

Measuring the benefits of lying in MARA under egalitarian social welfare*

Jonathan Carrero¹, Ismael Rodríguez^{1,2}, and Fernando Rubio^{1,2} †‡

Abstract

When some resources are to be distributed among a set of agents following egalitarian social welfare, the goal is to maximize the utility of the agent whose utility turns out to be minimal. In this context, agents can have an incentive to lie about their actual preferences, so that more valuable resources are assigned to them. In this paper we analyze this situation, and we present a practical study where genetic algorithms are used to assess the benefits of lying under different situations.

Keywords: Social Welfare; Genetic Algorithms; Complexity.

1 Introduction

Computer Science and Economics have been successfully applied to each other for decades. Economic concepts and mechanisms can be applied to deliver scarce resources in a computational environment (see e.g. [1, 2]), while Computer Science has been applied to identify the difficulty of developing several economic tasks (see e.g. [3, 4, 5, 6]).

Actually, some of the challenges one finds in the problem of delivering resources are the same regardless of the nature of agents (natural or artificial). Multiagent Resource Allocation (MARA) deals with the problem of delivering some resources to some agents, given their preferences over the resources. The definition of a particular MARA environment depends on many factors, and the first decision to be addressed is what the goal of the distribution is. For instance, the goal may be maximizing the revenue of the particular agent making

*This paper was published in 2020 IEEE International Conference on Systems, Man, and Cybernetics (SMC). The present version is the author's accepted manuscript. This work has been partially supported by projects TIN2015-67522-C3-3-R, PID2019-108528RB-C22, and by Comunidad de Madrid as part of the program S2018/TCS-4339 (BLOQUES-CM) co-funded by EIE Funds of the European Union.

†¹Dpto. Sistemas Informáticos y Computación, Facultad de Informática, Universidad Complutense de Madrid, 28040 Madrid, Spain. joncarre@ucm.es, isrodrig@sip.ucm.es, fernando@sip.ucm.es

‡²Instituto de Tecnologías del Conocimiento, Universidad Complutense de Madrid, 28040 Madrid, Spain.

the distribution (or *auctioneer*) by means of the payments done by bidders for the resources or bundles of resources (called *Winner Determination Problem* [7] in the latter case). Alternatively, if the goal is satisfying the preferences of the agents these resources are assigned to, other alternatives unfold depending on the kind of society satisfaction (*social welfare*) to be pursued. For instance, the goal may be maximizing the addition of utilities of all agents these resources are assigned to, which is the so-called *utilitarian* social welfare; or it may consist in maximizing the utility of the agent whose utility turns out to be minimal, which is called *egalitarian* social welfare (think, for instance, in distributing humanitarian resources after a disaster); or other criteria may be considered (see e.g. [8]). Regarding how agent preferences are denoted, agents may just assign a numerical value to each available resource and assume *additive* preferences (i.e. the utility of a bundle of resources is the addition of the utilities of all resources in the bundle). Alternatively, non-additive settings where bundles of resources can be given utilities different to the simple addition (higher or lower) can be denoted in several ways [9]. Regarding the process followed to achieve the desired delivery, we may assume that it is performed in one-shot by some party (say an auctioneer) that initially receives the preferences of all agents, or that the desired delivery is the result of iterating local interactions between agents according to some local rules [10].

Let us consider deliveries performed in one shot. Note that agents may have an incentive to communicate false preferences to the auctioneer in order to maximize their utility. For instance, if the utilitarian social welfare is the goal of the delivery, then agents may benefit by pretending the utility they get for each resource is higher than it is, as the addition of utilities of agents is maximized by giving resources to those agents wanting them the most. In order to cope with this problem, the Generalized Vickrey Auction (GVA) [11] introduces a payment mechanism to eliminate any incentive for communicating preferences different to the true ones. In this auction, the distribution of resources maximizing the addition of utilities is found and assigned to agents, but in addition each agent has to pay an amount equivalent to the *loss of utility* of the other agents due to the participation of the agent. More specifically, this payment is the subtraction of the utility of the others in a distribution where the agent did not exist, *minus* the utility of the others according to the actual distribution of resources. After taking this payment into account, the final utility of each agent turns out to be the addition of utility of *all* agents (due to the resources the agent receives and, on the other side, the payment *reduction* due to the utility of the other agents in the actual distribution) *minus* another term the agent has no control at: the payment due to the distribution that would be reached if the agent did not participate. Thus, to the extent that can be controlled by the agent, the utility of the agent (due to the resources it receives and due to the payment) is maximized by maximizing the addition of utility of *all* agents, which is actually maximized by communicating the true preferences. This property, usually called *strategy-proofness*, eliminates the incentive for strategic behavior in the GVA — although this mechanism has other drawbacks (see e.g. [12, 13], where problems related with cheating, disclosure of confidential information, etc. are discussed).

To the best of our knowledge, no similar “falsehood disincentivizing” mechanism has been found for the case of the egalitarian social welfare. In fact, the intrinsic discontinuities of the function to be maximized in this case hinder this goal, as the utility of the agent with minimum utility dramatically changes by introducing slight resource allocation changes. Apparently, we could hardly make the utility of agents coincide with the egalitarian social welfare by introducing (rational, not extremely distorting) payments, like GVA does with the utilitarian social welfare. Note that not disincentivizing the falsehood is a problem indeed: if the distribution mechanism just delivers the resources in such a way the egalitarian social welfare is maximized, then agents do have a trivial incentive for lying. In fact, an agent just has to do the opposite as we said for the utilitarian social welfare: by falsely *reducing* the values of the utilities given by resources, an agent increases its chances to be the agent with minimum utility, and in turn the agent whose utility is to be maximized by the distribution. Thus, agents can trivially increase their profit by underestimating their preferences.

