

On the Complexity of Learning to Cooperate in Populations of Socially Rational Agents

Bandyopadhyay, Saptarashmi; Çelikok, Mustafa Mert; Loftin, Robert

Publication date
2025

Document Version
Submitted manuscript

Published in
Proceedings of the 24th International Conference on Autonomous Agents and Multiagent Systems, AAMAS 2025

Citation (APA)

Bandyopadhyay, S., Çelikok, M. M., & Loftin, R. (2025). On the Complexity of Learning to Cooperate in Populations of Socially Rational Agents. In Y. Vorobeychik, S. Das, & A. Nowe (Eds.), *Proceedings of the 24th International Conference on Autonomous Agents and Multiagent Systems, AAMAS 2025* (pp. 233-241). (Proceedings of the International Joint Conference on Autonomous Agents and Multiagent Systems, AAMAS). International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS).

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

On the Complexity of Learning to Cooperate with Populations of Socially Rational Agents

Robert Loftin

Department of Computer Science
University of Sheffield
Sheffield, S10 2TN, UK
r.loftin@sheffield.ac.uk

Saptarashmi Bandyopadhyay

Department of Computer Science
University of Maryland
College Park, MD 20742, USA
saptab1@umd.edu

Mustafa Mert Çelikok

Department of Intelligent Systems
Delft University of Technology
Delft, 2600 AA, The Netherlands
m.m.celikok@tudelft.nl

Abstract

Artificially intelligent agents deployed in the real-world will require the ability to reliably *cooperate* with humans (as well as other, heterogeneous AI agents). To provide formal guarantees of successful cooperation, we must make some assumptions about how partner agents could plausibly behave. Any realistic set of assumptions must account for the fact that other agents may be just as adaptable as our agent is. In this work, we consider the problem of cooperating with a *population* of agents in a finitely-repeated, two player general-sum matrix game with private utilities. Two natural assumptions in such settings are that: 1) all agents in the population are individually rational learners, and 2) when any two members of the population are paired together, with high-probability they will achieve at least the same utility as they would under some Pareto efficient equilibrium strategy. Our results first show that these assumptions alone are insufficient to ensure *zero-shot* cooperation with members of the target population. We therefore consider the problem of *learning* a strategy for cooperating with such a population using prior observations its members interacting with one another. We provide upper and lower bounds on the number of samples needed to learn an effective cooperation strategy. Most importantly, we show that these bounds can be much stronger than those arising from a "naive" reduction of the problem to one of imitation learning.

1 Introduction

In this work, we address the problem of learning to cooperate with a *socially intelligent* population of agents from observations interactions between members of this population. We study cooperation in finitely-repeated, two-player, general-sum matrix games with private payoffs. We say that a population of adaptive agents is socially intelligent if its members are (1) individually Hannan-consistent and (2) compatible in the sense that any pair of agents will perform nearly as well as some Pareto-optimal Nash equilibrium of the matrix game. We argue that this model of cooperation is more realistic than those that assume identical payoffs or public utilities. In real-world applications it is unlikely that independent agents will have identical utilities, or that they will provide complete information about their preferences or future behaviour to others. In the case of AI–AI cooperation,

agents developed by different companies will not have access to each other’s source-code, while in the case of human–AI cooperation, having the human fully describe their preferences or behaviour in advance may be infeasible. Therefore, the question we address in this work is: *Can we learn to cooperate with a socially intelligent population of agent by observing its members cooperate with each other?* We answer this question by providing upper and lower bounds on the sample complexity of learning good cooperation strategies.

If we make no assumptions about the target population, we can do little more than attempt to mimic observed behavior as closely as possible, reducing the problem to one of imitation learning. Unfortunately, the strategies of adaptive agents may depend on the full history of interaction, and so the sample complexity of imitation learning will grow exponentially in the length of the repeated game. Our main contribution is an upper-bound showing that, for partners drawn from a socially intelligent (consistent and compatible) population, we can learn to cooperate with far fewer samples than would be required by a pure imitation learning approach.

This result utilizes a class of what we refer to as *imitate-then-commit* strategies, which leverage the fact that the population is socially intelligent to achieve cooperation without perfect imitation. The key idea is that our agent only needs to learn to imitate a member of the target population long enough to for the average strategy to approximate a Pareto-efficient solution. Once such a strategy is identified, our agent can switch to a *coercive* strategy such that any Hannan-consistent partner will either continue to adhere to the current joint strategy, or else switch to a superior strategy, with either case corresponding to “successful” cooperation.

In section 2 formalize our repeated game setting, and provide background on external regret and Hannan-consistency. We also propose a definition of cooperative compatibility (Definition 2.2) that is closely related to the notion of compatibility used in [1]. In Section 2.3, we provide our novel definition of social intelligence, and describe a realistic class of agents that satisfy it. In Section 3 we formalize our learning problem as that of trying to minimize *altruistic regret*, which we argue is the most natural measure of successful cooperation in this setting. We also give lower bounds on its sample complexity under different sets of assumptions. Finally, in Section 4 we present an upper-bound on the number of samples needed to learn strategies that achieve small altruistic regret.

2 Preliminaries

Repeated bi-matrix games with private types. Let $i \in \{1, 2\}$ denote the agent index. We assume both agents have N pure strategies (henceforth “actions”). Let Θ denote the *finite* type space, where $\theta_1, \theta_2 \in \Theta$ denote the *private* types of the two agents, and $\theta = (\theta_1, \theta_2)$ denotes the joint type. We denote agent i ’s payoff matrix as $G(\theta_i) \in \mathbb{R}^{N \times N}$, and let $G(\theta) = [G(\theta_1), G(\theta_2)]^\top$ denote the bi-matrix game parameterized by θ (with agent 1 as the row player). In a single *episode*, the agents play $G(\theta)$ for a fixed number of stages $0 < T < \infty$. We let a_t^1 and a_t^2 denote the actions chosen by agents 1 and 2 in stage $0 < t \leq T$. For mixed strategies $\sigma, \sigma' \in \Delta(N)$, we let $G(\sigma, \sigma'; \theta) = \sigma^\top G(\theta) \sigma'$. We overload a_t^1 and a_t^2 to also denote the mixed strategies that assign all probability mass to actions a_t^1 and a_t^2 , such that $G(a_t^1, a_t^2; \theta_1)$ and $G(a_t^1, a_t^2; \theta_2)$ are agent 1 and 2’s respective payoffs at stage t . We also assume that for all $\theta \in \Theta$, $G_{ij}(\theta) \in [0, 1]$, $\forall i, j \in [N]$.

Let $\mathcal{H}_t = (N \times N)^t$ be the set of histories of length t (with $\mathcal{H}_0 = \{\emptyset\}$), and let $\mathcal{H}_{\leq t} = \bigcup_{s=0}^t \mathcal{H}_s$ be the set of all histories of length at most t . The strategy space Π for an agent is then the space of mappings $\pi : \Theta \times \mathcal{H}_{\leq T-1} \mapsto \Delta(N)$, where $\Delta(N)$ is the set of probability distributions over the action set $[N]$. As a functional, a strategy π maps each type θ to a *behavioral strategy* [2, Chapter 5.2.2] that maps histories of play to action distributions, such that $a_t^i \sim \pi_i(\theta_i, h_{t-1})$. We denote agent i ’s expected total payoff for following strategy π against π' as

$$M_i(\pi, \pi'; \theta, \theta') = \mathbb{E} \left[\sum_{t=1}^T G(a_t^i, a_t^{-i}; \theta_i) \middle| \pi_i = \pi, \pi_{-i} = \pi', \theta_i = \theta, \theta_{-i} = \theta' \right], \quad (1)$$

where the expectation is taken over the actions a_t^i and a_t^{-i} sampled from the agents’ strategies.

