
Learning Pareto-Optimal Policies in Low-Rank Cooperative Markov Games

Abhimanyu Dubey¹ Alex Pentland¹

Abstract

We study cooperative multi-agent reinforcement learning in episodic Markov games with n agents. While the global state and action spaces typically grow exponentially in n in this setting, in this paper, we introduce a novel framework that combines function approximation and a graphical dependence structure that restricts the decision space to $o(dn)$ for each agent, where d is the ambient dimensionality of the problem. We present a multi-agent value iteration algorithm that, under mild assumptions, provably recovers the set of Pareto-optimal policies, with finite-sample guarantees on the incurred regret. Furthermore, we demonstrate that our algorithm is *no-regret* even when there are only $\mathcal{O}(1)$ episodes with communication, providing a scalable and provably no-regret algorithm for multi-agent reinforcement learning with function approximation. Our work provides a tractable approach to multi-agent decision-making that is provably efficient and amenable to large-scale collaborative systems.

Cooperative multi-agent reinforcement learning (MARL) is becoming increasingly prevalent in applications such as robotics (Ding et al., 2020), power grid management (Yu et al., 2014), traffic control (Bazzan, 2009) and team games (Zhao et al., 2019). In this setting, a group of n agents, each with their own state and action spaces, interact simultaneously to maximize their cumulative rewards. The foundational challenge in these multi-agent environments (also known as multi-agent MDPs (Boutilier, 1996) or *cooperative* Markov games (Shapley, 1953)) is that despite having small individual state and action spaces, the joint space grows exponentially in n , introducing a curse of dimensionality that makes standard approaches intractable. Furthermore, designing a *globally* optimal policy is difficult owing to communication and computational constraints.

¹Media Lab and Institute for Data, Systems and Society, Massachusetts Institute of Technology. Correspondence to: Abhimanyu Dubey <dubeya@mit.edu>.

In single-agent *tabular* reinforcement learning (RL), algorithms exist that provably incur a regret over T episodes that scales as $\tilde{\mathcal{O}}(H\sqrt{|\mathcal{S}||\mathcal{A}|T})$ ¹, where $\mathcal{S}$ and $\mathcal{A}$ are the state and action spaces, respectively, and H denotes the length of each episode. Such settings normally are agnostic to the low-rank structure present in many environments, and recent work (Jin et al., 2020; Wang et al., 2020; Yang et al., 2020) has explored a *low-rank* linear formulation of MDPs, where the transition kernels and reward functions are assumed to be linear functions of a known d -dimensional feature of the state and action. Under this assumption, algorithms have been proposed that provably incur a regret of $\tilde{\mathcal{O}}(H^2\sqrt{d^3T})$, and when $d^3 \ll |\mathcal{S}||\mathcal{A}|$, this low-rank structure can be exploited effectively in many environments. Concurrently, the multi-agent RL literature has focused on establishing *local dependence structures* (Qu and Li, 2019; Qu et al., 2020), where the dynamics are assumed to be a function of only a subset of agents, effectively reducing the dependence on n from exponential to polynomial, providing *localized* algorithms with provable asymptotic convergence. This complements the approaches based on *factored* MDPs (Guestrin et al., 2002; 2001; Roth et al., 2007), where the rewards incurred by any agent is decomposed into a sum of several latent reward functions.

In this paper, we unify these two perspectives of low-rank function approximation and local dependence structures to present a scalable, provably efficient approach to cooperative multi-agent reinforcement learning. Specifically, we seek to answer the following open question - *can we design tractable, scalable and provably efficient cooperative MARL algorithms with function approximation?*

Contributions. We answer the question affirmatively under mild environmental conditions. First, we present a characterization of cooperative Markov games based on a graphical influence model, where a known (connected, undirected) graph G determines the structure of influence (i.e., an edge (i, j) exists in G if agents i and j influence each other). We extend the single-agent low-rank MDP (Jin et al., 2020) environment to multi-player MDPs and provide a set of weak assumptions, titled *clique-dominance*, that are sufficient to reduce the effective size of the joint state-action space from $\mathcal{O}((|\mathcal{S}||\mathcal{A}|)^n)$ to $o(dn)$, where d is the dimensionality of the

¹The $\tilde{\mathcal{O}}$ notation ignores polylogarithmic factors.

approximating function class. Next, we generalize the cooperative MARL objective from maximizing total reward to a broader class of *Pareto-optimal* policies, and characterize conditions in which this class of policies can be efficiently recovered by the *method of scalarization* (Knowles, 2006) by minimizing *Bayes regret*. Thirdly, we introduce **MultOVI**, a decentralized *vector-valued* optimistic value iteration algorithm that even under partial observability conditions, obtains a cumulative *Bayes regret* of $\tilde{O}(\theta(G)H^2\sqrt{d^3T})$ over T episodes, where $\theta(G)$ denotes the *clique covering number* of G . **MultOVI** runs in polynomial time and only requires a communication budget of $o(nd^2 \log T)$ rounds per agent, which can be much smaller for sparse G . This ensures that **MultOVI** is scalable to very large environments and adapts to the sparsity of influence as well. Furthermore, in contrast to the existing work in cooperative MARL that converges to the global optimal policy (i.e., maximizing total reward), **MultOVI** can, under mild conditions, recover any subset of policies in the Pareto frontier, additionally enabling *adaptive* load-balancing (Schaerf et al., 1994). Moreover, a direct corollary of our analysis also provides the first no-regret algorithm for multi-objective RL (Mossalam et al., 2016) with function approximation.

Organization. Section 2 presents assumptions about the Markov game considered. Section 3 presents our performance objective and recovery guarantees. Section 4 presents our algorithm and associated regret upper and lower bounds. We defer full proofs, a survey of related work, additional remarks, and experiments to the Appendix for brevity.

2. Preliminaries

Notation. We denote vectors by lowercase solid letters, i.e., $\mathbf{x}$, matrices by uppercase solid letters $\mathbf{X}$, and sets by calligraphic letters, i.e., $\mathcal{X}$, the ellipsoid norm of a vector $\mathbf{x}$ as $\|\mathbf{x}\|_{\mathbf{S}} = \sqrt{\mathbf{x}^\top \mathbf{S} \mathbf{x}}$ for some matrix $\mathbf{S}$. We denote the interval $a, \dots, b$ for $b \geq a$ by $[a, b]$ and as $[b]$ when $a = 1$.