Despite the apparent impossibility to fully eliminate the incentive for lying (i.e. sending false preferences) in this kind of distribution, it is easy to notice that some simple rule changes can make it harder to get profit by lying or, at least, to *certainly* get profit by lying. For instance, if all agents are forced to make the addition of their utilities for individual resources be equal to some fix constant, then agents cannot arbitrarily underestimate their preferences for *all* resources, as underestimating some resource implies the risky action of overestimating another. Even if we cannot *fully* disincentivize false preferences under the egalitarian social welfare, we can still measure to what extent some modifications of the rules do it—and check out if some of them virtually disincentivize false preferences *in practice*. Note that some ways of lying could have a higher probability of being profitable than others. Moreover, let us assume the agent willing to lie has an estimation of the preferences the other agents will communicate (regardless of whether they are true or not). Then, the optimal set of *false preferences* for this agent is the set such that, if it is communicated to the auctioneer *and* the other agents also send the preferences estimated by this agent, then the distribution of resources maximizing the egalitarian social welfare will also maximize its *true* utility—maybe beyond the utility achieved by sending the true preferences. Note that this so-called optimal set of false preferences could not be so good if the preferences sent by the other agents turn out to be different to those estimated by the agent. Actually, some deviation (i.e. estimation imprecision) is expected in any real setting, so the suitability of that alleged optimal set of false preferences might strongly depend on the degree of that deviation: maybe using these false preferences is better than saying the truth if that deviation is low, but it is not if the deviation is high. Moreover, some distribution mechanism rules (e.g. forcing the addition of utilities to be equal to some constant) may make the distributions for these optimal false preferences be more *sensitive* to those deviations—to the extent that, even under small and likely deviations, saying the truth is better than lying in general. In this case, we could say that the incentive for lying is *practically* eliminated under

those particular rules.

The goal of this paper will be investigating, by means of experiments, if we can disincentive lying in practice in the context of egalitarian social welfare. Experiments will be conducted where some agent lies according to some predefined lying strategies or by using the lie that would be optimal provided that the other agents behaved exactly as estimated by the agent. We will assume a basic egalitarian social welfare MARA setting where the preferences are additive, although our experiments could be easily extended to consider other more complex preferences. Actually, as maximizing the egalitarian social welfare is NP-hard, solving each instance of this optimization problem will require approximate algorithms (in particular, we will choose genetic algorithms, although any other metaheuristic could be used, for instance [14, 15, 16]). In a nutshell, our experiments will show that, if the addition of preferences is required to be equal to some given constant, then (a) no predefined lying strategy is profitable; and (b) using the optimal lie according to some estimation of the preferences of the other agents will be worse than saying the truth when that estimation is even slightly imprecise.

The rest of the paper is organized as follows. In the next section we formally introduce the problems under consideration. Section 3 reports our experimental results, and conclusions are given in Section 4.

2 Formal model

In this section we introduce the formal notions and problems under consideration in this paper.

Definition 1 Let $R = \{r_1, \dots, r_m\}$ be a set of resources and $A = \{A_1, \dots, A_n\}$ be a set of agents.

The set of possible allocations of R to A is the set $\mathcal{A} = \underbrace{A \times \dots \times A}_m$. Given $\alpha \in \mathcal{A}$ with $\alpha = (\alpha_1, \dots, \alpha_m)$, we say that the allocation α assigns each resource r_j to agent α_j for all $1 \leq j \leq m$. $\square$

A utility function will be a function receiving a distribution of resources and returning a real number. As usual, utility functions will be used to denote the preferences of agents for bundles of resources: if the utility function of an agent is u and $u(\alpha_1) > u(\alpha_2)$, then the agent prefers distribution α_1 over distribution α_2 .

Definition 2 A utility function is a function $u : \mathcal{A} \rightarrow \mathbb{R}^{\geq 0}$.

We say that u *depends only* on agent A_i if, for all $\alpha = (\alpha_1, \dots, \alpha_m) \in \mathcal{A}$ and $\beta = (\beta_1, \dots, \beta_m) \in \mathcal{A}$ fulfilling $\alpha_j = \beta_j = A_i$ or $\alpha_j \neq A_i \neq \beta_j$ for each $1 \leq j \leq m$, we have $u(\alpha) = u(\beta)$. In addition, we also say that u is *additive* for agent A_i if for all $\alpha = (\alpha_1, \dots, \alpha_m) \in \mathcal{A}$ we have $u(\alpha) = \sum_{\{j \mid \alpha_j = A_i\}} u(\underbrace{A_k, \dots, A_k}_{i-1}, A_i, \underbrace{A_k, \dots, A_k}_{m-i})$, where A_k is any arbitrary agent with $A_k \neq A_i$. $\square$

Note that, if u is additive for A_i , then we can simply denote u with a vector $P = (p_1, \dots, p_n)$ with $p_i \in \mathbb{R}^{\geq 0}$ representing the individual utilities added by receiving each resource. Formally, u is unambiguously constructed from P as follows: for all $\alpha = (\alpha_1, \dots, \alpha_m)$, $u(\alpha) = \sum_{\{j \mid \alpha_j = A_i\}} p_j$. Thus, an additive utility function u can be simply represented by that vector $P = (p_1, \dots, p_n)$. Given $r \in \mathbb{R}^{>0}$, we will say that an additive utility function u is *r-limited* if $\sum_{1 \leq i \leq n} p_i = r$ (in Section 3, experiments assuming and not assuming this constraint will be conducted).

Regarding non-additive utility functions, there are several ways to represent the preferences for *bundles* of resources: extensionally defining the value of all possible *subsets* of resources, defining values to be added or subtracted if specific k -combinations of resources are owned, etc. This representation choice affects the complexity of any optimization problem based on them [9]. Hereinafter, only additive utility functions will be considered in this paper.

