

Robust Networked Multiagent Optimization

Designing Agents to Repair Their Own Utility Functions

Philip N. Brown · Joshua H. Seaton ·
Jason R. Marden

Received: date / Accepted: date

Abstract We study settings in which autonomous agents are designed to optimize a given system-level objective. In typical approaches to this problem, each agent is endowed with a decision-making rule that specifies the agent's choice as a function of relevant information pertaining to the system's state. The choices of other agents in the system comprise a key component of this information. This paper considers a scenario in which the designed decision-making rules are not implementable in the realized system due to discrepancies between the anticipated and realized information available to the agents. The focus of this paper is to develop methods by which the agents can preserve system-level performance guarantees in these unanticipated scenarios through local and independent redesigns of their own decision-making rules. First, we show a general impossibility result which states that in general settings, there are no local redesign methodologies that can offer any preservation of system-level performance guarantees, even when the affected agents satisfy an inconsequentiality criterion. However, we then show that when system-level

This work is supported in part by ONR Grant #N00014-20-1-2359, AFOSR Grants #FA9550-20-1-0054 and #FA9550-21-1-0203, NSF Grant #ECCS-2013779, and the Army Research Lab through ARL DCIST CRA #W911NF-17-2-0181. A preliminary version of this work appeared in (Brown and Marden, 2019).

P. N. Brown
University of Colorado Colorado Springs
Dept. of Computer Science
E-mail: pbrown2@uccs.edu

J. H. Seaton
University of Colorado Colorado Springs
Dept. of Computer Science
E-mail: jseaton@uccs.edu

J. R. Marden
University of California Santa Barbara
Dept. of Electrical and Computer Engineering
E-mail: jrmarden@ece.ucsb.edu

objectives are submodular, local redesigns of utility functions do exist which allow nominal performance guarantees to degrade gracefully as information is denied to agents. That is, in these submodular settings, agents can adapt to informational inconsistencies independently without incurring much loss in terms of system-level performance.

Keywords Game Theory · Networked Learning · Distributed Optimization · Utility Design

1 Introduction

A multi-agent system can be viewed as a collection of decision-making entities each making local decisions in response to locally available information. These decision-making entities could be self-interested, e.g., drivers utilizing a transportation networks (Kleer and Schäfer, 2017), or could be computer controlled, e.g., unmanned aerial vehicles in a swarm (Ni et al., 2020). Regardless of the physical makeup of these decision-making entities, the central job of the system operator is to ensure that the emergent collective behavior is desirable relative to a performance metric of interest.

Game theory has received considerable attention as an overarching framework for the analysis and design of such multi-agent systems (Alós-Ferrer and Netzer, 2010; Blume, 1993, 1995; Pradelski and Young, 2012; Young, 1993). Clearly, game theory is relevant to the design and control of multi-agent systems when the system is comprised of self-interested decision-making entities (Collins et al., 2021; Ferguson et al., 2021; Vetta, 2002). Interestingly, game theory is also relevant for systems where the decision-making entities are computer controlled and a system operator is tasked with deriving local decision-making rules (or control strategies) for the agents in the systems (Wolpert et al., 1999). The connection to game theory is fueled by the fact that these control strategies can often be viewed through the lens of distributed learning, where agents are responding in a predetermined way to some intrinsic utility functions.

The connection between game theory and the design of engineered multi-agent systems has fueled a branch of research under the umbrella of mechanism design that focuses on the design of agent utility functions (Marden and Shamma, 2014; Yang et al., 2017). Here, a central objective of the system designer is to construct agent utility functions such that the resulting equilibria have desirable efficiency properties. It is often the case that the system designer can focus exclusively on the design of utility functions without explicitly considering the dynamics that will ultimately drive the collective behavior to such equilibria (de Jong and Uetz, 2020; Li and Marden, 2013). The reason for this decomposition stems from the fact that if the designed agent utility functions possess a desirable structure, e.g., form a weakly acyclic or potential game, then a system operator can appeal to off-the-shelf distributed learning algorithms that can be employed to drive the collective behavior to an equilibrium of interest, e.g., fictitious play or log-linear learning (Alós-Ferrer

and Netzer, 2010; Candogan et al., 2013; Fudenberg and Levine, 1998; Li et al., 2021). Accordingly, the design of distributed control algorithms for engineered multi-agent systems can be recast as the problem of designing utility functions for the agents such that the induced games have desirable equilibrium properties. Interestingly, recent results show that this game theoretic approach to distributed control design is without loss for various problems of interest (Paccagnan et al., 2020). That is, such game theoretic algorithms can actually match the performance bounds of the best distributed algorithms.

One of the central research themes associated with the problem of utility design in multiagent systems focuses on the well-studied efficiency metric termed the *price of anarchy* (Papadimitriou, 2001). The price of anarchy is a worst-case measure that seeks to bound the system-level cost associated with an equilibrium of interest, e.g., pure Nash equilibrium, when compared to the optimal system-level cost. With regard to multiagent system design, the price of anarchy serves as the competitive ratio of the distributed algorithm that results from merging a particular utility design with a given dynamic process that is guaranteed to drive the collective behavior to an equilibrium. Accordingly, there have been significant recent advances in deriving utility functions that optimize the price of anarchy (Ramaswamy Pillai et al., 2021).

This paper studies implementing such game theoretic approaches for the online coordination of multiagent systems, focusing in particular on the robustness of such approaches to informational deficiencies. Implementing a given distributed learning algorithm for online coordination of multi-agent systems typically requires an individual agent to have access to (i) the agent’s local utility function and (ii) the behavior of the other agents in the system. While knowledge of the local utility function is not necessarily a problem, the same need not hold true regarding knowledge of the behavior of the other agents in the system (Grimsman et al., 2020; Seaton and Brown, 2022). These informational deficiencies can be due to numerous issues including agent failures, adversarial interventions, and communication or sensing issues. Accordingly, the focus of this paper is on understanding how to revise these online learning algorithms to accommodate unplanned informational deficiencies while preserving the original equilibrium performance guarantees.

We begin by describing our model to ensure that our contributions are clear.

1.1 Model

A multiagent optimization problem has agent set $\mathcal{I} = \{1, \dots, n\}$ where each agent $i \in \mathcal{I}$ has a finite action set (or choice set) $\mathcal{A}_i$. The quality of each joint action profile $a = (a_1, \dots, a_n) \in \mathcal{A} = \mathcal{A}_1 \times \dots \times \mathcal{A}_n$ is expressed by a system-level objective function of the form $W : \mathcal{A} \rightarrow \mathbb{R}_{\geq 0}$. Further, each agent i is assigned a utility function $U_i : \mathcal{A} \rightarrow \mathbb{R}_{\geq 0}$ which will guide their decision-making process. We will refer to these utility functions as the *nominal* utility functions and will denote a multiagent optimization problem of the above form

		$\bar{\mathcal{A}}_2$					$\mathcal{A}_2$					$\mathcal{A}_2$				
		x	y	z			x	y	z			x	y	z		
$\bar{\mathcal{A}}_1$	x	$W_{xx} = 1$	$W_{xy} = 2$	$W_{xz} = 0$		x	$W_{xx} = 1$	$W_{xy} = 2$			x	$W_{xx} = 1$		$W_{xz} = 0$		
	y	$W_{yx} = 2$	$W_{yy} = 3$	$W_{yz} = 0$		y	$W_{yx} = 2$	$W_{yy} = 3$			y		$W_{yy} = 3$			
	z	$W_{zx} = 0$	$W_{zy} = 0$	$W_{zz} = 4$		z					z	$W_{zx} = 0$		$W_{zz} = 4$		
All Possible Actions						Realization #1						Realization #2				

Fig. 1 This figure highlights the inherent challenges associated with the utility design process for a simplified problem. To that end, suppose that each agent $i \in \mathcal{I} = \{1, 2\}$ has three potential actions, which we denote by $\bar{\mathcal{A}}_1 = \bar{\mathcal{A}}_2 = \{x, y, z\}$. Furthermore, the system-level objective (or welfare) associated with the nine possible joint actions is provided on the left figure. For simplicity, consider the common interest design where each agent is associated with a utility function equal to the global objective, e.g., $U_1(x, y) = U_2(x, y) = W_{xy}$. Once this utility structure is specified, the action set of each agent is realized and the designed utility functions are implemented. In Realization #1, the realized action sets are $\mathcal{A}_1 = \mathcal{A}_2 = \{x, y\}$ and there is a unique pure Nash equilibrium with the optimal total welfare of 3. However, in Realization #2, the realized action sets are now of the form $\mathcal{A}_1 = \mathcal{A}_2 = \{x, z\}$ and there are two pure Nash equilibria, (x, x) and (z, z) , with a welfare of 1 and 4 respectively. Hence, this particular utility design is unable to guarantee that any Nash equilibrium in any problem instance will have a welfare $> 25\%$ of the optimal welfare.

by the tuple $G = (\mathcal{I}, \mathcal{A}, \{U_i\}_{i \in \mathcal{I}}, W)$. We will denote a set of such multiagent optimization problems by $\mathcal{G}$.

The field of utility design focuses on how to construct these agent utility functions $U = (U_1, \dots, U_n)$ to achieve desirable equilibrium properties. During the design process, it is typically assumed that the system designer has limited knowledge regarding the specific structure of the multiagent optimization problem $(\mathcal{I}, \mathcal{A}, \cdot, W)$, and a common mode of uncertainty pertains to the structure of the agents' action sets. More formally, each agent $i \in \mathcal{I}$ is associated with a set of possible action choices $\bar{\mathcal{A}}_i$, and the system-level objective is defined over these possible action sets, i.e., $W : \bar{\mathcal{A}}_1 \times \dots \times \bar{\mathcal{A}}_n \rightarrow \mathbb{R}$. However, the specific realization of these action sets, $\mathcal{A}_i \subseteq \bar{\mathcal{A}}_i$ for all $i \in \mathcal{I}$, is not available to the system designer. The goal of a system designer is to design a utility function for each agent $i \in \mathcal{I}$ of the form $U_i : \bar{\mathcal{A}}_1 \times \dots \times \bar{\mathcal{A}}_n \rightarrow \mathbb{R}$ that leads to desirable collective behavior regardless of the specific realization of the agents' action sets. Note that a particular choice of utility functions $U = (U_1, \dots, U_n)$ coupled with a system objective function W induces a family of multiagent optimization problems of the form $\mathcal{G}_U = \{(\mathcal{I}, \mathcal{A}, \{U_i\}_{i \in \mathcal{I}}, W) : \mathcal{A} \subseteq \bar{\mathcal{A}}\}$, where we denote an admissible joint action set as $\mathcal{A} \subseteq \bar{\mathcal{A}}$ with the understanding that this means $\mathcal{A} = \mathcal{A}_1 \times \dots \times \mathcal{A}_n$ where $\mathcal{A}_i \subseteq \bar{\mathcal{A}}_i$ for each agent $i \in \mathcal{I}$. See Figure 1 for an illustration.