	A	B
A	2, 2	0, 0
B	0, 0	1, 1

(a) A fully-cooperative 2x2 matrix game.

	C	D
C	2, 2	0, 3
D	3, 0	1, 1

(b) The prisoner's dilemma game.

Table 1

2.1 Consistency

A natural criterion for rationality is that an agent should attempt to achieve a payoff nearly as large as the best response to its partner's average strategy, which we refer to as *consistency*. To account for the non-stationary behavior of other agents', we specifically consider *Hannan consistency* [3], which in our finite-time setting simply requires that an agent have bounded *external regret* over T stages. The external regret for agent i is defined as

$$R_i^{\text{ext}}(h; \theta) = \max_{a^i \in [N]} \sum_{t=1}^{|h|} \left\{ G(a^i, a_t^{-i}(h); \theta_i) - G(a_t^i(h), a_t^{-i}(h); \theta_i) \right\} \quad (2)$$

where $a_t^i(h)$ denotes the action i played at stage t within the history $h \in \mathcal{H}_{\leq T}$.

Definition 2.1 (Consistency). For $\delta, \epsilon, T > 0$, an agent $i \in \{1, 2\}$ is (δ, ϵ, T) -consistent if, for all types $\theta \in \Theta$, and any partner strategy, we have that $\frac{1}{T} R_i^{\text{ext}}(h_T; \theta) \leq \epsilon$ with probability at least $1 - \delta$.

We also define the *expected* external regret $\bar{R}_i^{\text{ext}}(h; \theta)$ by replacing the $a_t^i(h)$ (the action i played at stage t) with their full strategy $\pi^i(\theta, h_t)$. $R_i^{\text{ext}}(h; \theta)$ and $\bar{R}_i^{\text{ext}}(h; \theta)$ are related by the inequality

$$R_i^{\text{ext}}(h_t; \theta) \leq \bar{R}_i^{\text{ext}}(h_t; \theta) + \sqrt{\frac{T}{2} \ln \frac{1}{\delta}}, \quad (3)$$

which holds w.p. at least $1 - \delta$ for all $t \leq T$ simultaneously (this follows directly from [4, Lemma 4.1]). We therefore only need to bound $\bar{R}_i^{\text{ext}}(h_t; \theta)$ to provide high-probability regret bounds.

2.2 Cooperative compatibility

Even in a fully cooperative game, the fact that both agents are consistent does not guarantee that they will achieve an optimal outcome. In the 2×2 game in Table 1a for example, both (A, A) and (B, B) are Nash equilibria to which consistent agents could converge, but only (A, A) is optimal. In general-sum games, consistency may preclude Pareto-optimal outcomes, as in the classic prisoner's dilemma game (Table 1b), where the only outcome in which neither player incurs positive regret is (D, D) , which is Pareto-dominated by (C, C) . Therefore, similar to [1], we define successful cooperation in terms of the *Pareto-optimal Nash equilibria* (PONE) [5] of a game G .

Let $\mathcal{N}(G) \subseteq \Delta(N) \times \Delta(N)$ be the set of Nash equilibria (NE) of G . For a fully-cooperative game, $\mathcal{N}(G)$ will contain all globally optimal strategy profiles for G . It may, however, also contain joint strategies that are highly sub-optimal. Let $\mathcal{P}(G) \subseteq \mathcal{N}(G)$ denote the set of Pareto optimal Nash equilibria. In this work, we say that a strategy profile $\langle \sigma_1, \sigma_2 \rangle \in \mathcal{P}(G)$ if and only if $\langle \sigma_1, \sigma_2 \rangle \in \mathcal{N}(G)$, and there does not exist $\langle \sigma'_1, \sigma'_2 \rangle \in \mathcal{N}(G)$ such that $G(\sigma'_1, \sigma'_2; \theta_1) > G(\sigma_1, \sigma_2; \theta_1)$ and $G(\sigma'_2, \sigma'_1; \theta_2) > G(\sigma_2, \sigma_1; \theta_2)$. This means that $\langle \sigma_1, \sigma_2 \rangle$ is a PONE if it is a Nash equilibrium of G , and it is not *strongly* Pareto-dominated by any other Nash equilibrium of G . Intuitively, if two agents are individually consistent, and willing to cooperate with each other, their joint payoff profile should not be dominated by any PONE. We formalize this intuition as follows:

Definition 2.2 (Compatibility). For $\delta, \epsilon, T > 0$, two agents π^1 and π^2 are (δ, ϵ, T) -compatible if, when played together, for any joint type $\theta \in \Theta \times \Theta$, w.p. at least $1 - \delta$, $\exists \langle \sigma_1^*, \sigma_2^* \rangle \in \mathcal{P}(G(\theta))$ s.t.

$$\frac{1}{T} \sum_{t=1}^T G(\sigma_i^*, \sigma_{-i}^*; \theta_i) - G(a_t^i, a_t^{-i}; \theta_i) \leq \epsilon, \quad (4)$$

for both $i = 1$ and $i = 2$.

A pair of agents is compatible if, when paired together, with high-probability over their path of play h_T there will exist some PONE that does not ϵ -dominate their realized payoffs. Note that this definition is the approximate and finite-horizon version of the one provided in [1].

2.3 Socially intelligent agents

We argue that it is natural to model an existing population of cooperating agents as a set of approximately compatible, but otherwise heterogeneous agents. We therefore introduce the more general idea of a socially intelligent *class* of agents that are compatible with any other member of their class:

Definition 2.3 (Social Intelligence). A set C of agents forms a *socially intelligent class* w.r.t. Θ if, for some $\delta, \epsilon, T > 0$, each agent $\pi \in C$ is (δ, ϵ, T) -consistent for all $\theta \in \Theta$, and any two agents $\pi, \pi' \in C$ are (δ, ϵ, T) -compatible over all joint types Θ . An individual agent π is called *socially intelligent* if it forms a socially intelligent class $\{\pi\}$ with itself.

The Hannan consistency requirement ensures that any agent in the population always has bounded average regret, whereas the approximate compatibility means if both agents are from C , with high probability there will exist some PONE that does not ϵ -dominate their path of play. Below we describe a socially intelligent class based on a pre-agreed *coordination protocol*.