Cooperative Markov Games. We consider the simultaneous-move Markov game (Xie et al., 2020), which is an extension of an MDP to multiple agents, and is also known as a multi-agent MDP (Boutillier, 1996). A Markov game (MG) can be formally described as $\text{MG}(\mathcal{S}, \mathcal{M}, \mathcal{A}, H, \mathbb{P}, \mathbf{R})$, where the set of agents $\mathcal{M}$ is finite and countable with size n , the state and action spaces are factorized as $\mathcal{S} = \mathcal{S}_1 \times \mathcal{S}_2 \times \dots \times \mathcal{S}_n$ and $\mathcal{A} = \mathcal{A}_1 \times \mathcal{A}_2 \times \dots \times \mathcal{A}_n$, where $\mathcal{S}_\nu$ and $\mathcal{A}_\nu$ denote the individual state and action space for agent ν respectively. The transition matrix $\mathbb{P} = \{\mathbb{P}_h\}_{h \in [H]}$, $\mathbb{P}_h : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$ determines how the joint state evolves given an existing joint state-action, and the reward function $\mathbf{R} = \{\mathbf{r}_h\}_{h=1}^H$, $\mathbf{r}_h = \{r_{\nu,h}\}_{\nu=1}^n, r_{\nu,h} : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ denotes the reward obtained by each agent ν in the MG. We further denote, for any subset $\mathcal{Z} \subseteq \mathcal{M}$ of agents, the marginal transition

probability for the subset as $\mathbb{P}^{\mathcal{Z}} = \{\mathbb{P}_h^{\mathcal{Z}}\}_{h=1}^H$ such that $\mathbb{P}_h^{\mathcal{Z}} : \mathcal{S} \times \mathcal{A} \times (\prod_{i \in \mathcal{Z}} \mathcal{S}_i) \rightarrow [0, 1]$. Next, we consider a graphical model of influence in order to remove the exponential dependence on n (a generalization of prior work, e.g., Qu and Li (2019); Qu et al. (2020)), summarized below.

Assumption 1 (Local Influence). *Let $G = (\mathcal{M}, \mathcal{E})$ denote an undirected network of influence between agents in $\mathcal{M}$, i.e., $\mathcal{E}$ contains an edge (i, j) if the reward of agent i is a function of agent j (and vice-versa), and let $\mathcal{N}^+(\nu)$ denote, for any agent ν its neighborhood in G (including itself). Alternatively, this implies that the reward for any agent ν obeys $r_{\nu,h} = r_{\nu,h}(\tilde{\mathbf{x}}_\nu, \tilde{\mathbf{a}}_\nu)$ where $\tilde{\mathbf{x}}_\nu = \{x_j\}_{j \in \mathcal{N}^+(\nu)}$ and $\tilde{\mathbf{a}}_\nu = \{a_j\}_{j \in \mathcal{N}^+(\nu)}$ denote the local state-action for ν .*

Ref. Remark 4 in Appendix B. We are not quite yet equipped with a feasible learning model. This is evident as Assumption 1 in the worst case still leads to an exponential dependence on n , and combinatorial state-action spaces have been known to be intractable (Blondel and Tsitsiklis, 1999; Papadimitriou and Tsitsiklis, 1987). As a consequence, recent work has suggested additional conditions bounding the strength of interactions between agents to develop efficient policies (Qu and Li, 2019; Qu et al., 2020). We now describe a similar assumption to characterize dynamics.

Definition 1 (Clique covering number). *A k -clique cover $\mathcal{C} = \{C_1, \dots, C_k\}$ of any graph G is a partition of G into k non-overlapping subgraphs such that each subgraph $C_i, i \in [k]$ is strongly connected. The clique covering number $\theta(G)$ is the size of the smallest clique covering $\mathcal{C}^*$ of G .*

Assumption 2 (Clique-Dominant Dynamics). *For the network G defined in Assumption 1 let $\mathcal{C} = \cup_{l \in [k]} C_l$ be a known k -clique cover. For any $V \subseteq G$ and joint state-action pair $(\mathbf{x}, \mathbf{a})$, let $\mathbf{x}_V = \{\cup_{i \in V} x_i\}$ and $\mathbf{a}_V = \{\cup_{i \in V} a_i\}$ denote the joint state and action of all agents in V , and $\bar{\mathbf{x}}_V, \bar{\mathbf{a}}_V$ denote the joint state and action of all agents not in V . We assume that for each $C \in \mathcal{C}$ and $h \in [H]$ there exists an unknown kernel $\tilde{\mathbb{P}}_h^C : (\prod_{i \in C} \mathcal{S}_i) \times (\prod_{i \in C} \mathcal{A}_i) \times (\prod_{i \in C} \mathcal{S}_i) \rightarrow [0, 1]$, unknown functions $\{\tilde{r}_{\nu,h}\}_{\nu=1}^n$ and a known nondecreasing function $\varepsilon(\cdot) : [1, n] \rightarrow [0, 1]$ such that for any joint state-action $(\mathbf{x}, \mathbf{a}) = (\{\mathbf{x}_C, \bar{\mathbf{x}}_C\}, \{\mathbf{a}_C, \bar{\mathbf{a}}_C\})$, we have for any $\nu \in C$,*

$$|r_{\nu,h}(\mathbf{x}, \mathbf{a}) - \tilde{r}_{\nu,h}(\mathbf{x}_C, \mathbf{a}_C)| \leq \varepsilon(k) \text{ and },$$

$$\left\| \mathbb{P}_h^C(\cdot | \mathbf{x}, \mathbf{a}) - \tilde{\mathbb{P}}_h^C(\cdot | \mathbf{x}_C, \mathbf{a}_C) \right\|_{\text{TV}} \leq \varepsilon(k).$$

Remark 1 (Feasibility of Clique-Dominance). Assumption 2 assumes that if any group of agents C is strongly-connected (i.e., all influence each other), their joint information suffices to “approximately” explain the individual reward and joint marginal transition dynamics up to a factor ε for all agents in C . Naturally, for a smaller clique-covering, a lower approximation error ε can be expected. Similar assumptions for local regularity have been made in prior

work: Qu and Li (2019) introduce the (c, ρ) -exponential decay property, see Remarks 5, 6 in Appendix B.