When a utility design induces a learnable game structure (e.g., potential games or weakly acyclic games) and is coupled with a suitable distributed learning rule that guarantees convergence to an equilibrium (e.g., fictitious play or best response dynamics), the result is a distributed algorithm that provides a competitive ratio equal to the price of anarchy, which is defined as

$$\text{PoA}(\mathcal{G}_U) = \inf_{G \in \mathcal{G}_U} \left(\frac{\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}})}{\max_{a^{\text{opt}} \in G} W(a^{\text{opt}})} \right) \geq 0. \quad (1)$$

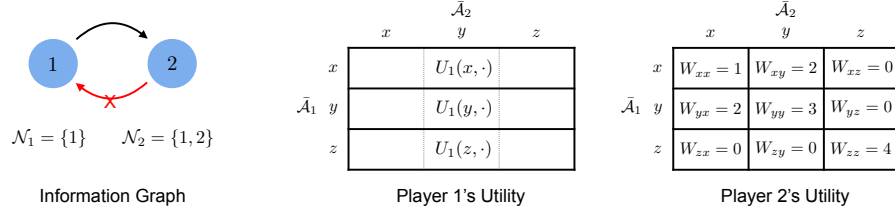


Fig. 2 This figure highlights the impact of informational deficiencies on the implementation of utility functions. Consider the particular design considered in Figure 1. Now suppose that Agent 1 loses the ability to observe the behavior of Agent 2, i.e., $\mathcal{N}_1 = \{1\}$ and $\mathcal{N}_2 = \{1, 2\}$. Since Agent 1 is unable to observe the choice of Agent 2, Agent 1 must construct a lower dimensional payoff function that associates a single payoff to each choice in the set $\mathcal{A}_1 = \{x, y, z\}$. The efficacy of this construction is ultimately gauged by the quality of the equilibria in this new game.

where $\text{PNE}(G) \subseteq \mathcal{A}$ denotes the set of pure Nash equilibria in the game G . There are several recent positive results that characterize the utility functions that optimize the price of anarchy for several interesting classes of systems (Ramaswamy Pillai et al., 2021), e.g., congestion games.

Implementing such an algorithm requires that each individual agent is able to evaluate its utility function, which may not be possible given potential informational deficiencies that may prevent agents from observing the action choices of other agents in the system (for instance, a faulty sensor or jamming by an adversary). We model the information available to the agents using an *information graph* $\mathcal{N} := \{\mathcal{N}_1, \dots, \mathcal{N}_{|\mathcal{I}|}\}$ where we let $\mathcal{N}_i \subseteq \mathcal{I}$, $i \in \mathcal{N}_i$, denote the set of agents whose actions can be observed by agent i . Throughout this paper, we write $|\mathcal{N}| := \sum_{i \in \mathcal{I}} |\mathcal{N}_i|$ to denote the number of *edges* in $\mathcal{N}$. We write $\mathcal{N}^c = \{\mathcal{N}_1^c, \dots, \mathcal{N}_{|\mathcal{I}|}^c\}$ to denote the *complement graph* of $\mathcal{N}$; that is, for each $i \in \mathcal{I}$, we write $\mathcal{N}_i^c := \mathcal{I} \setminus \mathcal{N}_i$ to denote the set of agents whose actions *cannot* be observed by Agent i . To account for these informational deficiencies, it is imperative that the nominal utility functions $U_i : \bar{\mathcal{A}} \rightarrow \mathbb{R}_{\geq 0}$ are replaced with modified lower-dimensional utility functions of the form $\mathcal{U}_i : \bar{\mathcal{A}}_{\mathcal{N}_i} \rightarrow \mathbb{R}_{\geq 0}$, where $\bar{\mathcal{A}}_{\mathcal{N}_i} = \prod_{i \in \mathcal{N}_i} \bar{\mathcal{A}}_i$ denotes the actions of agents visible to i . This modification will then allow the agents to follow the prescribed learning rule, albeit with these new modified utility functions that adhere to the realized information structure. See Figure 2 for an illustration.

This paper focuses on designing local mechanisms for constructing these new lower-dimensional utility functions $\mathcal{U}_1, \dots, \mathcal{U}_n$ that maintain some notion of equilibrium quality relative to the nominal utility functions U . The process by which agents adapt their utility functions to account for information deficiencies is given by a local adaptation rule, defined as follows:

Definition 1 Let $U_i : \bar{\mathcal{A}} \rightarrow \mathbb{R}$ be the nominal utility function of agent i . A *local adaptation rule* for agent i is a function f_i of the form $\mathcal{U}_i = f_i(U_i, \mathcal{N}_i)$. That is, f_i takes in a nominal utility function U_i and a local information

graph $\mathcal{N}_i$, and outputs a new lower-dimensional utility function of the form $\mathcal{U}_i : \bar{\mathcal{A}}_{\mathcal{N}_i} \rightarrow \mathbb{R}_{\geq 0}$.¹

When applying a local adaptation rule, agent i knows $\mathcal{N}_i$ (which agents it can observe) and U_i (its own full nominal utility function) but has no knowledge of the system objective W , or the utility functions or information sets $\mathcal{N}_j$ of other agents. We let $G_f(\mathcal{N})$ denote the locally adapted version of the game G through the adaptation rule f and information graph $\mathcal{N}$. We will sometimes omit highlighting the information graph, i.e., express $G_f(\mathcal{N})$ as merely G_f , when the dependence is clear.

1.2 Summary of Contributions

While there has been extensive work done in the area of utility design, c.f., (Paccagnan et al., 2020), to the best of our knowledge this is the first work to address the design of adaptation rules to accommodate unplanned informational deficiencies. Accordingly, to disentangle these two design elements, we fix the utility design as merely the common interest design $U_i = W$, as highlighted in Figures 1 and 2, and concentrate purely on the design of local adaptation rules $f = (f_1, \dots, f_n)$ that gracefully preserve efficiency guarantees of the resulting equilibria. While we do not explicitly consider situation where $U_i \neq W$ in this manuscript, many of the results immediately extend to this setting as well.

Our first set of results are largely negative, stating that in general it is impossible for any local adaptation rule f to preserve any level of optimality associated with the resulting pure Nash equilibria for any degree of informational losses. We state these results here less formally, and refer readers to Section 2 for the formal statements of Proposition 2 and Theorem 1.

Contribution 1 *Let $\mathcal{G}$ be a sufficiently rich family of multiagent optimization problems with the common interest utility design. Then for any information graph $\mathcal{N}$ that is missing one directed edge,² we have*

$$\sup_f \inf_{G \in \mathcal{G}} \left(\frac{\max_{a_f^{\text{ne}} \in \text{PNE}(G_f)} W(a_f^{\text{ne}})}{\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}})} \right) = 0. \quad (2)$$

The first result above highlights the fragility of distributed approaches to informational deficiencies in multiagent optimization problems. Note that the metric introduced in (2) is optimistic, comparing the best equilibrium in the adapted game G_f to the worst equilibrium in the nominal game G , and that this optimism bias does not help preserve efficiency guarantees. Furthermore,

¹ Our forthcoming results impose an extra consistency condition on local adaptation rules; see Definition 2 in Section 2 for details.

² That is, one agent lacks the ability to observe a single other agent, but every other agent can observe all agents. Formally, the complement graph of $\mathcal{N}$ contains exactly one edge: $|\mathcal{N}^c| = 1$.

in Section 2 we also demonstrate that this negative result holds even when the affected agents i, j satisfy an inconsequentiality criterion in which the behavior of agent j has minimal impact on the utility of agent i . Finally, we mention that the formal statement of this result in Theorem 1 is actually stronger than the statement above, since there we show that there exist pathological multiagent optimization problems on which *every* local adaptation rule fails to preserve equilibrium performance guarantees.

While the above result paints a negative picture regarding the availability of adaptation rules to account for informational deficiencies, our next set of results is more positive and states that certain structures of multiagent optimization problems do provide agents with opportunities to preserve equilibrium quality despite informational deficiencies. These positive results pertain to the class of multiagent optimization problems with submodular system-level objective functions, i.e., objective functions that exhibit a property of diminishing returns. In this case, we establish a lower bound on equilibrium quality degradation that is parameterized by $|\mathcal{N}^c|$, the number of edges missing from the information graph; this bound is depicted as well in Figure 3.