Coordination protocols For a type space Θ , we first define a function $s(\theta) \in \mathcal{P}(G(\theta))$ that maps from each joint type θ to a strategy profile in $\mathcal{P}(G(\theta))$. We can think of $s(\theta)$ as a common “convention” the agents in C have settled upon. Since we assume private types, members of C do not know each other’s type at the beginning of their interaction. If any type $\theta \in \Theta$ can be communicated to others in a sequence of $k < T$ actions, then agents in C can agree on a coordination protocol similar to a handshake. Let the protocol be a map $\kappa(\theta)$ from types to a history-dependent policy. Then, at the beginning of each interaction, both agents will play κ for k -steps in order to communicate their types. After coordinating with each other, the agents play $s((\theta_i, \theta_{-i}))$ for the remaining $T - k$ steps. The agents must still ensure their partner does not deviate from $s((\theta_i, \theta_{-i}))$ for safety against adversarial “imposters”. Since playing a PONE jointly will lead to low regret for both, if i ’s regret exceeds a certain threshold, this would indicate $-i$ is deviating from s significantly. The threshold can be chosen by the aid of the following lemma,

Lemma 2.4. For any $\delta, T > 0$, if both players follow strategy $s(\theta)$ at each stage, then with probability at least $1 - \delta$ we have

$$\bar{R}_i^{\text{ext}}(h_t; \theta_i) \leq \sqrt{2T \ln \frac{2}{\delta}} \quad \text{and} \quad R_i^{\text{ext}}(h_t; \theta_i) \leq 2\sqrt{2T \ln \frac{4}{\delta}}, \quad (5)$$

which follows from an application of the Azuma-Hoeffding inequality (shown in Appendix A.1). Then the question is what safe strategy should the i fall back into, if the rule is triggered. We base the fallback strategy on the *multiplicative weights* [6] update rule, defined as:

$$s_{\text{mw},k}^i(h_t; \theta_i) \propto s_{\text{mw},k}^i(h_{t-1}; \theta_i) \exp \left(-\eta G(k, a_{t-1}^{-i}(h); \theta_i) \right) \quad (6)$$

for $k \in N$, where $s_{\text{mw}}^i(h_0; \theta_i)$ is the uniform strategy. Define $\pi^{\text{mw},T}$ as the agent that plays $s_{\text{mw}}^i(h_t; \theta_i)$ with learning rate $\eta = \sqrt{8 \ln(N/T)}$. The expected external regret of $\pi^{\text{mw},T}$ is bounded as

$$\bar{R}_i^{\text{ext}}(h_T; \theta_i) \leq \sqrt{\frac{T}{2} \ln N} \quad (7)$$

surely [4, Theorem 2.2]. We then define the agent’s overall strategy $\pi^{T,\epsilon}$ as follows:

1. In first k steps, play $\kappa(\theta_i)$.
2. If $-i$ ’s behaviour in h_k not compatible with $\kappa(\theta)$ for any $\theta \in \Theta$, switch to $\pi^{\text{mw},T}$ for all subsequent stages.
3. While $\bar{R}_i^{\text{ext}}(h_t; \theta_i) \leq k + \epsilon(T - k) - \sqrt{\frac{T-k}{2} \ln N} - 1$, play $s_i(\theta)$.
4. Otherwise, switch to $\pi^{\text{mw},T}$ for all subsequent stages.

The theorem below shows that agents that follow the social authentication strategy above form a socially intelligent class among themselves. All proofs have been deferred to appendix A.2.

Theorem 2.5. For any $\delta, T > k$, let $\epsilon_0 \geq \sqrt{\frac{2}{(T-k)} \ln \frac{2}{\delta}}$, and let $\epsilon_1 = \epsilon_0 + \sqrt{\frac{1}{2(T-k)} \ln N} + \frac{1}{(T-k)}$. Then for $\epsilon = \epsilon_1 + \sqrt{\frac{(T-k)}{2} \ln \frac{1}{\delta}}$, the π^{T,ϵ_1} is (δ, ϵ, T) -socially intelligent.

3 Learning to Cooperate

Going forward, we will assume that our agent (henceforth referred to as the “AI”) will take the role of agent 1, while the other agent (referred to as the “partner”) will be agent 2. Our goal is to choose a strategy for the AI that can cooperate with a partner drawn from some *target population* nearly as effectively as agents from this population cooperate with one another. For parametric game G , with type space Θ , we will let the target population be a set C of strategies forming an (δ, ϵ, T) -SI class w.r.t. Θ . Ideally, we would hope to choose an AI strategy π that can cooperate with C *without* any additional information the strategies in C . Looking at the coordination protocol example in Section 2.3, we can see that in many cases a population is likely to use arbitrary conventions to coordinate their behavior, and intuitively we would imagine cooperation to be impossible without prior knowledge of these conventions. (We make this intuition formal in Theorem 3.5).

We therefore consider the problem of learning an cooperative AI strategy using prior observations of members of the target population interacting with one another. We define a *social learning problem* by a tuple $\{G, \Theta, C, \rho, \mu\}$, where C is the target population (SI w.r.t. Θ), ρ is a distribution over C , while μ is a distribution over the joint type space $\Theta \times \Theta$. We can think of C as the set of possible strategies that any member of the target population might follow, while ρ is the frequency of those strategies within the population. To choose an AI strategy, we leverage a dataset $\mathcal{D} = \{(\theta_1^j, \theta_2^j, h_T^j) | j \in [n]\}$ covering n episodes of length T . In each episode j , two agents π_j^1 and π_j^2 are sampled independently from ρ , and played together under the joint type $\theta_j \sim \mu$. The AI observes the full history h_T^j , along with the agents’ types θ_1^j and θ_2^j . We denote a specific learning algorithm as a data conditioned strategy $\pi(\mathcal{D})$.

3.1 Altruistic Regret

We seek an AI strategy that minimizes the regret relative to some Pareto optimal solution to $G(\theta)$. Rather than minimizing regret in terms of the AI’s own payoffs, however, we seek to minimize *partner’s* relative to their (worst case) PONE in $G(\theta)$. We formalize this regret with the following definition:

Definition 3.1 (Altruistic Regret). Let the $(\sigma_i^*, \sigma_{-i}^*) \in \mathcal{P}(G_{-i}(\theta_{-i}))$ denote the PONE with the *lowest payoff* for the agent $-i$ where $i \in \{1, 2\}$. The altruistic regret of agent i is defined as

$$R_i^{\text{alt}}(h_T; \theta_{-i}) = \sum_{t=1}^T G(\sigma_i^*, \sigma_{-i}^*; \theta_{-i}) - G(a^i(h_t), a^{-i}(h_t); \theta_{-i}). \quad (8)$$

In practical cooperation tasks, we would expect outcomes that have low regret for the partner will have low regret for the AI as well.

The cooperation objective for the AI agent can then be formalized as minimising the altruistic regret. Unlike the definition suggests, the AI agent must know its own type as well. This is due to the fact that as seen in the coordination protocols example, if the AI fails to imitate a human of its type or fail to communicate its type correctly, the partner might switch to a safe strategy.

The goal for the AI is to minimize its *expected* altruistic regret over partners sampled from ρ and types sampled from μ . The following lemma shows that we can treat the problem of minimizing regret with respect to a heterogeneous population C as that of minimizing regret w.r.t. a single stochastic strategy.

Lemma 3.2. *Let C be a finite set of agents that are (δ, ϵ, T) -socially intelligent w.r.t. type space Θ , and let ρ be a distribution over C . There exists a mixed strategy $\bar{\rho}$ that forms an (δ, ϵ, T) -socially intelligent class, and which is equivalent to playing against partners sampled from ρ in expectation.*

Proof. In a perfect recall game, every behavioural strategy has an equivalent mixed strategy and vice-versa [7]. Thus ρ can equivalently be defined as a distribution over mixed strategies so that $\rho \in \Delta(\Delta(N))$. Then defining $\bar{\rho}(a) = \int_{\Delta(N)} \sigma(a) d\rho(\sigma)$ where $a \in [N]$ denotes a pure strategy (i.e. action) completes the proof.

In order to show the joint impact of consistency and compatibility on the learning problem, we discuss the cases where the population is either consistent or compatible, but not both.