Let us denote by $\mathcal{U}$ the set of possible additive utility functions. The utility functions of all agents in a tuple of agents A will be denoted by a tuple $U = (u_1, \dots, u_n) \in \mathcal{U}^n$, where each u_i is the utility function of agent A_i .

We define the optimization problem of distributing resources in such a way that the utility of the agent receiving less utility is maximized. In the next definition, the auxiliary term $\mathbf{eg}_{A,R,U}(\alpha)$ denotes the utility of the less satisfied agent when the distribution of resources is α .

Definition 3 Given $A = (A_1, \dots, A_n)$, $R = (r_1, \dots, r_m)$ and $U = (u_1, \dots, u_n)$, for all $\alpha \in \mathcal{A}$ let $\mathbf{eg}_{A,R,U}(\alpha) = \min\{u_j(\alpha) \mid 1 \leq j \leq n\}$. The egalitarian social welfare optimization problem consists in, given A, R, U , finding α maximizing $\mathbf{eg}_{A,R,U}(\alpha)$. We will denote the solution of the problem for A, R, U by $\mathbf{egsol}(A, R, U)$. $\square$

The NP-hardness of the previous problem is proved in [17], and it cannot be approximated in polynomial time within a factor of $\beta > 1/2$ unless $P = NP$ [18].

Proposition 1 *Let us consider a variant of the egalitarian social welfare optimization problem where an additional input $r \in \mathbb{R}^{\geq 0}$ is given and all considered utility functions must be r -limited. The resulting problem is also NP-hard.*

Proof. In [17], the NP-hardness of the problem in Definition 3 is proved by a simple polynomial reduction from NP-hard problem PARTITION into it. In this reduction, the constructed instance of that problem contains only two agents having the same utility function, which in turn shows that the egalitarian optimization problem is NP-hard even under the additional constraint of taking only instances with two agents whose utility functions coincide. This implies the NP-hardness of a variant of the problem in Definition 3 where only r -limited utility functions, for some $r \in \mathbb{R}$ given as problem input, can be used (i.e. when, for each agent, the addition of the valuations of all resources must be equal to a given constant r), as the aforementioned variant dealing with two agents having equal preferences (which is NP-hard) can be trivially polynomially reduced to

that r -limited problem: note that the addition of preferences of both agents are necessarily equal to some value v , so we can just set $r = v$. $\square$

Let us call *auctioneer* to the third party responsible for, given the utility functions of all agents, finding the distribution of resources maximizing the egalitarian social welfare. Next we introduce the problem of finding the optimal fake utility function of a given agent, that is, the (potentially false) utility function letting that agent maximize its utility in an egalitarian social welfare distribution. That is, we search for the utility function such that, if agent A_i sends it to auctioneer *and* the utility functions sent to the auctioneer by the other agents are those actually estimated by agent A_i , then the utility received by agent A_i (with its *true* utility function) in the distribution maximizing the egalitarian social welfare is maximized. In the next definition, term $\mathbf{fake}_{A,R,U,i}(f)$ denotes the utility achieved by agent A_i if the utility function it sends to the auctioneer is f . Note that calculating this term implies solving the optimization problem given in Definition 3, so *maximizing* that term means performing, in turn, an optimization of a term whose calculation requires another optimization.

Definition 4 Given A, R, U as before and $i \in \mathbb{N}$, for all $f \in \mathcal{U}$ let $\mathbf{fake}_{A,R,U,i}(f) = u_i(\alpha_f)$ with $\alpha_f = \mathbf{egsol}(A, R, (u_1, \dots, u_{i-1}, f, u_{i+1}, \dots, u_n))$. The optimal fake utility problem consists in, given A, R, U , and i , finding f maximizing $\mathbf{fake}_{A,R,U,i}(f)$. We will denote the solution of the problem for A, R, U, i by $\mathbf{fakesol}(A, R, U, i)$. $\square$

Despite the fact that the previous problem intuitively consists in performing an optimization of a term whose calculation requires another optimization, we can see that solving it as it is defined is not difficult at all.

Proposition 2 *Let the utility functions of each agent A_j estimated by agent A_i (and the utility of agent A_i itself when $j = i$) be $P^j = (p_1^j, \dots, p_n^j)$. The optimal fake utility function for A_i is $P^i = (c \cdot p_1^i, \dots, c \cdot p_n^i)$ where $c \in \mathbb{R}^{>0}$ is any positive value such that $\sum_{1 \leq s \leq n} c \cdot p_s^i < p_l^k$ for all $k \neq i$ and l with $p_l^k > 0$.*

Proof. First, let us suppose that there exists some allocation of resources $\alpha \in \mathcal{A}$ such that $\mathbf{eg}_{A,R,U}(\alpha) > 0$ (i.e. there is some allocation of resources giving a non-null utility to all agents). We show that the optimal fake utility function for A_i is the proposed one (note that c always exists). These fake preferences of agent A_i are such that, even if the auctioneer gave all resources to A_i , then A_i would still obtain less utility than the utility reached by any other agent by receiving any resource this agent has a non-null preference for. Note that, in this case, the goal of the auctioneer (i.e. maximizing $\mathbf{eg}_{A,R,U}(\alpha)$) coincides with maximizing the *fake* utility of A_i (i.e. preferences P^i) *provided that* some non-null utility is given to the other agents (note that, if this constraint is impossible to achieve, then it contradicts our previous assumption that $\mathbf{eg}_{A,R,U}(\alpha) > 0$ for some α). On the other hand, note that the goal of A_i is maximizing its *true* utility provided that *the same* constraint holds (as allocations not giving a non-null utility to all agents will never be picked by the auctioneer). Note that maximizing $P^i = (p_1^i, \dots, p_n^i)$ subject to that constraint is equivalent to

maximizing $P^i = (c \cdot p_1^i, \dots, c \cdot p_n^i)$ subject to the same constraint. Thus, the optimal strategy for A_i consists in sending preferences P^i to the auctioneer.

