

All Stable Equilibria Have Improved Performance Guarantees in Submodular Maximization with Communication-Denied Agents

Joshua H. Seaton and Philip N. Brown

Abstract—This paper considers the robustness of game-theoretic approaches to distributed submodular maximization problems, which have been used to model a wide variety of applications such as competitive facility location, distributed sensor coverage, and routing in transportation networks. Recent work showed that in this class of games, if k agents suffer a technical fault and cannot observe the actions of other agents, Nash equilibria are still guaranteed to be within a factor of $k + 2$ of optimal. However, our paper shows that at a Nash equilibrium with a very low objective function value, the total payoffs of compromised agents are very close to the payoffs they would receive at an optimal allocation. At the extreme worst-case equilibria, all agents are perfectly indifferent between their equilibrium and optimal action; hence, the equilibria have low stability. Conversely, we show that if agents’ equilibrium payoffs are much higher than their optimal-allocation payoffs (i.e., the equilibrium is “stable”), then this ensures that the equilibrium must be of relatively high quality. To demonstrate how this phenomenon may be exploited algorithmically, we perform simulations using the log-linear learning algorithm and show that average performance on worst-case instances is far better even than our improved analytical guarantees.

Index Terms—Game Theory, Multiagent Control, Submodular Maximization

I. INTRODUCTION

GAME-THEORETIC approaches to distributed control of multiagent systems have been proposed for such diverse systems as distributed power generation, swarming of autonomous vehicles, network routing, and more [1]–[4]. Here, a system designer assigns a utility function to each agent and programs the agent with an algorithm to optimize its own utility function. By selecting the utility functions properly, a system designer can guarantee that the local algorithms drive the agents’ collective behavior to a state which is globally desirable. This approach leverages the broader literature on learning in games to arrive at generalizable convergence and performance guarantees [5], [6]. Furthermore, previous work [7], [8] provides background and a framework for analyzing games in which agents have limited knowledge of other agents’ strategies.

Applying these game-theoretic techniques to multiagent submodular maximization (a problem with applications in a wide variety of engineering contexts) yields a class of games known as *valid utility games* [9]–[11]. It has long been known

that in any valid utility game, every Nash equilibrium is always within a factor of 2 of optimal; it is said that the *price of anarchy* of valid utility games is $1/2$ [12].

However, recent work has suggested that the derived game-theoretic algorithms may lack robustness to unanticipated changes in problem structure [13]. Accordingly, recent years have seen a growing interest in the role played by communication and observation among agents in these approaches; for instance, by investigating optimal communication topologies [14] and studying the harm induced by communication or observation failures [15], [16]. This is also related to similar work on the robustness of other distributed algorithms for submodular maximization [17].

Building on this, the recent work in [18] identifies a class of tight worst-case guarantees for communication-denied multiagent submodular maximization problems. There, it is shown that if k agents cannot observe the actions or detect the contributions of other agents (we call these agents *compromised*), then the price of anarchy guarantee worsens from $1/2$ to $1/(k + 2)$. This bound is shown to be tight via finely-tuned worst-case problem instances; however, small perturbations to agent utility functions can easily destabilize the low-quality Nash equilibria and thereby greatly improve worst-case performance. In essence, the specific worst-case examples presented in [18] are *fragile*.

In the present paper, we show that this fragility is generic and is exhibited by all worst-case problem instances in this class of games; furthermore, we show a fundamental relationship between this fragility and the quality of equilibria.

First, we show the following general principle analytically: if an instance of a communication-denied valid utility game has a very low-quality Nash equilibrium, then at that equilibrium, the agents are collectively very *close* to selecting an optimal action profile: the low-quality equilibrium is not *stable*. Specifically, if agents switch from their equilibrium actions to system-optimal actions, their total payoff loss is very small. Conversely, if at equilibrium all agents strongly “prefer” their equilibrium actions to their optimal actions (i.e., the equilibrium is *stable*), the equilibrium must have relatively high quality. All low quality equilibria have low stability, and all highly stable equilibria are of high quality.

To make this notion precise, we write S to denote the *payoff stability* of a game, defined as the the total payoff loss of compromised agents at system-optimal actions relative to equilibrium. This stability is analogous to the magnitude of a perturbation to agent utility functions which would be sufficient to cause agents to deviate from a suboptimal equilibrium to an optimal action profile. Our main theoretical

This work was supported in part by the National Science Foundation under Grant ECCS-2013779.

The authors are with the University of Colorado Colorado Springs, CO 80918, USA. {jseaton, philip.brown}@uccs.edu

result states that if G is a game with $k \geq 1$ compromised agents, then the price of anarchy of G satisfies

$$\text{PoA}(G) \geq \frac{1}{2 + k - S}.$$

In other words, games in which all compromised agents strongly prefer their equilibrium actions (i.e., S is large) necessarily have equilibria that are relatively close to optimal. Here, S is a measure of the “stability” of worst-case equilibria: a very small S means that compromised agents experience a very small payoff loss if they switch from equilibrium to system-optimal actions.