3.2 Consistency without Compatibility

Assume that C consists of agents that are consistent but not necessarily compatible. The most general class in this case is the class of all no-external-regret learners (no-regret henceforth). It is a well-established result that the long-run average of no-regret learning converges to the set of coarse correlated equilibria. The question is whether the AI agent can learn to do better than a coarse correlated equilibrium when paired with a member of C , using only a dataset $\mathcal{D}$ that consists of histories of play for different CCEs.

Theorem 3.3. *There exists a consistent yet incompatible class of agents C such that even with an infinite amount of data, the AI cannot learn strategies that minimise altruistic regret.*

Proof. The proof follows from the theorem 3 of Monnot and Piliouras [8] which shows that given any coarse correlated equilibrium of a two-player normal-form game, there exists a pair of no-regret learners that would converge to it. Since C can be any subset of no-regret learners, we cannot exclude those who converge to inefficient CCE. If the class C contains only the agents that converge to Pareto-inefficient CCE, we cannot hope to learn optimal strategies from any dataset. Given an observed CCE z in the dataset, assume that the AI knows it is facing one of the two agents that generated z , but does not know their type explicitly. Using a Stackelberg argument similar to Brown et al. [9], we prove in appendix B.2 that the AI can compute and commit to a leader strategy such that the payoffs are never *strongly* Pareto-dominated by z . However even in this case, we cannot eliminate the possibility of it being weakly dominated.

Regardless of the dataset, in the online phase, the AI faces a new agent from C each time and does not know their type. We may hope to learn a classifier to quickly infer our partner's type online from their behaviour, assuming there exists a mapping from initial behaviour to types. However, since C consists only of no-regret learners guaranteed to converge to a CCE in self-play, they have no reason to initially communicate their types to each other.

3.3 Compatibility without consistency

Assume that the members of C are compatible but not consistent. We can construct such a class by using the coordination protocols example from section 2.3. Now, when agents from C successfully identify each other after the authentication phase, they proceed with playing the agreed-upon PONE. However, if at any moment they play the wrong action, there is no constraint on what strategy they will switch to. This setting is equivalent to the case considered by Loftin and Oliehoek [10] in their impossibility result. The members of C can employ grim-trigger strategies that forever punish the other agent, triggered by a mistake at any point. Even if we eliminate grim-trigger strategies, the impossibility result has proven that there still exists strategies the members of C can play once triggered, and make the other agent suffer regret arbitrarily close to $\frac{1}{2}$ with payoffs in $[0, 1]$. Since a single mistake during the online interaction can lead to partner playing strategies that yield linear regret, the outsider must learn to imitate at least one member of C perfectly from the dataset. Therefore the offline problem in this setting reduces to imitation learning, in particular the no-interaction case from Rajaraman et al. [11].

For each agent, the authentication protocol κ is equivalent to a history-dependent policy that they commit to playing in the first k time-steps. The lower-bound on the expected sub-optimality of the imitation learning from Rajaraman et al. [11] is based on the fact that the imitator cannot do better than uniformly random in unseen states. In the case of κ , states correspond to histories up to length k . Since every k -step history can be uniquely embedding a type, an unseen history means a high probability of making a mistake if paired with the corresponding type. Therefore, to avoid linear altruistic regret, the AI must observe at least $|\mathcal{H}_k|$ samples, where $\mathcal{H}_k$ is the set of all possible k -step histories.

Theorem 3.4. *Let M be the number of unique samples of k -step histories in the dataset. There exists a class of agents C with a k -step social authentication protocol such that to bound the probability of failing to authenticate, we need $M \geq \frac{N^{3k} - \delta N^{3k} - N^{2k}}{N-1}$ samples. Then for growing k , the sample complexity lower bound is $M = \Omega(N^{2k})$.*

Proof. Consider the coordination protocol example mentioned above. Let $h_k \in \mathcal{H}_k$ be missing from the dataset. When the AI is paired with the corresponding partner type, the probability of correctly authenticating is $\frac{1}{N^k}$, and thus authentication fails with probability $\frac{N^k - 1}{N^k}$. Assuming we face each

type uniformly randomly, if we have M unique samples, the probability of facing an unobserved history is $\frac{N^{2k}-M}{N^{2k}}$ since $|\mathcal{H}_k| = N^{2k}$. Then the probability of failing is $\frac{N^k-1}{N^k} \times \frac{N^{2k}-M}{N^{2k}} = 1 - \frac{M}{N^{2k}} - \frac{1}{N} + \frac{M}{N^{3k}}$. In order to bound this by δ , we need $M \geq \frac{N^{3k}-\delta N^{3k}-N^{2k}}{N-1}$ samples. Since k -steps need to embed each type uniquely, k grows with the size of the type space. For large k , the bound is dominated by N^{3k} , thus we have $M = \Omega(N^{3k})$ as k grows.

An immediate conclusion that follows from theorem 3.4 is that for the case of compatibility without consistency, this sample complexity is for bounding the probability of suffering linear regret. This is due to the fact that failing to authenticate can now lead to linear regret, since the partner can switch to arbitrary strategies.

3.4 Lower bound for socially intelligent populations

Theorem 3.5. *Let M denote the number of histories with unique first k -steps in dataset $\mathcal{D}$ generated by the members of a socially intelligent class $\mathcal{C}$. There exists a $\mathcal{C}$ where $R_i^{\text{alt}}(h_T; \theta_{-i}) = T$ with probability $\frac{N^k-1}{N^k} \times \frac{N^{2k}-M}{N^{2k}} = 1 - \frac{M}{N^{2k}} - \frac{1}{N} + \frac{M}{N^{3k}}$.*

Proof: Let $\mathcal{C}$ be a socially intelligent class of agents following a coordination protocol akin to the one described in section 2.3. The probability follows from the proof of theorem 3.4 as the probability of failing to authenticate. If the authentication fails, the partner switches to an arbitrary Hannan-consistent strategy. As stated in section 3.2, a consistent partner strategy may never communicate the partner's type. Without knowing the partner's type, the agent's worst-case average altruistic regret can be 1, since it cannot compute its true regret without the partner's type (see definition 3.1). Let there be two partner types $\theta_{-i} = \theta_2$ or θ_3 . If the agent i mistakenly assumes $\theta_{-i} = \theta_2$, its behaviour attempts to minimize $R_i^{\text{alt}}(h_T; \theta_2) = \sum_{t=1}^T G(\sigma_i^*, \sigma_{-i}^*; \theta_2) - G(a^i(h_t), a^{-i}(h_t); \theta_2)$. Meanwhile, the play of the partner will be a no-regret algorithm with respect to the external regret $R_{-i}^{\text{ext}}(h; \theta_3)$. Having no other constraints in the type space, there is nothing stopping us from constructing a Θ such that a strategy minimizing $R_{-i}^{\text{ext}}(h_T; \theta_3)$ ends up maximizing $R_i^{\text{alt}}(h_T; \theta_2)$. Imagine the ideal case of $R_{-i}^{\text{ext}}(h_T; \theta_3) = 0$ where $-i$ plays the fixed best action in hindsight a^* throughout h_T . Then the altruistic regret observed by i is $R_i^{\text{alt}}(h_T; \theta_2) = \sum_{t=1}^T G(\sigma_i^*, \sigma_{-i}^*; \theta_2) - G(a^i(h_t), a^{-i} = a^*; \theta_2)$. Let $G(a^i, a^*; \theta_2) = 0$ for all a^i . Then the altruistic regret is $\sum_{t=1}^T G(\sigma_i^*, \sigma_{-i}^*; \theta_2)$ which is T in the worst-case.