In the (degenerate) case where there does not exist α with $\mathbf{eg}_{A,R,U}(\alpha) > 0$, agent A_i has no control over the allocation of resources picked by the auctioneer, because all of them are equally good according to the goal of the auctioneer (as any allocation α yields $\mathbf{eg}_{A,R,U}(\alpha) = 0$). Thus, preferences P^i are as optimal for A_i as any other preferences. $\square$

Let us consider the problem given in Definition 4 under the additional constraint that all agent preferences are required to be r -limited for some $r \in \mathbb{R}^{>0}$ given as problem input. In this case solutions cannot be trivially constructed as before (note that P^i would not respect the constraint). Actually it seems that, in this case, the optimization of a term whose calculation requires another optimization *does* make the resulting problem much harder. We conjecture the resulting (decision) problem is Σ_2^P -complete.

3 Experiments

In this section we describe several experiments we have conducted to check the usefulness of lying in an egalitarian social welfare scenario. That is, we study what happens when an agent sends to the auctioneer a preference vector P different to its real preferences. From now on, we will denote by T (True) the vector of real preferences of a given agent for the available resources. On the other hand, we will denote by M (Modified) the modified preference vector, i.e. the one actually communicated to the auctioneer.

We will consider three stages in our experiments. First, we will analyze the usefulness of some predefined lying strategies (e.g. increase/decrease our preferences about the resources we value the most or the resources valued the most by other users, etc.); second, we will use genetic algorithms to dynamically search for the optimal false preferences for each specific problem instance (instead of composing them according to predefined rules as before); and third, we will analyse the performance of our lies when the information we have about the preferences of other users is not perfect.

According to the restrictions the agents have on the preferences they are allowed to communicate, for each of the aforementioned stages we consider two scenarios. In the first one, denoted “unlimited scenario,” each agent must establish a preference for each resource so that this value is within the interval $(0, 100)$. In the second case, denoted as “limited scenario,” in addition to the previous restriction, the sum of all the preferences is required to be equal to 100.

In all experiments, it is assumed that Agent 1 is the lying agent we are analysing. In each case, this agent will try to achieve the greatest possible utility by modifying its communicated preferences. We performed different sets of experiments with different numbers of agents and available resources, and we observed the same trends and patterns in all of them. Hence, for the sake of conciseness and specificity in the data exposition, we will only show the

Table 1: Real preferences of the agents in the unlimited scenario

	1	2	3	4	5	6	7	8	9	10
Agent 1	42.36	19.18	75.37	42.32	60.20	68.98	85.30	20.66	31.67	60.41
Agent 2	53.70	96.73	70.87	68.47	86.95	51.34	3.06	30.01	55.03	45.42
Agent 3	88.49	17.88	70.73	98.71	30.42	40.71	98.41	85.19	42.08	57.05
Agent 4	95.08	50.57	69.31	26.01	56.03	64.09	55.08	6.62	49.24	31.69

Table 2: Real preferences of the agents in the limited scenario

	1	2	3	4	5	6	7	8	9	10
Agent 1	17.67	12.58	4.35	12.00	5.47	0.77	14.57	3.49	16.55	12.504
Agent 2	1.927	6.09	19.66	18.17	10.51	10.75	2.42	3.31	23.04	4.07
Agent 3	16.95	12.24	11.64	7.33	7.63	12.57	9.87	8.73	11.33	1.66
Agent 4	14.41	2.20	16.78	1.03	13.15	9.70	5.18	5.46	17.19	14.85

results for the case study where ten resources are to be distributed among four agents. Table 1 contains the real preferences of the four agents in the unlimited scenario, whereas Table 2 contains the preferences of the agents in the limited case. As we said before, our experiments will analyse what happens when Agent 1 communicates different preferences. However, in any case, the usefulness obtained by Agent 1 will be computed by using its actual preferences, that is, those shown in the corresponding row of Table 1 and Table 2.

3.1 Testing generic lying strategies

In order to achieve a greater benefit, Agent 1 may lie when it communicates its preferences to the auctioneer. Note that, when Agent 1 communicates its lie, in general it does not certainly know if that lie will provide it with a greater benefit, since at that moment the distribution of resources has not been made yet. Next we investigate if Agent 1 has some kind of lying strategy to follow that will generally provide it with profitable lies.

The aim of our first set of experiments is to test a set of predefined lying strategies. For instance, is it a good strategy to always communicate lower (or greater) preferences for the item we are most interested in? Is it useful to follow the same strategy for the item we are less interested in (or for the items supposed to be less valued by other agents)?

We have tested different strategies. For each of them, we compare the results obtained when the agent does not lie and when the agent uses such strategy to lie up to different *lying levels* (higher levels mean higher deviations from the corresponding true values). Table 3 shows some examples of strategies that have been applied.

In order to calculate the benefit (or loss) Agent 1 would obtain by communicating its lie, first we must know the utility it obtains when it communicates

Table 3: Summary of basic lying strategies

Test	Description
1	Decrease the preference for random resources
2	Increase the preference for random resources
3	Increase preference for the resources most valued by other agents
4	Increase preference for the resources less valued by other agents
5	Increase preference for the resources most valued by itself
6	Increase preference for the resources less valued by itself
7	Decrease preference for the resources most valued by other agents
8	Decrease preference for the resources less valued by other agents
9	Decrease preference for the resources most valued by itself
10	Decrease preference for the resources less valued by itself

its real preferences. Since obtaining the optimal resource allocation according to the egalitarian social welfare (see Definition 3) is an NP-hard problem, we use a genetic algorithm to deal with this problem. The algorithm in charge of finding the solution to this problem will be called LLGA (*Lower-Level Genetic Algorithm*). In our examples, the population of LLGA will have 50 individuals, running 50 iterations and using tournament selection.

