
Hardness of Independent Learning and Sparse Equilibrium Computation in Markov Games

Dylan J. Foster¹ Noah Golowich² Sham M. Kakade³

Abstract

We consider the problem of decentralized multi-agent reinforcement learning in Markov games. A key question is whether there are algorithms that, when run independently by all agents, lead to no-regret for each player, analogous to celebrated results for no-regret learning in normal-form games. While recent work has shown that such algorithms exist for restricted settings (e.g., when regret is defined with respect to deviations to Markov policies), the question of whether independent no-regret learning can be achieved in the standard Markov game framework was open. We provide a decisive negative resolution to this problem, both from a computational and statistical perspective. We show that:

1. Under the assumption that PPAD-hard problems cannot be solved in polynomial time, there is no polynomial-time algorithm that attains no-regret in general-sum Markov games when executed independently by all players, even when the game is known to the algorithm designer and the number of players is a small constant.
2. When the game is unknown, no algorithm, efficient or otherwise, can achieve no-regret without observing exponentially many episodes in the number of players.

These results are proven via lower bounds for a simpler problem we refer to as `SPARSE CCE`, in which the goal is to compute a coarse correlated equilibrium that is “sparse” in the sense that it can be represented as a mixture of a small number of product policies.

1. Introduction

The framework of multi-agent reinforcement learning (MARL), which describes settings in which multiple agents interact in a dynamic environment, has played a key role in recent breakthroughs in artificial intelligence, including the development of agents that approach or surpass human performance in games such as Go (Silver et al., 2016), Poker (Brown & Sandholm, 2018), Stratego (Perolat et al., 2022), and Diplomacy (Kramar’et al., 2022; Bakhtin et al., 2022). MARL also shows promise for real-world multi-agent systems, including autonomous driving (Shalev-Shwartz et al., 2016), and cybersecurity (Malialis & Kudenko, 2015), and economic policy (Zheng et al., 2022). These applications, where reliability is critical, necessitate the development of algorithms that are practical and efficient, yet provide strong formal guarantees and robustness.

Multi-agent reinforcement learning is typically studied using the framework of Markov games (also known as stochastic games) (Shapley, 1953). In a Markov game, agents interact over a finite number of steps: at each step, each agent observes the state of the environment, takes an action, and observes a reward which depends on the current state as well as the other agents’ actions. Then the environment transitions to a new state as a function of the current state and the actions taken. An episode consists of a finite number of such steps, and agents interact over the course of multiple episodes, progressively learning new information about their environment. Markov games generalize the well-known model of Markov Decision Processes (MDPs) (Puterman, 1994), which describe the special case in which there is a single agent acting in a dynamic environment, and we wish to find a policy that maximizes its reward. By contrast, for Markov games, we typically aim to find a distribution over agents’ policies which constitutes some type of equilibrium.

1.1. Decentralized learning

In this paper, we focus on the problem of decentralized (or, independent) learning in Markov games. In decentralized MARL, each agent in the Markov game behaves independently, optimizing their policy myopically while treating the effects of the other agents as exogenous. Agents observe local information (in particular, their own actions and rewards), but do not observe the actions of the other agents directly. Decentralized learning enjoys a number of desirable properties, including

¹Microsoft Research ²Massachusetts Institute of Technology
³Harvard University. Correspondence to: Dylan J. Foster<dylanfos-
ter@microsoft.com>, Noah Golowich<nzg@mit.edu>, Sham M.
Kakade<sham@seas.harvard.edu>.

Proceedings of the 40th International Conference on Machine
Learning, Honolulu, Hawaii, USA. PMLR 202, 2023. Copyright
2023 by the author(s).

scalability, versatility, and practicality. The central question we consider is whether there exist decentralized learning algorithms which, when employed by all agents in a Markov game, lead them to play near-equilibrium strategies over time.

Decentralized equilibrium computation in MARL is not well understood theoretically, and algorithms with provable guarantees are scarce. To motivate the challenges and most salient issues, it will be helpful to contrast with the simpler problem of decentralized learning in normal-form games, which may be interpreted as Markov games with a single state. Much of the modern work on decentralized learning in normal-form games centers on no-regret learning, where agents select actions independently using online learning algorithms (Cesa-Bianchi & Lugosi, 2006) designed to minimize their regret (that is, the gap between realized payoffs and the payoff of the best fixed action in hindsight). In particular, a foundational result is that if each agent employs a no-regret learning strategy, then the average of the agents’ joint action distributions approaches a coarse correlated equilibrium (CCE) for the normal-form game (Cesa-Bianchi & Lugosi, 2006; Hannan, 1957; Blackwell, 1956). CCE is a natural relaxation of the foundational concept of Nash equilibrium, which has the downside of being intractable to compute. On the other hand, there are many efficient algorithms that can achieve vanishing regret in a normal-form game, even when opponents select their actions in an arbitrary, potentially adaptive fashion, and thus converge to a CCE (Vovk, 1990; Littlestone & Warmuth, 1994; Cesa-Bianchi et al., 1997; Hart & Mas-Colell, 2000; Syrgkanis et al., 2015).

This simple connection between no-regret learning and decentralized convergence to equilibria has been influential in game theory, leading to numerous lines of research including fast rates of convergence to equilibria (Syrgkanis et al., 2015; Chen & Peng, 2020; Daskalakis et al., 2021; Anagnostides et al., 2022), price of anarchy bounds for smooth games (Roughgarden, 2015), and lower bounds on query and communication complexity for equilibrium computation (Fearnley et al., 2013; Rubinstein, 2016; Babichenko & Rubinstein, 2017). Empirically, no-regret algorithms such as regret matching (Hart & Mas-Colell, 2000) and Hedge (Vovk, 1990; Littlestone & Warmuth, 1994; Cesa-Bianchi et al., 1997) have been used to compute equilibria that can achieve state-of-the-art performance in application domains such as Poker (Brown & Sandholm, 2018) and Diplomacy (Bakhtin et al., 2022). Motivated by these successes, we ask whether an analogous theory can be developed for Markov games. In particular:

Are there efficient algorithms
for no-regret learning in Markov games?

Challenges for no-regret learning. In spite of active research effort and many promising pieces of progress (Jin et al., 2021; Song et al., 2022; Mao & Basar, 2021; Daskalakis et al., 2022; Erez et al., 2022), no-regret learning guarantees for Markov games have been elusive. A barrier faced by naive algorithms

is that it is intractable to ensure no-regret against an arbitrary adversary, both computationally (Bai et al., 2020; Abbasi-Yadkori et al., 2013) and statistically (Liu et al., 2022; Kwon et al., 2021; Foster et al., 2022). Fortunately, many of the implications of no-regret learning (in particular, convergence to equilibria) do not require the algorithm to have sublinear regret against an arbitrary adversary, but rather only against other agents who are running the same algorithm independently. This observation has been influential in normal-form games, where the line of work on fast rates of convergence to equilibrium (Syrgkanis et al., 2015; Chen & Peng, 2020; Daskalakis et al., 2021; Anagnostides et al., 2022) holds only in this more restrictive setting. This motivates the following relaxation to our central question.

Problem 1.1. Is there an efficient algorithm that, when adopted by all agents in a Markov game and run independently, leads to sublinear regret for each individual agent?

Attempts to address Problem 1.1. Two recent lines of research have made progress toward addressing Problem 1.1 and related questions. In one direction, several recent papers have provided algorithms, including V-learning (Jin et al., 2021; Song et al., 2022; Mao & Basar, 2021) and SPoCMAR (Daskalakis et al., 2022), that do not achieve no-regret, but can nevertheless compute and then sample from a coarse correlated equilibrium in a Markov game in a (mostly) decentralized fashion, with the caveat that they require a shared source of random bits as a mechanism to coordinate. Notably, V-learning depends only mildly on the shared randomness: agents first play policies in a fully independent fashion (i.e., without shared randomness) according to a simple learning algorithm for T episodes, and use shared random bits only once learning finishes as part of a post-processing procedure to extract a CCE policy. A question left open by these works, is whether the sequence of policies played by the V-learning algorithm in the initial independent phase can itself guarantee each agent sublinear regret.

Most closely related to our work, Erez et al. (2022) recently showed that Problem 1.1 can be solved positively for a restricted setting in which regret for each agent is defined as the maximum gain in value they can achieve by deviating to a fixed Markov policy. Markov policies are those whose choice of action depends only on the current state as opposed to the entire history of interaction. This notion of deviation is restrictive because in general, even when the opponent plays a sequence of Markov policies, the best response will be non-Markov. In challenging settings that abound in practice, it is standard to consider non-Markov policies (Leibo et al., 2021; Agapiou et al., 2022), since they often achieve higher value than Markov policies; we provide a simple example in Proposition B.1. Thus, while a regret guarantee with respect to the class of Markov policies (as in (Erez et al., 2022)) is certainly interesting, it may be too weak in general, and it is of great interest to understand whether Problem 1.1 can be answered positively in the general setting.¹

¹We remark that the V-learning and SPoCMAR algorithms

We refer the reader to Appendix B.2 for further discussion.

1.2. Our contributions

We resolve Problem 1.1 in the negative, from both a computational and statistical perspective.

Computational hardness. We provide two computational lower bounds (Theorems 1.2 and 1.3) which show that under standard complexity-theoretic assumptions, there is no efficient algorithm that runs for a polynomial number of episodes and guarantees each agent non-trivial (“sublinear”) regret when used in tandem by all agents. Both results hold even if the Markov game is explicitly known to the algorithm designer; Theorem 1.3 is stronger and more general, but applies only to 3-player games, while Theorem 1.2 applies to 2-player games, but only for agents restricted to playing Markovian policies.

To state our first result, Theorem 1.2, we define a product Markov policy to be a joint policy in which players choose their actions independently according to Markov policies (see Sections 2 and 3 for formal definitions). Note that if all players use independent no-regret algorithms to choose Markov policies at each episode, then their joint play at each round is described by a product Markov policy, since any randomness in each player’s policy must be generated independently.

Theorem 1.2 (Informal version of Corollary 3.3). If $\text{PPAD} \not\leq \text{RP}$, then there is no polynomial-time algorithm that, given the description of a 2-player Markov game, outputs a sequence of joint product Markov policies which guarantees each agent sublinear regret.

Theorem 1.2 provides a decisive negative resolution to Problem 1.1 under the assumption that $\text{PPAD} \not\leq \text{RP}$,² which is standard in the theory of computational complexity (Papadimitriou, 1994).³ Beyond simply ruling out the existence of fully decentralized no-regret algorithms, it rules out existence of centralized algorithms that compute a sequence of product policies for which each agent has sublinear regret, even if such a sequence does not arise naturally as the result of agents independently following some learning algorithm. Salient implications include:

- Theorem 1.2 provides a separation between Markov games and normal-form games, since standard no-regret algorithms for normal-form games i) run in polynomial

mentioned above do learn equilibria that are robust to deviations to non-Markov policies, though they do not address Problem 1.1 since they do not have sublinear regret.

²Technically, the class we are denoting by P , namely of total search problems that have a deterministic polynomial-time algorithm, is sometimes denoted by FP , as it is a search problem. We ignore this distinction.

³ PPAD is the most well-studied complexity class in algorithmic game theory, and is widely believed to not admit polynomial time algorithms. Notably, the problem of computing a Nash equilibrium for normal-form games with two or more players is PPAD -complete (Daskalakis et al., 2009; Chen et al., 2006; Rubinstein, 2018).

time and ii) produce sequences of joint product policies that guarantee each agent sublinear regret. Notably, no-regret learning for normal-form games is efficient whenever the number of agents is polynomial, whereas Theorem 1.2 rules out polynomial-time algorithms for as few as two agents.

- A question left open by the work of Jin et al. (2021); Song et al. (2022); Mao & Basar (2021) was whether the sequence of policies played by the V-learning algorithm during its independent learning phase can guarantee each agent sublinear regret. Since V-learning plays product Markov policies during the independent phase and is computationally efficient, Theorem 1.2 implies that these policies do not enjoy sublinear regret (assuming $\text{PPAD} \not\leq \text{RP}$).

Our second result, Theorem 1.3, extends the guarantee of Theorem 1.2 to the more general setting in which agents can select arbitrary, potentially non-Markovian policies at each episode. This comes at the cost of only providing hardness for 3-player games as opposed to 2-player games, as well as relying on the slightly stronger complexity-theoretic assumption that $\text{PPAD} \not\leq \text{RP}$.⁴

Theorem 1.3 (Informal version of Corollary 4.4). If $\text{PPAD} \not\leq \text{RP}$, then there is no polynomial-time algorithm that, given the description of a 3-player Markov game, outputs a sequence of joint product general policies (i.e., potentially non-Markov) which guarantees each agent sublinear regret.

Statistical hardness. Theorems 1.2 and 1.3 rely on the widely-believed complexity theoretic assumption that PPAD -complete problems cannot be solved in (randomized) polynomial time. Such a restriction is inherent if we assume that the game is known to the algorithm designer. To avoid complexity-theoretic assumptions, we consider a setting in which the Markov game is unknown to the algorithm designer, and algorithms must learn about the game by executing policies (“querying”) and observing the resulting sequences of states, actions, and rewards. Our final result, Theorem 1.4, shows unconditionally that, for m -player Markov games whose parameters are unknown, any algorithm computing a no-regret sequence as in Theorem 1.3 requires a number of queries that is exponential in m .

Theorem 1.4 (Informal version of Theorem 5.2). Given query access to a m -player Markov game, no algorithm that makes fewer than $2^{\Omega(m)}$ queries can output a sequence of joint product policies which guarantees each agent sublinear regret.

Similar to our computational lower bounds, Theorem 1.4 goes far beyond decentralized algorithms, and rules out even centralized algorithms that compute a no-regret sequence by jointly controlling all players. The result provides another

⁴We use RP to denote the class of total search problems for which there exists a polynomial-time randomized algorithm which outputs a solution with probability at least $2/3$, and otherwise outputs “fail”.

4

Markov) policies, which select actions based on the entire sequence of states and actions observed up the current step. To streamline notation, for $i \in \mathcal{P}$, let $i;h$ denote the history of agent i 's states, actions, and reward up to step h . Let $H_{i;h}$ denote the space of all possible histories of agent i up to step h . For $i \in \mathcal{P}$, a randomized general (i.e., non-Markov) policy of agent i is a collection of mappings $\pi_{i;h} : H_{i;h} \rightarrow \Delta(\mathcal{A}_i)$ where $\pi_{i;h} : H_{i;h} \rightarrow \Delta(\mathcal{A}_i)$ is a mapping that takes the history observed by agent i up to step h and the current state and outputs a distribution over actions for agent i .

We denote by $\mathcal{P}_i^{\text{gen;rnd}}$ the space of randomized general policies of agent i , and further write $\mathcal{P}^{\text{gen;rnd}} = \prod_{i \in \mathcal{P}} \mathcal{P}_i^{\text{gen;rnd}}$ to denote the space of product general policies; note that $\text{markov} \in \mathcal{P}_i^{\text{gen;rnd}}$ and $\text{markov} \in \mathcal{P}^{\text{gen;rnd}}$. In particular, a policy $P \in \mathcal{P}^{\text{gen;rnd}}$ is specified by a collection $\pi_{i;h} : H_{i;h} \rightarrow \Delta(\mathcal{A}_i)$, where $\pi_{i;h} : H_{i;h} \rightarrow \Delta(\mathcal{A}_i)$. When agents play according to a general policy $P \in \mathcal{P}^{\text{gen;rnd}}$, at each step h , each agent, given the current state s_h and their history $i;h$, chooses to play an action $a_{i;h} \in \mathcal{A}_i$, independently from all other agents. For a policy $P \in \mathcal{P}^{\text{gen;rnd}}$, we let $\Pr_s$ and $\mathbb{E}_s$ denote the law and expectation operator for the trajectory when players select actions via $a_{i;h} \in \mathcal{A}_i$, and write π to denote the collection of policies of all agents but i , i.e., $\pi = \prod_{j \neq i} \pi_{j;h}$.

We will also consider distributions over product randomized general policies, namely elements of $\mathcal{P}^{\text{gen;rnd}}$.⁷ We will refer to elements of $\mathcal{P}^{\text{gen;rnd}}$ as distributional policies. To play a distributional policy $P \in \mathcal{P}^{\text{gen;rnd}}$, agents draw a randomized policy P (so that $P \in \mathcal{P}^{\text{gen;rnd}}$) and then play.

Value functions. For a general policy $P \in \mathcal{P}^{\text{gen;rnd}}$, we define the value function for agent $i \in \mathcal{P}$ as $V_i : E \rightarrow \mathbb{R}$, where $E = \prod_{h=1}^H \mathcal{R}_h \times \prod_{i \in \mathcal{P}} \mathcal{A}_i$; this represents the expected reward that agent i receives when each agent chooses their actions via $a_{i;h} \in \mathcal{A}_i$. For a distributional policy $P \in \mathcal{P}^{\text{gen;rnd}}$, we extend this notation by defining $V^P : E \rightarrow \mathbb{R}$.

2.3. Normal-form games

To motivate the solution concepts we consider for Markov games, let us first revisit the notion of normal-form games, which may be interpreted as Markov games with a single state. For $m, n \in \mathbb{N}$, an m -player n -action normal-form game G is specified by a tuple of m reward tensors $M_1, \dots, M_m \in \mathbb{R}^{n \times \dots \times n}$, where each tensor is of order m (i.e., has n^m entries). We will write $G = (M_1, \dots, M_m)$. We assume for simplicity that each player has the same number n of actions, and identify each player's action space with $\mathcal{R}_n$. Then

⁷When T is not a finite set, we take $\mathcal{P}^{\text{gen;rnd}}$ to be the set of Radon probability measures over T equipped with the Borel σ -algebra.

an action profile is specified by a $\mathbf{r} \in \mathcal{R}_n^m$; if each player acts according to $\mathbf{a}$, then the reward for player $i \in \mathcal{P}$ is given by $r_i(\mathbf{a})$. Our hardness results will use the standard notion of Nash equilibrium in normal-form games. We define the m -player n -NASH problem to be the problem of computing an approximate Nash equilibrium of a given m -player n -action normal-form game. (See Definition C.2 for a formal definition of ϵ -Nash equilibrium.) A celebrated result is that Nash equilibria are PPAD-hard to approximate, i.e., the 2-player n -NASH problem is PPAD-hard for any constant $c > 0$ (Daskalakis et al., 2009; Chen et al., 2006). We refer the reader to Section C.2 for further background on these concepts.