Contribution 2 *Let $\mathcal{G}$ be a sufficiently rich family of submodular multiagent optimization problems with the common interest utility design. Then for any information graph $\mathcal{N}$, we have*

$$\sup_f \inf_{G \in \mathcal{G}} \left(\frac{\min_{a_f^{\text{ne}} \in \text{NE}(G_f)} W(a_f^{\text{ne}})}{\max_{a^{\text{ne}} \in \text{NE}(G)} W(a^{\text{ne}})} \right) \geq \frac{1}{2 + \lfloor |\mathcal{N}^c|/2 \rfloor}. \quad (3)$$

Here, we write $\text{NE}(G)$ to denote the set of all Nash equilibria (not merely pure). Note that the metric introduced in (3) is far more pessimistic than its

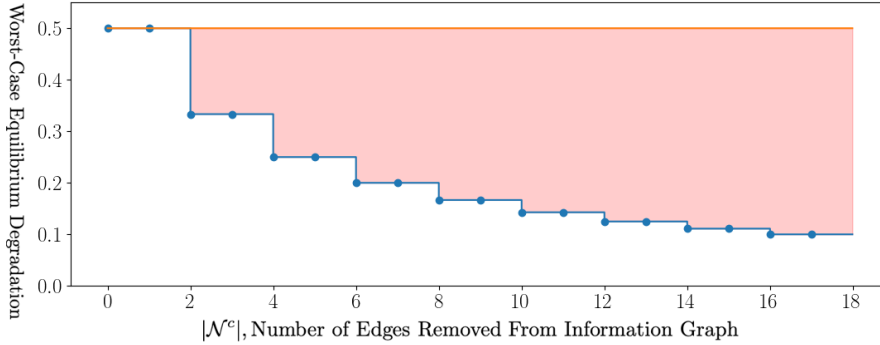


Fig. 3 Quality metric from Contribution 1 as a function of the number of edges “missing” from $\mathcal{N}$; for formal statements see Theorem 2 and Corollary 2 in Section 3. In general, the more edges missing, the worse our achievable guarantee, but there is a wide range of possible levels of degradation. Upper orange trace indicates that if the edges are chosen carefully, it is possible to remove a substantial number of edges with no impact on worst-case guarantees (see Corollary 1 in Section 3). Blue trace and discs indicate a provable lower bound (possibly loose) on worst-case equilibrium degradation as a function of the number of edges missing from $\mathcal{N}$ (Corollary 2, Section 3). Pink shading indicates possible range of values of; the actual value depends on which edges are missing in $\mathcal{N}$.

counterpart in (2), comparing the worst equilibrium in the adapted game G_f to the best equilibrium in the nominal game G . Nonetheless, this result indicates that important classes of multiagent optimization problems in $\mathcal{G}$ have a degree of robustness to informational deficiencies, since it means that even the worst adapted equilibria cannot be arbitrarily worse than the nominal equilibria, and this worst-case guarantee degrades gracefully as edges are removed from the information graph $\mathcal{N}$. To provide some intuition, the pathological games which enable Contribution 1 have highly fragile equilibria due to a lack of alignment between the actions of different agents (i.e., individual local agent best responses can be made to point “away” from system-optimal action profiles). However, submodularity enforces a degree of agent action alignment and this appears to be part of why positive performance guarantees are possible in this case.

Figure 3 depicts the possible bounds provided by (3) as a function of the number of directed edges missing from $\mathcal{N}$; note that one consequence of Contribution 2 is that removing a single edge from $\mathcal{N}$ cannot degrade worst-case equilibrium guarantees relative to the nominal case. Also, note that the edge-based bound given in (3) is a consequence of a somewhat tighter bound which we show in Theorem 2.

Finally, note that this paper considers two sources of uncertainty: first, the agents’ inability to observe other agents’ actions; second, the agents’ uncertainty over the realized action sets. We mention here that both types of uncertainty are required to obtain the negative results in this paper. Specifically, if agents cannot reliably observe each others’ actions but *do* know the realized action sets with certainty, then each agent could simply compute the system optimal action profile and play its corresponding action. However, if agents can always observe each others’ actions, then none of the pathologies in this paper are possible, even if agents have uncertainty over realized action sets.

2 Contribution 1: General multiagent optimization problems are fragile

Our first set of contributions focuses on designing local adaptation rules for general multiagent optimization problems, and asks if any adaptation rule can preserve system-level performance guarantees in the presence of informational inconsistencies. To pose the problem rigorously, we consider local adaptation rules that take the following form for some *payoff evaluator* f , for agent i and joint action $a_{\mathcal{N}_i} = \{a_j\}_{j \in \mathcal{N}_i}$:

$$\mathcal{U}_i(a_{\mathcal{N}_i}) = f(\{U_i(a') : a' \in \bar{\mathcal{A}}, a'_{\mathcal{N}_i} = a_{\mathcal{N}_i}\}). \quad (4)$$

That is, each agent is given a function $f(\cdot)$ which aggregates the payoffs which could *possibly* be obtained and are consistent with the information available to the agent into one single proxy payoff. In the interest of generality, we allow any evaluator f which satisfies the following definition:

Definition 2 For any possible utility functions U_i and U_i^* , let

$$S = \{U_i(a') : a' \in \bar{\mathcal{A}}, a'_{\mathcal{N}_i} = a_{\mathcal{N}_i}\} \quad \text{and} \quad S^* = \{U_i^*(a') : a' \in \bar{\mathcal{A}}, a'_{\mathcal{N}_i} = a_{\mathcal{N}_i}\}.$$

If $S = (s_1, s_2, \dots, s_k)$ and $S^* = (s_1^*, s_2^*, \dots, s_k^*)$ can be indexed such that $s_i > s_i^*$ for each i , then a function $f : \mathbb{R}^k \rightarrow \mathbb{R}$ is called an *acceptable evaluator* if $f(S) > f(S^*)$.

If Agent i uses acceptable evaluator f to compute proxy payoffs for the case when Agent j 's action is unobservable, we say that Agent i *applies* f to Agent j . Proposition 1 gives a partial list of evaluators which are acceptable by Definition 2; its proof is included in the Appendix.

Proposition 1 *The following functions are acceptable evaluators:*

1. *Sum*: $f_{\text{sum}}(S) = \sum_{s \in S} s$
2. *Maximum element*: $f_{\text{max}}(S) = \max_{s \in S} s$
3. *Minimum element*: $f_{\text{min}}(S) = \min_{s \in S} s$
4. *Mean*: $f_{\text{mean}}(S) = \frac{1}{|S|} \sum_{s \in S} s$

Throughout, we denote the set of all acceptable evaluators by $\mathcal{F}$.

Note that some evaluators have an appealing intuitive explanation. For instance, f_{min} assigns proxy utility functions to agent i which implicitly assume that the unobserved agents are selecting actions *adversarially*, acting to minimize the payoffs of agent i . On the other hand, the f_{max} evaluator assigns proxy utility functions to agent i which implicitly assume that unobserved agents are selecting actions to maximize the payoffs of agent i . Furthermore, in the case of common interest games, the f_{max} evaluator implicitly models unobservable agents as simply playing best response actions.

2.1 Unrestricted games are fragile.

First, we show that if no restriction is placed on which type of game is under consideration, a single missing edge in the information graph $\mathcal{N}$ can easily and catastrophically degrade performance. The following proposition assumes that Agent 1 loses information about the action choice of Agent 2, and thus must compute proxy payoffs for the case when Agent 2's action is unobservable.

Proposition 2 *Let $\mathcal{G}$ denote the set of all multiagent optimization problems, and let $\mathcal{N}$ be such that Agent 2 cannot observe the action choice of Agent 1. For any acceptable evaluator f that Agent 2 applies to Agent 1, it holds that*

$$\sup_{f \in \mathcal{F}} \inf_{G \in \mathcal{G}} \left(\frac{\max_{a_f^{\text{ne}} \in \text{PNE}(G_f)} W(a_f^{\text{ne}})}{\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}})} \right) = 0. \quad (5)$$

Here, by showing that the “optimistic” metric obtains a value of 0, we see that in general games, no evaluator can prevent pure Nash equilibria from having arbitrarily low quality when even one edge is missing from $\mathcal{N}$. The proof of Proposition 2 is given in the Appendix.

2.2 Even inconsequential agents matter.

In light of Proposition 2, we now ask whether adding additional structure to the class of optimization problems can confer some resilience. Accordingly, we now study systems in which individual unobservable agents are only weakly important to the system objective. We introduce the following notion of weak interrelation: we say Agent j is “inconsequential” if it can never cause a large change in the system objective value by changing actions. We make this notion precise in Definition 3:

Definition 3 Agent i is ϵ -inconsequential if for all $a_{-i} \in \bar{\mathcal{A}}_{-i}$, and all $a_i, a'_i \in \bar{\mathcal{A}}_i$,

$$\frac{|W(a_i, a_{-i}) - W(a'_i, a_{-i})|}{\max_{a \in \bar{\mathcal{A}}} W(a)} \leq \epsilon. \quad (6)$$

Now, let $\mathcal{G}_\epsilon$ denote the class of games such that for each $G \in \mathcal{G}_\epsilon$, we have that Agent 1 is no more than ϵ -inconsequential. Inconsequentiality gives that for each game $G \in \mathcal{G}_\epsilon$, we are assured that even if some other agent(s) cannot observe Agent 1’s action, a unilateral deviation by Agent 1 can have only a small impact on the system objective. One might hope that this property could help avoid the severe pathologies of Proposition 2.

Unfortunately, Theorem 1 demonstrates that whenever $\epsilon > 0$, no acceptable evaluator can prevent the loss of a single information edge from causing significant harm to the emergent behavior in the problem, even if that edge is associated with observing an inconsequential agent.

Theorem 1 For any $\epsilon > 0$, let $\mathcal{G}_\epsilon$ be the set of multiagent optimization problems as defined above, and let information graph $\mathcal{N}$ satisfy $|\mathcal{N}^c| = 1$ and be such that $1 \notin \mathcal{N}_2$ (that is, Agent 2 cannot observe the action choices of Agent 1).³ Then it holds that

$$\inf_{G \in \mathcal{G}_\epsilon} \sup_{f \in \mathcal{F}} \left(\frac{\max_{a_f^{\text{ne}} \in \text{PNE}(G_f)} W(a_f^{\text{ne}})}{\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}})} \right) = 0. \quad (7)$$

We provide the proof of Theorem 1 in the Appendix.

This theorem indicates a fundamental fragility in these systems, even when the class of systems has carefully been selected to perform well. As we detail in the proof of Theorem 1, the statement in (7) is obtained by demonstrating a family of games which all have unique Nash equilibria, and whose adapted games are weakly acyclic with unique Nash equilibria. That is, it is possible to reach the Nash equilibrium in each adapted game from any action profile via a finite sequence of unilateral payoff-improving moves.