4 Upper bound for socially intelligent populations

A key idea behind this work is that against a socially intelligent target population, rather than trying to imitate a member of the population perfectly throughout the entire episode, the AI only needs to imitate them long enough to learn about its partner's private type. Once it has this information, the AI can leverage the fact that the partner's strategy is consistent against *any* strategy, and try to "coerce" the human partner into playing a strategy that minimizes the altruistic regret. We will refer to such strategies as *imitate-then-commit* (IC) strategies, which use the previous observations $\mathcal{D}$ to learn an imitation strategy to follow over the first $\tilde{T} < T$ steps of the interaction. In this section we provide an upper bound on the altruistic regret of a specific (IC) strategy, as a function of the number of episodes in $\mathcal{D}$, subject to the following assumptions:

Assumption 4.1. For $\delta, \epsilon > 0$, and some $\tilde{T} < T$, we have that

1. ρ is (δ, ϵ, T) -consistent.
2. ρ is $(\delta, \epsilon, \tilde{T})$ -compatible.

Imitation learning. Under an imitate-then-commit strategy, the sample complexity is defined entirely by the number of episodes the AI needs to observe to learn a good $\tilde{T}$ -step imitation policy. Fortunately, imitation learning is a well-studied problem, and we can largely leverage existing complexity bounds. The one caveat is that in this setting we need bounds on the total variation distance between the distribution over the partial history $h_{\tilde{T}}$ under the population strategy ρ , and that under the learned strategy. Given the dataset $\mathcal{D}$, we define the imitation strategy $\hat{\pi}_{\tilde{T}}^1(\mathcal{D})$ such that

$\hat{\pi}_{\tilde{T}}^1(h; \theta, \mathcal{D})$ is the empirical distribution over agent 1's actions for each history-type pair (h, θ) occurring in $\mathcal{D}$, while $\hat{\pi}_{\tilde{T}}^1(h; \theta, \mathcal{D})$ is the uniform distribution over N for $(h, \theta) \notin \mathcal{D}$. We then define the *marginal* strategy $\hat{\pi}_{\tilde{T}}^1$, which can be implemented by sampling a dataset $\mathcal{D}$, and then following the imitation strategy defined by $\mathcal{D}$ for the next $\tilde{T}$ steps. We then have the following bound on the distribution of $h_{\tilde{T}}$ under the imitation strategy:

Lemma 4.2. *Let $p_{\tilde{T}}$ be the distribution over partial histories $h_{\tilde{T}}$ under the population strategy ρ , and let $\hat{p}_{\tilde{T}}$ be their distribution under $\hat{\pi}_{\tilde{T}}^1$. We have that*

$$\|p_{\tilde{T}} - \hat{p}_{\tilde{T}}\|_{TV} \leq \min \left\{ \tilde{T}, \frac{N^{2(\tilde{T}+1)} |\Theta| \tilde{T}^2 \log(K)}{K} \right\}, \quad (9)$$

where $K = |\mathcal{D}|$

This bound follows directly from that of [11] via Lemma 1 of [12] (see Appendix B.1 for full proof).

Imitate-then-commit strategy. For history $h_{\tilde{T}} \in \mathcal{H}_{\tilde{T}}$, we let $\hat{z}(h_{\tilde{T}}) \in \Delta(N \times N)$ denote the empirical *joint* strategy played up to and including step $\tilde{T}$. We show that, using $\hat{z}(h_{\tilde{T}})$, it is possible to construct a *mixture* ν over mixed strategies $x \in \Delta(N)$ that, in expectation over ν , the partner's payoff under their best response to $x \sim \nu$ will be at least as large as their payoff under $\hat{z}(h_{\tilde{T}})$. The corresponding IC strategy will operate as follows:

1. Sample $\mathcal{D}$ and compute the imitation strategy $\hat{\pi}_{\tilde{T}}^1(\mathcal{D})$.
2. Play $\hat{\pi}_{\tilde{T}}^1(\mathcal{D})$ for the first $\tilde{T}$ steps, and observe $h_{\tilde{T}}$.
3. Compute a suitable mixture ν from $\hat{z}(h_{\tilde{T}})$, and sample $x \sim \nu$
4. Sample actions from x for the remaining $T - \tilde{T}$ steps

We then have the following upper bound on the altruistic regret achievable with an imitate-then-commit strategy:

Theorem 4.3. *Given that Assumption 4.1 holds for ρ , there exists a data-dependent strategy $\pi^{IC}(\mathcal{D})$ such that when played by the AI as agent 2, the altruistic regret satisfies*

$$E \left[R_1^{alt}(h_T, \theta_2) \right] \leq 2\delta + \delta(K) + \left(2 \frac{T - \tilde{T}}{T} + 1 \right) \epsilon, \quad (10)$$

where $K = |\mathcal{D}|$ and $\delta(K)$ is defined as

$$\delta(K) = \min \left\{ \tilde{T}, \frac{N^{2(\tilde{T}+1)} |\Theta| \tilde{T}^2 \log(K)}{K} \right\} \quad (11)$$

and where the expectation is taken over h_T , θ , and $\mathcal{D}$.

Proof sketch: By Lemma 4.2, we can learn an imitation strategy such that the corresponding distribution over $h_{\tilde{T}}$ and $\hat{z}(h_{\tilde{T}})$ is close to that under ρ in self-play. As ρ is compatible, both agents' payoffs under $\hat{z}(h_{\tilde{T}})$ must be close to those under *some* PONE. Finally, we can construct a mixture ν for agent 1 such that agent 2's payoffs under its (approximate) best-response are almost as large as those under $\hat{z}(h_{\tilde{T}})$ (see Appendix B.2).

5 Related Work

Our work is closely related to the previous targeted learning model [1, 13, 14], which defines similar compatibility and consistency criteria. The notion of targeted optimality [15] include convergence to learning an approximately best response in a multi-agent model with high probability in a tractable number of steps against a population of memory-bounded adaptive agents. The main difference with our work is that targeted learning only requires consistency against a specific target class of partners, which generally would not include the agent itself, or other adaptive agents. We require

socially intelligent agents to be consistent against all possible partner strategies. We also require that cooperation and consistent learning occur over a fixed time horizon T , rather than asymptotically. These differences mean that a hypothetical “universally cooperative” agent might be able to leverage the consistency of its partner to achieve cooperation without a prearranged convention. Socially intelligent agents can be modeled as individually rational learners [16] to achieve Pareto-efficient joint behavior. Our research builds on this work by considering a learning setting where the agent when paired with any member of the population will achieve at least the same utility with high probability as the Pareto-efficient approach.

The problem of training agents to be able to cooperate with previously unseen partners is sometimes referred to as *ad hoc teamwork* [17, 18] or *zero-shot coordination* [19], especially in the context of multiagent reinforcement learning. Many approaches in reinforcement learning train cooperative policies that are *robust* to possible strategies that a human or an AI agent can follow [20]. A lot of these methods build a “population” of partner strategies and maximize the diversity of this population in order to train the AI’s policy against it [21, 22]. Other approaches assume that there is no prior coordination between the agents [19] to learn rational joint strategies while estimating the agents’ mutual uncertainty about one-another’s strategies [23]. Ad-hoc multiagent coordination can be helpful to learn cooperation among AI agents with the “other-play” algorithm [19] that finds such a strategy as a solution to the corresponding *label free coordination* problem [23]. A possible approach to solve these problems can be self-play [24] where the agent can optimize themselves by playing with past iterations of themselves in order to estimate the strategies of unseen partners. However, the “self-play” approach can learn cooperative strategies which can “over-fit” [25] to one another in the population of agents. A key goal of Ad hoc coordination (teamwork) and aligned research in zero-shot coordination work has been to avoid this type of overfitting [26]. Our problem domain is closely related to both *ad hoc teamwork* or *zero-shot coordination*, since we consider training an agent to cooperate with previously unseen partners, and assume no control over the partner. Even though population-based training approaches to ad hoc teamwork are common, they focus on fully cooperative environments such as Dec-POMDPs, where the main issue is creating a diverse enough population to train with [27]. We consider partners that are self-interested, and do not assume identical payoffs.