2.4. Markov games: Equilibria and no-regret

We now turn our focus back to Markov games, and introduce the main solution concepts we consider, as well as the notion of no-regret. Since computing Nash equilibria is intractable even for normal-form games, much of the work on efficient equilibrium computation has focused on alternative notions of equilibrium, notably coarse correlated equilibria.

For a distributional policy $P \in \mathcal{P}^{\text{gen;rnd}}$ and a randomized policy π_i of player i , we let π_i^P denote the distributional policy which is given by the distribution of π_i^P for P (and π_i denotes the marginal of π_i^P on all players but i). For $P \in \mathcal{P}^{\text{gen;rnd}}$, we write π_i^P to denote the policy given by π_i^P . Let us fix a Markov game G , which in particular determines the players' value functions V .

Definition 2.1 (Coarse correlated equilibrium). For $\epsilon > 0$, a distributional policy $P \in \mathcal{P}^{\text{gen;rnd}}$ is defined to be an ϵ -coarse correlated equilibrium (CCE) if for each $i \in \mathcal{P}$, it holds that $\max_{\pi_i \in \Delta(\mathcal{A}_i)} \mathbb{E}_{\pi_i^P} [V_i(\pi_i, \pi_{-i}^P)] \geq \mathbb{E}_{\pi_i^P} [V_i(\pi_i^P, \pi_{-i}^P)] - \epsilon$.

Coarse correlated equilibria can be computed efficiently for both normal-form games and Markov games, and are fundamentally connected to the notion of no-regret and independent learning, which we now introduce.

Regret. For a policy $P \in \mathcal{P}^{\text{gen;rnd}}$, we denote the distributional policy which puts all its mass on P by $\mathbb{I}_P \in \mathcal{P}^{\text{gen;rnd}}$. Thus $\mathbb{I}_P^T$ denotes the distributional policy which randomizes uniformly over the $\mathcal{P}^{\text{gen;rnd}}$. We define regret as follows.

Definition 2.2 (Regret). Consider a sequence of policies $\pi_1, \dots, \pi_T \in \mathcal{P}^{\text{gen;rnd}}$. For $i \in \mathcal{P}$, the regret of agent i with respect to this sequence is defined as:

$$\text{Reg}_{i,T}(\pi_1, \dots, \pi_T) = \max_{\pi_i \in \Delta(\mathcal{A}_i)} \mathbb{E}_{\pi_i^T} [V_i(\pi_i, \pi_{-i}^T)] - \mathbb{E}_{\pi_i^T} [V_i(\pi_i^T, \pi_{-i}^T)] \quad (1)$$

It is immediate from the above definitions that a sequence of policies $\pi_1, \dots, \pi_T \in \mathcal{P}^{\text{gen;rnd}}$ satisfies $\text{Reg}_{i,T}(\pi_1, \dots, \pi_T) \leq \epsilon$ if and only if

T if and only if the distributional policy $\pi_{p,q}^T$ is an -CCE (stated formally in Fact C.1 in the appendix).

No-regret learning. A standard approach to decentralized equilibrium computation, which exploits Fact C.1, is to select $p^{1q}, \dots, p^{Tq}$ using independent no-regret learning algorithms. A no-regret learning algorithm for player i selects $p_i^{gen;rnd}$ based on the realized trajectories $p^{1q}, \dots, p^{Tq}$ that player i observes over the course of play,⁸ but with no knowledge of $p^{1q}, \dots, p^{Tq}$, so as to ensure that no-regret is achieved: $\text{Reg}_{i,T}(p^{1q}, \dots, p^{Tq})/T$. If each player i uses their own, independent no-regret learning algorithm, this approach yields product policies $p^{1q}, \dots, p^{Tq}$, and the uniform average of the $p^{1q}, \dots, p^{Tq}$ yields a CCE as long as all of the players can keep their regret small.⁹

For the special case of normal-form games, there are several efficient algorithms, which—when run independently—ensure that each player’s regret after T episodes is bounded above by Op_T (that is Op_1), even when the other players’ actions are chosen adversarially.

3. Lower bound for Markovian algorithms

In this section we prove Theorem 1.2 (restated formally below as Theorem 3.2), establishing that in two-player Markov games, there is no computationally efficient algorithm that computes a sequence $p^{1q}, \dots, p^{Tq}$ of product Markov policies so that each player has small regret under this sequence. This section serves as a warm-up for our results in Section 4, which remove the assumption that $p^{1q}, \dots, p^{Tq}$ are Markovian.

3.1. SPARSEMARKOV CCE and computational model

As discussed in the introduction, our lower bounds for no-regret learning are a consequence of lower bounds for the SPARSE CCE problem. In what follows, we formalize this problem (specifically, the Markovian variant, which we refer to as SPARSEMARKOV CCE), as well as our computational model.

Description length for Markov games (constant m). Given a Markov game G , we let pG denote the maximum number of bits needed to describe any of the rewards $R_{i,h;p;a,q}$ or transition probabilities $P_{h,p^1;s;a,q}$ in binary.¹⁰ We define $|G| : \max_T S; \max_{i,p,rms} A_i; H; pG$. The interpretation of $|G|$ depends on the number of players m : If m is a constant (as will be the case in the current section and Section 4), then

⁸An alternative model allows for player i to have knowledge of the previous joint policies $p^{1q}, \dots, p^{(t-1)q}$, when selecting p^{tq} .

⁹In Appendix B, we discuss the implications of relaxing the stipulation that $p^{1q}, \dots, p^{Tq}$ be product policies (for example, by allowing the use of shared randomness, as in V-learning). In short, allowing $p^{1q}, \dots, p^{Tq}$ to be non-product essentially trivializes the problem.

¹⁰We emphasize that pG is defined as the maximum number of bits required by any particular $p;s;a,q$ pair, not the total number of bits required for all $p;s;a,q$ pairs.

$|G|$ should be interpreted as the description length of the game G , up to polynomial factors. In particular, for constant m , the game G can be described using $|G|^{\text{Op}1q}$ bits. In Section 5, we discuss the interpretation of $|G|$ when m is large.

The SPARSEMARKOV CCE problem. From Fact C.1, we know that the problem of computing a sequence $p^{1q}, \dots, p^{Tq}$ of joint product Markov policies for which each player has at most T regret is equivalent to computing a sequence $p^{1q}, \dots, p^{Tq}$ for which the uniform mixture forms an -approximate CCE. We define $pT;q$ -SPARSEMARKOV CCE as the computational problem of computing such a CCE directly.

Definition 3.1 (SPARSEMARKOV CCE problem). For an m -player Markov game G and parameters T, P, N and $i, 0$ (which may depend on the size of the game G), $pT;q$ -SPARSEMARKOV CCE is the problem of finding a sequence $p^{1q}, \dots, p^{Tq}$, with each p^{tq} Markov, such that the distributional policy $\pi_{p,q}^T$ is an -CCE of G (or equivalently, such that for all i, P, rms , $\text{Reg}_{i,T}(p^{1q}, \dots, p^{Tq})/T$).

Decentralized learning algorithms naturally lead to solutions to the SPARSEMARKOV CCE problem. In particular, consider any decentralized protocol which runs for T episodes, where at each timestep t , each player i chooses a Markov policy p_i^{tq} to play, without knowledge of the other players’ policies p_j^{tq} (but possibly using the history); any strategy in which players independently run online learning algorithms falls under this protocol. If each player experiences overall regret at most T , then the sequence $p^{1q}, \dots, p^{Tq}$ is a solution to the $pT;q$ -SPARSEMARKOV CCE problem. However, one might expect the $pT;q$ -SPARSEMARKOV CCE problem to be much easier than decentralized learning, since it allows for algorithms that produce $p^{1q}, \dots, p^{Tq}$ satisfying the constraints of Definition 3.1 in a centralized manner. The main result of this section, Theorem 3.2, rules out the existence of any efficient algorithms, including centralized ones, that solve the SPARSEMARKOV CCE problem.

Before moving on, let us give a sense for what sort of scaling one should expect for the parameters T and i in the $pT;q$ -SPARSEMARKOV CCE problem. First, we note that there always exists a solution to the $p1;0q$ -SPARSEMARKOV CCE problem in a Markov game, which is given by a (Markov) Nash equilibrium of the game; of course, Nash equilibria are intractable to compute in general.¹¹ For the special case of normal-form games (where there is only a single state, and $H=1$), no-regret learning (e.g., Hedge) yields a computationally efficient solution to the $pT;\text{Op}1$ -SPARSEMARKOV CCE problem, where the $\text{Op}1$ hides a $\max_i \log |A_i|$ factor. Refined convergence guarantees of Daskalakis et al. (2021); Anagnostides et al. (2022) improve upon this result, and yield an efficient solution

¹¹Such a Nash equilibrium can be seen to exist by using backwards induction to specify the player’s joint distribution of play at each state at steps $H; H-1; \dots; 1$.

to the $pT; q\text{-SPARSEMARKOVCCCE}$ problem.

3.2. Main result

Theorem 3.2. There is a constant $C_0 \geq 1$ so that the following holds. Let n, P, N be given, and let T, P, N and $\epsilon \geq 0$ satisfy $T \leq \exp^{2n^{1/2}} \{2^{5q}\}$. Suppose there is an algorithm that, given the description of any 2-player Markov game G with $|G| \leq n$, solves the $pT; q\text{-SPARSEMARKOVCCCE}$ problem in time U , for some $U \leq P, N$. Then, for each n, P, N , the 2-player $pn^{1/2}u; 4q\text{-NASH}$ problem (Definition C.2) can be solved in time $pnTUq^{C_0}$.

We emphasize that the range $T \leq \exp^{2n^{1/2}} \{2^{5q}\}$ ruled out by Theorem 3.2 is the most natural parameter regime, since the runtime of any decentralized algorithm which runs for T episodes and produces a solution to the SPARSEMARKOVCCCE problem is at least linear in T . Using that 2-player $pn; q\text{-NASH}$ is PPAD-complete for n^c (for any $c \geq 0$) (Daskalakis et al., 2009; Chen et al., 2006; Rubinstein, 2018), we obtain the following corollary.

Corollary 3.3 (SPARSEMARKOVCCCE is PPAD-complete). For any constant $C \geq 4$, if there is an algorithm which, given the description of a 2-player Markov game G , solves the $p|G|^C; |G|^C q\text{-SPARSEMARKOVCCCE}$ problem in time $\text{poly}(|G|, q)$, then PPAD $\leq P$.

The condition $C \geq 4$ in Corollary 3.3 is set to ensure that $|G|^C \leq \exp^{2n^{1/2}} \{2^{5q}\}$ for sufficiently large $|G|$, so as to satisfy the condition of Theorem 3.2. Corollary 3.3 rules out the existence of a polynomial-time algorithm that solves the SPARSEMARKOVCCCE problem with accuracy polynomially small and T polynomially large in $|G|$.

Proof overview. The proof of Theorem 3.2 is based on a reduction, which shows that any algorithm that efficiently solves the $pT; q\text{-SPARSEMARKOVCCCE}$ problem, for T not too large, can be used to efficiently compute an approximate Nash equilibrium of any given normal-form game. In particular, fix n_0, P, N , and let a 2-player normal form game G with n_0 actions be given.

We construct a Markov game $G \text{ GpGq}$ with horizon $H \leq n_0$ and action sets identical to those of the game G , i.e., $A_1 = A_2 = n_0$.

The state space of G consists of n^2 states, which are indexed by joint action profiles; the transitions are defined so that the value of the state at step h encodes the action profile taken by the agents at step $h-1$.¹² At each state of G , the reward functions are given by the payoff matrices of G , scaled down by a factor of $1/H$ (which ensures that the rewards received at each step belong to $[0, 1/H]$). In particular, the rewards and transitions out of a given state do not depend on the identity of the state, and so G can be thought of as a repeated game where G is played H times. The formal definition of G is given in Definition D.3.

Fix any algorithm for the SPARSEMARKOVCCCE prob-

lem, and recall that for each step h and state s for G , $p_h^{tq} p_{s,q}$ denotes the joint action distribution taken in s at step h for the sequence of $p^{1q}, \dots, p^{Tq}$ produced by the algorithm. The bulk of the proof of Theorem 3.2 consists of proving a key technical result, Lemma D.4, which states that if $p^{1q}, \dots, p^{Tq}$ indeed solves $pT; q\text{-SPARSEMARKOVCCCE}$, then there exists some tuple $p_h; s; t_q$ such that $p_h^{tq} p_{s,q}$ is an approximate Nash equilibrium for G . With this established, it follows that we can find a Nash equilibrium efficiently by simply trying all $H \cdot T$ choices for $p_h; s; t_q$.

To prove Lemma D.4, we reason as follows. Assume that $\tau = \frac{1}{T} \sum_{t=1}^T p_{t,q} \in \mathcal{P}^{\text{gen}; \text{rnd}}_q$ is an ϵ -CCE. If, by contradiction, none of the distributions $p_h^{tq} p_{s,q}$ are approximate Nash equilibria for G , then it must be the case that for each t , one of the players has a profitable deviation in G with respect to the product strategy $p_h^{tq} p_{s,q}$, at least for a constant fraction of the tuples $p_h; s; t_q$. We will argue that if this were to be the case, it would imply that there exists a non-Markov deviation policy for at least one player i in Definition 2.1, meaning that τ is not in fact an ϵ -CCE.

To sketch the idea, recall that to draw a trajectory from τ , we first draw a random index $t \in [1, T]$ uniformly at random, and then execute p^{tq} for an episode. We will show (roughly) that for each player i , it is possible to compute a non-Markov deviation policy σ_i which, under the draw of a trajectory from τ , can “infer” the value of the index t within the first few steps of the episode. Then policy σ_i then, at each state s and step h after the first few steps, play a best response to their opponent’s portion of the strategy $p^{tq} p_{s,q}$. If, for each possible value of t , none of the distributions $p_h^{tq} p_{s,q}$ are approximate Nash equilibria of G , this means that at least one of the players i can significantly increase their value in G over that of τ by playing σ_i , which contradicts the assumption that τ is an ϵ -CCE.

It remains to explain how we can construct a non-Markov policy σ_i which “infers” the value of t . Unfortunately, exactly inferring the value of t in the fashion described above is impossible: for instance, if there are t_1, t_2 so that $p^{t_1q} \neq p^{t_2q}$, then clearly it is impossible to distinguish between the cases $t = t_1$ and $t = t_2$. Nevertheless, by using the fact that each player observes the full joint action profile played at each step h , we can construct a non-Markov policy which employs Vovk’s aggregating algorithm for online density estimation (Vovk, 1990; Cesa-

Bianchi & Lugosi, 2006) in order to compute a distribution which is close to $p_h^{tq} p_{s,q}$ for most $h \in [1, H]$.¹³ This guarantee is stated formally in an abstract setting in Proposition D.2, and is instantiated in the proof of Theorem 3.2 in ((5)). As we show in Section D.2, approximating $p_h^{tq} p_{s,q}$ as we have described is sufficient to carry out the reasoning from the previous paragraph.

¹³Vovk’s aggregating algorithm is essentially the exponential weights algorithm with the logarithmic loss. A detailed background for the algorithm is provided in Section D.1.

¹²For technical reasons, this only is the case for even values of h ; we discuss further details in the full proof in Section D.2.

4. Lower bound for non-Markov algorithms

In this section, we prove Theorem 1.3 (restated formally below as Theorem 4.3), which strengthens Theorem 3.2 by allowing the sequence $p^{1q}, \dots, p^{Tq}$ of product policies to be non-Markovian. This additional strength comes at the cost of our lower bound only applying to 3-player Markov games (as opposed to Theorem 3.2, which applied to 2-player games).

4.1. SPARSECCE problem and computational model

To formalize the computational model for the SPARSECCE problem, we must first describe how the non-Markov product policies $p^{1q}, p^{2q}, \dots, p^{Tq}$ are represented. Recall that a non-Markov policy $p_i^{\text{gen};\text{rnd}}$ is, by definition, a mapping from agent i 's history and current state to a distribution over their next action. Since there are exponentially many possible histories, it is information-theoretically impossible to express an arbitrary policy in $p_i^{\text{gen};\text{rnd}}$ with polynomially many bits. As our focus is on computing a sequence of such policies p^{1q} in polynomial time, certainly a prerequisite is that p^{1q} can be expressed in polynomial space. Thus, we adopt the representational assumption, stated formally in Definition 4.1, that each of the policies $p_i^{\text{gen};\text{rnd}}$ is described by a bounded-size circuit that can compute the conditional distribution of each next action given the history. This assumption is satisfied by essentially all empirical and theoretical work concerning non-Markov policies (e.g., (Leibo et al., 2021; Agapiou et al., 2022; Jin et al., 2021; Song et al., 2022)).

Definition 4.1 (Computable policy). Given a m -player Markov game G and $N \in \mathbb{N}$, we say that a policy $p_i^{\text{gen};\text{rnd}}$ is N -computable if for each $h \in \mathcal{H}_i$, there is a circuit of size N that, ¹⁴ on input $p_{i,h1}; s \in \mathcal{S}$, outputs the distribution $p_{i,h1}; s \in \mathcal{S}$. A policy $p_{1,q}, \dots, p_{m,q}$ is N -computable if each constituent policy p_i is.

Our lower bound applies to algorithms that produce sequences $p^{1q}, \dots, p^{Tq}$ for which each p^{1q} is N -computable, where the value N is taken to be polynomial in the description length of the game G . For example, Markov policies whose probabilities can be expressed with b bits are $\text{OpHSA}_{b,q}$ -computable for each player i , since one can simply store each of the probabilities $p_{i,h1}; a_{i,h}q$, each of which takes b bits to represent.

The SPARSECCE problem. SPARSECCE is the problem of computing a sequence of non-Markov product policies $p^{1q}, \dots, p^{Tq}$ such that the uniform mixture forms an ϵ -approximate CCE. The problem generalizes SPARSE-MARKOV CCE (Definition 3.1) by relaxing the condition that the policies p^{1q} be Markov.

¹⁴For concreteness, we suppose that “circuit” means “boolean circuit” as in Definition 6.1 of (Arora & Barak, 2006), where probabilities are represented in binary. The precise model of computation we use does not matter, though, and we could equally assume that the policies may be computed by Turing machines that terminate after N steps.

Definition 4.2 (SPARSECCE Problem). For an m -player Markov game G and parameters $T; N \in \mathbb{N}$ and $\epsilon > 0$ (which may depend on the size of the game G), $p^{Tq}; Nq$ -SPARSECCE is the problem of finding a sequence $p^{1q}, \dots, p^{Tq}$ $p^{\text{gen};\text{rnd}}$, with each p^{1q} being N -computable, such that the distributional policy $\frac{1}{T} \sum_{t=1}^T p^{tq}$ is an ϵ -CCE for G (equivalently, such that for all $i \in \{1, \dots, m\}$, $\text{Reg}_i(p^{1q}, \dots, p^{Tq}) \leq \epsilon$).