Setting. The game proceeds as follows. In each episode $t \in [T]$ each agent ν fixes a policy $\pi_\nu(t) = \{\pi_\nu^h(t)\}_{h=1}^H$ in a (joint) initial state $\mathbf{x}_1(t) = \{x_\nu^1(t)\}_{\nu=1}^n$ picked arbitrarily by the environment. For each step $h \in [H]$ of the episode, each agent observes the local state $\tilde{\mathbf{x}}_\nu^h(t)$, selects an individual action $a_\nu^h(t) \sim \pi_\nu^h(\cdot | \tilde{\mathbf{x}}_\nu^h(t))$ (collectively the joint action $\mathbf{a}_h(t) = \{a_\nu^h(t)\}_{\nu=1}^n$), and obtains a reward $r_\nu^h(\tilde{\mathbf{x}}_\nu^h(t), \tilde{\mathbf{a}}_\nu^h(t))$ (collectively the joint reward $\mathbf{r}_h(\mathbf{x}_h(t), \mathbf{a}_h(t)) = \{r_\nu^h(\tilde{\mathbf{x}}_\nu^h(t), \tilde{\mathbf{a}}_\nu^h(t))\}_{\nu=1}^n$). All agents transition subsequently to a new joint state $\mathbf{x}_{h+1}(t) = \{x_\nu^{h+1}(t)\}_{\nu=1}^n$ sampled according to $\mathbb{P}_h(\cdot | \mathbf{x}_h(t), \mathbf{a}_h(t))$. The episode terminates at step $H + 1$ where all agents receive no reward. The agents can then (optionally) communicate by sharing messages to neighbors in G after each episode. Let $\pi = \{\pi_\nu\}_{\nu=1}^n$ denote a joint policy for all n agents. We can define the vector-valued value function over all joint states $\mathbf{x} \in \mathcal{S}$ for a policy π and step h as $\mathbf{V}_h^\pi(\mathbf{x}) \triangleq \mathbb{E}_\pi \left[\sum_{i=h}^H \mathbf{r}_i(\mathbf{x}_i, \mathbf{a}_i) \mid \mathbf{x}_h = \mathbf{x} \right]$. Analogously, we define the vector-valued Q -function for a policy π and any $\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}, h \in [H]$, $\mathbf{Q}_h^\pi(\mathbf{x}, \mathbf{a}) \triangleq r_h(\mathbf{x}, \mathbf{a}) + \mathbb{E}_\pi \left[\sum_{i=h+1}^H \mathbf{r}_i(\mathbf{x}_i, \mathbf{a}_i) \mid \mathbf{x}_h = \mathbf{x}, \mathbf{a}_h = \mathbf{a} \right]$.

3. Beyond Team-Average Rewards

Cooperative MARL focuses primarily on global objectives, most commonly that of *team-average* reward. While this objective is indeed valid in many environments, we aim to recover the richer class of *Pareto-optimal* objectives (Buchanan, 1962). Formally,

Definition 2 (Pareto optimality, Paria et al. (2020)). A policy π Pareto-dominates policy π' iff $\mathbf{V}_1^\pi(\mathbf{x}) \succeq \mathbf{V}_1^{\pi'}(\mathbf{x}) \forall \mathbf{x} \in \mathcal{S}$. A policy is Pareto-optimal if it is not Pareto-dominated by another policy. We denote the set of all policies by Π , and the Pareto-optimal policies by Π^* .

It is evident that *joint* policies that maximize any agent's individual reward as well as the average reward are all elements of Π^* . The motivation to consider recovering the Pareto frontier is indeed derived from applications, e.g., in EV charging (Marinescu et al., 2014), smart grids (Chiu et al., 2019), and workflow optimization (Wang et al., 2019).

Random Scalarizations. To recover Π^* , our approach is to utilize the method of *random scalarizations* (Knowles, 2006). The key idea in the method of scalarization is to observe that if the Pareto frontier Π^* is convex, then there is a bijective mapping of each policy in Π^* to the optimal policy of a *scalarized* MDP. Consider a *scalarization function* $\mathbf{s}_v(\mathbf{x}) = \mathbf{v}^\top \mathbf{x} : \mathbb{R}^n \rightarrow \mathbb{R}$ parameterized by $\mathbf{v}$ belonging to the set $\Upsilon \subseteq \Delta^n$ (unit simplex in n dimensions). We then have the *scalarized* value function $V_{v,h}^\pi(\mathbf{x}) : \mathcal{S} \rightarrow \mathbb{R}$ and

Q -function $Q_{v,h}^\pi : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ for some joint policy π as $V_{v,h}^\pi(\mathbf{x}) \triangleq \mathbf{s}_v(\mathbf{V}_h^\pi(\mathbf{x})) = \mathbf{v}^\top \mathbf{V}_h^\pi(\mathbf{x})$, and $Q_{v,h}^\pi(\mathbf{x}, \mathbf{a}) \triangleq \mathbf{s}_v(\mathbf{Q}_h^\pi(\mathbf{x}, \mathbf{a})) = \mathbf{v}^\top \mathbf{Q}_h^\pi(\mathbf{x}, \mathbf{a})$. Since both $\mathcal{A} = \prod_i \mathcal{A}_i$ and H are finite, there exists an optimal multi-agent policy for any fixed scalarization $\mathbf{v}$, which gives the value $V_{v,h}^* = \sup_{\pi \in \Pi} V_{v,h}^\pi(\mathbf{x})$ for all $\mathbf{x} \in \mathcal{S}$ and $h \in [H]$. This policy coincides with the optimal policy for an MDP over the space $\mathcal{S} \times \mathcal{A}$, defined as follows.

Proposition 1. For the scalarized value function given above, the Bellman optimality conditions are given as, for all $h \in [H], \mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}, \mathbf{v} \in \Upsilon$, $Q_{v,h}^*(\mathbf{x}, \mathbf{a}) = \mathbf{s}_v(\mathbf{r}_h(\mathbf{x}, \mathbf{a}) + \mathbb{P}_h V_{v,h}^*(\mathbf{x}, \mathbf{a}))$, $V_{v,h}^*(\mathbf{x}) = \max_{\mathbf{a} \in \mathcal{A}} Q_{v,h}^*(\mathbf{x}, \mathbf{a})$, and $V_{v,H+1}^*(\mathbf{x}) = 0$.

The optimal policy for any fixed $\mathbf{v}$ is given by the greedy policy with respect to the Bellman-optimal scalarized Q -values. We denote this (unique) optimal policy by π_v^* . The next result claims that by “projecting” a cooperative Markov game to an MDP via scalarization, one can recover a policy on the Pareto frontier. Indeed, when the set Π^* is convex, then the set of policies $\Pi_\Upsilon^* = \{\pi_v^* | \mathbf{v} \in \Delta^n\}$ spans Π^* , and one can recover Π^* by simply learning Π_Υ^* .

Theorem 1. For any Markov game with finite $\mathcal{A}$ and H , $\Pi_\Upsilon^* \subseteq \Pi^*$. If Π^* is convex, $\Pi_\Upsilon^* = \Pi^*$.