The above bound holds for all valid utility games, but we also show that if all agents are assigned the specific *marginal cost utility function* that the worst-case bound improves slightly to $1/(1 + k - S)$.

Our main theoretical result thus suggests a simple algorithmic approach to mitigate the harm of observation or communication failures: program all agents to select randomly among their high-payoff actions, using a payoff-sensitive algorithm such as log-linear learning [19]. The fact that all stable equilibria are high-quality suggests that this approach will typically not harm high-quality equilibria. Conversely, the fact that all low-quality equilibria are very unstable suggests that under this approach, agents may frequently select their system-optimal actions, leading to large performance gains. In Section IV we provide empirical evidence for this using the log-linear learning algorithm. Our simulations confirm that log-linear learning can provide substantial improvements relative to worst-case guarantees, with the greatest improvements coming for the game instances with the worst nominal performance guarantees.

II. MODEL

A. Distributed Submodular Maximization: Valid Utility Games

A multiagent resource allocation game has agent set $N = \{1, \dots, n\}$ and a finite set of resources $\mathcal{R}$. Each agent $i \in N$ has a set of admissible *actions* given by a subset of resources: $\mathcal{A}_i \subseteq 2^{\mathcal{R}}$. For convenience, we assume that $\emptyset \in \mathcal{A}$ for every i . We denote the *joint action space* by $\mathcal{A} := \mathcal{A}_1 \times \dots \times \mathcal{A}_n$. We write $a \in \mathcal{A}$ to denote a joint action and a_i to denote the resource selected by agent i in a . All set operations on action profiles proceed componentwise, e.g. $a \subset a'$ means $a_i \subset a'_i$ for all $i \in N$. For a subset of agents $J \subseteq N$, we write a_J to denote the action profile a restricted only to agents in J , and similarly we write a_{-i} to mean $a_{N \setminus \{i\}}$. With this notation, we will sometimes write an action profile a as (a_i, a_{-i}) or $(a_J, a_{N \setminus J})$ when we wish to highlight the actions of a particular agent or set of agents.

Agents are tasked with maximizing the value of an objective function $W : \mathcal{A} \rightarrow [0, \infty)$. To accomplish this goal, we assume that a system designer has endowed each agent with a *utility function* $U_i : \mathcal{A} \rightarrow \mathbb{R}$. Note that the agent utility functions are subjective representations of the agent’s objective; any information available to the agent about the choices of other agents must be represented by this utility function. Since each agent is an individual optimizer attempting to maximize

its own utility function, this formulation induces a non-cooperative game among the agents. In this work we consider utility functions which satisfy the *valid utility game* criteria of [12]:

Definition 1: A Valid Utility Game is a multiagent optimization problem satisfying the following three conditions:

- 1) W is submodular, nondecreasing, and normalized,
- 2) $U_i(a_i, a_{-i}) \geq W(a_i, a_{-i}) - W(\emptyset, a_{-i})$
- 3) $\sum_i U_i(a_i, a_{-i}) \leq W(a_i, a_{-i})$

The objective function W is *submodular* if for all $i \in N$, a_i , and $a_{-i}, a'_{-i} \in \mathcal{A}_{-i}$, whenever $a_{-i} \subseteq a'_{-i}$ it holds that

$$W(a_i, a_{-i}) - W(a_{-i}) \geq W(a_i, a'_{-i}) - W(a'_{-i}), \quad (1)$$

In other words, W exhibits a generalized form of decreasing marginal returns. We say that W is *nondecreasing* if $W(a) \leq W(a')$ for all $a \subseteq a'$, and *normalized* if $W(\emptyset) = 0$. Many utility functions are known which yield valid utility games. In particular, we often refer to the *marginal contribution* utility function defined formally as

$$U_i(a_i, a_{-i}) := W(a_i, a_{-i}) - W(\emptyset, a_{-i}). \quad (2)$$

Under marginal contribution, agent i computes the payoff of an action as the difference between the system objective function with agent i present and what it would be if agent i were not present. It is known that marginal-contribution utility functions render all system-optimal action profiles as Nash equilibria [20].

B. Compromised Agents

Recent work [18] has extended this model to focus on the performance of a multiagent system in which some subset of agents suffer a fault which compromises their ability to observe the actions of other agents in the system. We write $K \subseteq N$ to denote the set of compromised agents. A compromised agent $i \in K$ receives no information about the choices made by other agents; this models jamming by an adversary, a persistent communication failure, or a failed sensor. We call an agent which is not compromised *normal*.