4.2. Main result

Our main theorem for this section, Theorem 4.3, shows that for appropriate values of T , ϵ , and N , solving the $p^{Tq}; Nq$ -SPARSECCE problem is at least as hard as computing Nash equilibria in normal-form games.

Theorem 4.3. Fix $n \in \mathbb{N}$, and let $T; N \in \mathbb{N}$, and $\epsilon > 0$ satisfy $1/T \leq \exp(-\frac{\epsilon}{16})$. Suppose there exists an algorithm that, given the description of any 3-player Markov game G with $|G| \leq n$, solves the $p^{Tq}; Nq$ -SPARSECCE problem in time U , for some $U \in \mathbb{N}$. Then, for any $\epsilon > 0$, the 2-player $\text{ptn}(2u; 50q)$ -NASH problem can be solved in randomized time $\text{pnTNU} \log p_1(q^{C_0})$ with failure probability ϵ , where $C_0 > 0$ is an absolute constant.

By analogy to Corollary 3.3, we obtain the following immediate consequence.

Corollary 4.4 (SPARSECCE is hard under PPA + RP). For any $C > 4$, if there is an algorithm which, given the description of a 3-player Markov game G , solves the $p|G|^C; |G|^{\frac{1}{C}}; |G|^C$ -SPARSECCE problem in time $\text{poly}(|G|q)$, then PPA + RP.

Proof overview for Theorem 4.3. The proof of Theorem 4.3 has a similar high-level structure to that of Theorem 3.2: given an m -player normal-form G , we define an $\text{pm}_{1,q}$ -player Markov game G_{pGq} which has $n_0 = \text{tn}(\mu)$ actions per player and horizon $H = n_0$. The key difference in the proof of Theorem 4.3 is the structure of the players' reward functions. To motivate this difference and the addition of an $\text{pm}_{1,q}$ -th player, we explain why the proof of Theorem 3.2 fails to extend: a sequence $p^{1q}, \dots, p^{Tq}$ can hypothetically solve the SPARSECCE problem by attempting to punish any one player's deviation policy, and thus avoid having to compute a Nash equilibrium of G . In particular, if player i plays according to the policy $\hat{p}_i$ that we described in Section 3.2, then other players $j \neq i$ can use the non-Markov property of p^{1q} to adjust their choice of actions in later rounds to decrease player i 's value.

This behavior is reminiscent of “tit-for-tat” strategies which are used to establish the folk theorem in the theory of repeated games (Maskin & Fudenberg, 1986). The folk theorem describes how Nash equilibria are more numerous in repeated games than in single-shot normal form games. As it turns out, the folk theorem does not yield to worst-case speedups in repeated games, when the number of players is at least 3. Indeed, Borgs et al. (2008) gave an “anti-folk theorem”, showing that computing Nash equilibria in $\text{pm}_{1,q}$ -player

repeated games is PPAD-hard for $m \geq 2$, via a reduction to m -player normal-form games. We adapt their reduction to our setting: roughly speaking, this approach adds an $m-1$ -th player whose actions represent potential deviations for each of the m players. The structure of the rewards ensures that if π^T is an ϵ -CCE, then for some policy π of the $m-1$ -th player, the first m players will play an approximate Nash of G with constant probability, under a trajectory drawn from the joint policy $\pi \times \pi^T$. Thus, in order to efficiently find a Nash (see Algorithm 2), we need to simulate the policy $\pi \times \pi^T$, which involves running Vovk's algorithm. This approach is in contrast to the proof of Theorem 3.2, which used Vovk's algorithm as an ingredient in the proof but not in the Nash computation algorithm.

Two-player games. One intriguing question we leave open is whether the SPARSECCE problem remains hard for two-player Markov games. Interestingly, as shown by Littman & Stone (2005), there is a polynomial time algorithm to find an exact Nash equilibrium for the special case of repeated two-player normal-form games. Though their result only applies in the infinite-horizon setting, it is possible to extend their results to the finite-horizon setting, which rules out naive approaches to extend the proof of Theorem 4.3 and Corollary 4.4 to two players.

5. Multi-player games: lower bounds

In this section we present Theorem 1.4 (restated formally below as Theorem 5.2), which gives a statistical lower bound for the SPARSECCE problem. The lower bound applies to any algorithm, regardless of computational cost, that accesses the underlying Markov game through a generative model.

Definition 5.1 (Generative model). For an m -player Markov game $G = (S; H; p; A_1, \dots, A_m; R_1, \dots, R_m)$, a generative model oracle is defined as follows: given a query described by a tuple $\langle \pi; s; a_1, \dots, a_m \rangle$, the oracle returns the distribution $P_{\pi, s; a_1, \dots, a_m}$ and the tuple of rewards $\langle R_1; \dots; R_m \rangle$.

From the perspective of lower bounds, the assumption that the algorithm has access to a generative model is quite reasonable, as it encompasses most standard access models in RL, including the online access model, in which the algorithm repeatedly queries a policy and observes a trajectory drawn from it, as well as the local access generative model used in from (Yin et al., 2022; Weisz et al., 2021). We remark that it is slightly more standard to assume that queries to the generative model only return a sample from the distribution $P_{\pi, s; a_1, \dots, a_m}$ as opposed to the distribution itself (Kakade, 2003; Kearns et al., 1999), but since our goal is to prove lower bounds, the notion in Definition 5.1 only makes our results stronger.

To state our main result, we recall the definition $|G| = \max\{S; \max_i |A_i|; H; p; G\}$. In the present section, we consider the setting where the number of players m is large. Here, $|G|$ does not necessarily correspond to the

description length for G , and should be interpreted, roughly speaking, as a measure of the description complexity of G $|G|$ with respect to decentralized learning algorithms. In particular, from the perspective of an individual agent implementing a decentralized learning algorithm, their sample complexity should depend only on the size of their individual action set (as well as the global parameters $S; H; p; G$), as opposed to the size of the joint action set, which grows exponentially in m ; the former is captured by $|G|$, while the latter is not. Indeed, a key advantage shared by much prior work on decentralized RL (Jin et al., 2021; Song et al., 2022; Mao & Basar, 2021; Daskalakis et al., 2022) is their avoidance of the curse of multi-agents, which describes the situation where an algorithm has sample and computational costs that scale exponentially in m .

Our main result for this section, Theorem 5.2, states that for m -player Markov games, exponentially many generative model queries (in m) are necessary to produce a solution to the $pT; Nq$ -SPARSECCE problem, unless T is exponential in m .

Theorem 5.2. Let $m \geq 2$ be given. There are constants $c; \epsilon > 0$ so that the following holds. Suppose there is an algorithm B which, given access to a generative model for a $m-1$ -player Markov game G with $|G| \leq 2^m$, solves the $pT; \{p10mq; Nq$ -SPARSECCE problem for G for some T satisfying $1/T \leq \epsilon \exp(-cm)$, and any $N \geq N$. Then B must make at least 2^{cm} queries to the generative model.

Theorem 5.2 establishes that there are m -player Markov games, where the number of states, actions per player, and horizon are bounded by $\text{poly}(mq)$, but any algorithm with regret $\text{opt}(mq)$ must make 2^{cm} queries (via Fact C.1). In particular, if there are $\text{poly}(mq)$ queries per episode, as is standard in the online simulator model where a trajectory is drawn from the policy π^t at each episode t , then $T \geq 2^{cm}$.

2^{cm} episodes are required to have regret $\text{opt}(mq)$. This is in stark contrast to the setting of normal-form games, where even for the case of bandit feedback (which is a special case of the generative model setting), standard no-regret algorithms have the property that each player's regret scales as $\text{Opt}(T) \leq n \sqrt{T}$ (i.e., independently of m), where n denotes the number of actions per player (Lattimore & Szepesvari, 2020). As with our computational lower bounds, Theorem 5.2 is not limited to decentralized algorithms, and also rules out centralized algorithms which have access to a generative model.

Acknowledgements

This work was performed in part while Noah Golowich was an intern at Microsoft Research. Noah Golowich is supported at MIT by a Fannie & John Hertz Foundation Fellowship and an NSF Graduate Fellowship. This work has been made possible in part by a gift from the Chan Zuckerberg Initiative Foundation to establish the Kempner Institute for the Study of Natural and Artificial Intelligence. Sham Kakade acknowledges funding from the Office of Naval Research under award N00014-22-1-2377 and the National Science Foundation Grant under award #CCF-2212841.

References

- Abbasi Yadkori, Y., Bartlett, P. L., Kanade, V., Seldin, Y., and Szepesvari, C. Online learning in markov decision processes with adversarially chosen transition probability distributions. In Burges, C., Bottou, L., Welling, M., Ghahramani, Z., and Weinberger, K. (eds.), *Advances in Neural Information Processing Systems*, volume 26. Curran Associates, Inc., 2013. URL <https://proceedings.neurips.cc/paper/2013/file/4f284803bd0966cc24fa8683a34afc6e-Paper.pdf>.
- Agapiou, J. P., Vezhnevets, A. S., Duñez-Guzmán, E. A., Matyas, J., Mao, Y., Sunehag, P., Köster, R., Madhushani, U., Kopparapu, K., Comanescu, R., Strouse, D., Johanson, M. B., Singh, S., Haas, J., Mordatch, I., Mobbs, D., and Leibo, J. Z. Melting pot 2.0, 2022. URL <https://arxiv.org/abs/2211.13746>.
- Anagnostides, I., Farina, G., Kroer, C., Lee, C.-W., Luo, H., and Sandholm, T. Uncoupled learning dynamics with $\mathcal{O}(\log t)$ swap regret in multiplayer games. In Oh, A. H., Agarwal, A., Belgrave, D., and Cho, K. (eds.), *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=CZwh1XdAhNv>.
- Arora, S. and Barak, B. *Computational Complexity: A Modern Approach*. Cambridge University Press, 2006. ISBN 978-0-521-42426-4.
- Babichenko, Y. Query complexity of approximate nash equilibria. *J. ACM*, 63(4), oct 2016. ISSN 0004-5411.
- Babichenko, Y. and Rubinstein, A. Communication complexity of approximate nash equilibria. In *Proceedings of the 49th Annual ACM SIGACT Symposium on Theory of Computing, STOC 2017*, pp. 878–889, New York, NY, USA, 2017. Association for Computing Machinery. ISBN 9781450345286. doi: 10.1145/3055399.3055407. URL <https://doi.org/10.1145/3055399.3055407>.
- Bai, Y., Jin, C., and Yu, T. Near-optimal reinforcement learning with self-play. In *Proceedings of the 34th International Conference on Neural Information Processing Systems, NIPS’20*, Red Hook, NY, USA, 2020. Curran Associates Inc. ISBN 9781713829546.
- Bakhtin, A., Brown, N., Dinan, E., Farina, G., Flaherty, C., Fried, D., Goff, A., Gray, J., Hu, H., Jacob, A. P., Komeili, M., Konath, K., Kwon, M., Lerer, A., Lewis, M., Miller, A. H., Mitts, S., Renduchintala, A., Roller, S., Rowe, D., Shi, W., Spisak, J., Wei, A., Wu, D., Zhang, H., and Zijlstra, M. Human-level play in the game of iĉdiplomacyiĉ by combining language models with strategic reasoning. *Science*, 378(6624):1067–1074, 2022. doi: 10.1126/science.ade9097. URL <https://www.science.org/doi/abs/10.1126/science.ade9097>.
- Blackwell, D. An analog of the minimax theorem for vector payoffs. *Pacific Journal of Mathematics*, 6:1–8, 1956.
- Borgs, C., Chayes, J., Immorlica, N., Kalai, A. T., Mirrokni, V., and Papadimitriou, C. The myth of the folk theorem. In *Proceedings of the fortieth annual ACM symposium on Theory of computing*, pp. 365–372, 2008.
- Brown, G. W. Some notes on computation of games solutions. 1949.
- Brown, N. and Sandholm, T. Superhuman ai for heads-up no-limit poker: Libratus beats top professionals. *Science*, 359(6374):418–424, 2018. doi: 10.1126/science.aao1733. URL <https://www.science.org/doi/abs/10.1126/science.aao1733>.
- Cesa-Bianchi, N. and Lugosi, G. *Prediction, Learning, and Games*. Cambridge University Press, New York, NY, USA, 2006. ISBN 0521841089.
- Cesa-Bianchi, N., Freund, Y., Haussler, D., Helmbold, D. P., Schapire, R. E., and Warmuth, M. K. How to use expert advice. *Journal of the ACM*, 44(3):427–485, 1997.
- Chen, X. and Peng, B. Hedging in games: Faster convergence of external and swap regrets. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 18990–18999. Curran Associates, Inc., 2020. URL <https://proceedings.neurips.cc/paper/2020/file/db346ccb62d491029b590bbb0f5c412-Paper.pdf>.
- Chen, X., Deng, X., and Teng, S.-h. Computing nash equilibria: Approximation and smoothed complexity. In *2006 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS’06)*, pp. 603–612, 2006.
- Chen, X., Cheng, Y., and Tang, B. Well-Supported vs. Approximate Nash Equilibria: Query Complexity of Large Games. In Papadimitriou, C. H. (ed.), *8th Innovations in Theoretical Computer Science Conference (ITCS 2017)*, volume 67 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pp. 57:1–57:9, Dagstuhl, Germany, 2017. Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik. ISBN 978-3-95977-029-3. doi: 10.4230/LIPIcs.ITCS.2017.57. URL <http://drops.dagstuhl.de/opus/volltexte/2017/8163>.
- Daskalakis, C., Goldberg, P. W., and Papadimitriou, C. H. The complexity of computing a nash equilibrium. *SIAM Journal on Computing*, 39(1):195–259, 2009.
- Daskalakis, C., Golowich, N., and Zhang, K. The complexity of markov equilibrium in stochastic games, 2022. URL <https://arxiv.org/abs/2204.03991>.
- Daskalakis, C. C., Fishelson, M., and Golowich, N. Near-optimal no-regret learning in general

- games. In Beygelzimer, A., Dauphin, Y., Liang, P., and Vaughan, J. W. (eds.), *Advances in Neural Information Processing Systems*, 2021. URL <https://openreview.net/forum?id=cVwc7IHWEWi>.
- Erez, L., Lancewicki, T., Sherman, U., Koren, T., and Mansour, Y. Regret minimization and convergence to equilibria in general-sum markov games, 2022. URL <https://arxiv.org/abs/2207.14211>.
- Fearnley, J., Gairing, M., Goldberg, P., and Savani, R. Learning equilibria of games via payoff queries. In *Proceedings of the Fourteenth ACM Conference on Electronic Commerce, EC '13*, pp. 397–414, New York, NY, USA, 2013. Association for Computing Machinery. ISBN 9781450319621. doi: 10.1145/2492002.2482558. URL <https://doi.org/10.1145/2492002.2482558>.
- Foster, D. J., Rakhlin, A., Sekhari, A., and Sridharan, K. On the complexity of adversarial decision making. *arXiv preprint arXiv:2206.13063*, 2022.
- Fudenberg, D., Levine, D., and Maskin, E. The folk theorem with imperfect public information. *Econometrica*, 62(5): 997–1039, 1994. ISSN 00129682, 14680262.
- Hannan, J. Approximation to Bayes risk in repeated play. *Contributions to the Theory of Games*, 3:97–139, 1957.
- Hart, S. and Mas-Colell, A. A simple adaptive procedure leading to correlated equilibrium. *Econometrica*, 68(5):1127–1150, 2000. doi: <https://doi.org/10.1111/1468-0262.00153>. URL <https://onlinelibrary.wiley.com/doi/abs/10.1111/1468-0262.00153>.
- Jin, C., Krishnamurthy, A., Simchowitz, M., and Yu, T. Reward-free exploration for reinforcement learning. In *International Conference on Machine Learning (ICML)*, pp. 4870–4879. PMLR, 2020.
- Jin, C., Liu, Q., Wang, Y., and Yu, T. V-learning—a simple, efficient, decentralized algorithm for multiagent RL. *arXiv preprint arXiv:2110.14555*, 2021.
- Jin, Y., Muthukumar, V., and Sidford, A. The complexity of infinite-horizon general-sum stochastic games, 2022.
- Kakade, S. M. On the sample complexity of reinforcement learning, 2003.
- Kearns, M., Mansour, Y., and Ng, A. Y. A sparse sampling algorithm for near-optimal planning in large markov decision processes. In *Proceedings of the 16th International Joint Conference on Artificial Intelligence - Volume 2, IJCAI'99*, pp. 1324–1331, San Francisco, CA, USA, 1999. Morgan Kaufmann Publishers Inc.
- Kramár, J., Eccles, T., Gemp, I., Tacchetti, A., McKee, K. R., Malinowski, M., Graepel, T., and Bachrach, Y. Negotiation and honesty in artificial intelligence methods for the board game of Diplomacy. *Nature Communications*, 13(1):7214, December 2022. ISSN 2041-1723. doi: 10.1038/s41467-022-34473-5. URL <https://www.nature.com/articles/s41467-022-34473-5>. Number: 1 Publisher: Nature Publishing Group.
- Kwon, J., Efroni, Y., Caramanis, C., and Mannor, S. RL for latent MDPs: Regret guarantees and a lower bound. *Advances in Neural Information Processing Systems*, 34, 2021.
- Lattimore, T. and Szepesvári, C. *Bandit algorithms*. Cambridge University Press, 2020.
- Leibo, J. Z., Dueñez-Guzman, E. A., Vezhnevets, A., Agapiou, J. P., Sunehag, P., Koster, R., Matyas, J., Beattie, C., Mordatch, I., and Graepel, T. Scalable evaluation of multi-agent reinforcement learning with melting pot. In Meila, M. and Zhang, T. (eds.), *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pp. 6187–6199. PMLR, 18–24 Jul 2021. URL <https://proceedings.mlr.press/v139/leibo21a.html>.
- Littlestone, N. and Warmuth, M. K. The weighted majority algorithm. *Information and computation*, 108(2):212–261, 1994.
- Littman, M. L. and Stone, P. A polynomial-time Nash equilibrium algorithm for repeated games. *Decision Support Systems*, 39:55–66, 2005.
- Liu, Q., Wang, Y., and Jin, C. Learning Markov games with adversarial opponents: Efficient algorithms and fundamental limits. In Chaudhuri, K., Jegelka, S., Song, L., Szepesvari, C., Niu, G., and Sabato, S. (eds.), *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pp. 14036–14053. PMLR, 17–23 Jul 2022. URL <https://proceedings.mlr.press/v162/liu22r.html>.
- Malialis, K. and Kudenko, D. Distributed response to network intrusions using multiagent reinforcement learning. *Engineering Applications of Artificial Intelligence*, 41:270–284, 2015. ISSN 0952-1976. doi: <https://doi.org/10.1016/j.engappai.2015.01.013>. URL <https://www.sciencedirect.com/science/article/pii/S095219761500024X>.
- Mao, W. and Basar, T. Provably efficient reinforcement learning in decentralized general-sum markov games. *CoRR*, abs/2110.05682, 2021. URL <https://arxiv.org/abs/2110.05682>.
- Maskin, E. and Fudenberg, D. The folk theorem in repeated games with discounting or with incomplete information.