Remark 2 (Limits of Scalarization). This approach suffers from the drawback that convexity assumptions on the scalarization function limit algorithms to only recover policies within the convex regions of Π^* (Vamplew et al., 2008), which is exact when Π^* is convex. Subsequently, our algorithm is limited in this sense as it relies on scalarizations, however, we leave the extension to non-convex regions as future work, and assume Π^* to be convex for simplicity.

Bayes Regret. As mentioned earlier, in many applications, we may require learning policies that prioritize an agent over others. Hence, we consider a general notion of *Bayes regret*. Our objective is to approximate Π^* by learning a set of T policies $\hat{\Pi}_T$ that minimize the Bayes regret, given by,

$$\mathfrak{R}_B(T) \triangleq \mathbb{E}_{\mathbf{v} \sim p_\Upsilon} \left[\max_{\mathbf{x} \in \mathcal{S}} \left[V_{v,1}^*(\mathbf{x}) - \max_{\pi \in \hat{\Pi}_T} V_{v,1}^\pi(\mathbf{x}) \right] \right]. \quad (1)$$

Here p_Υ is a distribution over Υ that characterizes the nature of policies we wish to recover. For example, if we set p_Υ as the uniform distribution over Δ^n then we can expect the policies recovered to prioritize all agents equally². The advantage of minimizing Bayes regret can be understood as follows. For any $\mathbf{v} \in \Upsilon$, if $\pi_v^* \in \hat{\Pi}_T$, then the regret incurred is 0. Hence, by collecting policies that minimize Bayes regret, we are effectively searching for policies

²One may consider minimizing the regret for a fixed scalarization $\mathbf{v}' = \mathbb{E}_{p_\Upsilon} \mathbf{v}$, however, that will also recover only one policy in Π^* , whereas we desire to capture regions of Π^* .

that span dense regions of Π^* (assuming convexity, see Remark 2). Consider now the *cumulative regret*:

$$\mathfrak{R}_C(T) \triangleq \sum_{t \in [T]} \mathbb{E}_{\mathbf{v}_t \sim p_{\Upsilon}} \left[\max_{\mathbf{x} \in \mathcal{S}} [V_{\mathbf{v}_t, 1}^*(\mathbf{x}) - V_{\mathbf{v}_t, 1}^{\pi_t}(\mathbf{x})] \right]. \quad (2)$$

Where $\mathbf{v}_1, \dots, \mathbf{v}_T \sim p_{\Upsilon}$ are sampled i.i.d. from p_{Υ} , and π_t refers to the joint policy at episode t . Under suitable conditions on $\mathfrak{s}$ and Υ , we can bound the two quantities.

Proposition 2. *For $\mathfrak{s}$ that is Lipschitz and bounded Υ , we have that $\mathfrak{R}_B(T) \leq \frac{1}{T} \mathfrak{R}_C(T) + o(1)$.*

4. An Efficient Algorithm

We now present our algorithm **MultOVI** (Multiagent Optimistic Value Iteration) that provides a polynomial sample complexity for environments with low-rank structure.

Assumption 3 (Clique-dominant Linear Markov Game). *Let $\mathcal{C}$ be a clique covering of G , and for any clique $C \in \mathcal{C}$, let $\mathcal{S}_C = \prod_{\nu \in C} \mathcal{S}_{\nu}$ and $\mathcal{A}_C = \prod_{\nu \in C} \mathcal{A}_{\nu}$ denote the joint state and action space of agents within C . A Markov Game $MG(\mathcal{S}, \mathcal{M}, \mathcal{A}, H, \mathbb{P}, \mathbf{R})$ is a clique-dominant linear Markov game if (a) it is clique-dominant (i.e., obeys Assumption 2), and (b) for every $C \in \mathcal{C}$, $h \in [H]$, for a set of $|C| + 1$ features $\{\phi_{\nu}\}_{\nu \in C}, \phi_{\nu} : \mathcal{S}_C \times \mathcal{A}_C \rightarrow \mathbb{R}^d$ and $\psi_C : \mathcal{S}_C \times \mathcal{A}_C \rightarrow \mathbb{R}^d$, there exist d unknown measures $\mu_h^C(\cdot) = \{\mu_{1,h}^C(\cdot), \dots, \mu_{|C|,h}^C(\cdot)\}$ over $\mathcal{S}_C$ and an unknown vector $\theta_h^C \in \mathbb{R}^d$ such that $\forall (\mathbf{x}, \mathbf{a}) \in \mathcal{S}_C \times \mathcal{A}_C$ and $\nu \in C$,*

$$\begin{aligned} \tilde{\mathbb{P}}_h^C(\cdot | \mathbf{x}, \mathbf{a}) &= \langle \psi_C(\mathbf{x}, \mathbf{a}), \mu_h^C(\cdot) \rangle, \text{ and} \\ \tilde{r}_{\nu, h}(\mathbf{x}, \mathbf{a}) &= \langle \phi_{\nu}(\mathbf{x}, \mathbf{a}), \theta_h^C \rangle. \end{aligned}$$

We denote the overall clique feature vector as $\Phi_C(\cdot) \in \mathbb{R}^{d \times |C|}$, where, for any $\mathbf{x} \in \mathcal{S}_C, \mathbf{a} \in \mathcal{A}_C$, $\Phi_C(\mathbf{x}, \mathbf{a}) = [[\phi_1(\mathbf{x}, \mathbf{a}), \psi_C(\mathbf{x}, \mathbf{a})]^\top, \dots, [\phi_{|C|}(\mathbf{x}, \mathbf{a}), \psi_C(\mathbf{x}, \mathbf{a})]^\top]^\top$, and the overall approximate clique reward $\tilde{r}_h^C(\mathbf{x}, \mathbf{a}) = [\tilde{r}_{1,h}(\mathbf{x}, \mathbf{a}), \dots, \tilde{r}_{|C|,h}(\mathbf{x}, \mathbf{a})]^\top$. Under this representation, we have that for any $\mathbf{x} \in \mathcal{S}_C, \mathbf{a} \in \mathcal{A}_C, h \in [H]$,

$$\begin{aligned} \tilde{r}_h^C(\mathbf{x}, \mathbf{a}) &= \Phi_C(\mathbf{x}, \mathbf{a})^\top \begin{bmatrix} \theta_h^C \\ \mathbf{0}_d \end{bmatrix}, \text{ and,} \\ \mathbf{1}_{|C|} \cdot \tilde{\mathbb{P}}_h^C(\cdot | \mathbf{x}, \mathbf{a}) &= \Phi_C(\mathbf{x}, \mathbf{a})^\top \begin{bmatrix} \mathbf{0}_d \\ \mu_h^C(\cdot) \end{bmatrix}. \end{aligned}$$

Assume WLOG $\forall C \in \mathcal{C}, \|\Phi_C(\mathbf{x}, \mathbf{a})\| \leq \sqrt{|C|} \forall (\mathbf{x}, \mathbf{a}) \in \mathcal{S}_C \times \mathcal{A}_C, \|\theta_h^C\| \leq \sqrt{d}, \|\mu_h^C(\mathcal{S}_C)\| \leq \sqrt{d}$.

