $A_v = \{a_i | O(g_i) \cap O(\mathcal{P}_{A_v}^U) \neq \emptyset\}$ (a fixed point). In the above example, such a grouping mechanism yields three groups: $A_1 = \{a_1\}$, $A_2 = \{a_2, a_3, a_5\}$, $A_3 = \{a_4, a_6\}$.

This grouping mechanism can be refined by taking into consideration only top level goals that share equivalent subgoals, or top level goals with conflicting subgoals (this could be done using a set of consistency axioms as shown in the previous section) such that $A_v = \{a_i | g_i \ominus \mathcal{P}_{A_v}^U \neq \emptyset\}$, where $p_1 \ominus p_2$ (where p_1 and p_2 are sets of predicates) stands for $(p_1 \cap p_2) \cup \{\tilde{p} | \tilde{p} \in p_1 \text{ and } \tilde{p} \text{ is inconsistent with } p_2\}$. Using this partition we get $A_1 = \{a_1\}$, $A_2 = \{a_2, a_3, a_5\}$ (since g_5 is inconsistent with g_2 , and g_3 shares $On(G, W)$ with g_2), $A_3 = \{a_4\}$, and $A_4 = \{a_6\}$.

A further refinement to the above approach is to also take into consideration the *actual plan* needed for achieving each goal, such that agents with interfering goals share the same voting group. The purpose of such a grouping mechanism is to allow agents with conflicting plans to "argue" about the plan to be carried out. Similarly, it allows agents whose plans share segments to vote for states that yield possibly cooperative plans. If, for example, we momentarily ignore the scheduling problem, and take into consideration plans that share mutual *Move* actions, we get two voting groups: $A_1 = \{a_1, a_2, a_3, a_4, a_5\}$, $A_2 = \{a_6\}$.⁵

Along with the added complexity of having to form voting groups, this solution to the tax waste problem might impose the need for an extensive bookkeeping mechanism, where each agents' debts are recorded. This would allow an agent to pay his debts by performing actions in pursuit of others' goals later on. The commitment to future action can remove the need

⁵ a_1 is in the same group as a_5 since $Move(W, 4)$ serves their mutual interests, and a_4 shares $Move(G, 3)$ with a_2 and a_3 . If scheduling is to be considered we would have only one group, since a_6 's $Move(Y, 1)$ may conflict with a_4 's $Move(R, 1)$.

for agents to share an explicit common currency.

Work Distribution

The Clarke tax mechanism assumes that the financing of any choice is equally divided among the voting agents.⁶ Since each agent declares its true "willingness to pay" for each alternative, one may be tempted to conclude that the agents' contribution to the creation of the chosen state should be based upon this willingness. Unfortunately, this does not maintain truth-telling as a dominant strategy. If an agent's stated preference is used to decide how to share the burden of a plan, the agents have an incentive to lie. To operate correctly, the share of work must be defined *a priori*.

There are several ways to determine each agent's share of work. The most convenient is to distribute the work equally among all members of the voting group, such that each agent has to contribute his share to the overall activity.

Another approach is to let agents vote on the work distribution. Instead of each s_k , the state generator has to generate a set S_K of alternatives. Each member of this set denotes a distinct state and work distribution. There are two drawbacks to this kind of procedure. First, the set of alternatives explodes combinatorially (instead of M we have, in the general case, $\sum_{k=1}^M N^{C(s_0 \rightsquigarrow s_k)}$ states). Second, if each action costs the same to all agents, and they are indifferent with regard to which specific action they take, then all $s_k \in S_K$ will get the same score, and one will have to be chosen arbitrarily (if $c(i, k)$ is the cost a_i has to pay as specified in s_k , then $W(i, K) = V(i, K) - c(i, k)$, and the score of each s_k is $(\sum_i V(i, K)) - C(s_0 \rightsquigarrow s_k)$).

A more desirable and just way to apportion work is to set each agent's share in direct relation to how fully the state in question satisfies his goal (as given to the candidate generator). One such measure could, for example, be $g_i \cap s_k$ (the actual portion of the goal satisfied by the state), or $C(a_i, s_0 \rightsquigarrow g_i \cap s_k)$ (the cost needed for a_i to accomplish the part of the state that he really wanted) or $C(a_i, s_0 \rightsquigarrow g_i) - C(a_i, s_k \rightsquigarrow g_i)$ (how much the state improves a_i 's position with respect to his goal).

Unfortunately, even though the work distribution has been set *a priori*, the Clarke mechanism fails. Realizing that his cost share is based upon such considerations, an agent is given an incentive to declare a false goal to the candidate generator. If the agent can ascertain the other agents' goals, he can benefit by declaring a goal that guarantees his participation in the voting group (depending on the grouping method) but very different from what he believes will be the group consensus. Thus, he may hope that the state generator will generate some states that he favors, while his predefined share of the work (based on his declared goal)

⁶ In classical voting theory the financing question is not addressed.

would be minimal or nonexistent.