- Econometrica, 53(3):533–554, 1986. Reprinted in A. Rubinstein (ed.), *Game Theory in Economics*, London: Edward Elgar, 1995. Also reprinted in D. Fudenberg and D. Levine (eds.), *A Long-Run Collaboration on Games with Long-Run Patient Players*, World Scientific Publishers, 2009, pp. 209–230.
- Nash, J. Non-cooperative games. *Annals of Mathematics*, 54 (2):286–295, 1951. ISSN 0003486X.
- Papadimitriou, C. H. On the complexity of the parity argument and other inefficient proofs of existence. *J. Comput. Syst. Sci.*, 48(3):498–532, 1994.
- Perolat, J., Vylder, B. D., Hennes, D., Tarassov, E., Strub, F., de Boer, V., Muller, P., Connor, J. T., Burch, N., Anthony, T., McAleer, S., Elie, R., Cen, S. H., Wang, Z., Gruslys, A., Malysheva, A., Khan, M., Ozair, S., Timbers, F., Pohlen, T., Eccles, T., Rowland, M., Lanctot, M., Lepia, J.-B., Piot, B., Omidshafiei, S., Lockhart, E., Sifre, L., Beauguerlange, N., Munos, R., Silver, D., Singh, S., Hassabis, D., and Tuyls, K. Mastering the game of stratego with model-free multiagent reinforcement learning. *Science*, 378(6623):990–996, 2022. doi: 10.1126/science.add4679. URL <https://www.science.org/doi/abs/10.1126/science.add4679>.
- Puterman, M. *Markov Decision Processes*. John Wiley & Sons, Ltd, 1 edition, 1994. URL <http://onlinelibrary.wiley.com/doi/10.1002/9780470316887>.
- Roughgarden, T. Intrinsic robustness of the price of anarchy. *Journal of the ACM*, 2015.
- Rubinstein, A. Settling the complexity of computing approximate two-player Nash equilibria. In *Annual Symposium on Foundations of Computer Science (FOCS)*, pp. 258–265. IEEE, 2016.
- Rubinstein, A. Inapproximability of Nash equilibrium. *SIAM Journal on Computing*, 47(3):917–959, 2018.
- Shalev-Shwartz, S., Shammah, S., and Shashua, A. Safe, multi-agent, reinforcement learning for autonomous driving. *CoRR*, abs/1610.03295, 2016. URL <http://arxiv.org/abs/1610.03295>.
- Shapley, L. *Stochastic Games*. PNAS, 1953.
- Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., Van Den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., et al. Mastering the game of go with deep neural networks and tree search. *nature*, 529(7587):484, 2016.
- Song, Z., Mei, S., and Bai, Y. When can we learn general-sum markov games with a large number of players sample-efficiently? In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=6MmiSOHJHR>.
- Syrkanis, V., Agarwal, A., Luo, H., and Schapire, R. E. Fast convergence of regularized learning in games. In *Advances in Neural Information Processing Systems (NIPS)*, pp. 2989–2997, 2015.
- Vovk, V. Aggregating strategies. *Proc. of Computational Learning Theory*, 1990, 1990.
- Weisz, G., Amortila, P., Janzer, B., Abbasi-Yadkori, Y., Jiang, N., and Szepesvari, C. On query-efficient planning in mdps under linear realizability of the optimal state-value function. In Belkin, M. and Kpotufe, S. (eds.), *Proceedings of Thirty Fourth Conference on Learning Theory*, volume 134 of *Proceedings of Machine Learning Research*, pp. 4355–4385. PMLR, 15–19 Aug 2021.
- Yin, D., Hao, B., Abbasi-Yadkori, Y., Lazić, N., and Szepesvári, C. Efficient local planning with linear function approximation. In Dasgupta, S. and Haghtalab, N. (eds.), *Proceedings of The 33rd International Conference on Algorithmic Learning Theory*, volume 167 of *Proceedings of Machine Learning Research*, pp. 1165–1192. PMLR, 29 Mar–01 Apr 2022.
- Zhan, W., Lee, J. D., and Yang, Z. Decentralized optimistic hyperpolicy mirror descent: Provably no-regret learning in markov games. *arXiv preprint arXiv:2206.01588*, 2022.
- Zheng, S., Trott, A., Srinivasa, S., Parkes, D. C., and Socher, R. The ai economist: Taxation policy design via two-level deep multiagent reinforcement learning. *Science Advances*, 8(18):eabk2607, 2022. doi: 10.1126/sciadv.abk2607. URL <https://www.science.org/doi/abs/10.1126/sciadv.abk2607>.

Contents of Appendix

I	Additional results and discussion	14
A	Tighter computational lower bounds under ETH for PPAD	14
B	Discussion and interpretation	15
B.1	Comparison to V-learning	15
B.2	No-regret learning against Markov deviations	16
B.3	On the role of shared randomness	17
B.4	Comparison to lower bounds for finding stationary CCE	17
B.5	Proof of Proposition B.1	17
II	Proofs	17
C	Additional preliminaries	18
C.1	Additional preliminaries for Markov games	18
C.2	Nash equilibria and computational hardness	18
C.3	Query complexity of Nash equilibria	19
D	Proofs of lower bounds for SPARSEMARKOVCCCE (Section 3)	19
D.1	Preliminaries: Online density estimation	19
D.2	Proof of Theorem 3.2	21
E	Proofs of lower bounds for SPARSECCCE (Sections 4 and 5)	24
E.1	Proof of Theorem E.1	25
E.2	Remarks on bit complexity of the rewards	32
F	Equivalence between $p_j^{\text{gen;rnd}}$ and $p_j^{\text{gen;det}}$	33
F.1	Proofs of the equivalence	34

Part I

Additional results and discussion

A. Tighter computational lower bounds under ETH for **PPAD**

Recall that Corollary 3.3 states that if **PPAD** P , then there is no constant $C \leq 4$ and $\text{poly}(|G|)^q$ -time algorithm which solves the $p(|G|^C; |G|^{1/C})$ -**SPARSEMARKOVCCCE** problem for any 2-player Markov game G . Using a stronger complexity-theoretic assumption, the Exponential Time Hypothesis for **PPAD** (Rubinstein, 2016), we can obtain a hardness result which rules out

efficient algorithms even when 1) the accuracy is constant, as opposed to being $|G|^{1/C}$, and 2) T is quasipolynomially large, as opposed to only being of polynomial size, i.e., $|G|^C$.

Corollary A.1 (ETH-hardness of SPARSEMARKOVCCCE). There is a constant $\epsilon_0 > 0$ such that if there exists an algorithm that solves the $p|G|^{\epsilon_0 \log |G|}; \epsilon_0 q$ -SPARSEMARKOVCCCE problem in $|G|^{\epsilon_0 \log |G|}$ time, then the Exponential Time Hypothesis for PPAD fails to hold.

Corollary A.1 is an immediate consequence of Theorem 3.2 and the fact that for some absolute constant $\epsilon_0 > 0$, there are no polynomial-time algorithms for computing ϵ_0 -Nash equilibria in 2-player normal-form games under the Exponential Time Hypothesis for PPAD (as shown in (Rubinstein, 2016)).

B. Discussion and interpretation

Theorems 3.2, 4.3, and 5.2 present barriers—both computational and statistical—toward developing efficient decentralized no-regret guarantees for multi-agent reinforcement learning. We emphasize that no-regret algorithms are the only known approach for obtaining fully decentralized learning algorithms (i.e., those which do not rely even on shared randomness) in normal-form games, and it seems unlikely that a substantially different approach would work in Markov games. Thus, these lower bounds for finding subexponential-length sequences of policies with the no-regret property represent a significant obstacle for fully decentralized multi-agent reinforcement learning. Moreover, these results rule out even the prospect of developing efficient centralized algorithms that produce no-regret sequences of policies, i.e., those which “resemble” independent learning. In this section, we compare our lower bounds with recent upper bounds for decentralized learning in Markov games, and explain how to reconcile these results.

B.1. Comparison to V-learning

The V-learning algorithm (Jin et al., 2021; Song et al., 2022; Mao & Basar, 2021) is a polynomial-time decentralized learning algorithm that proceeds in two phases. In the first phase, the m agents interact over the course of K episodes in a decentralized fashion, playing product Markov policies $p^{1q}, \dots, p^{Kq}$ p^{markov} . In the second phase, the agents use data gathered during the first phase to produce a distributional policy $p p^{\text{gen;rnd}} q$, which we refer to as the output policy of V-learning. As discussed in Section 1, one implication of Theorem 3.2 is that the first phase of V-learning cannot guarantee each agent sublinear regret. Indeed if K is of polynomial size (and PPAD P), this follows because a bound of the form $\text{Reg}_{i;K} p^{1q}, \dots, p^{Kq} q / K$ for all i implies that $p^{1q}, \dots, p^{Kq} q$ solves the $pK; q$ -SPARSEMARKOVCCCE problem.

The output policy $p p^{\text{gen;rnd}} q$ produced by V-learning is an approximate CCE (per Definition 2.1), and it is natural to ask how many product policies it takes to represent as a uniform mixture (that is, whether solves the $p p^{\text{gen;rnd}} q$ -SPARSEMARKOVCCCE problem for a reasonable value of T). First, recall that V-learning requires $K \text{ poly}(H; S; \max_i A_i) q^2$ episodes to ensure that p is an ϵ -CCE. It is straightforward to show that p can be expressed as a non-uniform mixture of at most K^{KHS-1} policies in gen;rnd (we prove this fact in detail below). By discretizing the non-uniform mixture, one can equivalently represent it as uniform mixture of $O(p^{1/q} K^{KHS-1})$ product policies, up to ϵ error. Recalling the value of K , we conclude that we can express as p a uniform mixture of $T \exp(O(p^{1/q} \text{poly}(H; S; \max_i A_i) q))$ product policies in gen;rnd . Note that the lower bound of Theorem 4.3 rules out the efficient computation of an ϵ -CCE represented as a uniform mixture of $T \leq \exp^2 \max(H; S; \max_i A_i) q$ efficiently computable policies in gen;rnd . Thus, in the regime where $1/q$ is polynomial in $H; S; \max_i A_i$, this upper bound on the sparsity of the policy p produced by V-learning matches that from Theorem 4.3, up to a polynomial in the exponent.

The sparsity of the output policy from V-learning. We now sketch a proof of the fact that the output policy p produced by V-learning can be expressed as a (non-uniform) average of K^{KHS-1} policies in gen;rnd , where K is the number of episodes in the algorithm’s initial phase. We adopt the notation and terminology from Jin et al. (2021).

Consider Algorithm 3 of Jin et al. (2021), which describes the second phase of V-learning, which produces the output policy p . We describe how to write p as a weighted average of a collection of product policies, each of which is indexed by a function $h : \mathcal{H} \times \mathcal{S} \times \mathcal{R} \times \mathcal{K} \rightarrow \mathbb{R}$ and a parameter $k_0 \in \mathcal{R} \times \mathcal{K}$: in particular, we will write $p = \sum_{h, k_0} w_{k_0} p^{h, k_0}$, where $w_{k_0} \in \mathbb{R}_{\geq 0}$ are mixing weights summing to 1 and $p^{h, k_0} \in \mathcal{P}^{\text{gen;rnd}}$. The number of tuples $pk_0; q$ is K^{1-KHS} .

We define the mixing weight allocated w_{k_0} to any tuple $pk_0; q$ to be:

$$\frac{1}{K} \prod_{h, k_0} \frac{1}{N_h} \prod_{p \in \mathcal{P}^{\text{gen;rnd}}} \frac{1}{N_p} \prod_{s, k_0} \frac{1}{N_s} \prod_{r, k_0} \frac{1}{N_r} \prod_{k, k_0} \frac{1}{N_k} \prod_{q, k_0} \frac{1}{N_q}$$

where $N_h \in \mathbb{N}$ and $N_{k_0} \in \mathbb{N}$ (for $h \in \mathcal{H}$ and $k_0 \in \mathcal{R} \times \mathcal{K}$) are defined as in (Jin et al., 2021).

Next, for each k_0 , we define $k; P_0^{\text{gen;rnd}}$ to be the following policy: it maintains a parameter k P r k s over the first h / H steps of the episode (as in Algorithm 3 of (Jin et al., 2021)), but upon reaching state s at step h , given the present value of k P r k s, sets $i; p; s; k; q$, and updates $k \leftarrow k^i p; s; q$, and then samples an action $a^i p; s; q$ (where $k^i p; s; q; p; s; q$ are defined in (Jin et al., 2021)). Since the mixing weights w_k ; defined above exactly simulate the random draws of the parameter k in Line 1 and the parameters i in Algorithm 3, Line 4 of (Jin et al., 2021), it follows that the distributional policy defined by Algorithm 3 of (Jin et al., 2021) is equal to $k_0; w_{k_0; k_0}; P P^{\text{gen;rnd}} q$.

B.2. No-regret learning against Markov deviations

As discussed in Section 1, Erez et al. (2022) showed the existence of a learning algorithm with the property that if each agent plays it independently for T episodes, then no player can achieve regret more than $\text{Oppolypm;H;S;max}_i A_i q T^{3/4} q$ by deviating to any fixed Markov policy. This notion of regret corresponds to, in the context of Definition 2.2, replacing $\max_{P^{\text{gen;rnd}}}$ with the smaller quantity $\max_{P^{\text{markov}}}$. Thus, the result of Erez et al. (2022) applies to a weaker notion of regret than that of the SPARSECCE problem, and so does not contradict any of our lower bounds. One may wonder which of these two notions of regret (namely, best possible gain via deviation to a Markov versus non-Markov policy) is the “right” one. We do not believe that there is a definitive answer to this question, but we remark that in many empirical applications of multi-agent reinforcement learning it is standard to consider non-Markov policies (Leibo et al., 2021; Agapiou et al., 2022). Furthermore, as shown in the proposition below, there are extremely simple games, e.g., of constant size, in which Markov deviations lead to “vacuous” behavior: in particular, all Markov policies have the same (suboptimal) value but the best non-Markov policy has much greater value:

Proposition B.1. There is a 2-player, 2-action, 1-state Markov game with horizon 2 and a non-Markov policy $P_2^{\text{gen;rnd}}$ for player 2 so that for all $P_1^{\text{markov}}, V^{1,2} \{1, 2\}$ yet $\max_{P_2^{\text{gen;rnd}}} V^{1,2} \{3, 4\}$.

The proof of Proposition B.1 is provided in Section B.5 below.

Other recent work has also proved no-regret guarantees with respect to deviations to restricted policy classes. In particular, Zhan et al. (2022) studies a setting in which each agent i is allowed to play policies in an arbitrary restricted policy class $P_i^{\text{gen;rnd}}$ in each episode, and regret is measured with respect to deviations to any policy in $P_i^{\text{gen;rnd}}$. Zhan et al. (2022) introduces an algorithm, DORIS, with the property that when all agents play it independently, each agent i experiences regret $O(\text{polypm;A;S;H} q \sqrt{T \log |P_i^{\text{gen;rnd}}|})$ to their respective class $P_i^{\text{gen;rnd}}$.

DORIS is not computationally efficient, since it involves performing exponential weights over the class $P_i^{\text{gen;rnd}}$, which requires space complexity $|P_i^{\text{gen;rnd}}|$. Nonetheless, one can compare the statistical guarantees the algorithm provides to our own results. Let $P_i^{\text{markov;det}} \in P_i^{\text{markov;det}}$ denote the set of deterministic Markov policies of agent i , namely sequences $p_i; 1; \dots; p_i; H; q$ so that $p_i; h : S \rightarrow A_i$. In the case that $P_i^{\text{markov;det}} = P_i^{\text{gen;rnd}}$, we have $\log |P_i^{\text{gen;rnd}}| \leq \text{OpSH log A}_i q$, which means that DORIS obtains no-regret against Markov deviations when m is constant, comparable to Erez et al. (2022).¹⁶ However, we are interested in the setting in which each player’s regret is measured with respect to all deviations in $P_i^{\text{gen;rnd}}$ (equivalently, $P_i^{\text{gen;det}}$). Accordingly, if we take $P_i^{\text{gen;det}} \in P_i^{\text{gen;rnd}}$,¹⁷ then $\log |P_i^{\text{gen;det}}| \leq \text{OpSH log A}_i q^{H-1}$, meaning that DORIS does not imply any sort of sample-efficient guarantee, even for $m = 2$.

Finally, we remark that the algorithm DORIS (Zhan et al., 2022), as well as the similar algorithm OPMD from earlier work of Liu et al. (2022), obtains the same regret bound stated above even when the opponents are controlled by (possibly adaptive) adversaries. However, this guarantee crucially relies on the fact that any agent implementing DORIS must observe the policies played by opponents following each episode; this feature is the reason that the regret bound of DORIS does not contradict the exponential lower bound of Liu et al. (2022) for no-regret learning against an adversarial opponent. As a result of being restricted to this “revealed-policy” setting, DORIS is not a fully decentralized algorithm in the sense we consider in this paper.

¹⁵Note that in the tabular setting, the sample complexity of DORIS (Corollary 1) scales with the size A of the joint action set, since each player’s value function class consists of the class of all functions $f : S \rightarrow \mathbb{R}$, which has Eluder dimension scaling with S , i.e., exponential in m .

¹⁶Erez et al. (2022) has the added bonus of computational efficiency, even for polynomially large m , though has the significant drawback of assuming that the Markov game is known.

¹⁷DORIS plays distributions over policies in $P_i^{\text{gen;det}}$ at each episode, whereas in our lower bounds we consider the setting where a policy in $P_i^{\text{gen;rnd}}$ is played each episode; Facts F.2 and F.3 shows that these two settings are essentially equivalent, in that any policy in $P_i^{\text{gen;rnd}}$ can be simulated by one in $P_i^{\text{gen;det}}$ and vice versa.