Essentially, this assumption requires that once we are provided a clique covering, and the Markov game obeys the *clique-dominance* property (Assumption 2), the approximate rewards $\tilde{r}_{\nu, h}$ are linear functions of a known feature vector ϕ_C evaluated on the joint state-action of the agents within its clique. Additionally, it assumes that the approximate marginal transitions $\tilde{\mathbb{P}}_h^C$ are linear functions of a known

feature ψ_C . This, in fact, is a straightforward extension of the single-agent linear MDP parameterization (see Assumption A in Jin et al. (2020)) to *clique-dominant* Markov games, see Remark 7 in Appendix B to compare.

4.1. Algorithm Design

The first step in our approach is to compute a k -clique covering $\mathcal{C}$ of the influence graph G . Recall that by Remark 6 that this can be done in polynomial time with a 1.25 approximation of $\mathcal{C}^*$. Since the game is clique-dominant (Assumption 2), we can learn k decentralized policies $\pi_1, \dots, \pi_k$, one corresponding to each clique of agents in $\mathcal{C}$ without incurring too much approximation error. Now, to motivate the design, we first observe that Assumption 3 implies that for each clique $C \in \mathcal{C}$, there exist a set of weights such that the scalarized Q -values for any parameter $\mathbf{v}_C$ are *almost* linear projections of the overall clique features $\Phi_C(\cdot)$, where the total error is no larger than $2H\epsilon(k)$.

Lemma 1 (Almost linear weights in Markov Games). *Under Assumption 3 for graph G with k cliques ordered from $1, \dots, k$, we have, for any fixed decentralized policy $\pi = \{\pi_1, \dots, \pi_k\}$ and $\mathbf{v} = \{\mathbf{v}_1, \dots, \mathbf{v}_k\} \in \Upsilon$, there exist weights $\{\mathbf{w}_{\mathbf{v}, h}^{\pi_\tau}\}_{h=1, \tau=1}^{H, k}$ such that $\left| Q_{\mathbf{v}, h}^{\pi}(\mathbf{x}, \mathbf{a}) - \sum_{\tau=1}^k \mathbf{v}_\tau^\top \Phi_\tau(\mathbf{x}_\tau, \mathbf{a}_\tau)^\top \mathbf{w}_{\mathbf{v}, h}^{\pi_\tau} \right| \leq 2H\epsilon(k) \forall (\mathbf{x}, \mathbf{a}, h)$, where $\|\mathbf{w}_{\mathbf{v}, h}^{\pi_\tau}\|_2 \leq 2H\sqrt{d}$.*

Armed with this observation, we design a policy using *vector-valued* linear least-squares regression, as the optimal policy is only at most $2H\epsilon$ away from the best least-squares fit. In a nutshell, our approach can be summed up in two steps: (a) first, we approximate the Pareto frontier Π^* with the set of policies Π_{Υ}^* recoverable by scalarization (see Remark 2), (b) next, we empirically approximate Π_{Υ}^* with a collection of T policies (one for each episode), such that the *Bayes Regret* is minimized (Proposition 2). In each episode $t \in [T]$, we sample a scalarization parameter $\mathbf{v}_t \sim p_{\Upsilon}$, and run k *vector-valued* decentralized linear least-squares regressions to approximate the optimal policy $\pi_{\mathbf{v}_t}^*$ with k policies $\pi_1(t), \dots, \pi_k(t)$ such that the resulting Q -values overestimate $Q_{\mathbf{v}_t, h}^*$ with high probability. Then, each agent in clique τ selects the corresponding greedy action with respect to $\pi_\tau(t)$. This approach is carried out via *value iteration with optimism*, as described below.

We describe the policy for any clique $C \in \mathcal{C}$ of size n_C . For any scalarization $\mathbf{v}(t) \in \mathbb{R}^n$, the n_C values corresponding to agents in $\mathcal{C}$ is denoted by $\mathbf{v}_C(t)$. Now, consider the MDP $\tilde{\text{MDP}}_C$ formed by scalarizing the Markov game corresponding to the approximate rewards $\tilde{r}_h^C$ and transition dynamics $\tilde{\mathbb{P}}_h^C$ with the parameter $\mathbf{v}_{C, t}$ (i.e., the reward function in $\tilde{\text{MDP}}_C$ is given by $\mathbf{v}_C(t)^\top \mathbf{r}_h^C$, transition by $\tilde{\mathbb{P}}_h$ and state-action spaces as $\mathcal{S}_C$ and $\mathcal{A}_C$ respectively). For each clique C , we will use value iteration to recover the optimal

$$\begin{aligned}
 V_C^{h+1}(t)(\mathbf{x}) &\leftarrow \arg \max_{\mathbf{a} \in \mathcal{A}} \left[\mathbf{v}_C(t)^\top \left(\mathbf{Q}_C^{h+1}(t)^\top \Phi_C(\mathbf{x}, \mathbf{a}) \right) \right] \quad \forall \mathbf{x} \in \mathcal{S}_C, \\
 \hat{\mathbf{Q}}_C^h(t) &\leftarrow \arg \min_{\mathbf{w} \in \mathbb{R}^d} \left[\sum_{\tau \in [s_t^C]} \left\| \mathbf{y}_C^h(\tau) - \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))^\top \mathbf{w} \right\|_2^2 + \lambda \|\mathbf{w}\|_2^2 \right], \\
 Q_C^h(t)(\mathbf{x}, \mathbf{a}) &\leftarrow \mathbf{v}_C(t)^\top \left(\hat{\mathbf{Q}}_C^h(t)^\top \Phi_C(\mathbf{x}, \mathbf{a}) \right) + \beta_C^h(t) \cdot \left\| \Phi_C(\mathbf{x}, \mathbf{a})^\top \Lambda_C^h(t)^{-1} \Phi_C(\mathbf{x}, \mathbf{a}) \right\|_2.
 \end{aligned}$$

Figure 1. Vector-valued least squares regression update rule for MulttOVI.