Finally, note the order of precedence of the inf and sup operators in (7): this formulation explicitly states that even if f is selected with *knowledge of the specific multiagent problem*, that it cannot preserve any efficiency guarantees.

³ Here, note that we highlight the specific identities of these agents simply to ensure that it is the inconsequential agent (Agent 1, per the definition of $\mathcal{G}_\epsilon$) that cannot be observed. Naturally, the agents could be indexed in any way and the result would still hold.

3 Contribution 2: Submodular multiagent optimization is resilient

In this section, we consider the well-studied case of submodular multiagent optimization problems, defined as follows: Let S be a finite set of elements and let W be a set-based function of the form $W : 2^S \rightarrow \mathbb{R}$. Function W is called *submodular* if for any sets $Q \subseteq R \subseteq S$ and any element $s \in S$ it holds that

$$W(Q \cup \{s\}) - W(Q) \geq W(R \cup \{s\}) - W(R), \quad (8)$$

i.e., the marginal benefit of adding the element s to the set Q is higher than the marginal benefit adding the element s to the superset R . In all our results pertaining to submodular functions, we also assume the functions to be *nondecreasing*, that is, for sets $Q \subseteq R \subseteq S$, $W(Q) \leq W(R)$. Furthermore, without loss of generality, we assume submodular functions to be *normalized*, or $W(\emptyset) = 0$. When cast in this paper's multiagent framework, agents' action sets are constrained to be of the form $\mathcal{A}_i \subseteq 2^S$; i.e., each agent's pure actions are subsets of some ground set S . We also assume that $\emptyset \in \mathcal{A}_i$ for all agents i ; that is, each agent can in essence choose not to participate.

Submodular maximization is a well-studied topic in a variety of fields due to its application in many common engineering problems, e.g., information gathering (Krause et al., 2008), influence in social networks (Kempe et al., 2003), object detection (Barinova et al., 2012), and document summarization (Lin and Bilmes, 2011), among others. Recent work has also focused on identifying centralized algorithms to solve submodular maximization problems in resilient ways (Tzoumas et al., 2018) or subject to informational constraints (Gharsifard and Smith, 2018). It is well-known that many multiagent optimization problems with submodular objective functions have nontrivial worst-case equilibrium quality guarantees. In particular, when $U_i = W$ the price of anarchy of a submodular multiagent optimization problem is known to be at least $1/2$ (Vetta, 2002).

Our first question is this: do there exist any local adaptation rules for submodular problems which provide efficiency guarantees that are close to the nominal price of anarchy of $1/2$? In answer, Theorem 2 indicates that in this setting, there exist payoff evaluators which give efficiency guarantees that gracefully degrade from the nominal $1/2$ as the information graph becomes increasingly disconnected. Furthermore, Theorem 2 provides a lower bound on the efficiency degradation which depends on the structure of the information graph.

Theorem 2 applies to all Nash equilibria; accordingly, we write a *mixed strategy* for player i as $a_i \in \Delta(\mathcal{A}_i)$, where $\Delta(\mathcal{A}_i)$ denotes the standard probability simplex over $\mathcal{A}_i$. For simplicity, we write $\Delta(\mathcal{A}) := \prod_{i \in \mathcal{I}} \Delta(\mathcal{A}_i)$ and $\Delta(\mathcal{A}_{-i}) := \prod_{j \neq i} \Delta(\mathcal{A}_j)$. If players are selecting mixed strategies, their utility functions are evaluated in expectation over the joint probability distribution of play in the standard way, so that $U : \Delta(\mathcal{A}) \rightarrow \mathbb{R}$; similarly, for any $a \in \Delta(\mathcal{A})$, we define the extended system objective $W(a)$ as the expected value of W under the distribution a .

Agent i 's *best response set* for an strategy profile $a_{-i} \in \Delta(\mathcal{A}_{-i})$ is defined as $B_i(a_{-i}) := \arg \max_{a_i \in \Delta(\mathcal{A}_i)} U_i(a_i, a_{-i})$. A strategy profile $a^{\text{ne}} \in \Delta(\mathcal{A})$ is a *Nash equilibrium* if for each agent i , $a_i^{\text{ne}} \in B_i(a_{-i}^{\text{ne}})$. We write the set of all Nash equilibria of multiagent optimization problem G as $\text{NE}(G)$.

We say that a graph $\mathcal{N}$ is *acyclic* if it contains no cycles; i.e., if every directed path in $\mathcal{N}$ has finite length. We write $\text{MF}(\mathcal{N})$ to denote the smallest set of agents whose removal (along with associated edges) renders $\mathcal{N}$ acyclic. In graph-theoretic terms, $\text{MF}(\mathcal{N})$ is the *minimum feedback vertex set* of graph $\mathcal{N}$.

We define the *disconnection factor* $\delta(\mathcal{N})$ as the cardinality of the minimum feedback vertex set of the complement graph: $\delta(\mathcal{N}) := |\text{MF}(\mathcal{N}^c)|$. Note that $\delta(\mathcal{N})$ expresses the “disconnectedness” of the information graph $\mathcal{N}$. For instance, if $\mathcal{N}$ is a complete graph (all agents can observe all agents), $\delta(\mathcal{N}) = 0$; if $\mathcal{N}$ is empty (no agent can observe any other agent), then $\delta(\mathcal{N}) = n - 1$. Furthermore, $\delta(\mathcal{N})$ is monotone in the following sense: if $\mathcal{N} \subseteq \mathcal{N}'$ are information graphs over $\mathcal{I}$, then $\delta(\mathcal{N}) \geq \delta(\mathcal{N}')$; adding an edge to a graph can only decrease its disconnection factor. We are now prepared to state our result:

Theorem 2 *Let $\mathcal{G}_{\text{SM}}$ be the family of multiagent optimization problems with nondecreasing, normalized, submodular system-level objective functions. For any information graph $\mathcal{N}$,*

$$\sup_{f \in \mathcal{F}} \inf_{G \in \mathcal{G}_{\text{SM}}} \left(\frac{\min_{a_f^{\text{ne}} \in \text{NE}(G_f)} W(a_f^{\text{ne}})}{\max_{a^{\text{ne}} \in \text{NE}(G)} W(a^{\text{ne}})} \right) \geq \frac{1}{2 + \delta(\mathcal{N})}. \quad (9)$$

Furthermore, a local adaptation rule f which achieves this bound is the minimum payoff evaluator $f_{\min}$, item 3 in Proposition 1.

Note that in a common interest game, the denominator in (9) resolves to the system-optimal objective value, and the numerator captures the worst-case degradation possible in the game's equilibria due to information losses. Thus, (9) represents a bound on the degree to which information losses can harm the emergent behavior of the multiagent system.

As we proceed to prove Theorem 2, we first highlight the structure of the payoff evaluator which gives (9). First, the lower bound in (9) is obtained using the *minimum* evaluator (see Proposition 1, item 3). While the structure of the *optimal* evaluator remains an open question, Theorem 2 demonstrates that choosing the evaluator correctly ensures that the performance guarantee associated with locally-adapted utility functions degrades gracefully as a function of the disconnection factor of the graph. The minimum evaluator has an appealing intuitive interpretation when the objective W is submodular and nondecreasing. That is, the minimum evaluator is equivalent to assuming that unobservable agents are selecting the $\emptyset$ action. For information graph with $\mathcal{N}_i$, the minimum evaluator generates a proxy utility function equal to

$$\begin{aligned} \mathcal{U}_i(a) &= W(a_{\mathcal{N}_i}) \\ &\leq W(a). \end{aligned} \quad (10)$$

Thus, Theorem 2 provides that the simple policy of “if I cannot see you, I will assume that you are not present” yields nontrivial equilibrium performance guarantees. We now proceed with the proof.

3.1 Proofs for Theorem 2

The proof of Theorem 2 proceeds by developing a sequence of inequalities which bound $W(a^{\text{opt}})$ using various properties of submodular problems. To facilitate these arguments, throughout this proof we replace the agents’ utility functions with *marginal-cost* utility functions as follows. For agent i , let the *marginal-cost* utility function be given by

$$\bar{U}_i(a) := W(a_i, a_{-i}) - W(a_{-i}). \quad (11)$$

Note that an action profile a^{ne} is a Nash equilibrium for marginal-cost utility functions $\bar{U}$ if and only if it is a Nash equilibrium for common-interest utility functions considered in this paper; this is because for any $a_i, a'_i \in \mathcal{A}_i$, it holds that

$$\bar{U}_i(a_i, a_{-i}) - \bar{U}_i(a'_i, a_{-i}) = W(a_i, a_{-i}) - W(a'_i, a_{-i}).$$

This equivalence holds for the adapted utility functions as well when considering the minimum evaluator. That is, Agent i ’s adapted marginal-cost utility function is simply given by

$$\bar{\mathcal{U}}_i(a) = W(a_{\mathcal{N}_i}) - W(a_{\mathcal{N}_i} \setminus a_i), \quad (12)$$

which (just as in (10)) is equivalent to assuming that unobservable agents are selecting $\emptyset$.

Thus, throughout the proofs for Theorem 2, without loss of generality we define all Nash equilibria with reference to the marginal-cost utility functions $\bar{U}$ defined in (11) and their adapted versions $\bar{\mathcal{U}}$ defined in (12). We begin with the following lemma.

Lemma 1 *Let $G \in \mathcal{G}_{\text{SM}}$ be a submodular multiagent optimization problem. Suppose that the associated information graph $\mathcal{N}$ is such that for some subset of agents $N \subseteq \mathcal{I}$, the subgraph of $\mathcal{N}$ associated with agents in N contains a complete DAG⁴. Then if agents apply the minimum evaluator (equivalently, they evaluate their nominal utility functions with unobservable agents selecting $\emptyset$), the following holds for any mixed strategy $a \in \Delta(\mathcal{A})$:*

$$\sum_{i \in N} \bar{\mathcal{U}}_i(a) \leq W(a). \quad (13)$$

⁴ A complete DAG is a directed acyclic graph such that the addition of any edge would create a cycle.