B.3. On the role of shared randomness

A key assumption in our lower bounds for no-regret learning is that each of the joint policies $p^{1q}, \dots, p^{Tq}$ produced by the algorithm is a product policy; such an assumption is natural, since it subsumes independent learning protocols in which each agent i selects p^{iq} without knowledge of p^{iq} . Compared to general (stochastic) joint policies, product policies have the desirable property that, to sample a trajectory from $p^{1q}, \dots, p^{Tq}$, the agents do not require access to shared randomness. In particular, each agent i can independently sample its action from p^{iq} at each of the h steps of the episode. It is natural to ask how the situation changes if we allow the agents to use shared random bits when sampling from their policies, which corresponds to allowing $p^{1q}, \dots, p^{Tq}$ to be non-product policies. In this case, V-learning yields a positive result via a standard “batch-to-online” conversion: by applying the first phase of V-learning during the first $T^{2/3}$ episodes and playing trajectories sampled i.i.d. from the output policy produced by V-learning during the remaining $T - T^{2/3}$ episodes (which requires shared randomness), it is straightforward to see that a regret bound of order $\text{poly}(H; S; \max_i A_i) T^{2/3}$ can be obtained. Similar remarks apply to SPoCMAR (Daskalakis et al., 2022), which can obtain a slightly worse regret bound of order $\text{poly}(H; S; \max_i A_i) T^{3/4}$ in the same fashion. In fact, the batch-to-online conversion approach gives a generic solution for the setting in which shared randomness is available. That is, the assumption of shared randomness eliminates any distinction between no-regret algorithms and (non-sparse) equilibrium computation algorithms, modulo slight loss in rates. For this reason, the shared randomness assumption is too strong to develop any sort of distinct theory of no-regret learning.

B.4. Comparison to lower bounds for finding stationary CCE

A separate line of work Daskalakis et al. (2022); Jin et al. (2022) has recently shown PPAD-hardness for the problem of finding stationary Markov CCE in infinite-horizon discounted stochastic games. These results are incomparable with our own: stationary Markov CCE are not sparse (in the sense of Definition 3.1), whereas we do not require stationarity of policies (as is standard in the finite-horizon setting).

B.5. Proof of Proposition B.1

Below we prove Proposition B.1.

Proof of Proposition B.1. We construct the claimed Markov game G as follows. The single state is denoted by s ; as there is only a single state, the transitions are trivial. We denote each player’s action space as $A_1 \times A_2 = \{1, 2\}$. The rewards to player 1 are given as follows: for all $a_1, a_2 \in A$,

$$R_{1,1}(s; a_1, a_2) = \begin{cases} 1 & \text{if } a_1 = 1 \\ 0 & \text{otherwise} \end{cases}; \quad R_{1,2}(s; a_1, a_2) = \begin{cases} 1 & \text{if } a_1 = a_2 \\ 0 & \text{otherwise} \end{cases}.$$

We allow the rewards of player 2 to be arbitrary; they do not affect the proof in any way.

We let $p_{2,1;2,2}^{\text{gen;rnd}}$ be the policy which plays a uniformly random action at step 1 and then plays the same action at step 2: formally, $p_{2,1} = \text{Unif}(A_2)$, and $p_{2,2}(a_2; a_1, r_1, q) = p_{2,1}(a_2)$. Then for any Markov policy p_1^{Markov} of player 1, we must have $V_1(p_1, p_2) = \mathbb{E}_{1,2} [r_1(a_1, a_2)] = \frac{1}{2}$, which means that $V_1 = \frac{1}{2}$.

On the other hand, any general (non-Markov) policy $p_1^{\text{gen;rnd}}$ which satisfies

$$\begin{aligned} p_{1,2}(a_1, a_2; r_1, q) &= \begin{cases} 1 & \text{if } a_1 = 1 \\ 0 & \text{otherwise} \end{cases} \\ p_{1,1}(a_1, a_2; r_1, q) &= \begin{cases} 1 & \text{if } a_1 = 2 \\ 0 & \text{otherwise} \end{cases} \end{aligned}$$

has $V_1 = \frac{1}{2}$.

□

Part II

Proofs

C. Additional preliminaries

C.1. Additional preliminaries for Markov games

Deterministic policies. It will be helpful to introduce notation for deterministic general (non-Markov) policies, which correspond to the special case of randomized policies where each policy $\pi_{i,h}$ exclusively maps to singleton distributions. In particular, a deterministic general policy of agent i is a collection of mappings $\pi_{i,1}; \dots; \pi_{i,H} : \mathcal{H}_i \rightarrow \mathcal{S}^{\tilde{N}_i}$. We denote by $\mathcal{P}^{\text{gen;det}}_i$ the space of deterministic general policies of agent i , and further write $\mathcal{P}^{\text{gen;det}}_1, \dots, \mathcal{P}^{\text{gen;det}}_m$ to denote the space of joint deterministic policies. We use the convention throughout that deterministic policies are denoted by the letter π , whereas randomized policies are denoted by μ .

Additional facts on regret and CCE. The following facts regarding deterministic policies and the definition of coarse correlated equilibria and regret are well-known:

- In the context of Definition 2.1 (defining an ϵ -CCE), the maximizing policy π^1 can always be chosen to be deterministic, so $\pi \in \mathcal{P}^{\text{gen;rnd}}_i$ is an ϵ -CCE if and only if $\max_{\pi \in \mathcal{P}^{\text{gen;det}}_i} V_i^{\pi} \geq V_i^{\pi^1} - \epsilon$.
- In the context of (1) in the definition of regret, the maximum over $\pi \in \mathcal{P}^{\text{gen;rnd}}_i$ is always achieved by a deterministic general policy, so we have $\text{Reg}_{i,T} = \max_{\pi \in \mathcal{P}^{\text{gen;det}}_i} \sum_{t=1}^T V_i^{\pi} - \sum_{t=1}^T V_i^{\pi_t}$.

Next, the following standard result shows that the uniform average of any no-regret sequence forms an approximate coarse correlated equilibrium.

Fact C.1 (No-regret is equivalent to CCE). Suppose that a sequence of policies $\pi^1; \dots; \pi^T \in \mathcal{P}^{\text{gen;rnd}}_i$ satisfies $\text{Reg}_{i,T} \leq \epsilon$ for each $i \in [m]$. Then the uniform average of these T policies, namely the distributional policy $\bar{\pi} = \frac{1}{T} \sum_{t=1}^T \pi_t$, is an ϵ -CCE.

Likewise if a sequence of policies $\pi^1; \dots; \pi^T \in \mathcal{P}^{\text{gen;rnd}}_i$ has the property that the distributional policy $\bar{\pi} = \frac{1}{T} \sum_{t=1}^T \pi_t$ is an ϵ -CCE, then we have $\text{Reg}_{i,T} \leq \epsilon T$ for all $i \in [m]$.

Fact C.1 is an immediate consequence of Definitions 2.1 and 2.2.

C.2. Nash equilibria and computational hardness.

The most foundational and well known solution concept for normal-form games is the Nash equilibrium (Nash, 1951).

Definition C.2 (pn;q-NASH problem). For a normal-form game $G = (M_1; \dots; M_m)$ and $\epsilon > 0$, a product distribution $\mu \in \prod_{i=1}^m \Delta(\mathcal{S}_i)$ is said to be an ϵ -Nash equilibrium for G if for all $i \in [m]$,

$$\max_{a_i \in \mathcal{S}_i} E_{\mu} [r_i(a_i, \mathbf{a}_{-i})] \geq E_{\mu} [r_i(\mathbf{a}_i, \mathbf{a}_{-i})] - \epsilon$$

We define the m -player pn;q-NASH problem to be the problem of computing an ϵ -Nash equilibrium of a given m -player n -action normal-form game.¹⁸

Informally, μ is an ϵ -Nash equilibrium if no player i can gain more than ϵ in reward by deviating to a single fixed action a_i^1 , while all other players randomly choose their actions according to μ . Despite the intuitive appeal of Nash equilibria, they are intractable to compute: for any $\epsilon > 0$, it is PPAD-hard to solve the pn;q-NASH problem, namely, to compute ϵ -approximate Nash

¹⁸One must also take care to specify the bit complexity of representing a normal-form game. We assume that the payoffs of any normal-form game given as an instance to the pn;q-NASH problem can each be expressed with $\max(n, m)$ bits; this assumption is without loss of generality as long as $\epsilon \geq 2^{-\max(n, m)}$ (which it will be for us).

equilibria in 2-player n -action normal-form games (Daskalakis et al., 2009; Chen et al., 2006; Rubinstein, 2018). We recall that the complexity class PPAD consists of all total search problems which have a polynomial-time reduction to the End-of-The-Line (EOTL) problem. PPAD is the most well-studied complexity class in algorithmic game theory, and it is widely believed that $\text{PPAD} \neq \text{P}$. We refer the reader to (Daskalakis et al., 2009; Chen et al., 2006; Rubinstein, 2018; Papadimitriou, 1994) for further background on the class PPAD and the EOTL problem.

C.3. Query complexity of Nash equilibria

Our statistical lower bound for the SPARSECCE problem in Theorem 5.2 relies on existing query complexity lower bounds for computing approximate Nash equilibria in m -player normal-form games. We first review the query complexity model for normal-form games.

Oracle model for normal-form games. For $m, n \in \mathbb{N}$, consider an m -player n -action normal form game G , specified by payoff tensors $M_1, \dots, M_m$. Since the tensors $M_1, \dots, M_m$ contain a total of $m n^m$ real-valued payoffs, in the setting when m is large, it is unrealistic to assume that an algorithm is given the full payoff tensors as input. Therefore, prior work on computing equilibria in such games has studied the setting in which the algorithm makes adaptive oracle queries to the payoff tensors.

In particular, the algorithm, which is allowed to be randomized, has access to a payoff oracle O_G for the game G , which works as follows. At each time step, the algorithm can choose to specify an action profile $a \in [n]^m$ and then query O_G at the action profile a . The oracle O_G then returns the payoffs $p_{M_1 a}, \dots, p_{M_m a}$ for each player if the action profile a is played.

Query complexity lower bound for approximate Nash equilibrium. The following theorem gives a lower bound on the number of queries any randomized algorithm needs to make to compute an approximate Nash equilibrium in an m -player game.

Theorem C.3 (Corollary 4.5 of (Rubinstein, 2016)). There is a constant $\epsilon > 0$ so that any randomized algorithm which solves the ϵ -NASH problem for m -player normal-form games with probability at least $2/3$ must use at least $2^{\Omega(m)}$ payoff queries.

We remark that (Babichenko, 2016; Chen et al., 2017) provide similar, though quantitatively weaker, lower bounds to that in Theorem C.3. We also emphasize that the lower bound of Theorem C.3 applies to any algorithm, i.e., including those which require extremely large computation time.

D. Proofs of lower bounds for SPARSEMARKOVCCCE (Section 3)

D.1. Preliminaries: Online density estimation

Our proof makes use of tools for online learning with the logarithmic loss, also known as conditional density estimation. In particular, we use a variant of the exponential weights algorithm known as Vovk's aggregating algorithm in the context of density estimation (Vovk, 1990; Cesa-Bianchi & Lugosi, 2006). We consider the following setting with two players, a Learner and Nature. Furthermore, there is a set Y , called the outcome space, and a set X , called the context space; for our applications it suffices to assume Y and X are finite. For some $T \in \mathbb{N}$, there are T time steps $t = 1, 2, \dots, T$. At each time step $t \in [T]$:

- Nature reveals a context $x^{(t)} \in X$;
- Having seen the context $x^{(t)}$, the learner predicts a distribution $\hat{p}^{(t)} \in \Delta(Y)$;
- Nature chooses an outcome $y^{(t)} \in Y$, and the learner suffers loss $-\log \hat{p}^{(t)}(y^{(t)})$.

For each $t \in [T]$, we let $H^{(t)} = (x^{(1)}, y^{(1)}; \dots; x^{(t)}, y^{(t)})$ denote the history of interaction up to step t ; we emphasize that each context $x^{(t)}$ may be chosen adaptively as a function of $H^{(t-1)}$. Let $\mathcal{F}^{(t)}$ denote the sigma-algebra generated by $H^{(t)}$. We measure performance in terms of regret against a set I of experts, also known as the expert setting. Each expert $i \in I$ consists of a function $p_i : X \rightarrow \Delta(Y)$. The regret of an algorithm against the expert class I when it receives contexts $x^{(1)}; \dots; x^{(T)}$ and observes outcomes $y^{(1)}; \dots; y^{(T)}$ is defined as

$$\text{Reg}_{I;T} = \sum_{t=1}^T -\log \hat{p}^{(t)}(y^{(t)}) - \min_{i \in I} \sum_{t=1}^T -\log p_i(x^{(t)})(y^{(t)})$$

Note that the learner can observe the expert predictions $p_i(x^{(t)})$ and use them to make its own prediction at each round t .

D.2. Proof of Theorem 3.2

Proof of Theorem 3.2. Fix $n \in \mathbb{N}$, which we recall represents an upper bound on the description length of the Markov game. Assume that we are given an algorithm B that solves the $pT; q$ -SPARSEMARKOVCCCE problem for Markov games G satisfying $|G|/n$ in time U . We proceed to describe an algorithm which solves the 2-player $pn^{1/2}\{2u; 4q$ -NASH problem in time $pnTUq^{C_0}$, as long as $T \leq \exp^2 n^{1/2}\{2^5 q$. First, define $n_0 := \lfloor n^{1/2}\{2u$, and consider an arbitrary 2-player n_0 -action normal form G , which is specified by payoff matrices $M_1; M_2 \in \mathbb{R}^{n_0 \times n_0}$, so that all entries of the game can be written in binary using at most n_0 bits (recall, per footnote 18, that we may assume that the entries of an instance of $pn_0; 4q$ -NASH can be specified with n_0 bits). Based on G , we construct a 2-player Markov game $G : GpGq$ as follows:

Definition D.3. We define the game $GpGq$ to consist of the tuple $GpGq; pS; H; pA; q; pR; q; pI; q; pR; q; q$, where:

- The horizon of G is $H = 2\lfloor n_0\{2u$ (i.e., the largest even number at most n_0).
- Let A_{n_0} ; the action spaces of the 2 agents are given by $A_1, A_2 \subseteq \mathbb{R}^S$.
- There are a total of $A^2 - 1$ states: in particular, there is a state $s_{pa_1; a_2}$ for each $pa_1; a_2 \in \mathbb{R}^{A^2}$, as well as a distinguished state s , so we have:

$$S = \{s_{pa_1; a_2} : pa_1; a_2 \in \mathbb{R}^{A^2}\} \cup \{s\}$$

- For all odd $h \in \mathbb{R}^H$, the reward to agents $j \in \mathbb{R}^2$ given that the action profile $pa_1; a_2$ is played at step h is given by $R_{j,h}(s; pa_1; a_2) = \frac{1}{H} \sum_{s \in S} M_j(s; a_1; a_2)$, for all $s \in S$. All agents receive 0 reward at even steps $h \in \mathbb{R}^H$.
- At odd steps $h \in \mathbb{R}^H$, if actions $a_1; a_2 \in \mathbb{R}^{A^2}$ are taken, the game transitions to the state $s_{pa_1; a_2}$. At even steps $h \in \mathbb{R}^H$, the game always transitions to the state s .
- The initial state (i.e., at step $h=1$) is s (i.e., is a singleton distribution supported on s).

It is evident that this construction takes polynomial time, and satisfies $|G|/A^2 \leq 1/n^2 \leq 1/n$. We will now show by applying the algorithm B to G , we can efficiently compute 4-approximate Nash equilibrium for the original game G . To do so, we appeal to Algorithm 1.

Algorithm 1 Algorithm to compute Nash equilibrium used in proof of Theorem 3.2.

- 1: Input: 2-player, n_0 -action normal form game G .
 - 2: Construct the 2-player Markov game $GpGq$ per Definition D.3, which satisfies $|G|/n$.
 - 3: Call the algorithm B on the game G , which produces a sequence $p^{1q}; \dots; p^{Tq}$, where each $p^{tq} \in \mathbb{P}^{\text{markov}}$.
 - 4: for $t \in \mathbb{R}^T$ and odd $h \in \mathbb{R}^H$ do
 - 5: if $p^{tq} \in \mathbb{P}^{\text{markov}}$ is a $p_4; nq$ -Nash equilibrium of G : then 6:
 - return p^{tq} .
 - 7: end if
 - 8: end for
 - 9: if the for loop terminates without returning: return fail.
-

Algorithm 1 proceeds as follows. First, it constructs the 2-player Markov game $GpGq$ as defined above, and calls the algorithm B , which returns a sequence $p^{1q}; \dots; p^{Tq} \in \mathbb{P}^{\text{markov}}$ of product Markov policies with the property that the average $\frac{1}{T} \sum_{t=1}^T I_{p^{tq}}$ is an -CCE of G . It then enumerates over the distributions $p^{tq} \in \mathbb{P}^{\text{markov}}$ for each $t \in \mathbb{R}^T$ and $h \in \mathbb{R}^H$ odd, and checks whether each one is a 4-approximate Nash equilibrium of G . If so, the algorithm outputs such a Nash equilibrium, and otherwise, it fails. The proof of Theorem 3.2 is thus completed by the following lemma, which states that as long as $\frac{1}{T} \sum_{t=1}^T I_{p^{tq}}$ is an -CCE of G , Algorithm 1 never fails.

Lemma D.4 (Correctness of Algorithm 1). Consider the normal form game G and the Markov game $GpGq$ as constructed above, which has horizon H . For any $\epsilon \in (0, 1]$, $T \in \mathbb{N}$, if $T \leq \exp^2 n^{1/2}\{2^8 q$ and $p^{1q}; \dots; p^{Tq} \in \mathbb{P}^{\text{markov}}$ are product Markov policies so that $\frac{1}{T} \sum_{t=1}^T I_{p^{tq}}$ is an ϵ -CCE of G , then there is some odd $h \in \mathbb{R}^H$ and $t \in \mathbb{R}^T$ so that p^{tq} is an ϵ -Nash equilibrium of G .

The proof of Lemma D.4 is given below. Applying Lemma D.4 with $\epsilon = 4$ (which is a valid application since $T \exp(n_0 p_4 q^2) \{2^8 q\}$ by our assumption on T), yields that Algorithm 1 always finds a 4-Nash equilibrium of the n_0 -action normal form game G , thus solving the given instance of the $p_{n_0, 4q}$ -NASH problem. Furthermore, it is straightforward to see that Algorithm 1 runs in time $O(p n T q^{C_0} / p U n T q^{C_0})$, for some constant $C_0 \in \mathbb{N}$.