policy for this $\widetilde{\text{MDP}}_C$ (let us call it $\tilde{\pi}_C^*(t)$). The algorithm is a distributed variant of least-squares value iteration with UCB exploration. Following Proposition 1, the key idea is to make sure that each agent in C acts according to the joint policy that is aiming to mimic $\tilde{\pi}_C^*(t)$. Therefore, we must ensure that the local estimate for the *joint* policy obtained by any agent must be identical, such that the *joint* action is in accordance with $\tilde{\pi}_C^*(t)$. To achieve this we will obtain the approximated (scalar) Q -values for $\tilde{\pi}_C^*(t)$ by recursively applying the Bellman equation and solving the resulting equations via a *vector-valued* regression. Since the policy variables are designed to be identical each agent in C at all times, we describe the procedure for any agent within C . For any episode t , let us assume that the last round of synchronization between agents in C occurred at time s_t^C . Each agent within the clique C obtains an *identical* sequence of value functions $\{Q_C^h(t)\}_{h \in [H]}$ by iteratively performing linear least-squares ridge regression from the history available from the previous s_t^C episodes by first learning a vector Q -function $\hat{\mathbf{Q}}_C^h(t)$ over $\mathbb{R}^{n_C}$, which is scalarized by $\mathbf{v}_C(t)$ to obtain the Q -values as $Q_C^h(t) = \mathbf{v}_C(t)^\top \hat{\mathbf{Q}}_C^h(t)$. Each agent m first sets $\hat{\mathbf{Q}}_C^{H+1}(t)$ to be a zero vector in $\mathbb{R}^{n_C}$, and for any $h \in [H]$, solves the following sequence of regressions to obtain Q -values described in Figure 1. Where the last equation holds for any $(\mathbf{x}, \mathbf{a}) \in (\mathcal{S}_C \times \mathcal{A}_C)$ and the targets $\mathbf{y}_C^h(\tau) = \mathbf{r}_C^h(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) + \mathbf{1}_{n_C} \cdot V_C^{h+1}(t)(\mathbf{x}_C^{h+1}(\tau))$, $\mathbf{1}_{n_C}$ denotes the all-ones vector in $\mathbb{R}^{n_C}$, $\beta_C^h(t)$ is selected such that with high probability the estimated Q -values overestimate the require Q -values, and $\Lambda_C^h(t)$ is described subsequently. Once all of these quantities are computed, each agent $\nu \in C$ selects the action $a_\nu^h(t) = [\arg \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(t)(\mathbf{x}_C^h(t), \mathbf{a})]_\nu$ for each $h \in [H]$. Hence, the joint clique action $\mathbf{a}_C^h(t) = \{a_\nu^h(t)\}_{\nu=1}^{n_C} = \arg \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(\mathbf{x}_C^h(t), \mathbf{a})$. Observe that while the computation of the policy is decentralized, the policies executed for all agents $\nu \in C$ coincide at all times by the modeling assumption and the periodic synchronizations between agents. We now present the closed form of $\hat{\mathbf{Q}}_C^h(t)$. Consider the contraction $\mathbf{z}_C^h(\tau) = (\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))$ and the map $\hat{\Phi}_C^h(t) : \mathbb{R}^d \rightarrow \mathbb{R}^{n_C}$ such that for any $\boldsymbol{\theta} \in \mathbb{R}^d$,

$$\hat{\Phi}_C^h(t)\boldsymbol{\theta} \triangleq [(\Phi_C(\mathbf{z}_C^h(1))^\top \boldsymbol{\theta})^\top, \dots, (\Phi_C(\mathbf{z}_C^h(t))^\top \boldsymbol{\theta})^\top]^\top.$$

Now, consider $\Lambda_C^h(t) = \Phi_C^h(s_t^C)^\top \Phi_C^h(s_t^C) + \lambda \mathbf{I}_d \in \mathbb{R}^{d \times d}$, and $\mathbf{U}_C^h(t) = \sum_{\tau=1}^{s_t^C} \Phi_C(\mathbf{z}_C^h(\tau)) \mathbf{y}_C^h(\tau)$. Then, we have by a multi-task concentration (see Appendix B of Chowdhury and Gopalan (2020)),

$$\hat{\mathbf{Q}}_C^h(t)(\mathbf{x}, \mathbf{a}) = \Phi_C(\mathbf{x}, \mathbf{a})^\top \Lambda_C^h(t)^{-1} \mathbf{U}_C^h(t).$$

The algorithm is presented in Algorithm 1. The algorithm is essentially learning k multi-agent policies by solving a vector-valued regression, one for each clique in the covering $\mathcal{C}$, such that the group of agents in each clique can learn the approximate clique-based MG (ref. Assumption 2).

4.2. Regret Analysis

Theorem 2. *Algorithm 1 on a game with n agents satisfying Assumptions 1, 2, 3 with error ε_* , approximate clique covering $\hat{\mathcal{C}}$, and $\kappa \cdot dH \cdot \theta(G)$ rounds of communication for some $\kappa > 1$, obtains, with probability at least $1 - \alpha$, regret:*

$$\tilde{\mathcal{O}} \left(\theta(G) \cdot d^{\frac{3}{2}} H^2 (2T \cdot n_{\max})^{\frac{2}{\kappa}} \left(\sqrt{T \log \left(\frac{1}{\alpha} \right)} + 2T n_{\max} \cdot \varepsilon_* \right) \right).$$

Where $\theta(G)$ denotes the clique covering number of G , and $n_{\max}$ is the size of the largest clique in $\hat{\mathcal{C}}$.

Remark 3 (Regret Bound). *The above Theorem suggests that our algorithm is no-regret (up to a factor ε_*), as long as there are a sufficient (logarithmic) rounds of communication. Further, if the clique-dominance error $\varepsilon_* = o(T^{-\gamma})$ for $\gamma > 0$ then the algorithm is no-regret regardless. Additionally, we see that MulttOVI can be simulated on a single agent with n objectives, where $S = 1$ and G is complete, which provides, to the best of our knowledge, the first no-regret algorithm for multi-objective reinforcement learning (Mossalam et al., 2016). We present further clarifications, discussions and lower bounds in Remarks 8, 9, 10 and 11 in Section B of the Appendix.*

Conclusion. We presented the first (to the best of our knowledge) no-regret algorithm for partially-observable cooperative Markov games, with competitive experimental performance (experiments deferred to Appendix for brevity). We generalize several concepts in the cooperative MARL literature, and we believe our results will be important for further work in cooperative MARL.