On the other hand, we also have to solve the same problem when actually using the corresponding lie, which is given by the preference vector M . Note that doing so will potentially provide a different solution. To know the benefit or loss that Agent 1 obtains by using its lie, it is enough to subtract the utility it obtains with the second solution (i.e. when it is using its vector M) minus the utility it obtains with the initial solution (i.e. when it is using its vector T). If the result of such subtraction is positive, then it is worth communicating the lie to the auctioneer.

Let us start with the most simple scenario, that is, the unlimited case. In the following figures, the y-axis represents the profit (or loss) obtained by Agent 1, whereas the x-axis represents the degree of application of the strategy (i.e. the amount of increase/decrease in which the communicated preferences differ from the preferences it actually has for the resources). These degrees of application are represented by 100 progressive increases/decreases.

Figure 1 shows the results of applying our first two tests, that is, decreasing and increasing, respectively, the preferences concerning some random resources. As it can be seen, when we start modifying a little bit the preferences of some of the resources, none of the strategies (decreasing or increasing preferences) seems to work. However, increasing preferences is clearly worse, as some negative peaks start appearing near from the beginning. These peaks in the graph are due to the amount of resources we are working with. Let us recall that we only have ten resources and that varying just one of them in a solution may make results suffer great variations (if more resources were considered, then these peaks would

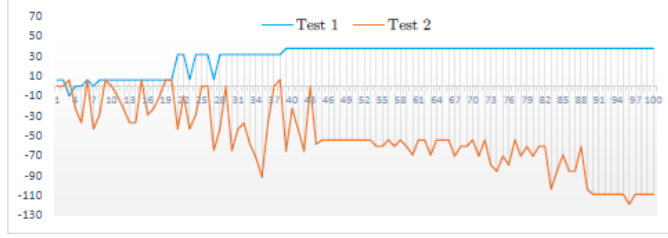


Figure 1: Results obtained with generic lying strategies in the unlimited scenario: decreasing (test 1) or increasing (test2) preferences

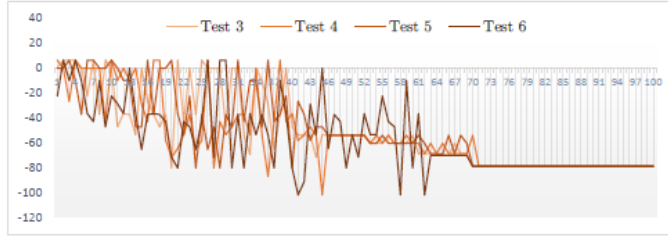


Figure 2: Results obtained with generic lying strategies in the unlimited scenario: different strategies increasing preferences

soften and we would see much more progressive variations). At the same time, as Agent 1 keeps modifying its preferences more, we observe how the strategy followed in Test 2 implies greater losses of utility. However, during Test 1 the utility tends to stabilize. These strategies suggest that communicating a lie with higher preferences with respect to the preference vector T seems to be a bad strategy for Agent 1 to follow. On the other hand, as Test 1 lies more about the preferences with respect to the preference vector T , the utility obtained is considerably increased, which provides a benefit to Agent 1.

Let us look at some more tests where the preference for resources is increased. Figure 2 shows that these tests do not provide any benefit to Agent 1. It seems that none of the lies in which higher preferences are communicated to the auctioneer provide benefit, which reinforces the results of Test 2 that we saw earlier.

On the other hand, Figure 3 shows some strategies in which Agent 1 communicates preferences that are lower than its real preferences. As we intuited, all strategies that involve decreasing preferences work, and they also become better as we consider a greater amount of resources. As we see, some of these strategies are better than others, especially Test 10, where Agent 1 decreases preferences for the resources it values less.

These last strategies make us wonder if the utility achieved during Test 10 is the highest possible. Despite being very close, it is not the maximum. By decreasing the preferences for all resources, a slightly better result is obtained.

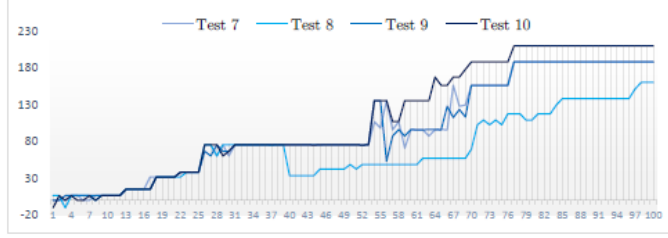


Figure 3: Results obtained with generic lying strategies in the unlimited scenario: different strategies decreasing preferences

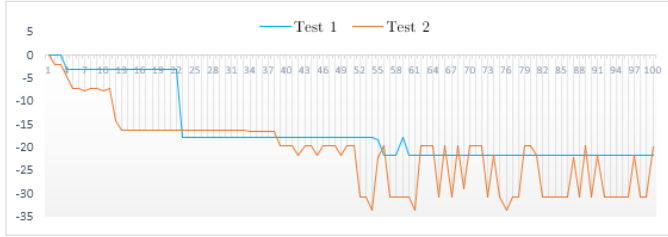


Figure 4: Results obtained with generic lying strategies in the limited scenario: decreasing (test 1) or increasing (test2) preferences

As we will see later, this last test verifies the best strategy that Agent 1 has to lie (accordingly to Proposition 2).