□

Proof of Lemma D.4. Consider a sequence of product Markov policies $\pi^{p_1 q}, \dots, \pi^{p_T q}$ with the property that the average $\frac{1}{T} \sum_{t=1}^T \mathbb{I}_{\pi^{p_t q}}$ is an $p_{n_0, 4q}$ -CCE of G . For all odd $h \in \mathcal{H}$ and $j \in \mathcal{P}$, let $\pi_{j,h}^{p_t q}$ be the distribution played under $\pi^{p_t q}$ by player j at step h (at the unique state s with positive probability of being reached at step h). For odd h , we have $\pi_{j,h}^{p_t q} = \pi_{j,h}^{p_{t-1} q}$, and our goal is to show that for some odd $h \in \mathcal{H}$ and $t \in \mathcal{T}$, $\pi_{j,h}^{p_t q}$ is an ϵ -Nash equilibrium of G . To proceed, suppose for the sake of contradiction that this is not the case.

Let us write $\mathcal{O}_H = \{h \in \mathcal{H} : h \text{ odd}\}$ to denote the set of odd-numbered steps, and $\mathcal{E}_H = \mathcal{H} \setminus \mathcal{O}_H$ to denote the set of even-numbered steps. Let $\mathcal{H}_0 = \mathcal{O}_H \cup \mathcal{E}_H$. We first note that for $j \in \mathcal{P}$, agent j 's value under the mixture policy is given as follows:

$$V_j = \frac{1}{T} \sum_{h \in \mathcal{H}_0} \mathbb{E}_{\pi_{j,h}^{p_t q}} [r_{j,h}(a_1, a_2, s)]$$

For each $j \in \mathcal{P}$, we will derive a contradiction by constructing a (non-Markov) deviation policy for player j in G , denoted $\pi_j^{\text{gen}, \det}$, which will give player j a significant gain in value against the policy π . To do so, we need to specify $\pi_{j,h}^{\text{gen}, \det}(s_h) \in \mathcal{A}_j$, for all $j \in \mathcal{P}$, $h \in \mathcal{H}$, and $s_h \in \mathcal{S}$; note that we may restrict our attention only to histories $j; h-1$ that occur with positive probability under the transitions of G .

Fix any $h_0 \in \mathcal{H}$, $j; h_0 \in \mathcal{H}$, and $s_{h_0} \in \mathcal{S}$. If $j; h_0$ occurs with positive probability under the transitions of G , then for each $h \in \mathcal{O}_H$, h_0-1 and both $j^1 \in \mathcal{P}$, the action played by agent j^1 at step h is determined by $j; h_0-1$. Namely, if the state at step $h-1$ of $j; h_0-1$ is $s_{p_{j,h_0-1} q}$, then player j^1 played action a^1 at step h . So, for each $h \in \mathcal{O}_H$ with h_0-1 , we may define $p_{j,h_0-1; h} q$ as the action profile played at step h , which is a measurable function of $j; h_0-1$. With this in mind, we define $\pi_{j,h_0-1; h} q$ by applying Vovk's aggregating algorithm (Proposition D.2) as follows.

1. If h_0 is even, play an arbitrary action (note that the actions at even-numbered steps have no influence on the transitions or rewards).
2. If h_0 is odd, define $\pi_{j,h_0} q$ by $\pi_{j,h_0} q = \mathbb{E}_{\pi_{j,h_0} q} [r_{j,h_0}(a_1, a_2, s)]$, where $\pi_{j,h_0} q$ is defined as follows: for $t \in \mathcal{T}$,

$$\pi_{j,h_0} q = \frac{\exp \left(\sum_{h=h_0}^t \log \frac{1}{\pi_{j,h} q} \right)}{\sum_{h=h_0}^t \exp \left(\sum_{h=h_0}^t \log \frac{1}{\pi_{j,h} q} \right)}$$

Note that $\pi_{j,h_0} q$ is a function of $j; h_0-1$ via the action profiles $p_{j,h_0-1; h} q$; to simplify notation, we suppress this dependence.

3. Then for any state $s_{h_0} \in \mathcal{S}$, define $\pi_{j,h_0} q$ to be a best response to $\pi_{j,h_0} q$ namely

$$\pi_{j,h_0} q = \arg \max_{a_j \in \mathcal{A}_j} \mathbb{E}_{\pi_{j,h_0} q} [r_{j,h_0}(a_j, a_{-j}, s)] \quad (4)$$

Note that, for odd h_0 , the distribution $\pi_{j,h_0} q$ defined above can be viewed as an application of Vovk's online aggregation algorithm at step $p_{j,h_0-1} q$ in the following setting: the number of steps (T , in the notation of Proposition D.2; note that T plays a different role in the present proof) is $H_0 \in \mathbb{N}$, the context space is $\mathcal{O}_H$, and the outcome space is $\mathcal{A}_j$.¹⁹ There are T experts $\pi^{p_1 q}, \dots, \pi^{p_T q}$ (i.e., we have $\mathbb{I}_{\pi^{p_t q}} = \mathbb{I}_{\pi^{p_t q}}(s_{p_{j,h_0-1} q})$), whose predictions on a context $h \in \mathcal{O}_H$ are defined as follows: the expert $\pi^{p_t q}$ predicts

¹⁹Here j denotes the index of the player who is not j .

$p^{ptq}_{j,h} q_{j,h}$. Then, the distribution $q_{j,h}$ is obtained by updating the aggregation algorithm with the context-observation pairs $p_{j,h}; a_{j,h} q$, for odd values of $h \in H_0$.

We next analyze the value of $V_j^{j,h}$ for j Pr2s to show that the deviation strategy we have defined indeed obtains significant gain. To do so, recall that this value represents the payoff for player j under the process in which we draw an index t PrT's uniformly at random, then for each step h PrHs, player j plays according to p_j and player j plays according to p^{ptq}_j . (In particular, at odd-numbered steps, player j plays according to $p_{j,h}$.) We recall that E_j rs denotes the expectation under this process. We let $j;h1$ $P H_{j;h1}$ denote the random variable which is the history observed by player j in this setup, i.e., when the policy played is j , and let $tpa_{1,h}; a_{2,h} q$ denote the action profiles for odd rounds, which are a measurable function of each player's trajectory.

We apply Proposition D.2 with the time horizon as H_0 , and with the set of experts set to $I : p^{p1q}; \dots; p^{ptq} u$ as defined above. The context sequence the sequence of increasing values of $h P O_H$, and for each $h P O_H$, the outcome at step $ph \in 1q\{2$ (for which the context is h) is distributed as $a_{j,h} p^{ptq}_{j,h} q$ conditioned on t , which in particular satisfies the realizability assumption stated in Proposition D.2. Then, since (as remarked above), the distributions $q_{j,h}$, for $h P O_H$, are exactly the predictions made by Vovk's aggregating algorithm, Proposition D.2 gives that²⁰

$$E_j \left[\sum_{h \in H_0} D_{TV}(p_{j,h}; p_j, q_{j,h}) \right] \leq \sum_{h \in H_0} D_{TV}(p_{j,h}; p^{ptq}_{j,h} q_{j,h}) \leq H_0 \log T : \quad (5)$$

Recall that we have assumed for the sake of contradiction that $p^{ptq}_{j,h}$ is not an α -Nash equilibrium of G for each h PrHs and t PrT's. Consider a fixed draw of the random variable t PrT's defined above. Then it holds that for j Pr2s and h PrHs, defining

$$o_{j;h} : \max_{a_j \in A_j} E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] - E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] \quad (6)$$

we have $o_{j;h} \geq 0$. Consider any j Pr2s, $h P O_H$, and a history $j;h1$ $P H_{j;h1}$ of agent j up to step $h1$ (conditioned on t). Let us write $j;h : D_{TV}(p_{j,h}; p_j, q_{j,h})$; note that $j;h$ is a function of $j;h1$, through its dependence on $q_{j,h}$. We have, by the definition of $j;h$ $p_{j,h}$ in (4) and the definition of $j;h$,

$$E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] - E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] \leq E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] - E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] \quad (7)$$

Combining (6) and (7), we get that for any fixed $h P O_H$, j Pr2s, and $j;h1$ $P H_{j;h1}$,

$$E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] - E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] \leq E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] - E_{j,h}^a [r(p_{j,h} q_{j,h}, a_j, a_{-j})] \quad (8)$$

²⁰In fact, Proposition D.2 implies that a similar bound holds uniformly for each possible realization of t , but (5) suffices for our purposes.

Averaging over the draw of $tPrTs$, which we recall is chosen uniformly, we see that

$$\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \quad (9)$$

$$\begin{aligned} & \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \\ & \leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \end{aligned} \quad (10)$$

$$\leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \leq \frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right] \quad (11)$$

where (9) follows from the definition $\frac{1}{T} \sum_{t=1}^T \mathbb{E}_{\mathbf{p}^{ptq}} \left[\sum_{j \in \mathcal{H}} V_j^{ptq} - V_j^{Pr2s} \right]$, (10) follows from (8), and (11) uses (5). As long as T is large enough, the this expression is bounded below by $\frac{1}{4T}$, meaning that $\mathbf{p}^{ptq}$ is not an ϵ_0 -approximate CCE. This completes the contradiction. $\square$

E. Proofs of lower bounds for SP^A RSECCCE (Sections 4 and 5)

In this section we prove our computational lower bounds for solving the SP^A RSECCCE problem with $m \geq 3$ players (Theorem 4.3 and Corollary 4.4), as well as our statistical lower bound for solving the SP^A RSECCCE problem with a general number m of players (Theorem 5.2).

Both theorems are proven as consequences of a more general result given in Theorem E.1 below, which reduces the NASH problem in m -player normal-form games to the SP^A RSECCCE problem in $m-1$ -player Markov games. In more detail, the theorem shows that (a) if an algorithm for SP^A RSECCCE makes few calls to a generative model oracle, then we get an algorithm for the NASH problem with few calls to a payoff oracle (see Section C.3 for background on the payoff oracle for the NASH problem), and (b) if the algorithm for SP^A RSECCCE is computationally efficient, then so is the algorithm for the NASH problem.

Theorem E.1. There is a constant $C_0 \geq 0$ so that the following holds. Consider $n, m \geq 1$, and suppose $T, N, Q \geq 1$ and $\epsilon_0 > 0$ satisfy $\frac{1}{T} \exp\left(-\frac{2\epsilon_0^2 m}{m^2}\right) \geq \frac{1}{2}$. Suppose there is an algorithm B which, given a generative model oracle for a $m-1$ -player Markov game G with $|G| \leq n$, solves the $pT; N$ - SP^A RSECCCE problem for G using Q generative model oracle queries. Then the following conclusions hold:

- For any $\epsilon_0 > 0$, the m -player $pT; N$ -NASH problem for any normal-form game G can be solved, with failure probability ϵ_0 , using at most $C_0 p Q \log p \log n$ queries to a payoff oracle O_G for G .
- If the algorithm B additionally runs in time U for some $U \leq N$, then the algorithm solving NASH from the previous bullet point runs in time $p n m T N U \log p \log n$.

Theorem 4.3 follows directly from Theorem E.1 by taking $m=2$.

Proof of Theorem 4.3. Suppose there is an algorithm which, given the description of any 3-player Markov game G with $|G| \leq n$, solves the $pT; N$ - SP^A RSECCCE problem in time U . Such an algorithm immediately yields an algorithm which can solve the $pT; N$ - SP^A RSECCCE problem in time $U \leq |G|^{O(1)}$ using only a generative model oracle, since the exact description of the Markov game can be obtained with $HS|A|/HS \max_i A_i q^3 / |G|^5$ queries to the generative model (across all $ph; s; aq$ tuples). We can now solve the problem of computing a $\frac{1}{2}$ -Nash equilibrium of a given 2-player $n \times n$ normal form game G as follows. We simply apply the algorithm of Theorem E.1 with $m=2$, noting that the oracle O_G in the theorem statement can be implemented by reading the corresponding bits of input of the input game G . The second bullet point yields that this algorithm

takes time $p n T N U \log p \{1/q\}^{C_0}$, for some constant C_0 . Furthermore, the assumption $T \leq \exp p^2 n \{\mu\} m^2 q$ of Theorem E.1 is implied by the assumption that $T \leq \exp p^2 n \{16q\}$ of Theorem 4.3. $\square$

In a similar manner, Theorem 5.2 follows from Theorem E.1 by applying Theorem C.3, which states that there is no randomized algorithm that finds approximate Nash equilibria of m -player, 2-action normal form games in time $2^{O(pmq)}$.

Proof of Theorem 5.2. Let ϵ be the constant from Theorem C.3, and consider any $m \geq 3$. Suppose there is an algorithm which, for any m -player Markov game G with $|G| \leq 2m^6$, makes Q oracle queries to a generative model oracle for G , and solves the $pT; \epsilon$ -NASH problem for G for some $T; N \geq N$ so that $T \leq \exp p c m q$, for a sufficiently small absolute constant c . Then, by Theorem E.1 with $\epsilon \leq p10mq$ and $n = m^6$ (which ensures that $T \leq \exp p p_0 \{p10mq\}^2 \ln \{\mu\} m^2 q$ as long as c is sufficiently small), there is an algorithm which solves the $p m^5; \epsilon$ -NASH problem—and thus the $p2; \epsilon$ -NASH problem—for $p m 1q$ -player games with failure probability $1/3$, using $O(pmq)$ queries to a payoff oracle. But by Theorem C.3, any such algorithm requires $Q \geq 2^{O(pmq)}$ queries to a payoff oracle. It follows that $Q \geq 2^{O(pmq)}$, as desired. $\square$

E.1. Proof of Theorem E.1

Proof of Theorem E.1. Fix any $m \geq 2$, $n \geq N$. Suppose we are given an algorithm B that solves the $p m 1q$ -player $pT; \epsilon$ -NASH problem for Markov games G satisfying $|G| \leq n$, running in time U and using at most Q generative model queries. We proceed to describe an algorithm which solves the m -player $p n \{\mu\} 16 p m 1q$ -NASH problem using $C_0 p Q \log p \{1/q\}^{C_0}$ queries to a payoff oracle, and running in time $p n m T N U \log p \{1/q\}^{C_0}$, where ϵ represents the failure probability. Define $n_0 = \lceil n \{\mu\} \rceil$, and assume we are given an arbitrary m -player n_0 -action normal form G , which is specified by payoff matrices $M_1, \dots, M_m \in \mathbb{R}^{n_0 \times n_0}$. We assume that all entries of each of the matrices M_j have only the most significant $\max\{n_0, \lceil \log \{1/\epsilon\} \rceil$ bits nonzero; this assumption is without loss of generality, since by truncating the utilities to satisfy this assumption, we change all payoffs by at most ϵ , which degrades the quality of any approximate equilibrium by at most 2ϵ (in addition, we have $\lceil \log \{1/\epsilon\} \rceil \leq n_0$ since we have assumed $1/T \leq \exp p^2 n_0 \{m^2 q\}$). We assume $\epsilon \leq 1/2$ without loss of generality. Based on G , we construct an $p m 1q$ -player Markov game $G: GpGq$ as follows.

Definition E.2. We define the Markov game $GpGq$ as the tuple $GpGqpS; H; pA_1 q; \dots; pA_m q; P; pR_1 q; \dots; pR_m q$, where:

- The horizon of G is chosen to be the power of 2 satisfying $n_0/H \leq 2n_0$.
- Let $A: n_0$. The action spaces of agents $1; 2; \dots; m$ are given by $A_1 \times \dots \times A_m$. The action space of agent $m+1$ is

$$A_{m+1} = \{p_j; a_j q : j \in \{1, \dots, m\}; a_j \in A_j\};$$

so that $|A_{m+1}| = |A|/n$.

We write $A = \prod_{j=1}^m A_j$ to denote the joint action space of the first m agents, and $\bar{A} = \prod_{j=1}^m A_j$ to denote the joint action space of all agents.

- There is a single state, denoted by s , i.e., $S = \{s\}$ (in particular, μ is a singleton distribution supported on s).
- For all $h \in \mathbb{R}^H$, the reward for agent $j \in \{1, \dots, m\}$, given an action profile $a = p a_1 q; \dots; a_m q$ at the unique state s , is as follows: writing $a = p j^1 q; a^1 q$, we have

$$R_{j,h}(p s; a q) = \frac{1}{H} \sum_{h=1}^H \text{encpaq} \quad (12)$$

where $R_{j,h}(p s; a q)$ is defined per the k-bitzer construction of (Borgs et al., 2008):

$$R_{j,h}(p s; a q) = \begin{cases} 0 & \text{if } j \in \{1, \dots, m\} \text{ and } a_j \neq j^1 \\ \frac{1}{H} \sum_{h=1}^H p M_j q_{a_1, \dots, a_m} & \text{if } j = j^1 \\ \frac{1}{H} \sum_{h=1}^H p M_j q_{a_1, \dots, a^1, \dots, a_m} & \text{if } j = m+1 \end{cases} \quad (13)$$

In (12) above, $\text{encpaq} \in \{0, 1\}^H$ is the binary representation of a binary encoding of the action profile a . In particular, if the binary encoding of a is $p b_1 q; \dots; b_N q$, with $b_i \in \{0, 1\}$, then $\text{encpaq} = \sum_{i=1}^N 2^{i-1} b_i$. Note that encpaq takes $N = O(p m \log n_0 q / O(p m \log n q))$ bits to specify.

Algorithm 2 Algorithm to compute Nash equilibrium used in proof of Theorem E.1.