References

- Y. Abbasi-Yadkori, D. Pál, and C. Szepesvári. Improved algorithms for linear stochastic bandits. In *Advances in Neural Information Processing Systems*, pages 2312–2320, 2011.
- A. Agarwal, H. Luo, B. Neyshabur, and R. E. Schapire. Corraling a band of bandit algorithms. In *Conference on Learning Theory*, pages 12–38. PMLR, 2017.
- A.-L. Barabasi. The origin of bursts and heavy tails in human dynamics. *Nature*, 435(7039):207, 2005.
- A. L. Bazzan. Opportunities for multiagent systems and multiagent reinforcement learning in traffic control. *Autonomous Agents and Multi-Agent Systems*, 18(3):342, 2009.
- V. D. Blondel and J. N. Tsitsiklis. Complexity of stability and controllability of elementary hybrid systems. *Automatica*, 35(3):479–489, 1999.
- C. Boutilier. Planning, learning and coordination in multiagent decision processes. Citeseer, 1996.
- J. M. Buchanan. The relevance of pareto optimality. *Journal of conflict resolution*, 6(4):341–354, 1962.
- G. D. Canas and L. Rosasco. Learning probability measures with respect to optimal transport metrics. *arXiv preprint arXiv:1209.1077*, 2012.
- M. R. Cerioli, L. Faria, T. O. Ferreira, C. A. Martinhon, F. Protti, and B. Reed. Partition into cliques for cubic graphs: Planar case, complexity and approximation. *Discrete Applied Mathematics*, 156(12):2270–2278, 2008.
- W.-Y. Chiu, J.-T. Hsieh, and C.-M. Chen. Pareto optimal demand response based on energy costs and load factor in smart grid. *IEEE Transactions on Industrial Informatics*, 16(3):1811–1822, 2019.
- S. R. Chowdhury and A. Gopalan. No-regret algorithms for multi-task bayesian optimization. *arXiv preprint arXiv:2008.08885*, 2020.
- G. Ding, J. J. Koh, K. Merckaert, B. Vanderborght, M. M. Nicotra, C. Heckman, A. Roncone, and L. Chen. Distributed reinforcement learning for cooperative multi-robot object manipulation. *arXiv preprint arXiv:2003.09540*, 2020.
- A. Ghosh, S. R. Chowdhury, and A. Gopalan. Misspecified linear bandits. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 31, 2017.
- M. Grötschel, L. Lovász, and A. Schrijver. Stable sets in graphs. In *Geometric Algorithms and Combinatorial Optimization*, pages 272–303. Springer, 1988.
- H. Gu, X. Guo, X. Wei, and R. Xu. Q-learning for mean-field controls. *arXiv preprint arXiv:2002.04131*, 2020.
- C. Guestrin, D. Koller, and R. Parr. Multiagent planning with factored mdps. In *NIPS*, volume 1, pages 1523–1530, 2001.
- C. Guestrin, S. Venkataraman, and D. Koller. Context-specific multiagent coordination and planning with factored mdps. In *AAAI/IAAI*, pages 253–259, 2002.
- E. Hillel, Z. Karnin, T. Koren, R. Lempel, and O. Somekh. Distributed exploration in multi-armed bandits. *arXiv preprint arXiv:1311.0800*, 2013.
- C. Jin, Z. Yang, Z. Wang, and M. I. Jordan. Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pages 2137–2143. PMLR, 2020.
- S. Kar, J. M. Moura, and H. V. Poor. Qd-learning: A collaborative distributed strategy for multi-agent reinforcement learning through consensus innovations. *IEEE Transactions on Signal Processing*, 61(7):1848–1862, 2013.
- R. M. Karp. Reducibility among combinatorial problems. In *Complexity of computer computations*, pages 85–103. Springer, 1972.
- J. Knowles. Parego: A hybrid algorithm with on-line landscape approximation for expensive multiobjective optimization problems. *IEEE Transactions on Evolutionary Computation*, 10(1):50–66, 2006.
- M. Lauer and M. Riedmiller. An algorithm for distributed reinforcement learning in cooperative multi-agent systems. In *In Proceedings of the Seventeenth International Conference on Machine Learning*. Citeseer, 2000.
- A. Marinescu, I. Dusparic, A. Taylor, V. Cahill, and S. Clarke. Decentralised multi-agent reinforcement learning for dynamic and uncertain environments. *arXiv preprint arXiv:1409.4561*, 2014.
- M. Molloy. The list chromatic number of graphs with small clique number. *Journal of Combinatorial Theory, Series B*, 134:264–284, 2019.
- H. Mossalam, Y. M. Assael, D. M. Roijers, and S. Whiteson. Multi-objective deep reinforcement learning. *arXiv preprint arXiv:1610.02707*, 2016.
- A. Pacchiano, M. Phan, Y. Abbasi-Yadkori, A. Rao, J. Zimmert, T. Lattimore, and C. Szepesvari. Model selection in contextual stochastic bandit problems. *arXiv preprint arXiv:2003.01704*, 2020.