Proof Without loss of generality, assume that the agents are indexed in increasing order of out-degree; or $|\mathcal{N}_i| < |\mathcal{N}_{i+1}|$. For each agent $i \in N$, let $\mathcal{N}_i^* := \{1, \dots, i-1\}$ (and define $\mathcal{N}_1^* = \emptyset$). Note that if $(\mathcal{N}_i)_{i \in N}$ contains a complete DAG, this implies that $(\mathcal{N}_i^*)_{i \in N} \subseteq (\mathcal{N}_i)_{i \in N}$. Then, assuming $|N| = n$,

$$\begin{aligned} \sum_{i \in N} \bar{U}_i(a) &= \sum_{i \in N} [W(a_{\mathcal{N}_i}) - W(a_{\mathcal{N}_i} \setminus a_i)] \\ &\leq \sum_{i=1}^n [W(a_{\mathcal{N}_i^*}) - W(a_{\mathcal{N}_i^*} \setminus a_i)] \\ &= W(a_{\mathcal{N}_n^*}) - W(\emptyset) \\ &\leq W(a). \end{aligned} \tag{14}$$

The first inequality follows simply from the submodularity and monotonicity of W and (12), yielding a telescoping sum; the second inequality follows from the monotonicity and normalization of W . $\square$

Proof of Theorem 2: Let $G \in \mathcal{G}_{\text{SM}}$ be a submodular maximization problem with nondecreasing and normalized objective function W . Let $a^{\text{opt}} \in \Delta(\mathcal{A})$ be an optimal solution to W , and let $a^{\text{ne}} \in \Delta(\mathcal{A})$ be a Nash equilibrium associated with proxy utility functions computed using the minimum evaluator $f_{\min}$ yielding utility functions (12). Let $M := \text{MF}(\mathcal{N}^c)$ be a minimum feedback vertex set of $\mathcal{N}^c$ (so that $\delta(\mathcal{N}) = |M|$). Then

$$\begin{aligned} W(a^{\text{opt}}) &\leq W(a^{\text{ne}}) + \sum_{i=1}^n \bar{U}_i(a^{\text{ne}}) \\ &= W(a^{\text{ne}}) + \sum_{i \in M} \bar{U}_i(a^{\text{ne}}) + \sum_{j \in \mathcal{I} \setminus M} \bar{U}_j(a^{\text{ne}}) \\ &\leq (1 + |M|)W(a^{\text{ne}}) + \sum_{j \in \mathcal{I} \setminus M} \bar{U}_j(a^{\text{ne}}) \\ &\leq (2 + |M|)W(a^{\text{ne}}). \end{aligned} \tag{15}$$

The first inequality is borrowed from (Vetta, 2002, Proof of Theorem 3); it is a consequence of utility functions computed by (12) and the definition of Nash equilibrium. The second inequality follows from the fact that by submodularity, monotonicity, the form of proxy utility functions from (10), and $\bar{U}_i(a) \leq W(a_i) \leq W(a)$. The last inequality follows from Lemma 1; note that the information (sub)graph of $\mathcal{I} \setminus M$ must contain a complete DAG, since its associated complement graph $\mathcal{N}^c$ is acyclic. The desired result is obtained by observing that by definition, $\delta(\mathcal{N}) = |M|$. $\square$

3.2 The Role of Graph Structure in Equilibrium Quality

While the lower bound in Theorem 2 holds for any information graph $\mathcal{N}$, we note that the relationship between a graph $\mathcal{N}$ and its disconnection factor $\delta(\mathcal{N})$

may not be immediately obvious. Indeed, identifying the minimum feedback vertex set of an arbitrary undirected graph is NP-Hard (Fomin et al., 2006). However, it is quite simple to parse for certain types of graph. For instance, a simple consequence of Theorem 2 is that if the complement graph $\mathcal{N}^c$ is acyclic, the nominal efficiency guarantee of $1/2$ is preserved:

Corollary 1 *Let $\mathcal{G}_{\text{SM}}$ denote the set of multiagent optimization problems with monotone, normalized, submodular objective functions, and let $\mathcal{N}_a$ be such that its associated complement graph $\mathcal{N}_a^c$ is acyclic. Then the optimal evaluator preserves nominal efficiency guarantees:*

$$\sup_{f \in \mathcal{F}} \inf_{G \in \mathcal{G}_{\text{SM}}} \left(\frac{\min_{a_f^{\text{ne}} \in \text{NE}(G_f)} W(a_f^{\text{ne}})}{\max_{a^{\text{ne}} \in \text{NE}(G)} W(a^{\text{ne}})} \right) = \frac{1}{2}. \quad (16)$$

Proof Note that if $\mathcal{N}^c$ is acyclic, $\delta(\mathcal{N}) = 0$ since no vertices are required to be removed to render $\mathcal{N}^c$ acyclic. Thus, (16) is immediate from (9) and from the known upper bound of $1/2$ for this class of problems. $\square$

We may also leverage straightforward bounds on the cardinality of minimum feedback vertex sets of directed graphs to reinterpret (9) directly in terms of $|\mathcal{N}^c|$ (i.e., the number of edges “missing” from $\mathcal{N}$).

Corollary 2 *Let $\mathcal{G}_{\text{SM}}$ denote the set of multiagent optimization problems with monotone, normalized, submodular objective functions. For any information graph $\mathcal{N}$,*

$$\sup_{f \in \mathcal{F}} \inf_{G \in \mathcal{G}_{\text{SM}}} \left(\frac{\min_{a_f^{\text{ne}} \in \text{NE}(G_f)} W(a_f^{\text{ne}})}{\max_{a^{\text{ne}} \in \text{NE}(G)} W(a^{\text{ne}})} \right) \geq \frac{1}{2 + \lfloor |\mathcal{N}^c|/2 \rfloor}. \quad (17)$$

Proof First, we must show for any graph $\mathcal{N}$ that $2\delta(\mathcal{N}) \leq |\mathcal{N}^c|$. Note that at least two unique edges must be incident on each vertex in $\text{MF}(\mathcal{N}^c)$; otherwise, $\text{MF}(\mathcal{N}^c)$ would not be minimal. Thus, $\mathcal{N}^c$ must contain at least twice as many edges as vertices in $\text{MF}(\mathcal{N}^c)$, or

$$2\delta(\mathcal{N}) = 2|\text{MF}(\mathcal{N}^c)| \leq e(\mathcal{N}^c). \quad (18)$$

Therefore (9), the integrality of $\delta(\mathcal{N})$, and (18) combine to give (17). $\square$

For a plot depicting the bound in (17), see Figure 3.

3.3 Empirical Comparison of Adaptation Rules

One example of a multiagent optimization problem that has received significant attention in the literature is known as the weighted set coverage problem. Here, there exists a collection of resources $\mathcal{R}$ where each resource $r \in \mathcal{R}$ is associated with a value $v_r \geq 0$. The goal of the weighted set coverage problem is to allocate agents over resources to maximize the total value of the covered

resources. To that end, we have a collection of agents $\mathcal{I}$ where each agent is associated with a given action set $\mathcal{A}_i \subseteq \bar{\mathcal{A}}_i = 2^{\mathcal{R}}$ that defines the set of permissible coverings. It is important to emphasize that agent i does not choose the action set $\mathcal{A}_i$; rather, the action set is chosen exogenously and the agent is tasked with choosing a particular covering choice $a_i \in \mathcal{A}_i$. The goal of the weighted set covering problem is to choose an admissible collective allocation $a = (a_1, \dots, a_n) \in \mathcal{A}$ that optimizes a system-level objective function of the form $W(a) = \sum_{r \in \bigcup_{i \in \mathcal{I}} \mathcal{A}_i} v_r$. Several different utility design methodologies have been considered for the maximum weighted set covering problem. The most natural design choice is that of common interest, where each agent $i \in \mathcal{I}$ is assigned a utility function of the form $U_i(a) = W(a)$ for any $a \in \bar{\mathcal{A}}$. This design choice ensures that the resulting game is an exact potential game and further that the price of anarchy is 0.5, meaning that such a choice ensures that all resulting pure Nash equilibria have performance within 50% of optimal irrespective of the specific weighted set covering problem (Vetta, 2002). Note that this design choice does not ensure that all Nash equilibria are optimal, as suboptimal Nash equilibrium can also exist as well.

It well-known that the weighted set cover objective is submodular, non-decreasing, and normalized; thus, weighted set cover problems satisfy the requirements of Theorem 2. Accordingly, we use this class of problems to perform an empirical comparison between the Minimum evaluator (which yields the bound in Theorem 2) and the Maximum evaluator (see Proposition 1, item 2). At first glance, in the weighted set cover problem, the Maximum evaluator seems to offer some hope for providing good performance under informational inconsistencies, since it biases agents toward resources that cannot be covered by other (unobserved) agents; thus, one might expect it to lead agents to avoid redundancies.

For concreteness, we first provide a simple example of a weighted set cover problem to illustrate the concept in Figure 4. This example has two agents $\mathcal{I} = \{1, 2\}$ and three resources $\mathcal{R} = \{a, b, c\}$ with values $v_a = v_c = 0.1$ and $v_b = 1$. Agent 1 can cover either resource a or b , and Agent 2 can cover either resource b or c ; i.e., $\mathcal{A}_1 = \{a, b\}$ and $\mathcal{A}_2 = \{b, c\}$. In this example, an optimal action profile has an objective value of 1.1, with one agent covering resource b and the other agent covering a different resource. If both agents can observe each others' actions, any optimal action profile is also a Nash equilibrium.

However, consider the scenario in which neither agent can observe the other's action; i.e., $\mathcal{N}_1 = \{1\}$ and $\mathcal{N}_2 = \{2\}$. In this case, the complement graph of $\mathcal{N}$ has a single directed cycle, so $\delta(\mathcal{N}) = 1$ and by Theorem 2, the Minimum payoff evaluator ensures that all Nash equilibria have objective value at least $1.1/3$.

Figure 4 depicts the payoff matrices corresponding to the Maximum and Minimum evaluators for this example problem; note that in this case the Maximum evaluator (bottom center) causes every action profile to be a Nash equilibrium (even the highly suboptimal (a, c) outcome), whereas the Minimum evaluator (bottom right) induces only one moderately-good equilibrium (b, b) .