- 1: Input:
- 2: Parameters $n; n_0; m; T; P, N, \{p_{j,h}^q, K, r_{j,h}^q\}_{j \in [m], h \in [n]}$.
- 3: An m -player, n_0 -action normal form game G , with utilities accessible by oracle O_G .
- 4: An algorithm B for computing approximate CCE of Markov games.
- 5: Call the algorithm B on the pm $1q$ -player Markov game $G_{p,q}$ constructed as in Definition E.2, which produces a sequence $p^{1q}, \dots, p^{Tq}$, where each $p^{tq} = p^{tq}_1, \dots, p^{tq}_{pm}$ with $p^{tq}_m \in \mathcal{P}^{gen;rnd}_m$. Here, we use the oracle O_G to simulate generative model oracle queries made by B .
- 6: Draw $tPrTs$ uniformly at random.
- 7: For each $j \in [m]$, initialize $j_{;0}$ to be an empty trajectory.
- 8: for $h \in [n]$: do
- 9: Set s_h (per the transitions of G).
- 10: For each $j \in [m]$, define $q_{j,h} : \mathcal{E}_{t,q_{j,h}} \rightarrow \mathcal{P}_{j,h}^{q_{j,h}}(S_h)$ where $q_{j,h} : \mathcal{P}^{PrTs}_j$ is defined as follows: for $tPrTs$,

$$q_{j,h}^{ptq} : \frac{\exp \left(\sum_{g \in [n]} \log \frac{1}{p_{j,g}^{ptq} | j_{;g-1}; s_g^q} \right)}{\sum_{t=1}^T \exp \left(\sum_{g \in [n]} \log \frac{1}{p_{j,g}^{ptq} | j_{;g-1}; s_g^q} \right)}$$
- 11: Draw K i.i.d. samples $a_h^1, \dots, a_h^K$ $j_{Prms} q_{j,h} \cdot p$.
- 12: For each $a^1 \in \mathcal{A}_{m-1}$, define $R_{m-1,h}^{pa^1q} : \mathcal{K} \rightarrow \mathbb{R}$ $R_{m-1,h}^{ps_h; pa^k; a^1q}$. Here, we use the oracle O_G to compute $R_{m-1,h}^{ps_h; pa^k; a^1q}$ for each tuple $pa^k; a^1q$.
- 13: For each $j \in [m]$, draw $a_{j,h} \in \mathcal{A}_{j,h}^{q_{j,h}}$.
- 14: Choose the action $a_{m-1,h}$ of player $m-1$ as follows: (Action $a_{m-1,h}$ corresponds to the action selected by the policy π_{m-1} of player $m-1$ defined within the proof of Lemma E.3; this policy is well-defined because the action profiles of all players $i \in [m]$ can be extracted from the lower-order bits of player $m-1$'s reward)

$$a_{m-1,h} : \argmax_{a^1 \in \mathcal{A}_{m-1}} R_{m-1,h}^{pa^1q} : \quad (14)$$
- 15: For each $j \in [m-1]$, let $r_{j,h} = R_{j,h}^{ps_h; pa_{1,h}; \dots; a_{m-1,h}^q}$.
- 16: Each player j constructs $j_{;h}$ by updating $j_{;h-1}$ with $ps_h; a_{j,h}; r_{j,h}^q$.
- 17: if $R_{m-1,h}^{pa_{m-1,h}^q} / 14pm$ $1q \{H$ then
- 18: return $q_h : j_{Prms} q_{j,h}$ as a candidate approximate Nash equilibrium for G .
- 19: end if
- 20: end for
- 21: if the for loop terminates without returning: return fail.

It is evident that this construction takes polynomial time and satisfies $|G|/mn_0/n$. Furthermore, it is clear that a single generative model oracle call for the Markov game G (per Definition 5.1) can be implemented using at most 2 calls to the oracle O_G for the normal-form game G . We will now show by applying the algorithm B to G , we can efficiently (in terms of runtime and oracle calls) compute a $16pm$ $1q$ -approximate Nash equilibrium for the original game G . To do so, we appeal to Algorithm 2.

Algorithm 2 proceeds as follows. First, it calls the algorithm B on the pm $1q$ -player Markov game $G_{p,q}$, using the oracle O_G to simulate B 's calls to the generative model oracle for G . By assumption, the algorithm B returns a sequence $p^{1q}, \dots, p^{Tq}$ of product policies of the form $p^{tq} = p^{tq}_1, \dots, p^{tq}_{pm}$, so that each $p^{tq}_m \in \mathcal{P}^{gen;rnd}_m$ is N -computable, and so that the average $\frac{1}{T} \sum_{t=1}^T p^{tq}$

is an ϵ -CCE of G . Next, Algorithm 2 samples a trajectory from G in which:

- Players $1, \dots, m$ each play according to a policy π^{tq} for an index t of T chosen uniformly at the start of the episode.
- Player $m+1$ plays according to a strategy that, at each step h of H , computes distributions $q_{j,h}$ representing its “belief” of what action each player j of m will play at step h (Line 10), and plays an approximate best response to the product of the strategies $\pi_{j,h}$, j of m (Line 14).

In order to avoid exponential dependence on the number of players m , when computing an approximate best response to $\prod_{j=1}^m q_{j,h}$ we draw $K = \lceil 4 \log m n_0 / \epsilon^2 \rceil$ (for $\{p_h^q\}$) samples from $\prod_{j=1}^m q_{j,h}$ and use these samples to compute the best response. In particular, letting $a_h^k \in A$ denote the k th sampled action profile, we construct a function $R_{m+1,h} : A_{m+1} \rightarrow \mathbb{R}$ in Lines 11 and 14 which, for each $a^{m+1} \in A_{m+1}$, is defined as the average over samples a_h^k of the realized payoffs $R_{m+1,h}(a^{m+1}, a_h^k)$; note that to compute the payoffs for each sample, Algorithm 2 needs only two oracle calls to O_G .

The following lemma, proven in the sequel, gives a correctness guarantee for Algorithm 2.

Lemma E.3 (Correctness of Algorithm 2). Given any m -player n_0 -action normal form game G , if the algorithm B solves the ϵ -Nash-SPARSECCE problem for the game G with T satisfying $T / \exp(n_0^2 m^2 \epsilon^2) \geq 1/3$, then Algorithm 2 outputs a ϵ -approximate Nash equilibrium of G with probability at least $1/3$, and otherwise fails.

The assumption that $T / \exp(n_0^2 m^2 \epsilon^2) \geq 1/3$ from the statement of Theorem E.1 yields that $T / \exp(n_0^2 m^2 \epsilon^2) \geq 1/3$, so Lemma E.3 yields that Algorithm 2 outputs a ϵ -Nash equilibrium of G with probability at least $1/3$ (and otherwise fails). By iterating Algorithm 2 for $\log(1/\epsilon)$ times, we may thus compute a ϵ -Nash equilibrium of G with failure probability 1.

We now analyze the oracle cost and computational cost of Algorithm 2. It takes $2Q$ oracle calls to O_G to simulate the Q generative model oracle calls of B , and therefore, if B runs in time U , then the call to B on Line 5, using oracle calls to O_G to simulate the generative model oracle calls, runs in time OpU . Next, the computations of $q_{j,h}$ (and thus $q_{j,h}$) in Line 10 can be performed in $pnmTn_0^{Op1q}$ time, the computation of $R_{m+1,h} : A_{m+1} \rightarrow \mathbb{R}$ in Line 14 requires time (and oracle calls to O_G) bounded above by $Op|A_{m+1}|Kq/pnm\log(1/\epsilon)$, constructing the actions $a_{j,h}$ (for j of m) in Lines 13 and 14 takes time $pNm n_0^{Op1q}$ (using the fact that the policies $\pi_{j,h}$ are n_0 -computable), and constructing the rewards $r_{j,h}$ on Line 15 requires another $2pm$ oracle calls to O_G . Altogether, Algorithm 2 requires $2Q + pnm\log(1/\epsilon)q^{C_0}$ oracle calls to O_G and, if B runs in time U , then Algorithm 2 takes time $pnmTNU\log(1/\epsilon)q^{C_0}$, for some absolute constant C_0 .

Remark E.4 (Bit complexity of exponential weights updates). In the above proof we have noted that $q_{j,h}$ (as defined in Line 10 of Algorithm 2) can be computed in time $pnmTn_0^{Op1q}$. A detail we do not handle formally is that, since the values of $q_{j,h}$ are in general irrational, only the $pnmTn_0^{Op1q}$ most significant bits of each real number $q_{j,h}$ can be computed in time $pnmTn_0^{Op1q}$. To give a truly polynomial-time implementation of Algorithm 2, one can compute only the $pnmTn_0^{Op1q}$ most significant bits of each distribution $q_{j,h}$, which is sufficient to approximate the true value of $q_{j,h}$ to within $\exp(-pnmTn_0^{Op1q})$ in total variation distance. Since $q_{j,h}$ only influences the subsequent execution of Algorithm 2 via the samples $a_h^1, \dots, a_h^K$ drawn in Line 11, by a union bound, the approximation of $q_{j,h}$ we have described perturbs the execution of the algorithm by at most $OpKq\exp(-pnmTn_0^{Op1q})$ in total variation distance. In particular, the correctness guarantee of Lemma E.3 still holds, with success probability at least $1/3 - \exp(-pnmTn_0^{Op1q}) \geq 1/4$.

It remains to prove Lemma E.3, which is the bulk of the proof of Theorem E.1.

Proof of Lemma E.3. We will establish the following two facts:

1. First, the choices of $a_{m+1,h}$ in Line 14 (i.e., Eq. 14) of Algorithm 2 correspond to a valid policy $\pi_{m+1}^{gen;rnd}$ for player $m+1$ (representing a strategy for deviating from the equilibrium), in that they can be expressed as a function of player $m+1$'s history, $p_{m+1,h}$, at each step h .
2. Second, we will show that, since G is an ϵ -CCE, the strategy $\pi_{m+1}^{gen;rnd}$ cannot lead to a large increase of value for player $m+1$, which will imply that Algorithm 2 must return a Nash equilibrium with high enough probability.

Defining π_i for $i \in \mathcal{P}$. We begin by constructing the policy $\pi_{i,1}$ described; for later use in the proof, it will be convenient to construct a collection of closely related policies $\pi_i^{\text{gen}, \text{rnd}}$ for $i \in \mathcal{P}$, also representing strategies for deviating from the equilibrium π .

Let $i \in \mathcal{P}$ be fixed. For $h \in \mathcal{H}$, the mapping $\pi_{i,h} : H_{i,h} \rightarrow \Sigma_{A_i}$ is defined as follows. Given a history $i_{:h-1}$ $\pi_{i,1}; a_{i,1}; r_{i,1}; \dots; s_{h-1}; a_{i,h-1}; r_{i,h-1} \in \mathcal{H}_{i,h-1}$ (we assume without loss of generality that $i_{:h-1}$ occurs with positive probability under some sequence of general policies) and a current state s_h , we define $\pi_{i,h} : \mathcal{H}_{i,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_i$ through the following process.

1. First, we claim that for all players $j \in \mathcal{P}$, it is possible to extract the trajectory $j_{:h-1}$ from the trajectory $i_{:h-1}$ of player i .
 - (a) Recall that for each $g \leq h$, from the definition in (12) and the function enc_{paq} , the bits following position $3\log_2 s$ of the reward $r_{i,g}$ given to player i at step g of the trajectory $i_{:g-1}$ encode an action profile $a_g \in \mathcal{A}$. Since $i_{:h-1}$ occurs with positive probability, this is precisely the action profile which was played by agents at step g . Note we also use here that by definition of the rewards $R_{j,h}; \text{ps}; a_q$ in (12), the component $R_{j,h}; \text{ps}; a_q$ of the reward only affects the first $2\log_2 s$ bits.
 - (b) For $g \leq h$ and $j \in \mathcal{P}$, define $r_{j,g} := R_{j,g}; \text{ps}; a_g$.
 - (c) For $j \in \mathcal{P}$, write $j_{:h-1} : \pi_{j,1}; a_{j,1}; r_{j,1}; \dots; s_{h-1}; a_{j,h-1}; r_{j,h-1}$; in particular, $j_{:h-1}$ is a deterministic function of $\pi_{i,h-1}; s_h$. (Note that, since $i_{:h-1}$ occurs with positive probability, the history $j_{:h-1}$ observed by player j up to step $h-1$ can be computed from it via Steps (a) and (b)). Going forward, for $g \leq h-1$, we let $j_{:g}$ denote the prefix of $j_{:h-1}$ up to step g .
2. Now, using that player i can compute all players' trajectories, for each $j \in \mathcal{P}$ we define

$$\theta_{j,h} : \mathcal{E}_{t,q,j,h} \rightarrow \mathcal{P}(\mathcal{A}_j); \quad (15)$$

where $q_{j,h} \in \mathcal{P}(\mathcal{T})$ is defined as follows: for $t \in \mathcal{T}$,

$$q_{j,h}(t) = \frac{\exp \left(\sum_{g=1}^h \log \left(\frac{1}{\text{pa}_{j,g} | j_{:g-1}; s_g} \right) \right)}{\sum_{t' \in \mathcal{T}} \exp \left(\sum_{g=1}^h \log \left(\frac{1}{\text{pa}_{j,g} | j_{:g-1}; s_g} \right) \right)} \quad (16)$$

Note that $\theta_{j,h}$ is a random variable which depends on the trajectory $\pi_{j,h-1}; s_h$ (which can be computed from $\pi_{i,h-1}; s_h$). In addition, the definition of $\theta_{j,h}$ (for each $j \in \mathcal{P}$) is exactly as is defined in Line 10 of Algorithm 2.

3. For $i \in \mathcal{P}$, define $\pi_{i,h} : \mathcal{H}_{i,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_i$ as follows:

$$\pi_{i,h} : \mathcal{H}_{i,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_i : \arg \max_{a \in \mathcal{A}_i} \mathbb{E}_{a \sim \pi_{i,h}} \left[R_{i,h}; \text{ps}; a; q \right] \quad (17)$$

For the case $i = 1$, define $\pi_{1,h} : \mathcal{H}_{1,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_1$ (implicitly) to be the following distribution over $a_{1,h} \in \mathcal{A}_1$: draw $a^1, \dots, a^K$ from $\mathcal{P}_{1,h}$, define $\mathbb{P}_{1,h} : \mathcal{A}_1 \rightarrow \mathcal{P}(\mathcal{A}_1)$ by $\mathbb{P}_{1,h}(a^k) = \frac{1}{K} R_{1,h}; \text{ps}; a^k; q$ for $a^k \in \mathcal{A}_1$, and finally set

$$\pi_{1,h} : \mathcal{H}_{1,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_1 : \arg \max_{a \in \mathcal{A}_1} \mathbb{E}_{a \sim \mathbb{P}_{1,h}} [R_{1,h}; \text{ps}; a; q] \quad (18)$$

Note that, for each choice of $\pi_{m-1,h-1}; s_h$, the distribution $\pi_{m-1,h} : \mathcal{H}_{m-1,h-1} \times \mathcal{S} \rightarrow \mathcal{A}_{m-1}$ as defined above coincides with the distribution of the action $a_{m-1,h}$ defined in Eq. 14 in Algorithm 2, when player $m-1$'s history is $\pi_{m-1,h-1}$ and the state at step h is s_h . The following lemma, for use later in the proof, bounds the approximation error incurred in sampling $a_1^1, \dots, a_h^K$ from $\mathcal{P}_{1,h}$.

Lemma E.5. Fix any $\pi_{m-1,h-1}; s_h \in \mathcal{H}_{m-1,h-1}$. With probability at least $1 - \epsilon$ over the draw of $a^1, \dots, a^K$ from $\mathcal{P}_{1,h}$, it holds that for all $a^1 \in \mathcal{A}_1$,

$$\mathbb{E}_{a \sim \pi_{1,h}} [R_{1,h}; \text{ps}; a; q] - \mathbb{E}_{a \sim \mathcal{P}_{1,h}} [R_{1,h}; \text{ps}; a; q] \leq \epsilon / H;$$

which implies in particular that with probability at least 1 over the draw of $a_{m-1;h}^{i-1}$ and $p_{m-1;h}^{i-1}$,

$$\max_{a_{m-1}^{i-1}} E_{a_{j;h} \sim \pi_{j;h}} [R_{m-1;h}(p_{j;h}; p_{a_1;h}; \dots; a_{m-1;h}^{i-1})] \geq \frac{2}{H} E_{a_{j;h} \sim \pi_{j;h}} [R_{m-1;h}(p_{j;h}; p_{a_1;h}; \dots; a_{m-1;h}^{i-1})] \quad (19)$$

It is immediate from our construction above that the following fact holds.

Lemma E.6. The joint distribution of $j_{j;h}$, for $j \in \mathcal{P} \setminus m$ and $h \in \mathcal{H}$ s, as computed by Algorithm 2, coincides with the distribution of $j_{j;h}$ in an episode of G when players follow the policy π_{m-1}^{i-1} .

Analyzing the distributions $\pi_{j;h}$. Fix any $i \in \mathcal{P} \setminus m$. We next prove some facts about the distributions $\pi_{j;h}$ defined above (as a function of $p_{i;h1}; S_h$) in the process of computing $\pi_{i;h}^{i-1}$.

For each $h \in \mathcal{H}$ s, consider any choice of $p_{i;h1}; S_h$. Note that for each $j \in \mathcal{P} \setminus m$, the distributions $\pi_{j;h}$ for $h \in \mathcal{H}$ s may be viewed as an application of Vovk's aggregating algorithm (Proposition D.2) in the following setting: the number of steps (T), in the context of Proposition D.2; note that T has a different meaning in the present proof) horizon is H , the context space is $\mathcal{H}_{j;h1}; S_h$, and the output space is $\mathcal{A}_j$. The expert set is $\{t^{p_{1q}}; \dots; t^{p_{Tq}}\}$ (which has $|J|$ elements), and the experts' predictions on a context $p_{j;h1}; S_h$ are defined via $\pi_{j;h1}; S_h$. Then for each $h \in \mathcal{H}$ s, the distribution $\pi_{j;h}$ is obtained by updating the aggregating algorithm with the context-observation pairs $p_{j;h1}; S_h$ for $h=1, 2, \dots, H$.

In more detail, fix any $t \in \mathcal{T}$ s and $j \in \mathcal{P} \setminus m$ with $j \neq i$. We may apply Proposition D.2 with the number of steps set to H , the set of experts as $\{t^{p_{1q}}; \dots; t^{p_{Tq}}\}$, and contexts and outcomes generated according to the distribution induced by running the policy π_{i-1}^{i-1} in the Markov game G as follows:

- For each $h \in \mathcal{H}$ s, we are given, at steps $h=1, \dots, H$, the actions $a_{k;h}$ rewards $r_{k;h}$ for all agents $k \in \mathcal{P} \setminus m$, as well as the states $s_1; \dots; s_h$.
 - For each $k \in \mathcal{P} \setminus m$, set $k;h1; p_{k;h1}; a_{k;h1}; r_{k;h1}; \dots; s_h; a_{k;h1}; r_{k;h1}$ to be agent k 's history. – The context fed to the aggregation algorithm at step h is $p_{j;h1}; S_h$.
 - The outcome at step h is given by $a_{j;h} \sim \pi_{j;h1}; S_h$; note that this choice satisfies the realizability assumption in Proposition D.2.
 - To aid in generating the next context at step $h+1$, choose $a_{k;h+1}; p_{k;h+1}; S_{h+1}$ for all $k \in \mathcal{P} \setminus m$ such that $j \neq k$ and $a_{i;h+1}; p_{i;h+1}; S_{h+1}$. Then set s_{h+1} to be the next state given the transitions of G and the action profile $a_h(p_{a_1;h}; \dots; p_{a_{m-1;h}})$.