We model this formally via a modified utility function for each compromised agent $i \in K$; we write this modified utility function as

$$\tilde{U}_i(a_i, a_{-i}) = U_i(a_i, \emptyset). \quad (3)$$

That is, a compromised agent simply assumes that the other (unobserved) agents are not present. For simplicity, we often write $\tilde{U}_i(a_i)$ to mean $\tilde{U}_i(a_i, a_{-i})$. The compromised agent’s lack of information means that the strategies of other agents cannot influence the payoff value of the compromised agent’s private utility function. Nevertheless, the behavior of the compromised agent remains a factor in the computation of the system objective function. That is, that compromised agents’ actions still contribute to the system-level objective W ; their activities remain valuable despite their inability to obtain information about other agents. We specify a game in this context by a tuple $G = (N, \mathcal{R}, \mathcal{A}, W, K, \{U_i\}_{i \in N})$.

With these definitions in hand, the core solution concept we consider is a *pure Nash equilibrium*, or an action profile from

which no agent has a unilateral incentive to deviate. Formally, a^{ne} is a Nash equilibrium if for every normal agent $i \in N \setminus K$ and every $a_i \in \mathcal{A}_i$ we have

$$U_i(a_i^{\text{ne}}, a_{-i}^{\text{ne}}) \geq U_i(a_i, a_{-i}^{\text{ne}}) \quad (4)$$

and for every compromised agent $j \in K$ and every $a_j \in \mathcal{A}_j$ we have

$$\tilde{U}_j(a_j^{\text{ne}}) \geq \tilde{U}_j(a_j). \quad (5)$$

For game G , we write the set of pure Nash equilibria of G as $\text{NE}(G)$.

C. Quality of Equilibria

In a distributed multiagent decision system modeled as a noncooperative game, a system operator is typically tasked with programming each agent with an algorithm that optimizes the agent's utility function; in many formulations, these algorithms are selected to guide the agents collectively to a Nash equilibrium. Accordingly, the quality of the game's Nash equilibria is a central concern. In this paper, we measure the worst-case quality of a game's equilibria using the measure known as the *price of anarchy*, defined as the worst-case ratio between the objective value of a Nash equilibrium and an optimal action profile, or

$$\text{PoA}(G) := \frac{\min_{a \in \text{NE}(G)} W(a)}{\max_{a' \in \mathcal{A}} W(a')} \leq 1. \quad (6)$$

It is known that for any game G as defined in this paper, it holds that if $K \subset N$, $\text{PoA}(G) \geq 1/(2 + |K|)$ [18], where $|K|$ is the number of elements in the set K .

D. Payoff Stability

The fundamental goal of this paper is to show that if a game has a very poor price of anarchy, that game's equilibria are "close" to an optimal allocation. To measure this closeness, we introduce the notion of *payoff stability*, defined as the total payoff loss of compromised agents at optimal action profiles relative to equilibrium. Let $a^{\text{ne}} \in \arg \min_{a \in \text{NE}(G)} W(a)$ be a worst-case Nash equilibrium¹ of G . Let $A^* = \arg \max_{a \in \mathcal{A}} W(a)$ be the set of optimal allocations. The payoff stability of game G is defined as

$$S(G) := \frac{\max_{a^* \in A^*} \sum_{i \in K} [\tilde{U}_i(a_i^{\text{ne}}) - \tilde{U}_i(a_i^*)]}{W(a^{\text{ne}})} \leq |K|. \quad (7)$$

When stability $S(G) = 0$, (7) shows that compromised agents are indifferent between their optimal and worst-case equilibrium choices. On the other hand, when compromised agents strongly prefer their worst-case equilibrium actions over their optimal actions, $S(G)$ is necessarily high. That is, $S(G)$ measures the level of satisfaction that compromised agents have with their worst-case equilibrium actions. Note that $\tilde{U}_i(a_i^{\text{ne}}) \leq W(a^{\text{ne}})$ by applying the second and third conditions of Definition 1. Therefore, since $\tilde{U}_i(a_i^{\text{ne}}) \geq \tilde{U}_i(a_i^*)$ (since a^{ne} is a Nash equilibrium), it holds that $0 \leq S(G) \leq$

$|K|$. For concreteness, we now provide Example 1 to illustrate payoff stability in the context of a specific game.

Example 1: An example instance of a single-selection weighted set cover game is illustrated in Figure 1. It consists of a set of four agents, $k = 3$ of which are compromised. Each resource has a nonnegative value, and the system objective is to maximize the total value of resources included in the union of agent actions. The agents are endowed with a marginal contribution utility function. The normal agent has as an admissible action a central resource with value 1. Each compromised agent has as admissible actions the central resource and another of value $1 - \frac{s}{k} = 0.95$, where $s = 0.15$. In the unique Nash equilibrium, all agents choose to cover the resource with value 1, yielding a system objective value of only 1. However, in the optimal allocation, each compromised agent covers the resource of value 0.95, which yields a system objective function value of 3.85. Thus, this game has a price of anarchy of $\frac{1}{3.85}$.

Note that each compromised agent would suffer a payoff loss of only 0.05 in switching from equilibrium to system-optimal actions; this intuitively indicates that this equilibrium has low stability. Indeed, this example instance has payoff stability $S(G) = 3(1 - 0.95) = 0.15$, which is quite small relative to its maximum possible value of 3. It can clearly be seen that the price of anarchy of this specific instance is $1/(1 + |K| - S(G)) = 1/3.85$, aligning with the analytical results of Theorem 2. Thus, we see that on this example the very poor price of anarchy is accompanied by a very low payoff stability: compromised agents can switch to their optimal actions and incur only a small payoff loss.