Conclusions and Future Work

Group voting mechanisms can feasibly allow autonomous agents to reach consensus. In designing such a mechanism, issues that need to be addressed include automatic generation of alternatives over which the group will vote, assessment by each agent of the worth of each alternative, incorporation of an effective "incentive mechanism" for truth-telling (e.g., the Clarke tax, which must be spent outside the voting group and thus necessitates having distinct voting groups), and distribution of the labor once consensus has been reached.

There remain important issues to be resolved. First, the Clarke tax mechanism can be manipulated by coalitions of agents; techniques must be devised to deal with this. There are also questions related to when and how payment of debts might be enforced among autonomous agents (e.g., is it possible that an agent might continually join voting groups to avoid paying his debts), and alternative (iterative) methods of forming voting groups. These issues, along with implementation of this mechanism, remain for future work.

Acknowledgments

This research has been partially supported by the Israel National Council for Research and Development (Grant 032-8284), and by the Center for Science Absorption, Office of Aliya Absorption, the State of Israel.

References

- Arrow, K. J. 1963. *Social Choice and Individual Values*. John Wiley, New York. First published in 1951.
- Clarke, E. H. 1971. Multipart pricing of public goods. *Public Choice* 11:17-33.
- Clarke, E. H. 1972. Multipart pricing of public goods: An example. In Muskin, S., editor, *Public Prices for Public Products*. Urban Inst., Washington.
- Conry, S. E.; Meyer, R. A.; and Lesser, V. R. 1988. Multistage negotiation in distributed planning. In Bond, Alan H. and Gasser, Les, editors, *Readings in Distributed Artificial Intelligence*. Morgan Kaufmann Publishers, San Mateo, California. 367-384.
- Dubins, L. 1977. Group decision devices. *American Mathematical Monthly* 84:350-356.
- Durfee, E. H. 1988. *Coordination of Distributed Problem Solvers*. Kluwer Academic Publishers, Boston.
- Ephrati, Eithan and Rosenschein, Jeffrey S. 1991. Voting and multi-agent consensus. Technical report, Computer Science Department, Hebrew University, Jerusalem, Israel. In preparation.
- Gibbard, A. 1973. Manipulation of voting schemes: a general result. *Econometrica* 41:587-602.
- Kraus, S. and Wilkenfeld, J. 1990. The function of time in cooperative negotiations: Extended abstract. In *Proceedings of the Tenth Workshop on Distributed Artificial Intelligence*, Bandera, Texas.
- Kreifelts, Thomas and von Martial, Frank 1990. A negotiation framework for autonomous agents. In *Proceedings of the Second European Workshop on Modeling Autonomous Agents in a Multi-Agent World*, Saint-Quentin en Yvelines, France. 169-182.
- Kuwabara, Kazuhiro and Lesser, Victor R. 1989. Extended protocol for multistage negotiation. In *Proceedings of the Ninth Workshop on Distributed Artificial Intelligence*, Rosario, Washington. 129-161.
- Laasri, Brigitte; Laasri, Hassan; and Lesser, Victor R. 1990. Negotiation and its role in cooperative distributed problem solving. In *Proceedings of the Tenth International Workshop on Distributed Artificial Intelligence*, Bandera, Texas.
- Malone, T. W.; Fikes, R. E.; and Howard, M. T. 1988. Enterprise: A market-like task scheduler for distributed computing environments. In Huberman, B. A., editor, *The Ecology of Computation*. North-Holland Publishing Company, Amsterdam. 177-205.
- Miller, M. S. and Drexler, K. E. 1988. Markets and computation: Agoric open systems. In Huberman, B. A., editor, *The Ecology of Computation*. North-Holland Publishing Company, Amsterdam. 133-176.
- Rosenschein, Jeffrey S. and Genesereth, Michael R. 1985. Deals among rational agents. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence*, Los Angeles, California. 91-99.
- Satterthwaite, N. A. 1975. Strategy-proofness and Arrow's conditions: existence and social welfare functions. *Journal of Economic Theory* 10:187-217.
- Smith, Reid G. 1978. *A Framework for Problem Solving in a Distributed Processing Environment*. Ph.D. Dissertation, Stanford University.
- Straffin, Philip D. Jr. 1980. *Topics in the Theory of Voting*. The UMAP Expository Monograph Series. Birkhäuser, Boston.
- Sycara, Katia P. 1988. Resolving goal conflicts via negotiation. In *Proceedings of the Seventh National Conference on Artificial Intelligence*, St. Paul, Minnesota. 245-250.
- Zlotkin, Gilad and Rosenschein, Jeffrey S. 1990a. Negotiation and conflict resolution in non-cooperative domains. In *Proceedings of The National Conference on Artificial Intelligence*, Boston, Massachusetts. 100-105.
- Zlotkin, Gilad and Rosenschein, Jeffrey S. 1990b. Negotiation and goal relaxation. In *Proceedings of The Workshop on Modelling Autonomous Agents in a Multi-Agent World*, Saint-Quentin en Yvelines, France. Office National D'Etudes et de Recherches Aerospatiales. 115-132.