By Proposition D.2, it follows that for any fixed $t \in \mathcal{T}$ s and $j \in \mathcal{P} \setminus m$ with $j \neq i$, under the process described above we have

$$E_{i-1}^{p_{1q}} \left[\sum_{h=1}^H D_{TV}(\pi_{j;h1}; S_h; \pi_{j;h}) \right] \leq \frac{a}{H} \log T \quad (20)$$

Analyzing the value of π_{m-1}^{i-1} . Next, using the development above, we show that if Algorithm 2 successfully computes a Nash equilibrium with constant probability (via π_{m-1}^{i-1}) whenever s is an ϵ -CCE. We first state the following claim, which is proven in the sequel by analyzing the values V_{i-1}^{i-1} for $i \in \mathcal{P} \setminus m$.

Lemma E.7. If s is an ϵ -CCE of G , then it holds that for all $i \in \mathcal{P} \setminus m$,

$$V_{i-1}^{i-1} \geq \frac{a}{m} \log T \cdot \frac{1}{H}$$

Note that in the game G , since for all $h \in \mathcal{H}$ s, $s \in \mathcal{S}$ and $a \in \mathcal{A}$, it holds that $\frac{1}{H} \sum_{h=1}^H R_{j;h}(p_{j;h}; a) \geq \frac{p_{m-1;h}^{i-1}}{H}$ (which holds since in (12), encpaq is multiplied by $\frac{1}{H} 2^{3 \log(1/\epsilon)}$), it follows that $\frac{1}{H} \sum_{h=1}^H V_{j-1}^{i-1} \geq \frac{p_{m-1;h}^{i-1}}{H}$. Thus, by Lemma E.7, we have $V_{m-1}^{i-1} \geq \frac{p_{m-1;h}^{i-1}}{H} \geq \frac{a}{m} \log T \cdot \frac{1}{H}$, and since s is an ϵ -CCE of G it follows that

$$V_{m-1}^{i-1} \geq \frac{a}{m} \log T \cdot \frac{1}{H} \geq \frac{a}{m} \log T \cdot \frac{1}{H} \quad (21)$$

To simplify notation, we will write $\mathbf{p}_h : \mathbf{p}_{j,h} \mathbf{q}_m$ in the below calculations, where we recall that each $\mathbf{q}_{j,h}$ is determined given the history up to step h , $\mathbf{p}_{j,h1}; \mathbf{S}_h \mathbf{q}$, as defined in (15) and (16). An action profile drawn from $\mathbf{q}_h$ is denoted as $\mathbf{a}_h$, with a P.A. We may now write $V_{1^{pm} 1q}$ as follows: —

$$\begin{aligned}
 & V_{m: 1^{pm} 1q}^m \\
 & E_{trTs}^H E_{h1}^{ptq: 1q} E_{p_{j,h}^{ptq} p_{j,h1} h q}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} E_{jPrms}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h m: 1;h p m} \\
 & \quad a: p_{a1;h}; \dots; a_{m: 1;h q} \\
 & \forall E_{trTs}^H E_{h1}^{ptq: 1q} E_{p_{j,h}^{ptq} p_{j,h1} h q}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} E_{jPrms}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} \\
 & \quad a: p_{a1;h}; \dots; a_{m: 1;h q} \\
 & \quad \frac{1}{H} D_{TV} p_{j,h}^{ptq} p_{j,h1} h q; S_h q; q_{j,h} p_{j,h1} h q \\
 & \forall E_{trTs}^H E_{h1}^{ptq: 1q} E_{p_{j,h}^{ptq} p_{j,h1} h q}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} \max_{PA} E_{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h q}^2 \frac{1}{H} \frac{pm}{H} \\
 & \quad \frac{1}{H} D_{TV} p_{j,h}^{ptq} p_{j,h1} h q; S_h q; q_{j,h} p_{j,h1} h q \\
 & \quad \forall \frac{1}{H} E_{trTs}^H E_{h1}^{ptq: 1q} E_{p_{j,h}^{ptq} p_{j,h1} h q}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} \max_{jPrms; a_{j,h} PA_j} E_{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h q}^2 \frac{m}{H} \frac{a}{H \log T^2};
 \end{aligned}$$

where:

- The first inequality follows from the fact that $R_{m: 1;h psh; aqs}$ takes values in $[1;H]$ and the fact that the total variation between product distributions is bounded above by the sum of total variation distances between each of the pairs of component distributions.
- The second inequality follows from the inequality (19) of Lemma E.5.
- The final equality follows from the definition of the rewards in (12) and (13), and by summing (20) over $jPrms$. We remark that the 2 term in the final line comes from the term $H^{2 \log 1/s} \text{encpaq}$ in (12).

Rearranging and using (21) as well as the fact that $\frac{2}{H} \{p_{6Hq}^2\} / (as/1\{2\})$, we get that

$$\begin{aligned}
 & E_{trTs}^H E_{h1}^{ptq: 1q} E_{p_{j,h}^{ptq} p_{j,h1} h q}^{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} \max_{jPrms; a_{j,h} PA_j} E_{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h q}^2 \\
 & \quad \frac{1}{2H pm: 1q pm: 1qm} \frac{a}{H \log T^2} \frac{1}{3H}:
 \end{aligned}$$

Since $\mathbf{p}_h$ is a product distribution a.s., we have that

$$\max_{jPrms; a_{j,h} PA_j} E_{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h q}^2 \geq 0:$$

Therefore, by Markov's inequality, with probability at least $1/2$ over the choice of t_{rTs} and the trajectories $\mathbf{p}_{j,h1}; \mathbf{S}_h \mathbf{q}_{m: 1^{pm} 1q}$ for $jPrms$ (which collectively determine $\mathbf{q}_h$), there is some $hPrHs$ so that

$$\max_{jPrms; a_{j,h} PA_j} E_{a_{j,h}^{ptq} p_{j,h} p_{j,h1} h q} R_{m: 1;h psh; aqs} a_{m: 1;h q}^2 \geq 10 pm: 1q 2 pm: 1qm \log p_{Tq} (H/12 pm: 1q); \quad (22)$$

where the final inequality follows as long as $H^2 \leq m^2 \log T$, i.e., $T / \exp^{-H} \frac{2}{m^2}$, which holds since $H \leq n_0$ and we have assumed that $T / \exp^2 n_0 \{m^2 q$.

Definition E.9 (Alternative construction to Definition E.2). Given an m -player, n_0 -action normal-form game G , we define the game G^1pGq as follows.

- The horizon of G is $H n_0$.
- Let A_{n_0} . The action spaces of agents $1, 2, \dots, m$ are given by A_1, A_m rAs. The action space of agent $m-1$ is A_{m-1}

$$tpj; a_j q : j \in P_{rms}; a_j \in A_j u;$$

so that $|A_{m-1}| = |A_m|/n$.

We write $A_{j=1}^m A_j$ to denote the joint action space of the first m agents, and $A_{j=1}^m A_j$ to denote the joint action space of all agents. Then $|A| = |A^m p m A q m A^m|^{1/n}$.

- The state space S is defined as follows. There are $|A|$ states, one for each action tuple $a \in A$. For each $a \in A$, we denote the corresponding state by s_a .
- For all $h \in P_{rH}$ s, the reward to agent $j \in P_{rm}$ is given action profile $a \in A$ at any state $s \in S$ is as follows: writing $a_{m-1} p j^1; a^1 q$,

$$R_{j;h} p s; a q : \begin{cases} \text{if } p M_j q_{a_1, \dots, a_m} p M_j q_{a^1, \dots, a^1} : j \in R_{tj}; m-1 u \\ \text{if } p M_j q_{a^1, \dots, a^1} p M_j q_{a_1, \dots, a_m} : j \in j^1 \\ \text{if } p M_j q_{a^1, \dots, a^1} p M_j q_{a_1, \dots, a_m} : j \in m-1 \end{cases} \quad (23)$$

- At each step $h \in P_{rH}$ s, if action profile $a \in A$ is taken, the game transitions to the state s_a .

Note that the number of states of G^1pGq is equal to $|A|^{1/n} m^{m-1} n_0$ and so $|G^1pGq| = m^{m-1} n_0$. As a result, if we were to use the game G^1pGq in place of $GpGq$ in the proof of Theorem E.1, we would need to define $n_0 : \ln^{1/(m-1)} \{m\}$ to ensure that $|G^1pGq|/n$, and so the condition $T \exp^{2 \ln \{m\}^2 q}$ would be replaced by $T \exp^{2 \ln^{1/(m-1)} \{m\}^2 q}$. This would only lead to a small quantitative degradation in the statement of Theorem 4.3, with the condition in the statement replaced by $T \exp^{c^2 n^{1/3} q}$ for some constant $c > 0$. However, it would render the statement of Theorem 5.2 essentially vacuous. For this reason, we opt to go with the approach of Definition E.2 as opposed to Definition E.9.

We expect that the construction of Definition E.2 can nevertheless still be modified to use $O(\log \{q\})$ bits to express each reward in the Markov game G . In particular, one could introduce stochastic transitions to encode in the state of the Markov game a small number of random bits of the full action profile played at each step. We leave such an approach for future work.

F. Equivalence between $p_j^{\text{gen;rnd}}$ and $p_j^{\text{gen;det}}$

In this section we consider an alternate definition of the space $p_j^{\text{gen;rnd}}$ of randomized general policies of player i , and show that it is equivalent to the one we gave in Section 2.

In particular, suppose we were to define a randomized general policy of agent i as a distribution over deterministic general policies of agent i : we write $p_i^{\text{gen;rnd}} : p_j^{\text{gen;det}}$ to denote the space of such distributions. Moreover, write $p_j^{\text{gen;rnd}} : p_j^{\text{gen;rnd}} p_j^{\text{gen;det}}$ to denote the space of product distributions over agents' deterministic policies. Our goal in this section is to show that policies in $p_j^{\text{gen;rnd}}$ are equivalent to those in $p_j^{\text{gen;rnd}}$ in the following sense: there is an embedding map $\text{Emb} : p_j^{\text{gen;rnd}} \rightarrow p_j^{\text{gen;rnd}}$, not depending on the Markov game, so that the distribution of a trajectory drawn from any $p \in p_j^{\text{gen;rnd}}$, for any Markov game, is the same as the distribution of a trajectory drawn from $\text{Emb} p$ (Fact F.2). Furthermore, Emb is surjective in the following sense: any policy $p \in p_j^{\text{gen;rnd}}$ produces trajectories that are distributed identically to those of $\text{Emb} p$ (and thus of p), for some $p \in p_j^{\text{gen;rnd}}$ (Fact F.3). In Definition F.1 below, we define Emb .

Definition F.1. For $j \in P_{rms}$ and $p_j^{\text{gen;rnd}}$, define $\text{Emb}_j p_j q p_j^{\text{gen;rnd}} p_j^{\text{gen;det}}$ to put the following amount of mass on each $p_j^{\text{gen;det}}$:

$$p \text{Emb}_j p_j q p_j q : \prod_{h=1}^H p_{j;h} p_{j;h;1} s_{h;1} q |_{j;h;1} s_{h;1} q \quad (24)$$

Furthermore, for $p_1, \dots, p_m \in \mathcal{P}^{\text{gen}, \text{rnd}}$, define $\text{Emb} p_1 q, \dots, \text{Emb} p_m q$.

Note that, in the special case that $p_j \in \mathcal{P}^{\text{gen}, \text{det}}_j$, $\text{Emb} p_j q$ is the point mass on j .

Fact F.2 (Embedding equivalence). Fix a m -player Markov game G and, arbitrary policies $p_j \in \mathcal{P}^{\text{gen}, \text{rnd}}_j$. Then a trajectory drawn from the product policy $p_1, \dots, p_m \in \mathcal{P}^{\text{gen}, \text{rnd}}$ is distributed identically to a trajectory drawn from $\text{Emb} p_1 q, \dots, \text{Emb} p_m q$.

The proof of Fact F.2 is provided in Section F.1. Next, we show that the mapping Emb is surjective in the following sense:

Fact F.3 (Right inverse of Emb). There is a mapping $\text{Fac} : \mathcal{P}^{\text{gen}, \text{rnd}} \rightarrow \tilde{\mathcal{N}}^{\text{gen}, \text{rnd}}$ so that for any Markov game G and any $r \in \mathcal{P}^{\text{gen}, \text{rnd}}$, the distribution of a trajectory drawn from r is identical to the distribution of a trajectory drawn from $\text{Emb} \text{Fac} r q$.

We will write $\text{Fac} p_1, \dots, p_m q : p_{\text{Fac} 1} p_1 q, \dots, p_{\text{Fac} m} p_m q$. Fact F.3 states that the policy $\text{Fac} r q$ maps, under Emb , to a policy in $\mathcal{P}^{\text{gen}, \text{rnd}}$ which is equivalent to r (in the sense that their trajectories are identically distributed for any Markov game).

An important consequence of Fact F.2 is that the expected reward (i.e., value) under any $p \in \mathcal{P}^{\text{gen}, \text{rnd}}$ is the same as that of $\text{Emb} p q$. Thus

given a Markov game, the induced normal-form game in which the players' pure action sets are $\mathcal{A}_1^{\text{gen}, \text{rnd}}, \dots, \mathcal{A}_m^{\text{gen}, \text{rnd}}$ is equivalent to the normal-form game in which the players' pure action sets are $\mathcal{A}_1^{\text{gen}, \text{det}}, \dots, \mathcal{A}_m^{\text{gen}, \text{det}}$, in the following sense: for any mixed strategy in the former, namely a product distributional policy $P \in \mathcal{P}^{\text{gen}, \text{rnd}}$, the policy $E_P \text{Emb} p q \in \mathcal{P}^{\text{gen}, \text{det}}$ is a mixed strategy in the latter which gives each player the same value as under P . (Note that $E_P \text{Emb} p q$ is indeed a product distribution since P is a product distribution and Emb factors into individual coordinates.) Furthermore, by Fact F.3, any distributional policy in $\mathcal{P}^{\text{gen}, \text{rnd}}$ arises in this manner, for some $P \in \mathcal{P}^{\text{gen}, \text{rnd}}$; in fact, P may be chosen to place all its mass on a single $p \in \mathcal{P}^{\text{gen}, \text{rnd}}$. Since Emb factors into individual coordinates, it follows that Emb yields a one-to-one mapping between the coarse correlated equilibria (or any other notion of equilibria, e.g., Nash equilibria or correlated equilibria) of these two normal-form games.

F.1. Proofs of the equivalence

Proof of Fact F.2. Consider any trajectory $p s_1; a_1; r_1; \dots; s_H; a_H; r_H q$ consisting of a sequence of H states and actions and rewards for each of the m agents. Assume that $r_{i,h} \in \mathcal{R}_{i,h} p s; a_h q$ for all i, h (as otherwise r has probability 0 under any policy). Write:

$$p : \prod_{h=1}^H p_{s_h} p_{s_h-1} | s_h; a_h q : h_1$$

Then the probability of observing γ under p is

$$p(\gamma) = \prod_{h=1}^H \prod_{j=1}^m p_{j,h} p_{j,h-1} | j; h_1; s_h q \quad (25)$$

where, per usual, $j, h_1 : p s_1; a_j; 1; r_j; 1; \dots; s_{h_1}; a_j; h_1; r_j; h_1 q$. Write $p_1, \dots, p_m \in \mathcal{P}^{\text{gen}, \text{rnd}}$. The probability of observing γ under p is

$$p(\gamma) = \prod_{j=1}^m p_{j,h_1} p_{j,h_1-1} | j; h_1; s_{h_1} q a_{j,h_1} q \quad (26)$$

It is now straightforward to see from the definition of $p_j q$ in (24) that the quantities in (25) and (26) are equal. $\square$

Proof of Fact F.3. Fix a policy $r_j \in \mathcal{P}^{\text{gen}, \text{rnd}}_j$. We define $\text{Fac}_j p_j q$ to be the policy $p_j \in \mathcal{P}^{\text{gen}, \text{rnd}}$, which is defined as follows: for $j, h_1 : p s_1; 1; a_j; 1; r_j; 1; \dots; s_{j,h_1}; a_j; h_1; r_j; h_1 q$, we have, for $a_j; h : p A_j$,

$$p_{j,h_1} p_{j,h_1-1} | j; h_1; s_{j,h_1} q a_{j,h} q = \frac{r_j^t p_j^{\text{gen}, \text{det}} : p_j; s_{j,h_1} q a_{j,h} q @g/hu}{r_j^t p_j^{\text{gen}, \text{det}} : p_j; s_{j,h_1} q a_{j,h} q @g/h1u}$$

If the denominator of the above expression is 0, then $p_{j,h1;s_h q}$ is defined to be an arbitrary distribution on $p_{A_j q}$. (For concreteness, let us say that it puts all its mass on a fixed action in A_j .) Furthermore, for $r \in P^{\text{gen}; \text{rnd}}_t$ define $\text{Fac}_{p,q} : p \mapsto \text{Fac}_1(p_1 q; \cdot; \cdot; \cdot; \text{Fac}_m(p_m q q) P_t^{\text{gen}; \text{rnd}})$.

Next, fix any $r = (r_1; \dots; r_m) \in P^{\text{gen}; \text{rnd}}_1$. Let $r \mapsto \text{Fac}_{p,q}$. By Fact F.2, it suffices to show that the distribution of trajectories under π is the same as the distribution of trajectories drawn from π .

So consider any trajectory $(s_1; a_1; r_1; \dots; s_H; a_H; r_H) q$ consisting of a sequence of H states and actions and rewards for each of the m agents. Assume that $r_{i,h} \in R_{i,h}; p_{s; a_h} q$ for all i, h (as otherwise π has probability 0 under any policy). Write:

$$p : \prod_{h=1}^H p_{s_h-1 | s_h; a_h q; h1}$$

Then the probability of observing π under π is

$$p(\pi) = \prod_{j=1}^m \prod_{h=1}^H p_{a_j; h | j; h1; s_h q; h1j1} \cdot \frac{\prod_{j=1}^m \prod_{h=1}^H r_j \cdot t_j \cdot p_e^{\text{gen}; d, t} : j p_{j; g; s_g q a_{j; g}} @g / hu}{\prod_{j=1}^m \prod_{h=1}^H t_j \cdot p_e^{\text{gen}; d, t} : j p_{j; g; s_g q a_{j; g}} @g / h1u p''_j t_j} \cdot p_e^{\text{gen}; d, t} : j p_{j; g; s_g q a_{j; g}} @g / Hu ;$$

which is equal to the probability of observing π under r . □