III. MAIN THEORETICAL RESULTS

Previous results [18] show that a game $G = (N, \mathcal{R}, \mathcal{A}, W, K, \{U_i\}_{i \in N})$ with $|K|$ compromised agents has $\text{PoA}(G) \geq \frac{1}{2 + |K|}$. In this paper we show that this bound is only achievable if all agents are indifferent between their equilibrium and optimal actions, and furthermore we show a fundamental relationship between the stability and quality of worst-case equilibria. Specifically, games with highly stable equilibria intrinsically have a more favorable price of anarchy. This effectively demonstrates that worst-case instances of communication-denied valid utility games are fragile in a strong sense: arbitrarily-small perturbations to their specifications can render their equilibria optimal.

Theorem 1: Let $G = (N, \mathcal{R}, \mathcal{A}, W, K, \{U_i\}_{i \in N})$ be a valid utility game with compromised agents and payoff stability $S(G)$. Then it holds that

$$\text{PoA}(G) \geq \frac{1}{2 + |K| - S(G)}. \quad (8)$$

Furthermore, this bound is essentially tight: that is, for any $k \geq 0$, $s \in [0, k]$, and $\epsilon > 0$, there exists a valid utility game G' with $|K| = k$ and $S(G') = s$ such that

$$\text{PoA}(G') \leq \frac{1}{2 + k - s - \epsilon}. \quad (9)$$

Note that (8) implies that the *only way* for performance $\text{PoA}(G)$ to be very poor (low) is for stability $S(G)$ to be very

¹This paper will not consider the trivial case in which $W(a^{\text{ne}}) = 0$, since in valid utility games this implies that $W(a) = 0$ for all $a \in \mathcal{A}$.

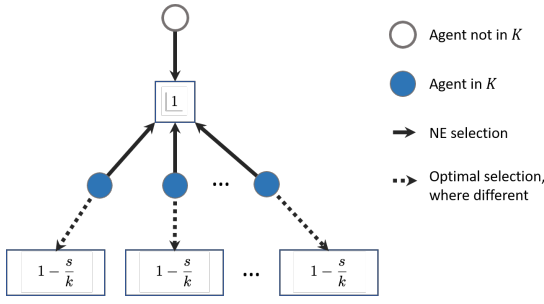


Fig. 1. Weighted set cover games used in simulations and to prove Theorem 2. It consists of a set of $n = k+1$ agents and $k+1$ resources. Every agent can access the central resource with value 1. Each of the k compromised agents can access its own “personal” resource of value $1 - s/k$, where $s \in [0, k]$. Thus, the payoff stability of the game is s .

low as well; that is, our result guarantees that a system with very bad performance also has agents which can be easily “persuaded” to switch to their optimal actions.

Proof: Let $a^{\text{ne}} \in \arg \min_{a \in \text{NE}(G)} W(a)$ be a worst-case Nash equilibrium of G . Let $a^{\text{opt}} \in \arg \max_{a \in A} W(a)$ be an optimal action profile that achieves the maximum in (7). In the following, we frequently write (a, a') to mean $(a_1 \cup a'_1, \dots, a_n \cup a'_n)$, and without loss we assume that W is defined on the extended action space $(2^{\mathcal{R}})^n$.

The proof proceeds via a sequence of inequalities; this sequence is adapted to our purposes from one introduced in [18]. We repeat it here for completeness, and highlight our modifications to make our novel contribution clear.

$$W(a^{\text{opt}}) \leq W(a^{\text{opt}}, a^{\text{ne}}) \quad (10)$$

$$\leq W(a^{\text{ne}}) + \sum_{i \in N} (W(a_i^{\text{opt}}, a_{j < i}^{\text{opt}}, a^{\text{ne}}) - W(a_{j < i}^{\text{opt}}, a^{\text{ne}})) \quad (11)$$

$$\leq W(a^{\text{ne}}) + \sum_{i \in N} (W(a_i^{\text{opt}}, a_{-i}^{\text{ne}}) - W(a_{-i}^{\text{ne}})) \quad (12)$$

$$\leq W(a^{\text{ne}}) + \sum_{i \notin K} U_i(a_i^{\text{opt}}, a_{-i}^{\text{ne}}) + \sum_{i \in K} W(a_i^{\text{opt}}) \quad (13)$$

$$\leq W(a^{\text{ne}}) + \sum_{i \notin K} U_i(a_i^{\text{ne}}, a_{-i}^{\text{ne}}) + \sum_{i \in K} \tilde{U}_i(a_i^{\text{opt}}) \quad (14)$$

$$\leq W(a^{\text{ne}}) + W(a^{\text{ne}}) + \sum_{i \in K} \tilde{U}_i(a_i^{\text{opt}}) \quad (15)$$