All in all, these tests clearly indicate that it is trivial to lie in the unlimited scenario, because a good strategy to follow by Agent 1 is to decrease its preferences for as many resources as possible. However, when we deal with the *limited* scenario, it is not possible to decrease the preferences concerning all the resources. Thus, each time a preference is lowered, it is necessary to increase the preference about another resource, so that the total addition keeps constant. Figure 4 shows the results obtained with the same strategies as before, but in the case of the limited scenario. If we look at it, we realize that the new constraint makes it very difficult to find a strategy that brings benefits to Agent 1. In fact, all of them have turned out to be unsuccessful strategies, producing a reduction in the satisfaction of Agent 1. Even when Agent 1 decreases the preferences for the resources it values the most (see Figure 6), it has to increase its preferences for other resources. Thus, the auctioneer can assign Agent 1 some resources which are not really very useful for that agent.

In other words: there does not seem to be a generic strategy to be followed by Agent 1 in the limited scenario so that it consistently improves its outcome when it lies. However, this does not necessarily imply that lying is completely discouraged—it only means that generic, predefined strategies do not seem to be useful. However, for each *specific* case study it could be possible to dynamically

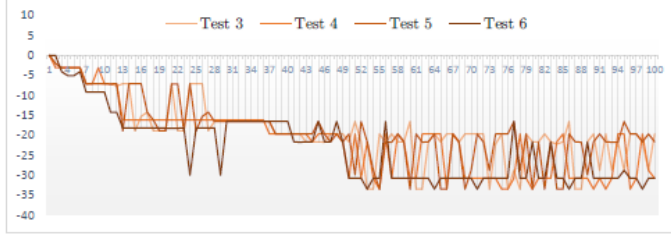


Figure 5: Results obtained with generic lying strategies in the limited scenario: different strategies increasing preferences

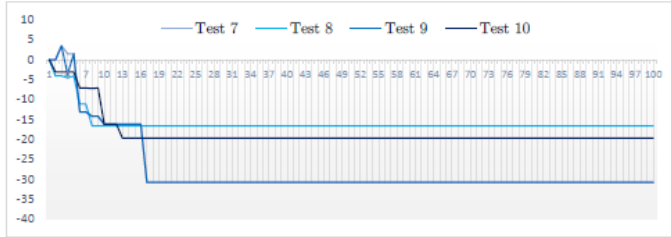


Figure 6: Results obtained with generic lying strategies in the limited scenario: different strategies decreasing preferences

find a tailored lie that provides a higher utility to Agent 1. In fact, in our next experiments we will analyse such a possibility.

3.2 Looking for the best lie

We have seen that generic lying strategies are useful in the unlimited scenario (in fact, reducing the valuation of all resources is a good strategy in any unlimited scenario), but we have not been able to find a good strategy in the case of the limited scenario.

Next we will try to construct *ad hoc* lies to deal with each specific configuration. In fact, now our aim is to find the *optimal* lie for each case study. That is, we want to find the best lie Agent 1 can communicate, i.e. the lie with which it would achieve the greatest possible utility. Obtaining this utility in the unlimited scenario is trivial: regardless of the number of agents and resources we have, Agent 1 achieves the highest utility when it proportionally underestimates the valuation of its own preferences so much as to make the auctioneer think Agent 1 will be the least satisfied agent in any distribution giving a non-null utility to all agents (see Proposition 2). In this case, Agent 1 will be assigned the $n-m$ resources that it values the most with its real preferences, where n is the number of available resources and m is the number of other agents. Note that giving some satisfaction to the other agents requires giving at least one resource to each agent. However, in order to maximize the utility received by

Agent 1, the rest of the agents must receive those resources Agent 1 values the least. By taking to the extreme the strategy of decreasing its preferences when lying, Agent 1 will achieve the remaining resources, which are those it values the most.

In order to find the best lie, we must solve the problem of finding the utility function f maximizing $\mathbf{fake}_{A,R,U,i}(f)$, where f denotes the false preferences to be sent to the auctioneer (see Definition 4). The resolution of this problem, as we already saw, involves making an optimization of an optimization. The genetic algorithm in charge of finding this best lie is called ULGA (*Upper-Level Genetic Algorithm*). In this algorithm, each individual of the population represents a possible lie of Agent 1. That is, an individual is a vector M representing a possible valuation of resources Agent 1 can communicate to the auctioneer. The fitness function used in ULGA to evaluate the quality of individuals (i.e. of candidate false preferences) requires evaluating the usefulness of the corresponding lie M to provide profit to Agent 1. Thus, for each individual M , computing its fitness firstly requires computing the optimal distribution of resources the auctioneer can perform according to the rules of egalitarian social welfare. Consequently, for each lie M in the population of ULGA, we use algorithm LLGA to compute the distribution of resources picked by the auctioneer when the preferences sent by Agent 1 are those given in M . The fitness of each possible individual (lie) M in ULGA is computed by adding the *real* values Agent 1 would get for each resource assigned to Agent 1 in the distribution of resources found by LLGA (note that those *real* values are those in vector T , not in M).

To check that our method obtains good solutions, we start with the most simple case, the unlimited scenario. Then, we will move to the most interesting case, the limited scenario.

The population of 50 individuals was evolved 3,000 times using tournament selection. Both the optimal lie (preferences) for Agent 1 and the reached distributions can be checked in Table 4 (in the distributions, numbers denote the agent each resource is assigned to). As it can be seen, all resources but three are assigned to Agent 1. Moreover, those three resources are the ones Agent 1 likes the least. Hence, ULGA actually found the solution providing the best distribution of resources for Agent 1, according to the obvious theoretical limit that at least one resource has to be assigned to each agent. Let us remark that ULGA does not need to decrease *forever* the preferences of Agent 1 for the resources in order to achieve continuously better lies. Actually, the best possible distribution of resources for Agent 1 is reached when its valuations are *low enough* (note that this experimental observation is consistent with Proposition 2).