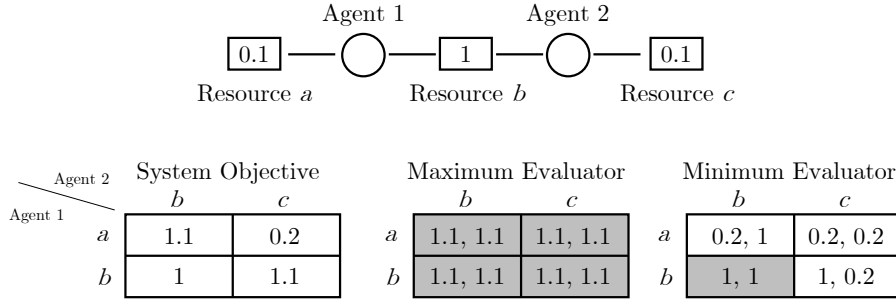


Fig. 4 Top: layout of game for illustrative example in text. Two agents (circles) each have access to two resources (boxes); resource values are indicated in the respective box. Bottom Left: system objective function with respect to agent selections; note that the objective is maximized when one agent selects b and the other agent selects a or c . Bottom Center/Right: If neither agent can observe the actions of the other, this is the payoff matrix induced by the Maximum and Minimum evaluators, respectively. Nash equilibria are shaded gray. Note that for the Maximum evaluator, all action profiles (even the highly suboptimal (a, c)) are Nash equilibria; however, the Minimum evaluator has only the nearly-optimal action profile (b, b) as a Nash equilibrium.

On this problem, note that the Maximum evaluator's worst Nash equilibrium is considerably worse than the lower bound ensured by the Minimum evaluator.

The results of our empirical study (depicted in Figure 5) corroborate this finding: the Minimum evaluator significantly outperforms the Maximum evaluator on random problems as well. To conduct this study, we generated 1200 random weighted set cover problems in the following way: each game was generated with 5 agents, 8 total resources, 3 randomly-selected resources available to each agent, and resource values selected uniformly at random from $[0, 1]$. For each randomly-generated game G , we computed the best nominal Nash equilibrium $a^{\text{ne}} \in \arg \max_{a \in \text{PNE}(G)} W(a)$. We then subjected the game to a sequence of edge removals; at each stage, we removed one uniformly-selected edge from $\mathcal{N}$, up to a total of 18 removed edges. For each information graph in the resulting sequence of graphs, we applied each of the considered evaluators and then computed the *worst* Nash equilibrium associated with that evaluator. Let G_f^k denote the game with k edges removed, subjected to the evaluator f ; then we may write the computed worst-case Nash equilibrium as $a_f^k \in \arg \min_{a \in \text{PNE}(G_f^k)} W(a)$ for each evaluator $f \in \{f_{\min}, f_{\max}\}$. Finally, for game G with k edges removed, we record the value of $W(a_f^k)/W(a^{\text{ne}})$ for each of the evaluators. Each (gray) trace in the plots in Figure 5 corresponds to this ratio evaluated for a single game parameterized by $|\mathcal{N}^c|$, the number of edges removed. The bold colored trace in each plot corresponds to the arithmetic mean of those traces.

Referring to Figure 5, note that the Minimum evaluator significantly outperforms the Maximum evaluator, and this advantage is the most pronounced on highly-disconnected graphs (large values of $|\mathcal{N}^c|$). First, Minimum average performance (blue trace) exceeds Maximum average performance (red trace) for all values of $|\mathcal{N}^c|$. Indeed, Minimum average performance is slightly higher

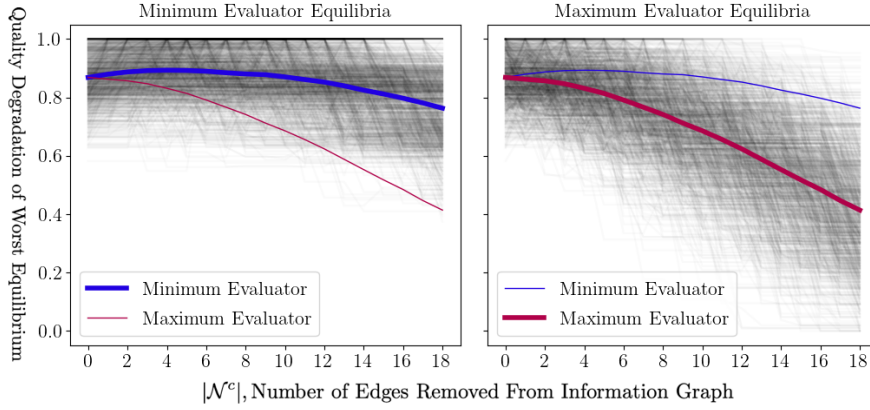


Fig. 5 Results of empirical investigation into the effect of local adaptation rule on equilibrium quality. To generate one (gray) trace on each plot a random weighted set cover game was generated with 5 agents, 8 total resources, 3 randomly-selected resources available to each agent, and resource values selected uniformly at random from $[0, 1]$. A single (gray) trace corresponds to a single random game under a sequence of removed edges; it records the objective value of the worst Nash equilibrium divided by the optimal objective value. The left plot was generated using the Minimum evaluator which is used to prove the lower bound in Contribution 2; and the blue trace in each plot reports the average performance of worst-case equilibria for the Minimum evaluator. The left plot was generated using the Maximum evaluator (see Proposition 1, item 2), and the red trace in each plot reports the average performance of worst-case equilibria for the Maximum evaluator. Note that for highly-disconnected graphs, the Minimum evaluator enjoys a substantial advantage over the Maximum evaluator.

for mostly-connected graphs than it is for fully-connected (complete) graphs. Second, note that for highly disconnected graphs, nearly all equilibria associated with Minimum (gray traces on the left plot) outperform Maximum’s average; conversely, nearly all equilibria associated with Maximum (gray traces on the right plot) *underperform* Minimum’s average. Finally, consider highly-disconnected information graphs in the plots in Figure 5, with approximately $|\mathcal{N}^c| > 14$. Even here, a setting in which almost no agent can observe other agents, the Minimum evaluator often induces equilibria which are as good as the *best* equilibria achievable when $\mathcal{N}$ is fully connected (i.e., their degradation reported on the plot is 1). However, in the same connectivity regime, the Maximum evaluator rarely or never induces equilibria of such high quality; practically all of its equilibria are worse than the best nominal equilibria, and some even obtain an objective value of 0.

While this empirical analysis is limited to a single class of simple submodular problems with a relatively low number of agents, it indicates that the problem of selecting a payoff evaluator is nontrivial, and that a misguided selection may lead to unnecessarily poor equilibrium performance.

4 Conclusions

This paper represents an initial study on the applicability of local methods for agents to adapt to informational inconsistencies in multiagent optimization, and suggests several open areas for future study. For example, it is unclear precisely what types of optimization problem might be subject to the fragility reported in Theorem 1, and future research could focus on identifying particular problem structures that render a problem susceptible to these issues. Another open question pertains to the optimal local adaptation rule for submodular problems. Theorem 2 shows a lower bound which is associated with the minimum evaluator and our empirical evidence suggests that the maximum evaluator is worse, but it is an open question whether any evaluators exist which can provably outperform the minimum evaluator in any sense.

5 Data Availability

Data sharing not applicable to this article as no datasets were generated or analysed during the current study.

References

- Alós-Ferrer C, Netzer N (2010) The logit-response dynamics. *Games and Economic Behavior* 68(2):413–427, DOI 10.1016/j.geb.2009.08.004, URL <http://dx.doi.org/10.1016/j.geb.2009.08.004>
- Barinova O, Lempitsky V, Kholi P (2012) On detection of multiple object instances using hough transforms. *IEEE Transactions on Pattern Analysis and Machine Intelligence* DOI 10.1109/TPAMI.2012.79
- Blume LE (1993) *The Statistical Mechanics of Strategic Interaction*. DOI 10.1006/game.1993.1023
- Blume LE (1995) The Statistical Mechanics of Best-Response Strategy Revision. *Games and Economic Behavior* 11(2):111–145
- Brown PN, Marden JR (2019) On the feasibility of local utility redesign for multiagent optimization. In: 18th European Control Conference, pp 3396–3401, URL <https://ieeexplore.ieee.org/document/8795902>
- Candogan O, Ozdaglar A, Parrilo PA (2013) Dynamics in near-potential games. *Games and Economic Behavior* 82:66–90, DOI 10.1016/j.geb.2013.07.001, URL <http://dx.doi.org/10.1016/j.geb.2013.07.001>, 1107.4386
- Collins BC, Hines L, Barboza G, Brown PN (2021) Robust stochastic stability in dynamic and reactive environments. In: 2021 60th IEEE Conference on Decision and Control (CDC), pp 1892–1897, DOI 10.1109/CDC45484.2021.9683342
- Ferguson BL, Brown PN, Marden JR (2021) The Effectiveness of Subsidies and Tolls in Congestion Games. *IEEE Transactions on Automatic Control* pp 1–