Finally, in the case of Hannan-consistent partners, our problem setting is closely related to strategizing against and learning to manipulate no-regret learners [28, 9]. This line of work studies whether an optimizer agent can achieve better payoff than CCE against no-regret learners by learning to enforce a Stackelberg equilibria on them. Their emphasis is on online learning and the optimizer’s payoff, while we focus on the offline setting and cooperation.

6 Conclusion

We provide formal guarantees for successful and reliable cooperation of AI agents with populations of socially intelligent rational agents. This is based on the assumptions that 1) agents in the population are individually rational, and 2) agents in the population when cooperating with another agent in the same group can achieve, at least the same utility that they would with respect to some Pareto efficient equilibrium strategy. We formalize the notion of consistency and cooperative compatibility of agents in two-player general-sum finitely-repeated bi-matrix games between the agents and the population with private type. Our theoretical guarantees are in the offline cooperation setting where the agent has to cooperate with unseen partners in the population to strategize against and manipulate no-regret policies for which we formalize the idea of altruistic regret. We prove that the assumptions on its own are insufficient to learn *zero-shot* cooperation with partners of the socially intelligent target population. We provide upper bounds on the sample complexity needed to learn a successful cooperation strategy along with lower bounds on when the multi-agent cooperation setting is needed with respect to the populations’ trajectories, the state space and the length of the learning episodes. The bounds in these settings of the agent actively querying the MDP without knowing the transition dynamics of the population or the agent observing the populations’ transition dynamics are much stronger than the bounds that can be derived by naively reducing the cooperation problem to one of reinforcement learning. These complexity analysis and formally proven bounds can be helpful to sustainably model the alignment problem of AI agents.

References

- [1] Rob Powers and Yoav Shoham. New criteria and a new algorithm for learning in multi-agent systems. *Advances in Neural Information Processing Systems*, 17, 2004.
- [2] Yoav Shoham and Kevin Leyton-Brown. *Multiagent systems: Algorithmic, game-theoretic, and logical foundations*. Cambridge University Press, 2008.
- [3] James Hannan. Approximation to Bayes risk in repeated play. *Contributions to the Theory of Games*, 3(2):97–140, 1957.
- [4] Nicolo Cesa-Bianchi and Gábor Lugosi. *Prediction, learning, and games*. Cambridge university press, 2006.
- [5] Andreu Mas-Colell, Michael Dennis Whinston, Jerry R Green, et al. *Microeconomic theory*, volume 1. Oxford university press New York, 1995.
- [6] Yoav Freund and Robert E Schapire. Adaptive game playing using multiplicative weights. *Games and Economic Behavior*, 29(1-2):79–103, 1999.
- [7] Robert J . Aumann. 28. *Mixed and Behavior Strategies in Infinite Extensive Games*, pages 627–650. Princeton University Press, Princeton, 1964. ISBN 9781400882014. doi: doi: 10.1515/9781400882014-029. URL <https://doi.org/10.1515/9781400882014-029>.
- [8] Barnabé Monnot and Georgios Piliouras. Limits and limitations of no-regret learning in games. *The Knowledge Engineering Review*, 32:e21, 2017.
- [9] William Brown, Jon Schneider, and Kiran Vodrahalli. Is learning in games good for the learners? *Advances in Neural Information Processing Systems*, 36, 2024.
- [10] Robert Loftin and Frans A Oliehoek. On the impossibility of learning to cooperate with adaptive partner strategies in repeated games. In *International Conference on Machine Learning*, pages 14197–14209. PMLR, 2022.
- [11] Nived Rajaraman, Lin Yang, Jiantao Jiao, and Kannan Ramchandran. Toward the fundamental limits of imitation learning. *Advances in Neural Information Processing Systems*, 33:2914–2924, 2020.
- [12] Kamil Ciosek. Imitation learning by reinforcement learning. In *International Conference on Learning Representations*, 2022.
- [13] Rob Powers and Yoav Shoham. Learning against opponents with bounded memory. In *The Nineteenth International Joint Conference on Artificial Intelligence*, pages 817–822, 2005.
- [14] Doran Chakraborty and Peter Stone. Convergence, targeted optimality and safety in multiagent learning. In *Proceedings of the Twenty-seventh International Conference on Machine Learning (ICML 2010)*, June 2010. URL <http://www.cs.utexas.edu/users/ai-lab?chakraborty:icml10>.
- [15] Doran Chakraborty and Peter Stone. Convergence, targeted optimality, and safety in multiagent learning. In *Proceedings of the 27th International Conference on International Conference on Machine Learning, ICML’10*, page 191–198, Madison, WI, USA, 2010. Omnipress. ISBN 9781605589077.
- [16] Robert Loftin, Mustafa Mert Çelikok, and Frans A. Oliehoek. Towards a unifying model of rationality in multiagent systems, 2023.
- [17] Peter Stone, Gal Kaminka, Sarit Kraus, and Jeffrey Rosenschein. Ad hoc autonomous agent teams: Collaboration without pre-coordination. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 24, pages 1504–1509, 2010.
- [18] Reuth Mirsky, Ignacio Carlucho, Arrasy Rahman, Elliot Fosong, William Macke, Mohan Sridharan, Peter Stone, and Stefano V Albrecht. A survey of ad hoc teamwork research. In *European conference on multi-agent systems*, pages 275–293. Springer, 2022.

- [19] Hengyuan Hu, Adam Lerer, Alex Peysakhovich, and Jakob Foerster. “other-play” for zero-shot coordination. In *International Conference on Machine Learning*, pages 4399–4410. PMLR, 2020.
- [20] Micah Carroll, Rohin Shah, Mark K Ho, Tom Griffiths, Sanjit Seshia, Pieter Abbeel, and Anca Dragan. On the utility of learning about humans for human-ai coordination. *Advances in Neural Information Processing Systems*, 32:5174–5185, 2019.
- [21] DJ Strouse, Kevin McKee, Matt Botvinick, Edward Hughes, and Richard Everett. Collaborating with humans without human data. *Advances in Neural Information Processing Systems*, 34:14502–14515, 2021.
- [22] Brandon Cui, Andrei Lupu, Samuel Sokota, Hengyuan Hu, David J Wu, and Jakob Nicolaus Foerster. Adversarial diversity in hanabi. In *The Eleventh International Conference on Learning Representations*, 2023.
- [23] Johannes Treutlein, Michael Dennis, Caspar Oesterheld, and Jakob Foerster. A new formalism, method and open issues for zero-shot coordination. In *International Conference on Machine Learning*, pages 10413–10423. PMLR, 2021.
- [24] Jaleh Zand, Jack Parker-Holder, and Stephen J. Roberts. On-the-fly strategy adaptation for ad-hoc agent coordination. In *Proceedings of the 21st International Conference on Autonomous Agents and Multiagent Systems*, AAMAS ’22, page 1771–1773, Richland, SC, 2022. International Foundation for Autonomous Agents and Multiagent Systems. ISBN 9781450392136.
- [25] DJ Strouse, Kevin McKee, Matt Botvinick, Edward Hughes, and Richard Everett. Collaborating with humans without human data. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 14502–14515. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/797134c3e42371bb4979a462eb2f042a-
- [26] Brandon Cui, Hengyuan Hu, Luis Pineda, and Jakob Foerster. K-level reasoning for zero-shot coordination in hanabi. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 8215–8228. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/4547dff5fd7604f18c8ee32cf3da41d7-
- [27] Muhammad Rahman, Jiaxun Cui, and Peter Stone. Minimum coverage sets for training robust ad hoc teamwork agents. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 17523–17530, 2024.
- [28] Yuan Deng, Jon Schneider, and Balasubramanian Sivan. Strategizing against no-regret learners. *Advances in neural information processing systems*, 32, 2019.
- [29] Wassily Hoeffding. Probability inequalities for sums of bounded random variables. *Journal of the American Statistical Association*, 58(301):13–30, 1963.
- [30] Bernhard von Stengel and Shmuel Zamir. Leadership with commitment to mixed strategies. 2004.