$$= 2W(a^{\text{ne}}) + \sum_{i \in K} \tilde{U}_i(a_i^{\text{ne}}) - S(G)W(a^{\text{ne}}) \quad (16)$$

$$\leq 2W(a^{\text{ne}}) + \sum_{i \in K} W(a_i^{\text{ne}}) - S(G)W(a^{\text{ne}}) \quad (17)$$

$$\leq (2 + |K| - S(G))W(a^{\text{ne}}), \quad (18)$$

where (10) is true since W is nondecreasing; (11) is true via telescoping; (12) is true by submodularity of W ; (13) holds since the original U_i satisfies 2) in Definition 1, and by submodularity of W ; (14) is true by definition of NE (2nd term) and by the utilities of the compromised agents (3rd term); (15) holds due to 3) in Definition 1. Our contribution begins with (16) which is substituted from the definition of payoff stability in (7); (17) is true by the definition of $\tilde{U}_i$ for agents in K ; and (18) is true since W is nondecreasing.

To see that (8) is essentially tight, let $k \geq 0$ and $s \in [0, k]$

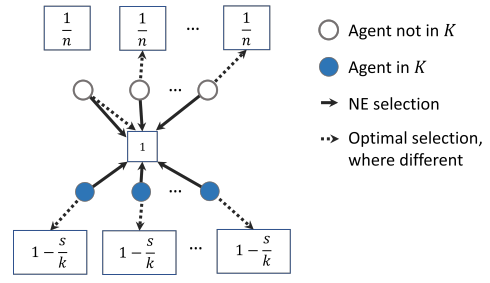


Fig. 2. This family of examples is used to show tightness of the proof of Theorem 1. An instance of a weighted set cover game, it consists of n agents, k of which are compromised. In a^{ne} , all agents select the central resource. In a^{opt} , one normal agent selects the central resource, while all other agents select their alternative resource. As $n \rightarrow \infty$, the PoA of this family of instances approaches $\frac{1}{k+2-s}$.

be given. We shall construct a valid utility game $G_{s,k}$ with $|K| = k$ and $S(G_{s,k}) = s$ with price of anarchy equal to (9). For reference, this example game is depicted in Figure 2.

The family of examples is that of *weighted set cover* games: each resource $r \in \mathcal{R}$ has value $v_r \geq 0$, and the system objective is to maximize the total value of resources included in the union of agents’ actions. Each agent has an *equal share* utility function: that is, if M is the set of agents selecting resource r , each agent $i \in M$ receives a payoff of $v_r/|M|$. It is well known that equal share utility functions satisfy Definition 1 [21].

Let there be $n+1$ resources and n agents; let $N = \{1, 2, \dots, n\}$ and let $K = \{1, 2, \dots, k\}$. Let resources $r_1, r_2, \dots, r_k$ have value $1 - s/k$ each, and let the resources $r_{k+1}, r_{k+2}, \dots, r_n$ have value $1/n$. Let resource r_{n+1} have value 1. For each agent $i \in N$, let $\mathcal{A}_i = \{r_i, r_{n+1}\}$. Thus, every agent i can select either its “personal” low-value resource r_i or the high-value “central” resource r_{n+1} .

Note that the action profile a^{ne} in which every agent $i \in N$ has $a_i^{\text{ne}} = r_{n+1}$ (all select the center resource) is a Nash equilibrium under equal share utility functions. Since only resource r_{n+1} is covered, we have $W(a^{\text{ne}}) = 1$.

Now, let a^* denote the action profile in which agent 1 selects central resource r_{n+1} and every other agent $i \in \{2, \dots, n\}$ selects its personal resource r_i . Here, every resource except one is covered, so $W(a^*) = 1 + \frac{n-k-1}{n} + k(1 - s/k) = 2 + k - s - \frac{k-1}{n}$. Thus, we have that

$$\text{PoA}(G_{s,k}) \leq \frac{W(a^{\text{ne}})}{W(a^*)} = \frac{1}{k+2-s-\frac{k-1}{n}}.$$

The upper bound in (9) can be obtained for any ϵ by making n sufficiently large. ■

Theorem 1 applies to all valid utility games, regardless of the specific utility functions chosen. However, we show in Theorem 2 that worst-case performance can be slightly improved by adopting the marginal-cost utility design:

Theorem 2: Given a valid utility game G with payoff stability $S(G)$ whose agents are endowed with the marginal contribution utility function, it holds that

$$\text{PoA}(G) \geq \min \left\{ \frac{1}{2}, \frac{1}{|K| + 1 - S(G)} \right\}. \quad (19)$$

Furthermore, this bound is tight: that is, for any $k \geq 0$ and $s \in [0, k]$, there exists a valid utility game G' with marginal

contribution utility functions with $|K| = k$ and $S(G') = s$ such that

$$\text{PoA}(G') = \min \left\{ \frac{1}{2}, \frac{1}{k+1-s} \right\}. \quad (20)$$

Proof: Let $G = (N, \mathcal{R}, \mathcal{A}, W, K, \{U_i\}_{i \in N})$ be a valid utility game with compromised agents with payoff stability $S(G)$ and marginal contribution utility functions as in (2); let a^{ne} denote a worst-case Nash equilibrium of G , and a^{opt} denote an optimal action profile for G that achieves the maximum in (7). First, note that if $S(G) > |K| - 1$, the minimum in (19) selects $1/2$; this result is known (even when $|K| = 0$) and we do not duplicate its proof here [12].