- 1, DOI 10.1109/TAC.2021.3088412, URL <http://arxiv.org/abs/2102.09655>
<https://ieeexplore.ieee.org/document/9451652/>, 2102.09655
- Fomin FV, Gaspers S, Pyatkin AV (2006) Finding a minimum feedback vertex set in time $\mathcal{O}(1.7548^n)$. In: International Workshop on Parameterized and Exact Computation, pp 184–191, URL https://link.springer.com/content/pdf/10.1007/11847250_17.pdf
- Fudenberg D, Levine D (1998) The Theory of Learning in Games. MIT Press, Cambridge, MA
- Gharesifard B, Smith SL (2018) Distributed Submodular Maximization with Limited Information. *IEEE Transactions on Control of Network Systems* 5(4):1635–1645, DOI 10.1109/TCNS.2017.2740625, 1706.04082
- Grimsman D, Seaton JH, Marden JR, Brown PN (2020) The Cost of Denied Observation in Multiagent Submodular Optimization. In: 2020 59th IEEE Conference on Decision and Control (CDC), IEEE, vol 2020-Decem, pp 1666–1671, DOI 10.1109/CDC42340.2020.9304054, URL <https://ieeexplore.ieee.org/document/9304054/>, 2009.05018
- de Jong J, Uetz M (2020) The quality of equilibria for set packing and throughput scheduling games. *International Journal of Game Theory* 49(1):321–344, DOI 10.1007/s00182-019-00693-1, URL <http://link.springer.com/10.1007/s00182-019-00693-1>
- Kempe D, Kleinberg J, Tardos É (2003) Maximizing the spread of influence through a social network. In: Proceedings of the ninth ACM SIGKDD international conference on Knowledge discovery and data mining - KDD '03, pp 137–146, DOI 10.1145/956755.956769, 0806.2034v2
- Kleer P, Schäfer G (2017) Path deviations outperform approximate stability in heterogeneous congestion games. In: International Symposium on Algorithmic Game Theory, pp 212–224, URL <http://arxiv.org/abs/1707.01278>, 1707.01278
- Krause A, McMahan HB, Guestrin C, Gupta A (2008) Robust Submodular Observation Selection. *Journal of Machine Learning Research* 9(Dec):2761–2801, DOI 10.1021/ma00200a025
- Li N, Marden JR (2013) Designing games for distributed optimization. *IEEE Journal on Selected Topics in Signal Processing* 7(2):230–242, DOI 10.1109/JSTSP.2013.2246511
- Li SHQ, Ratliff L, Acikmese B (2021) Disturbance Decoupling for Gradient-Based Multi-Agent Learning With Quadratic Costs. *IEEE Control Systems Letters* 5(1):223–228, DOI 10.1109/LCSYS.2020.3001240, URL <https://ieeexplore.ieee.org/document/9112258/>
- Lin H, Bilmes J (2011) A Class of Submodular Functions for Document Summarization. In: 49th Annual Meeting of the Association for Computational Linguistics, pp 510–520, DOI 10.1002/art.33407
- Marden JR, Shamma JS (2014) Game Theory and Distributed Control. In: Young H, Zamir S (eds) *Handbook of Game Theory Vol. 4*, Elsevier Science
- Ni J, Tang G, Mo Z, Cao W, Yang SX (2020) An Improved Potential Game Theory Based Method for Multi-UAV Cooperative Search. *IEEE Access* 8:47787–47796, DOI 10.1109/ACCESS.2020.2978853, URL

- <https://ieeexplore.ieee.org/document/9026958/>
- Paccagnan D, Chandan R, Marden JR (2020) Utility Design for Distributed Resource Allocation—Part I: Characterizing and Optimizing the Exact Price of Anarchy. *IEEE Transactions on Automatic Control* 65(11):4616–4631, DOI 10.1109/TAC.2019.2961995, URL <https://ieeexplore.ieee.org/document/8941319/>
- Papadimitriou C (2001) Algorithms, games, and the internet. In: *STOC '01: Proceedings of the thirty-third annual ACM symposium on Theory of computing*, pp 749–753, DOI <http://doi.acm.org/10.1145/380752.380883>
- Pradeliski B, Young HP (2012) Learning Efficient Nash Equilibria in Distributed Systems. *Games and Economic Behavior* 75(2):882–897
- Ramaswamy Pillai V, Paccagnan D, Marden JR (2021) Multiagent Maximum Coverage Problems: The Trade-off Between Anarchy and Stability. *IEEE Transactions on Automatic Control* pp 1–1, DOI 10.1109/TAC.2021.3069809, URL <https://ieeexplore.ieee.org/document/9393567/>
- Seaton JH, Brown PN (2022) All stable equilibria have improved performance guarantees in submodular maximization with communication-denied agents. *IEEE Control Systems Letters* pp 1–1, DOI 10.1109/LCSYS.2022.3166748
- Tzoumas V, Gatsis K, Jadbabaie A, Pappas GJ (2018) Resilient monotone submodular function maximization. 2017 IEEE 56th Annual Conference on Decision and Control, CDC 2017 2018-Janua(Cdc):1362–1367, DOI 10.1109/CDC.2017.8263844, 1703.07280
- Vetta A (2002) Nash equilibria in competitive societies, with applications to facility location, traffic routing and auctions. *The 43rd Annual IEEE Symposium on Foundations of Computer Science*, 2002 Proceedings pp 416–425, DOI 10.1109/SFCS.2002.1181966, URL <http://ieeexplore.ieee.org/lpdocs/epic03/wrapper.htm?arnumber=1181966>
- Wolpert DH, Tumer K, Frank J (1999) Using collective intelligence to route Internet traffic. *Proceedings of the 1998 conference on Advances in neural information processing systems II* pp 952–958, 9905004
- Yang C, Li J, Ejaz W, Anpalagan A, Guizani M (2017) Utility function design for strategic radio resource management games: An overview, taxonomy, and research challenges. *Transactions on Emerging Telecommunications Technologies* 28(5):e3128, DOI 10.1002/ett.3128, URL <https://onlinelibrary.wiley.com/doi/10.1002/ett.3128>
- Young HP (1993) The Evolution of Conventions. *Econometrica* 61(1):57–84, DOI 10.1007/s11229-006-9034-z

APPENDIX: Proofs for Section 2

First we present the proof for Proposition 1 characterizing several acceptable evaluators.

Proof of Proposition 1 To see that each satisfies Definition 2, let $S \in \mathbb{R}^k$ and $S' \in \mathbb{R}^k$ satisfy the assumption of Definition 2. Arrange S and S' in ascending

order and denote the i -th element of S and S' as s_i and s'_i , respectively so that $\min_{s \in S} = s_1$ and $\max_{s \in S} = s_k$. Thus, $f_{\text{sum}}(S) = \sum_{i=1}^k s_i > \sum_{i=1}^k s'_i > f_{\text{sum}}(S')$. Since f_{sum} satisfies Definition 2, it must be true that f_{mean} does as well. To see that f_{max} and f_{min} satisfy Definition 2, simply note that $f_{\text{max}}(S) = s_k > s'_k = f_{\text{max}}(S')$ and $f_{\text{min}}(S) = s_1 > s'_1 = f_{\text{min}}(S')$. $\square$

Next, we present the proof that general unrestricted multiagent optimization problems are fragile.

Proof of Proposition 2. Consider the multiagent optimization problem depicted in Figure 6, where $\epsilon > 0$ is a small positive constant. This problem has 3 agents with two possible actions each: $\bar{\mathcal{A}}_1 = \bar{\mathcal{A}}_2 = \bar{\mathcal{A}}_3 = \{1, 2\}$. Now, suppose that the realized problem has $\mathcal{A}_1 = \{1\}$; that is, Agent 1 does not have access to action 2 and thus selects the fixed action $a_1 = 1$ (i.e., selects the left-hand payoff matrix in Figure 6). Furthermore, suppose that the information graph is such that Agent 2 cannot observe the action choice of Agent 1 (formally: $\mathcal{N}_1 = \mathcal{N}_3 = \{1, 2, 3\}$ but $\mathcal{N}_2 = \{2, 3\}$ so that Agent 2 can observe only the choice of Agent 3).

Thus, the realized game is restricted simply to the left-hand payoff matrix in Figure 6, but Agent 2 lacks the ability to observe this. Since Agent 2 cannot observe Agent 1, Agent 2 applies an acceptable evaluator f to Agent 1. Agent 2's adapted utility function $\mathcal{U}_2$ is

		Agent 3	
		1	2
Agent 2	1	$f(\{0, 1\})$	$f(\{\epsilon, 1 - \epsilon\})$
	2	$f(\{\epsilon, 1 + \epsilon\})$	$f(\{2\epsilon, 1\})$

By the definition and established properties of acceptable evaluators, it must hold that $f(\{\epsilon, 1 + \epsilon\}) > f(\{0, 1\})$ and that $f(\{2\epsilon, 1\}) > f(\{\epsilon, 1 - \epsilon\})$. That is, Agent 2's adapted payoffs are such that $a_2 = 2$ becomes a strictly dominant strategy, leaving a unique dominant-strategy Nash equilibrium of $a_2 = 2$ and $a_3 = 2$ with an objective value of only 2ϵ . The result is proved by taking the limit as $\epsilon \rightarrow 0$. $\square$

$a_1 = 1$		Agent 3	
		1	2
Agent 2	1	<u>1</u>	<u>$1 - \epsilon$</u>
	2	<u>ϵ</u>	<u>2ϵ</u>

$a_1 = 2$		Agent 3	
		1	2
Agent 2	1	0	ϵ
	2	$1 + \epsilon$	1

Fig. 6 Payoff matrix for pathological game used to prove Proposition 2. As described in the proof, $\bar{\mathcal{A}}_1 = \{1, 2\}$, but that in the realized problem $\mathcal{A}_1 = \{1\}$. That is, Agent 1 lacks access to action 2 (that is, the action which selects the right-hand payoff matrix), so Agent 1 selects a fixed action $a_1 = 1$. Agent 2 cannot observe the action choice of Agent 1, so Agent 2 applies an acceptable evaluator to Agent 1. In the adapted game, whenever Agent 3 selects action $a_3 = k$, Agent 2's unique best response (according to the adapted payoffs) is $a_2 = 2$ as depicted by the underlines in the left-hand payoff matrix.

Finally, we present the proof of Theorem 1, showing that this fragility persists even in problems with a single ϵ -inconsequential agent which cannot be seen by a single other agent.

Proof of Theorem 1. We will construct a multiagent optimization problem $G \in \mathcal{G}_\epsilon$ in which Agent 1 is ϵ -inconsequential. In this problem, Agent 2 is unable to observe the action choices of Agent 1. However, despite Agent 1's inconsequentiality, regardless of which acceptable evaluator is applied by Agent 2, the problem's Nash equilibria are rendered arbitrarily poor. Furthermore, this problem has a unique Nash equilibrium in the adapted case (in which Agent 2 cannot observe Agent 1) which can be reached from any joint action via a finite sequence of payoff-improving unilateral deviations (i.e., the adapted game is weakly-acyclic under better replies; see (Marden and Shamma, 2014)).