Regarding the limited scenario, the population evolved the same number of times, but in this case the utility benefit is relatively lower, obtaining only a profit of a little bit more than 10. This is something we already suspected would happen, since none of the basic predefined strategies worked in our first experiments, as we saw in the previous section. Anyway, the experiments prove it is still profitable to lie in the limited scenario, although it is clearly harder to find a useful lie to be communicated to the auctioneer. Table 5 shows the

Table 4: Best lie obtained by ULGA in the unlimited scenario

	1	2	3	4	5	6	7	8	9	10
Real A1	42.36	19.18	75.37	42.32	60.20	68.98	85.30	20.66	31.67	60.41
Distribution	4	2	1	3	2	4	1	3	4	1
Lie A1	2.21	2.27	10.57	3.12	6.5	5.55	12.09	2.18	2.27	4.87
Distribution	1	4	1	1	1	1	1	3	2	1

Table 5: Best lie obtained by ULGA in the limited scenario

	1	2	3	4	5	6	7	8	9	10
Real A1	17.67	12.58	4.35	12.00	5.47	0.77	14.57	3.49	16.55	12.504
Distribution	1	3	2	2	4	3	4	3	1	4
Lie A1	13.53	9.52	14.35	9.86	6.1	8.83	9.03	5.43	12.2	11.11
Distribution	1	3	2	2	4	3	1	3	1	4

preference vector T of Agent 1 and the best lie found by the ULGA, as well as the corresponding distributions for each kind of preference. Let us remark that the best lie obtained is very different with respect to the real preferences. Moreover, there is not a clear general strategy to move from the real preferences to the best lie, as it depends on small adjustments related with the preferences of the rest of agents.

3.3 Analysing imprecise information

The previous experiments have shown that, even in the limited scenario, it is possible to find ad hoc lies making Agent 1 reach a higher profit if communicated to the auctioneer. However, the ad hoc solutions found by ULGA require that Agent 1 knows the preferences of the rest of agents, as they are needed to evaluate the usefulness of each lie when applying LLGA. Hence, a new question arises: is it still a good idea to use the lie obtained by ULGA when we are not completely sure as to whether the preferences of the rest of agents are those actually assumed by Agent 1 and introduced in the algorithm? That is, what happens when our predictions about the rest of agents are not good enough?

In this section we report some experiments conducted to analyze the aforementioned problem. We used ULGA to find the optimal lie Agent 1 has to communicate to the auctioneer, assuming the preferences of the other agents are those considered by Agent 1. Then, we evaluated the usefulness of this lie when such preferences are not correct. We evaluate different situations depending on the accuracy of the prediction of the preferences of the other agents assumed by Agent 1. In particular, we randomly generated the *real* preferences of the other agents from the preferences assumed by Agent 1 by applying a normal distribution whose average value is at the assumed preference. We generated several new instances where the standard deviation of such distribution is small, next other instances where it was slightly bigger, and so on for a variety of standard deviations.

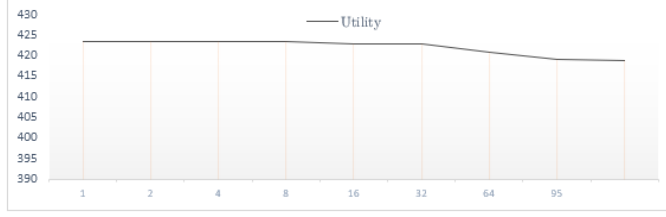


Figure 7: Lying in the unlimited scenario with imprecise information

In order to have a better view of how the utility of Agent 1 is affected by the imprecision, Figures 7 and 8 show the utility improvement obtained when using the lie found by ULGA in the context of imprecise information. The x -axis represents the standard deviation of the normal distributions used to randomly construct the actual preferences of the other agents from the corresponding assumed preferences.

For each new deviation, 1,000 alternative utility scenarios are created by using the corresponding deviation. For each of them, we computed the benefit of using the real preferences of Agent 1 (i.e. not lying) and the benefit Agent 1 obtains by using the lie provided by ULGA when it uses the preferences of the rest of agents assumed by Agent 1 (that is, the preferences before applying the normal distribution). Then, we compute the average behaviour of these 1,000 cases.

Again, let us firstly discuss the unlimited scenario. Figure 7 shows the obtained results. Specifically, the graph shows how the utility obtained by Agent 1 evolves when it uses its lie and, at the same time, we are increasing the deviation yielding the actual preferences of the other agents. As it can be seen, the line stays relatively horizontal as we modify the deviation until $\sigma = 32$, which implies that the deviation produces the same effect when applied on the different preferences. From this point on, the utility is slightly decreased. This means that Agent 1's lie continues to be efficient (at least, more efficient than telling the truth) even when the preferences of the other agents differ from those it estimated. Note that, even when we apply a normal distribution with $\sigma = 99.0$, the utility is still clearly above the utility obtained with the preference vector T of Agent 1 (which is 221.08 and can be calculated by using Table 4). This result shows that, independently of the estimation made by Agent 1 about the preferences of the other agents, it is always worth lying to the auctioneer.

If we move to the limited scenario, in Figure 8 we see very different results from those obtained in the unlimited scenario. This time, let us denote by A the total utility achieved by Agent 1 when using the lie provided by ULGA. In addition, let us denote by A' the total utility achieved by Agent 1 when it uses its real preference vector T . If we calculate the result of $A' - A$, we will see that after a certain deviation it is not worth for Agent 1 to communicate its lie to the auctioneer. Initially, the result of the subtraction keeps positive for small deviations. However, even with a relatively small deviation ($\sigma = 4$), the profit

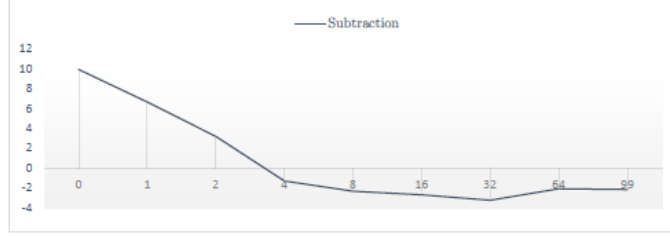


Figure 8: Lying in the limited scenario with imprecise information

obtained by Agent 1 when lying is worse than when saying the truth. This result shows that it is worth lying when the preferences of the other agents are very similar to those Agent 1 estimated but, otherwise, there is not incentive to try to lie. That is, according to the experimental results, in the limited scenario it is worth trying to lie only provided that the estimations about the preferences of the rest of agents are very close to their actual preferences.