Accordingly we focus on the case that $S(G) \leq |K| - 1$. In the following, we write (a, a') to mean $(a_1 \cup a'_1, \dots, a_n \cup a'_n)$, and without loss we assume that W is defined on the extended action space $(2^{\mathcal{R}})^n$. First, following a technique applied in [18], it can be shown that

$$\begin{aligned} W(a^{\text{opt}}) &\leq W(a^{\text{opt}}, a_K^{\text{ne}}) \\ &\leq 2W(a^{\text{ne}}) - W(a_K^{\text{ne}}) + \sum_{i \in K} W(a_i^{\text{opt}}). \end{aligned} \quad (21)$$

This approach involves explicitly modeling the agents in $N \setminus K$ as playing a “sub-game” among themselves, and leverages the known result that the price of anarchy of a game with no compromised agents is no worse than $1/2$ [12].

The proof then proceeds by a sequence of inequalities; an explanation for each follows.

$$\begin{aligned} W(a^{\text{opt}}) &\leq W(a^{\text{opt}}, a_K^{\text{ne}}) \\ &\leq 2W(a^{\text{ne}}) - W(a_K^{\text{ne}}) + \sum_{i \in K} W(a_i^{\text{opt}}) \end{aligned} \quad (22)$$

$$= 2W(a^{\text{ne}}) - W(a_K^{\text{ne}}) + \sum_{i \in K} \tilde{U}(a_i^{\text{ne}}) - S(G)W(a^{\text{ne}}) \quad (23)$$

$$\leq 2W(a^{\text{ne}}) - W(a_K^{\text{ne}}) + \sum_{i \in K} W(a_K^{\text{ne}}) - S(G)W(a^{\text{ne}}) \quad (24)$$

$$\leq 2W(a^{\text{ne}}) + (|K| - 1)W(a_K^{\text{ne}}) - S(G)W(a_K^{\text{ne}}) \quad (25)$$

$$\leq 2W(a^{\text{ne}}) + W(a^{\text{ne}})(|K| - 1 - S(G)). \quad (26)$$

Inequality (22) is due to (21); (23) is due to conditions 2) and 3) of Definition 1 and by the definition of payoff stability given in (7); (24) is due to conditions 2) and 3) of Definition 1; (25) is true since W is nondecreasing, and (26) is true since W is nondecreasing and $S(G) \leq |K| - 1$. Thus, when $S(G) \leq |K| - 1$, $W(a^{\text{opt}}) \leq (|K| + 1 - S(G))W(a^{\text{ne}})$, completing the proof of (19).

The tightness of (19) can be shown via a straightforward generalization of the phenomenon discussed in Example 1 using games of the form depicted in Figure 1. For reasons of space, we do not present it in full here. ■

IV. SIMULATIONS

In this section we present the results of multiple runs of a simulation² of weighted set cover games; these simulations

give direct empirical evidence that our core theoretical results from Theorems 1 and 2 can be exploited algorithmically to improve behavior in valid utility games with compromised agents. The simulation implements *log-linear learning*, a discrete-time asynchronous distributed learning algorithm. Log-linear learning operates in discrete steps at times $t_0, t_1, \dots$ resulting in a sequence of joint actions $a(0), a(1), \dots$ beginning with an arbitrary joint action $a(0)$. At each step, an agent is chosen uniformly at random to update its action. The updating agent i then selects an action a_i from its action set $\mathcal{A}_i$ with probability

$$p_i^{a_i}(t+1) = \frac{e^{U_i(a_i, a_{-i}(t))/T}}{\sum_{\tilde{a}_i \in \mathcal{A}_i} e^{U_i(\tilde{a}_i, a_{-i}(t))/T}}. \quad (27)$$

where T , called the *temperature*, is a parameter controlling agents’ randomness. Note that $T \rightarrow \infty$ results in agents playing uniformly at random. When $T \rightarrow 0$, agent actions are strictly best responses.

The structure of the simulated game is depicted and described in Figure 1. These weighted set cover games consist of a set of $n = k + 1$ agents and $k + 1$ resources. Every agent can select the “central” resource R_0 with value 1. Each of the k compromised agents can also select its own “personal” resource of value $1 - s/k$, where $s \in [0, k - 1]$. Thus, the payoff stability of the overall game is s , since it is always a Nash equilibrium for all agents to select R_0 , but the optimal action profile has every compromised agent selecting its own personal resource with value $1 - s/k$. To ensure that uniform randomization does not lead to spurious desirable outcomes, we also attach several “dummy actions” with 0 value to each of the compromised agents.