Acknowledgments and Disclosure of Funding

This work has been supported by the Hybrid Intelligence Center, <https://hybrid-intelligence-centre.nl>, grant number 024.004.022.

A Proofs for Section 2

A.1 Proof of Lemma 2.4

Here the joint type θ will be implicit. For $i \in \{1, 2\}$, we define V_t^i as

$$V_t^i = G_i(s_t^i, s_t^{-i}) - G_i(s_t^i, a_t^{-i}) \quad (12)$$

We can see that $\mathbb{E}[V_t^i | h_{t-1}] = 0$. We can then have that

$$\bar{R}_t^{\text{ext}} = \max_{a \in N} \sum_{r=1}^t \left\{ G_i(a, s_r^{-i}) - G_i(s_r^i, a_r^{-i}) \right\} \quad (13)$$

$$= \max_{a \in N} \sum_{r=1}^t \left\{ G_i(a, s_r^{-i}) - G_i(s_r^i, s_r^{-i}) + G_i(s_r^i, s_r^{-i}) - G_i(s_r^i, a_r^{-i}) \right\} \quad (14)$$

$$= \sum_{r=1}^t \left\{ G_i(s_r^i, s_r^{-i}) - G_i(s_r^i, a_r^{-i}) \right\} = \sum_{r=1}^t V_r^i \quad (15)$$

$$\leq \sqrt{\frac{2}{T} \ln \frac{1}{\delta}} \quad (16)$$

with probability $1 - \delta$ for all $t \leq T$ simultaneously.

This follows from the fact that $|V_t^i| \in [0, 1]$ and the “maximal” Azuma-Hoeffding inequality [29]. The second equality follows from the fact that $\langle s_t^i, s_t^{-i} \rangle = s(\theta)$ is a Nash equilibrium. The first bound of Lemma 2.4 follows from a union bound over the probability for both players, while the second bound combines this with Equation 3. $\square$

A.2 Proof of Theorem 2.5

Theorem A.1 (2.6). *For any $\delta, T > k$, let $\epsilon_0 \geq \sqrt{\frac{2}{(T-k)} \ln \frac{2}{\delta}}$, and let $\epsilon_1 = \epsilon_0 + \sqrt{\frac{1}{2(T-k)} \ln N} + \frac{1}{(T-k)}$. Then for $\epsilon = \epsilon_1 + \sqrt{\frac{(T-k)}{2} \ln \frac{1}{\delta}}$, the π^{T, ϵ_1} is (δ, ϵ, T) -socially intelligent.*

Proof. By the definition of ϵ_1 , π^{T, ϵ_1} will only deviate when playing with itself if at some point $k < t \leq T$ one player incurs an expected external regret of at least ϵ_0 , and by Lemma 2.4 that will occur with probability at most δ . Therefore, π^{T, ϵ_1} is (δ, ϵ_0, T) -compatible. We also have that the total expected external regret of the MW agent $\pi^{\text{mw}, T}$ is at most $\sqrt{(T/2) \ln N}$. This means that if π^{T, ϵ_1} switches at stage t , then the maximum possible expected external regret incurred by π^{T, ϵ_1} will be less than $\bar{R}_t^{\text{ext}}(h_t; \theta) + \sqrt{\frac{T}{2} \ln N}$. Since $\pi^{\text{mw}, T}$ will always switch just before this point is reached, its total expected regret will be less than ϵ_1 surely, and will be less than ϵ w.p. $1 - \delta$. As $\epsilon \geq \epsilon_0$, we have that the π^{T, ϵ_1} is (δ, ϵ, T) -socially intelligent.

B Proofs for Section 4

B.1 Proof of Lemma 4.2

We first apply Theorem 4.4 of [11], which states that, for episodic imitation learning over H -step trajectories, for any expert policy π^* we have

$$J(\pi^*) - \mathbb{E}_{\mathcal{D}} \left[J(\hat{\pi}_{\tilde{T}}^1(h; \theta, \mathcal{D})) \right] \leq \min \left\{ H, \frac{|S|H^2 \log(K)}{K} \right\}, \quad (17)$$

where S is the state space, with per-step rewards bounded in $[0, 1]$. We can model the interaction with ρ as a $\tilde{T}$ -step episodic MDP/R with $S = \mathcal{H}_{\leq \tilde{T}}$. Plugging in $H = \tilde{T}$, $|S| < N^{2(\tilde{T}+1)}$, and $\pi^* = \rho$ gives us

$$J(\rho) - \mathbb{E}_{\mathcal{D}} \left[J(\hat{\pi}_{\tilde{T}}^1(h; \theta, \mathcal{D})) \right] \leq \min \left\{ \tilde{T}, \frac{N^{2\tilde{T}} |\Theta| \tilde{T}^2 \log(N)}{K} \right\}. \quad (18)$$

This bound holds simultaneously for all possible reward functions bounded in $[0, 1]$. If we restrict the reward function r to be non-zero only for the terminal states $\mathcal{H}_{\tilde{T}}$, we have

$$J(\pi^*) - \mathbb{E}_{\mathcal{D}} \left[J(\hat{\pi}_{\tilde{T}}^1(h; \theta, \mathcal{D})) \right] = \mathbb{E}_{p_{\tilde{T}}} [r(h_{\tilde{T}})] - \mathbb{E}_{\hat{p}_{\tilde{T}}} [r(h_{\tilde{T}})], \quad (19)$$

using the definition of the marginal strategy $\hat{\pi}_{\tilde{T}}^1$. Finally, applying Lemma 1 of [12] gives us

$$\|p_{\tilde{T}} - \hat{p}_{\tilde{T}}\|_{\text{TV}} \leq \min \left\{ \tilde{T}, \frac{N^{2\tilde{T}} |\Theta| \tilde{T}^2 \log(N)}{K} \right\}, \quad (20)$$

the desired result. $\square$

B.2 Proof of Theorem 4.3

First, let $\tau^2(\theta)$, defined as

$$\tau^2(\theta) = \min_{(\sigma^1, \sigma^2) \in \mathcal{P}(G(\theta))} G(\sigma^2, \sigma^1; \theta^2), \quad (21)$$

denote agent 2's payoff under the worst possible payoff for a PONE of the game parameterized by joint type θ . Let $\mathcal{C}$ denote the event that

$$\tau^2(\theta) - \frac{1}{\tilde{T}} \sum_{t=1}^{\tilde{T}} G(a_t^2, a_t^1; \theta_2) \leq \epsilon \quad (22)$$