4 Conclusions and future work

The experimental results we have obtained in the unlimited scenario let us conclude that the best strategy to be used by Agent 1 is to lie by communicating a high underestimation of its preferences for the available resources (which is consistent with Proposition 2). This way, the auctioneer will be forced to assign the $n-m$ resources that it values the most to Agent 1, so that Agent 1 will achieve the greatest possible utility. On the other hand, the tests carried out in the limited scenario allow us to conclude that, in this case, none of the considered predefined strategies works. It seems that there is no strategy to be followed by Agent 1 other than simply testing many combinations by decreasing and increasing its preferences. In spite of this, the limited scenario does not succeed in completely eliminating the lie, because our ULGA algorithm manages to find fake preferences for Agent 1 providing a higher profit than sending the true preferences. Also, lying in the unlimited scenario remains the best option even when the estimation made by Agent 1 about the preferences of its mates differ greatly from the preferences these agents actually have. In other words: in the unlimited scenario, it is always worth lying. On the contrary, in the limited case there seems to be no clear strategy to lie. In fact, we can only find good lies to be communicated to the auctioneer provided that our estimations about the preferences of the rest of participants are really good. Otherwise, relatively small deviations from our estimations make our lies obtain results being worse than those achieved when revealing our real preferences.

Summarizing, although lying is not completely disincentivized in our limited scenario, we can conclude that, in practical terms, lying is not useful in such scenario, as we would need very detailed information about the preferences of

the rest of participants. Thus, a simple rule modification consisting in forcing the addition of all preferences be equal to some constant can be very useful in practical terms to avoid lying in the egalitarian social welfare.

References

- [1] R. Buyya, D. Abramson, and S. Venugopal, “The grid economy,” *Proceedings of the IEEE*, vol. 93, no. 3, pp. 698–714, 2005.
- [2] X. León, T. A. Trinh, and L. Navarro, “Using economic regulation to prevent resource congestion in large-scale shared infrastructures,” *Future Generation Computer Systems*, vol. 26, no. 4, pp. 599–607, 2010.
- [3] I. Rodríguez, F. Rubio, and P. Rabanal, “Automatic media planning: optimal advertisement placement problems,” in *2016 IEEE Congress on Evolutionary Computation (CEC)*. IEEE, 2016, pp. 5170–5177.
- [4] I. Rodríguez, P. Rabanal, and F. Rubio, “How to make a best-seller: Optimal product design problems,” *Applied Soft Computing*, vol. 55, pp. 178–196, 2017.
- [5] S.-H. Chen, M. Kaboudan, and Y.-R. Du, *The Oxford handbook of computational economics and finance*. Oxford University Press, 2018.
- [6] G. Monaco, L. Moscardelli, and Y. Velaj, “On the performance of stable outcomes in modified fractional hedonic games with egalitarian social welfare,” in *AAMAS’19*, 2019, pp. 873–881.
- [7] D. Lehmann, R. Müller, and T. Sandholm, “The winner determination problem,” *Combinatorial auctions*, pp. 297–318, 2006.
- [8] Y. Chevaleyre, P. Dunne, U. Endriss, J. Lang, M. Lemaître, N. Maudet, J. Padget, S. Phelps, J. Rodriguez-Aguilar, and P. Sousa, “Issues in multi-agent resource allocation,” *Informatica*, vol. 30, pp. 3–31, 2006.
- [9] N. Nguyen, T. Nguyen, M. Roos, and J. Rothe, “Computational complexity and approximability of social welfare optimization in multiagent resource allocation,” *Autonomous Agents and Multi-Agent Systems*, vol. 28, no. 2, pp. 256–289, 2014.
- [10] U. Endriss, N. Maudet, F. Sadri, and F. Toni, “Negotiating socially optimal allocations of resources,” *CoRR*, vol. abs/1109.6340, 2011.
- [11] L. M. Ausubel, “A generalized Vickrey auction,” *Univ. Maryland*, 1999.
- [12] T. W. Sandholm, “Limitations of the vickrey auction in computational multiagent systems,” in *ICMAS-96*, 1996, pp. 299–306.
- [13] M. H. Rothkopf, “13 reasons why the vickrey-clarke-groves process is not practical,” *Operations Research*, vol. 55, no. 2, pp. 191–197, 2007.

- [14] B. Xing and W.-J. Gao, “Innovative computational intelligence: a rough guide to 134 clever algorithms,” 2014.
- [15] P. Rabanal, I. Rodríguez, and F. Rubio, “Applications of river formation dynamics,” *Journal of computational science*, vol. 22, pp. 26–35, 2017.
- [16] F. Rubio and I. Rodríguez, “Water-based metaheuristics: How water dynamics can help us to solve np-hard problems,” *Complexity*, vol. 2019.
- [17] M. Roos and J. Rothe, “Complexity of social welfare optimization in multiagent resource allocation,” in *AAMAS’10*, 2010, pp. 641–648.
- [18] I. Bezáková and V. Dani, “Allocating indivisible goods,” *SIGecom Exch.*, vol. 5, no. 3, pp. 11–18, Apr. 2005.