The simulation proceeds by fixing a temperature T , initializing the agents to a random joint action and letting log-linear learning run for 200,000 iterations. After a burn-in of 50,000 iterations, the system objective function W is measured at each time step t ; the mean of these final 150,000 values are then recorded as the system’s empirical average performance. Figure 3 depicts these empirical averages across a range of temperatures with $k = 10$ and payoff stability $s = 1$.

At low temperatures, the agents’ selections approximate an asynchronous *best response* to other agents, so the average objective function value approximately achieves the theoretical price of anarchy of $1/11$. As the temperature increases, the compromised agents rapidly begin to select their “personal” actions, which dramatically improves the average objective value, rising to a maximum of approximately 0.456. At higher temperatures, agents act more randomly and begin to select the dummy actions, again harming the system objective.

Figure 4 then depicts the effect of payoff stability. The experiment depicted in Figure 3 is repeated for values of $s \in [0, k - 1]$, and the maximum and minimum values of the trace in Figure 3 are then plotted in Figure 4. The minimum value is represented by the red plus-signs, and the maximum value is represented by the dashed, orange line.

Clearly, the system objective value can be greatly increased when agents play with some randomness; this is due directly to the fact that at very poor equilibria, agents are nearly indifferent between their equilibrium actions and system-optimal

²The Python source code for the simulations is available at <https://github.com/descon-uccs/lcss-2022>.

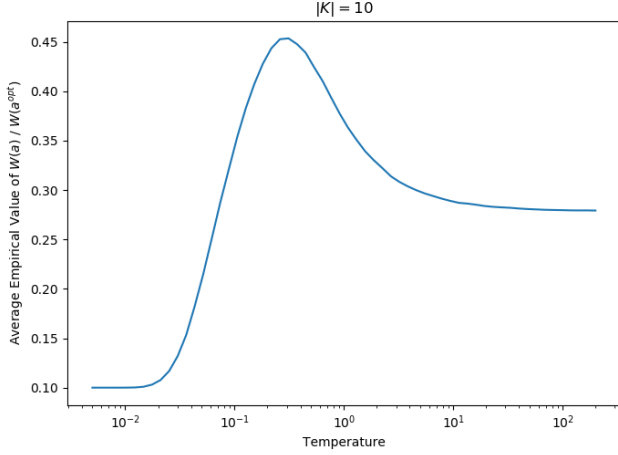


Fig. 3. Simulation across a range of temperatures with $|K| = 10$ and $S(G) = 1$, so $\text{PoA}(G) = \frac{1}{10}$. This is confirmed with the average empirical value of W being $\frac{1}{K}$ at very low temperatures. Note that a slight increase in temperature yields a large improvement in average performance, since agents begin to randomize and thus frequently select their system-optimal actions. The maximum and minimum values from multiple such simulations are plotted in Figure 4.

actions. Furthermore, note that the benefits of randomness are greatest when the payoff stability is lowest (i.e., when the nominal price of anarchy is the worst). This is due to the fact that at high-quality (and thus stable) equilibria, a small amount of randomness is unlikely to cause agents to deviate from their (relatively good) equilibrium actions.

V. CONCLUSION

This paper demonstrates that worst-case equilibrium behavior is fragile and that the previously-known theoretical lower bound is only achieved on finely-tuned games. Furthermore, our analysis and simulations indicate that payoff-sensitive probabilistic techniques such as log-linear learning can likely be used to compensate for communication failures and other types of information deprivation.

REFERENCES

- [1] W. Saad, Z. Han, H. Poor, and T. Basar, "Game-Theoretic Methods for the Smart Grid: An Overview of Microgrid Systems, Demand-Side Management, and Smart Grid Communications," *IEEE Signal Processing Magazine*, vol. 29, pp. 86–105, sep 2012.
- [2] Y. Xu, J. Wang, Q. Wu, A. Anpalagan, and Y. D. Yao, "Opportunistic spectrum access in cognitive radio networks: Global optimization using local interaction games," *IEEE Journal on Selected Topics in Signal Processing*, vol. 6, no. 2, pp. 180–194, 2012.
- [3] I. Kordonis, M. M. Dessouky, and P. A. Ioannou, "Mechanisms for Cooperative Freight Routing: Incentivizing Individual Participation," *IEEE Transactions on Intelligent Transportation Systems*, pp. 1–12, 2019.
- [4] B. L. Ferguson, P. N. Brown, and J. R. Marden, "Carrots or Sticks? The Effectiveness of Subsidies and Tolls in Congestion Games," in *2020 American Control Conference*, pp. 5370–5375, 2020.
- [5] J. R. Marden, H. P. Young, G. Arslan, and J. S. Shamma, "Payoff based dynamics for multi-player weakly acyclic games," *SIAM Journal on Control and Optimization*, vol. 48, no. 1, pp. 373–396, 2009.
- [6] R. Gopalakrishnan, J. R. Marden, and A. Wierman, "An architectural view of game theoretic control," *Performance Evaluation Review*, vol. 38, no. 3, pp. 31–36, 2010.
- [7] L. Su and N. H. Vaidya, "Byzantine-resilient multiagent optimization," *IEEE Transactions on Automatic Control*, vol. 66, no. 5, pp. 2227–2233, 2020.