Because ρ is $(\delta, \epsilon, \tilde{T})$ -compatible, we have that $\Pr_{\rho}\{\mathcal{C}\} \geq 1 - \delta$. For $\delta(K)$ defined as

$$\delta(K) = \min \left\{ \tilde{T}, \frac{N^{2(\tilde{T}+1)} |\Theta| \tilde{T}^2 \log(K)}{K} \right\}, \quad (23)$$

Lemma 4.2 also gives us $\Pr_{\hat{\pi}^1, \rho}\{\mathcal{C}\} \geq 1 - \delta - \delta(K)$. We therefore have that

$$\mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=1}^T G(a_t^2, a_t^1; \theta_2) \right] \geq \mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] - T(\delta + \delta(K)) \quad (24)$$

$$= \mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=1}^{\tilde{T}} G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] + \mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=\tilde{T}+1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] - T(\delta + \delta(K)) \quad (25)$$

$$\geq \mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=\tilde{T}+1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] + T(\tau^2(\theta) - \epsilon - \delta - \delta(K)) \quad (26)$$

We therefore need to lower-bound the term

$$\mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=\tilde{T}+1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] \quad (27)$$

This will be the expected payoff given the strategy $x \sim \nu$ the AI commits to for the remaining $T - \tilde{T}$ steps. The idea now is that we can construct a mixture ν over strategies that the AI can commit to for the remaining $T - \tilde{T}$ steps such that the partner's payoff under their (approximate) best-response will be nearly as good as that under $\hat{z}(h_{\tilde{T}})$.

Let $G(z; \theta^2) = \sum_{i \in M} \sum_{j \in M} z_{i,j} G(j, i; \theta^2)$ be agent 2's expected payoff under z . For any joint strategy z , we can construct ν such that if the AI commits to strategies sampled from ν , the partner will have the same information about the AI's probably actions as they would given their "recommended" action under $\hat{z}(h_{\tilde{T}})$. We build on the construction used by von Stengel and Zamir [30]. For any joint strategy z , we let $z_j = \sum_{i \in N} z_{ij}$ denote the *marginal* probability that the column player (agent 2) plays j under z . For all $j \in N$ such that $z_j > 0$, we define x_j as the *conditional* distribution over the row-player (agent 1's) actions given that the column player plays j , such that $x_j(i) = \frac{z_{ij}}{z_j}$. We then define ν as the strategy that commits to each x_j with probability z_j .

We can show that, when the partner plays a best-response to $x \sim \nu$, their payoff will be no worse than under z itself. We first construct a *response function* r_z such that when agent 2 responds to

$x \sim \nu$ with $r_z(x)$, its expected payoff equals $G(z; \theta^2)$. Let $S = \{j \in N : z_j > 0\}$, and partition S into $\mathcal{P}$ such that, for each $P \in \mathcal{P}$, we have $x_j = x_l$ for all $j, l \in P$. For each $P \in \mathcal{P}$, we then define the strategy y_P such that

$$y_P(j) = \frac{z_j}{\sum_{l \in P} z_l} \quad (28)$$

for each $j \in P$, with $y_P(j) = 0$ for $j \notin P$. (Note that if z corresponds to some *uncorrelated* strategy $\langle x, y \rangle$, then $P = N$ and $y_P = y$.) Finally, for $j \in S$, we define $P(j)$ as the partition containing j , and define r_z such that $r_z(x_j) = x_{P(j)}$. We leave r_z undefined for x where $\mu(x) = 0$. Now let x_P be the common conditional strategy for all $j \in P$, and let $z_P = \sum_{j \in P} z_j$. We then have that

$$\mathbb{E}_{x \sim \nu} G(r_z(x), x; \theta^2) = \sum_{j \in S} z_j \left[x_j^\top G(\theta^2)^\top r_z(x_j) \right] \quad (29)$$

$$= \sum_{P \in \mathcal{P}} z_P \left[x_P^\top G(\theta^2)^\top y_P \right] \quad (30)$$

$$= \sum_{P \in \mathcal{P}} z_P \left(\sum_{i \in N} \sum_{j \in N} x_P^\top(i) y_P(j) G(\theta^2)_{ij}^\top \right) \quad (31)$$

$$= \sum_{P \in \mathcal{P}} z_P \left(\sum_{i \in N} \sum_{j \in N} \Pr_z\{i|P\} \Pr_z\{j|P\} G(\theta^2)_{ij}^\top \right) \quad (32)$$

$$= \sum_{P \in \mathcal{P}} z_P \left(\sum_{i \in N} \sum_{j \in N} \Pr_z\{i, j|P\} G(\theta^2)_{ij}^\top \right) \quad (33)$$

$$= \sum_{i \in N} \sum_{j \in N} z_{ij} G(\theta^2)_{ij}^\top = G(z; \theta^2) \quad (34)$$

where we have used the fact that i and j are independent given that $j \in P$. Next, we have that for any best-response function r^* , we have

$$\begin{aligned} G(z; \theta^2) &= \mathbb{E}_{x \sim \nu} G(r_z(x), x; \theta^2) \\ &= \mathbb{E}_{x \sim \mu} [x^\top G(\theta^2)^\top r_z(x)] \\ &\leq \mathbb{E}_{x \sim \mu} \left[\max_{y \in \Delta(N)} x^\top G(\theta^2)^\top y \right] \\ &= \mathbb{E}_{x \sim \mu} [x^\top G(\theta^2)^\top r^*(x)] \\ &= \mathbb{E}_{x \sim \nu} G(r^*(x), x; \theta^2) \end{aligned} \quad (35)$$

Therefore, so long as the partner plays a best-response to the AIs chosen strategy, they will achieve at least the same payoff (in expectation) as they would under the strategy z from which ν was computed. Note however that ρ will be (approximately) consistent over the full T steps, not just the last $T - \tilde{T}$. Define $\alpha = \frac{\tilde{T}}{T}$ and $\beta = \frac{T - \tilde{T}}{T}$, and let z^1 be agent 1's marginal strategy under z . With probability $1 - \delta$, ρ will play an ϵ -best-response to the mixture $\alpha \hat{z}(h_{\tilde{T}})^1 - \beta x$, with $x \sim \nu$.

Let $\mathcal{C}'$ be the event that ρ is ϵ -consistent over T steps. We then have that

$$\mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=\tilde{T}+1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C} \right] \geq \mathbb{E}_{\hat{\pi}^1, \rho} \left[\sum_{t=\tilde{T}+1}^T G(a_t^2, a_t^1; \theta_2) | \mathcal{C}, \mathcal{C}' \right] - T\delta \quad (36)$$

$$\geq (T - \tilde{T}) \left(\tau^2(\theta) - 2\epsilon \right) - T\delta \quad (37)$$

Finally, dividing by T and subtracting from $\tau^2(\boldsymbol{\theta})$, we get

$$\mathbb{E} \left[R^{\text{alt}_1}(h_T, \theta_2) \right] \leq 2\delta + \delta(K) + \left(2\frac{T - \tilde{T}}{T} + 1 \right) \epsilon \quad (38)$$

the desired result.

□