Consider the multiagent optimization problem depicted in Figures 7 and 8, defined for small positive constant $\delta > 0$. This problem has 3 agents; $\bar{\mathcal{A}}_1 = \{1, 2\}$, and $\bar{\mathcal{A}}_2 = \bar{\mathcal{A}}_3 = \{1, 2, \dots, K\}$ where $K = \lfloor 1/\delta - 3 \rfloor$. First, note that Agent 1 is 3δ -inconsequential. Since $\max_{a \in \bar{\mathcal{A}}} W(a) = 1$, verifying 3δ -inconsequentiality reduces to verifying that a unilateral deviation by Agent 1 changes the objective value by no more than 3δ for any action profile; this is easily verified via Figure 7 by comparing like cells in the left ($a_1 = 1$) and right ($a_1 = 2$) matrices and via Figure 8 by comparing like cells in the upper ($a_1 = 1$) and lower ($a_1 = 2$) matrices.

Now, consider a system realization in which $\mathcal{A}_1 = \{1\}$; that is, Agent 1 only has access to a single action (i.e., the action which yields the left matrix in Figure 7 and the upper matrix in Figure 8). This effectively reduces the game to a two-player game between Agent 2 and Agent 3 with the common interest utility function depicted on the left in Figure 7 and the upper matrix in Figure 8). Throughout, we will write an action profile in the form (a_2, a_3) . As can readily be verified via the payoff matrices, this game has two pure Nash equilibria: $(1, 2)$ and $(2, 1)$ which each have an objective value of $1 - 2\delta$. Note also that no other action profile is a pure Nash equilibrium: in each action profile with $W(a) = 0$, one or both agents can deviate to obtain a nonzero payoff. In action profile (K, K) , Agent 2 can deviate to $a_2 = K - 1$ for a payoff improvement, and in every other action profile with nonzero payoffs, Agent 3 can deviate to $a_3 = 1$ (the left-most column) to obtain a payoff improvement. Thus, we have that

$$\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}}) \geq 1 - 2\delta. \quad (19)$$

Now, consider an information graph in which Agent 1 and Agent 3 can observe all agents' action choices, but Agent 2 cannot observe Agent 1. That is, $\mathcal{N}_1 = \mathcal{N}_3 = \{1, 2, 3\}$, but $\mathcal{N}_2 = \{2, 3\}$. Since Agent 2 cannot observe the action choice of Agent 1, Agent 2 applies an acceptable evaluator f to Agent 1. Now, suppose Agent 3 selects action $a_3 = k \in \{1, 2, \dots, K\}$. Agent 2's unique best response (given the adapted utility function) is to select $a_2 = k$. To see this, first let Agent 3's action satisfy $a_3 = k > 1$. Considering Figure 8, note

$W(\cdot), a_1 = 1$		Agent 3				$W(\cdot), a_1 = 2$		Agent 3			
		1	2	3	4			1	2	3	4
Agent 2	1	<u>$1 - 3\delta$</u>	$1 - 2\delta$	0	0	Agent 2	1	<u>1</u>	$1 - 5\delta$	0	0
	2	$1 - 2\delta$	<u>$1 - 4\delta$</u>	$1 - 3\delta$	0		2	$1 - 5\delta$	$1 - \delta$	$1 - 6\delta$	0
	3	$1 - 3\delta$	0	<u>$1 - 5\delta$</u>	$1 - 4\delta$		3	$1 - 6\delta$	0	<u>$1 - 2\delta$</u>	$1 - 7\delta$
	4	$1 - 4\delta$	0	0	<u>$1 - 6\delta$</u>		4	$1 - 7\delta$	0	0	<u>$1 - 3\delta$</u>

Fig. 7 Payoff matrix for pathological game used to prove Theorem 1; Agent 2 and 3 actions $\{1, 2, 3, 4\}$ are depicted. Note that $\bar{\mathcal{A}}_1 = \{1, 2\}$, but that in the realized problem $\mathcal{A}_1 = \{1\}$. That is, Agent 1 lacks access to action 2 (that is, the action which selects the right-hand payoff matrix), so Agent 1 selects a fixed action $a_1 = 1$. Agent 2 cannot observe the action choice of Agent 1, so Agent 2 applies an acceptable evaluator to Agent 1. In the adapted game, whenever Agent 3 selects action $a_3 = k$, Agent 2's unique best response is $a_1 = k$ as depicted by the underlines in the left-hand payoff matrix.

that Agent 2's adapted payoffs are

$$\mathcal{U}_2(k - 1, k) = f(\{1 - (k + 3)\delta, 1 - k\delta\})$$

and

$$\mathcal{U}_2(k, k) = f(\{1 - (k + 2)\delta, 1 - (k - 1)\delta\}).$$

Hence, by the properties of acceptable evaluators, it must hold that $\mathcal{U}_2(k, k) > \mathcal{U}_2(k - 1, k)$ whenever $\delta > 0$. Second, let Agent 3 select action $a_3 = 1$ (the left-most column in Figure 7). Here again it can be verified that Agent 2's best response is $a_2 = 1$, using a similar argument to above.

Finally, we show that the action profile (K, K) is a Nash equilibrium of the adapted game. Since $K = \lfloor 1/\delta - 3 \rfloor$, it holds that

$$\begin{aligned} W(K, K) &= 1 - (K + 2)\delta \\ &= 1 - (\lfloor 1/\delta - 3 \rfloor + 2)\delta \\ &\geq 1 - (1/\delta - 3 + 2)\delta \\ &= \delta. \end{aligned} \tag{20}$$

By previous arguments it is thus a best response for Agent 2, and any deviation by Agent 3 would yield a payoff of 0. Finally the objective value of (K, K) is small (obtained from Figure 9):

$$\begin{aligned} W(K, K) &= 1 - (K + 2)\delta \\ &= 1 - (\lfloor 1/\delta - 3 \rfloor + 2)\delta \\ &< 1 - (1/\delta - 4 + 2)\delta \\ &= 2\delta. \end{aligned} \tag{21}$$

Thus, we have that in the adapted multiagent optimization problem G_f ,

$$\max_{a_f^{\text{ne}} \in \text{PNE}(G_f)} W(a_f^{\text{ne}}) \leq 2\delta. \tag{22}$$

$$W(\cdot), a_1 = 1$$

		Agent 3			
		1	...	<u>k</u>	<u>$k+1$</u>
Agent 2	$k-1$	$1 - (k-1)\delta$	...	$1 - k\delta$	0
	k	$1 - k\delta$	...	<u>$1 - (k+2)\delta$</u>	<u>$1 - (k+1)\delta$</u>
	$k+1$	$1 - (k+1)\delta$	...	0	<u>$1 - (k+3)\delta$</u>

$$W(\cdot), a_1 = 2$$

		Agent 3			
		1	...	k	$k+1$
Agent 2	$k-1$	$1 - (k+2)\delta$	...	$1 - (k+3)\delta$	0
	k	$1 - (k+3)\delta$	...	$1 - (k-1)\delta$	$1 - (k+4)\delta$
	$k+1$	$1 - (k+4)\delta$	...	0	$1 - k\delta$

Fig. 8 Generic modular payoff matrix block for pathological game used to prove Theorem 1. Note that $\bar{\mathcal{A}}_1 = \{1, 2\}$, but that in the realized problem $\mathcal{A}_1 = \{1\}$. That is, Agent 1 lacks access to action 2 (that is, the action which selects the lower payoff matrix), so Agent 1 selects a fixed action $a_1 = 1$. Agent 2 cannot observe the action choice of Agent 1, so Agent 2 applies an acceptable evaluator to Agent 1. In the adapted game, whenever Agent 3 selects action $a_3 = k$, this leads Agent 2 to best-respond with action $a_1 = k$ as well as depicted by the underlines in the upper payoff matrix.

Combining (19) and (22), we have that for any acceptable evaluator $f \in \mathcal{F}$ and any $\delta > 0$, a multiagent problem G exists for which

$$\frac{\max_{a_f^{\text{ne}} \in \text{PNE}(G_f)} W(a_f^{\text{ne}})}{\min_{a^{\text{ne}} \in \text{PNE}(G)} W(a^{\text{ne}})} \leq \frac{2\delta}{1 - 2\delta}. \quad (23)$$

The proof of Theorem 1 is obtained in the limit by taking $\delta \rightarrow 0$. $\square$

$$W(\cdot), a_1 = 1$$

		Agent 3			
		1	$\dots$	$K-1$	K
Agent 2	$K-2$	$1 - (K-2)\delta$	$\dots$	$1 - (K-1)\delta$	0
	$K-1$	$1 - (K-1)\delta$	$\dots$	$1 - (K+1)\delta$	$1 - K\delta$
	K	0	$\dots$	0	$1 - (K+2)\delta$

$$W(\cdot), a_1 = 2$$

		Agent 3			
		1	$\dots$	$K-1$	K
Agent 2	$K-2$	$1 - (K+1)\delta$	$\dots$	$1 - (K+2)\delta$	0
	$K-1$	$1 - (K+2)\delta$	$\dots$	$1 - (K-2)\delta$	$1 - (K+3)\delta$
	K	0	$\dots$	0	$1 - (K-1)\delta$

Fig. 9 Portion of payoff matrix for pathological game used to prove Theorem 1; Agent 2 and 3 actions $\{K-1, K\}$ are depicted to show the pathological adapted Nash equilibrium at $a_2 = K$ and $a_3 = K$. Note that $\mathcal{A}_1 = \{1, 2\}$, but that in the realized problem $\mathcal{A}_1 = \{1\}$. That is, Agent 1 lacks access to action 2 (that is, the action which selects the lower payoff matrix), so Agent 1 selects a fixed action $a_1 = 1$. Agent 2 cannot observe the action choice of Agent 1, so Agent 2 applies an acceptable evaluator to Agent 1. In the adapted game, whenever Agent 3 selects action $a_3 = k$, this leads Agent 2 to best-respond with action $a_2 = k$ as well as depicted by the underlines in the upper payoff matrix.