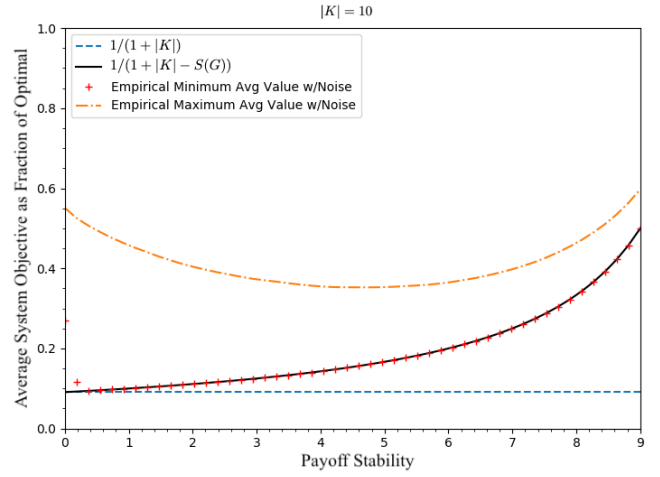


Fig. 4. Simulations of games constructed according to the model in Figure 1 with $|K| = 10$. The orange dashed line represents the maximum average value of W across a range of temperatures (see Figure 3). The red plus-signs represent the minimum average value. Note that this minimum trace follows the analytical prediction of $\frac{1}{k+1-s}$ when agents are endowed with marginal contribution utility functions; this is because the minimum average objective value is obtained for low temperatures (a best response process) and thus is a Nash equilibrium.

- [8] B. Ghahesifard and S. L. Smith, "On distributed submodular maximization with limited information," in *2016 American Control Conference (ACC)*, pp. 1048–1053, IEEE, 2016.
- [9] D. Kempe, J. Kleinberg, and É. Tardos, "Maximizing the spread of influence through a social network," in *9th ACM conference on Knowledge discovery and data mining*, pp. 137–146, 2003.
- [10] O. Barinova, V. Lempitsky, and P. Kholi, "On detection of multiple object instances using hough transforms," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2012.
- [11] H. Lin and J. Bilmes, "A Class of Submodular Functions for Document Summarization," in *49th Annual Meeting of the Association for Computational Linguistics*, pp. 510–520, 2011.
- [12] A. Vetta, "Nash equilibria in competitive societies, with applications to facility location, traffic routing and auctions," in *The 43rd Annual IEEE Symposium on Foundations of Computer Science, 2002. Proceedings.*, pp. 416–425, 2002.
- [13] P. N. Brown, H. P. Borowski, and J. R. Marden, "Projecting Network Games on to Sparse Graphs," in *52nd Asilomar Conference on Signals, Systems, and Computers*, pp. 307–309, 2018.
- [14] D. Grimsman, M. S. Ali, J. P. Hespanha, and J. R. Marden, "The Impact of Information in Greedy Submodular Maximization," *IEEE Transactions on Control of Network Systems*, vol. 6, no. 4, pp. 1334–1343, 2018.
- [15] P. N. Brown and J. R. Marden, "On the feasibility of local utility redesign for multiagent optimization," in *18th European Control Conference*, pp. 3396–3401, 2019.
- [16] H. Jaleel, W. Abbas, and J. S. Shamma, "Robustness Of Stochastic Learning Dynamics To Player Heterogeneity In Games," in *IEEE 58th Conference on Decision and Control*, pp. 5002–5007, 2019.
- [17] L. Zhou, V. Tzoumas, G. J. Pappas, and P. Tokekar, "Distributed Attack-Robust Submodular Maximization for Multi-Robot Planning," in *2020 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 2479–2485, IEEE, may 2020.
- [18] D. Grimsman, J. H. Seaton, J. R. Marden, and P. N. Brown, "The Cost of Denied Observation in Multiagent Submodular Optimization," in *2020 59th IEEE Conference on Decision and Control (CDC)*, pp. 1666–1671, IEEE, dec 2020.
- [19] C. Alós-Ferrer and N. Netzer, "The logit-response dynamics," *Games and Economic Behavior*, vol. 68, no. 2, pp. 413–427, 2010.
- [20] D. H. Wolpert and K. Tumer, "Optimal Payoff Functions for Members of Collectives," *Advances in Complex Systems*, vol. 04, no. 02n03, pp. 265–279, 2001.
- [21] M. Phillips and J. R. Marden, "Design tradeoffs in concave cost-sharing games," *IEEE Transactions on Automatic Control*, vol. 63, no. 7, pp. 2242–2247, 2018.