- C. H. Papadimitriou and J. N. Tsitsiklis. The complexity of markov decision processes. *Mathematics of operations research*, 12(3):441–450, 1987.
- B. Paria, K. Kandasamy, and B. Póczos. A flexible framework for multi-objective bayesian optimization using random scalarizations. In *Uncertainty in Artificial Intelligence*, pages 766–776. PMLR, 2020.
- G. Qu and N. Li. Exploiting fast decaying and locality in multi-agent mdp with tree dependence structure. In *2019 IEEE 58th Conference on Decision and Control (CDC)*, pages 6479–6486. IEEE, 2019.
- G. Qu, Y. Lin, A. Wierman, and N. Li. Scalable multi-agent reinforcement learning for networked systems with average reward. *arXiv preprint arXiv:2006.06626*, 2020.
- M. Roth, R. Simmons, and M. Veloso. Exploiting factored representations for decentralized execution in multiagent teams. In *Proceedings of the 6th international joint conference on Autonomous agents and multiagent systems*, pages 1–7, 2007.
- A. Schaerf, Y. Shoham, and M. Tennenholtz. Adaptive load balancing: A study in multi-agent learning. *Journal of artificial intelligence research*, 2:475–500, 1994.
- D. Shah, V. Somani, Q. Xie, and Z. Xu. On reinforcement learning for turn-based zero-sum markov games. *arXiv preprint arXiv:2002.10620*, 2020.
- L. S. Shapley. Stochastic games. *Proceedings of the national academy of sciences*, 39(10):1095–1100, 1953.
- H. Thadakamaila, U. N. Raghavan, S. Kumara, and R. Albert. Survivability of multiagent-based supply networks: a topological perspect. *IEEE Intelligent Systems*, 19(5): 24–31, 2004.
- P. Vamplew, J. Yearwood, R. Dazeley, and A. Berry. On the limitations of scalarisation for multi-objective reinforcement learning of pareto fronts. In *Australasian joint conference on artificial intelligence*, pages 372–378. Springer, 2008.
- R. Wang, R. Salakhutdinov, and L. F. Yang. Provably efficient reinforcement learning with general value function approximation. *arXiv preprint arXiv:2005.10804*, 2020.
- X. F. Wang and T. Sandholm. Learning near-pareto-optimal conventions in polynomial time. 2003.
- Y. Wang, H. Liu, W. Zheng, Y. Xia, Y. Li, P. Chen, K. Guo, and H. Xie. Multi-objective workflow scheduling with deep-q-network-based multi-agent reinforcement learning. *IEEE Access*, 7:39974–39982, 2019.
- Q. Xie, Y. Chen, Z. Wang, and Z. Yang. Learning zero-sum simultaneous-move markov games using function approximation and correlated equilibrium. In *Conference on Learning Theory*, pages 3674–3682. PMLR, 2020.
- L. Yang and M. Wang. Reinforcement learning in feature space: Matrix bandit, kernels, and regret bound. In *International Conference on Machine Learning*, pages 10746–10756. PMLR, 2020.
- Z. Yang, C. Jin, Z. Wang, M. Wang, and M. I. Jordan. Provably efficient reinforcement learning with kernel and neural function approximations. In H. Larochelle, M. Ranzato, R. Hadsell, M. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/9fa04f87c9138de23e92582b4ce549ec-Abstract.html>.
- T. Yoshikawa. Decomposition of dynamic team decision problems. *IEEE Transactions on Automatic Control*, 23(4):627–632, 1978.
- T. Yu, H. Wang, B. Zhou, K. Chan, and J. Tang. Multi-agent correlated equilibrium $q(\lambda)$ learning for coordinated smart generation control of interconnected power grids. *IEEE transactions on power systems*, 30(4):1669–1679, 2014.
- K. Zhang, Z. Yang, and T. Basar. Networked multi-agent reinforcement learning in continuous spaces. In *2018 IEEE Conference on Decision and Control (CDC)*, pages 2771–2776. IEEE, 2018a.
- K. Zhang, Z. Yang, H. Liu, T. Zhang, and T. Basar. Fully decentralized multi-agent reinforcement learning with networked agents. In *International Conference on Machine Learning*, pages 5872–5881. PMLR, 2018b.
- K. Zhang, Z. Yang, and T. Başar. Multi-agent reinforcement learning: A selective overview of theories and algorithms. *arXiv preprint arXiv:1911.10635*, 2019.
- Y. Zhao, I. Borovikov, J. Rupert, C. Somers, and A. Beirami. On multi-agent learning in team sports games. *arXiv preprint arXiv:1906.10124*, 2019.
- Y. Zhao, Y. Tian, J. D. Lee, and S. S. Du. Provably efficient policy gradient methods for two-player zero-sum markov games. *arXiv preprint arXiv:2102.08903*, 2021.

A. Related Work

It is difficult to summarize the rich literature on cooperative multi-agent reinforcement learning, being examined by various perspectives from the AI (Lauer and Riedmiller, 2000; Boutilier, 1996), control (Yoshikawa, 1978; Wang and Sandholm, 2003) and statistical learning communities (Xie et al., 2020). While there has been extensive recent work on provably efficient algorithms for *competitive* multiplayer RL (Xie et al., 2020; Zhao et al., 2021; Shah et al., 2020), our work is placed in the *cooperative* MARL setting, with the objective being to efficiently find *globally* optimal policies, where recent work has focused on *locality* assumptions in order to reduce the policy search space (Qu and Li, 2019; Qu et al., 2020). However, the more general *heterogeneous* reward setting considered in our work, where each agent may have unique rewards, corresponds to the *team average* games studied previously (Kar et al., 2013; Zhang et al., 2018b;a). While some of these approaches do provide tractable algorithms that are decentralized and convergent, none provide finite-time regret guarantees, and moreover, focus only on maximizing the *team average* reward. In our paper, however, we study a more general form of regret in order to recover a set of policies on the *Pareto frontier*. For a detailed overview of algorithms in cooperative MARL, we refer the readers to the illuminating survey by Zhang et al. (2019). Our work builds on the increasingly relevant line of work in (single-agent) reinforcement learning with function approximation (Wang et al., 2020; Yang et al., 2020; Yang and Wang, 2020; Jin et al., 2020), however, our environment suffers from several additional challenges not present in single-agent settings, such as communication costs, scalability issues and decentralized multi-agent planning, which are the key contributions of this paper.

B. Omitted Remarks

Remark 4 (Feasibility of Local Influence). Networked influence assumptions similar to Assumption 1 have been explored extensively in the literature (Gu et al., 2020; Qu and Li, 2019; Qu et al., 2020; Guestrin et al., 2001), and is commonly present in many real-world environments such as supply-chain networks (Thadakamaila et al., 2004) and social networks (Barabasi, 2005). However, in contrast to prior work, which assume the individual reward functions to be functions only of the *local* state and action, we consider a broader model where even *local* rewards are functions of the neighborhood.

Remark 5 (Comparison with Exponential Decay). Compared to the (c, ρ) -decay of (Qu and Li, 2019), our assumptions are both weaker and stronger in some aspects. First, we do not require any knowledge of *pairwise* interactions, and make assumptions at the subgraph level, and second, we do not require an exponential decay: simply an upper bound on the error suffices. Consequently, our guarantee only utilizes local neighborhoods (i.e., agents at distance 1), whereas (c, ρ) -exponential decay utilizes *all* interactions. In this regard, we remark that our clique-dominance assumption can incorporate further neighbors by partitioning the κ -power of G and introducing state-action communication between agents (as any agent can only observe its neighbors, hence information about distant neighbors must be communicated), which we omit for simplicity.

Remark 6 (Complexity of clique covering). Assumption 2 requires a clique covering of G , which is NP-hard (Karp, 1972), however, for special cases, can be found in polynomial time (e.g., triangle-free graphs (Molloy, 2019) and perfect graphs (Grötschel et al., 1988)). Cerioli et al. (2008) provide a polynomial-time algorithm that gives a 1.25 approximation of the minimal clique covering, therefore, we can replace $\mathcal{C}^*$ with an approximate covering $\hat{\mathcal{C}}$ such that $|\hat{\mathcal{C}}| \leq 1.25|\mathcal{C}^*|$ for any G in our approach.

Remark 7 (Multi-agent modeling assumptions). In contrast to the typical linear MDP assumption, here we model the rewards and dynamics for each clique of agents separately, each with d linear dimensions each. In the single-agent setting, identical assumptions on the reward and transition kernels will lead to a model with complexity d , whereas in our formulation we have a complexity of $2d$, implying that our fomulation incurs an overhead of $2\sqrt{2}$ in the regret if applied to the single-agent setting, compared to the model presented in Jin et al. (2020). Furthermore, observe that in the *fully-cooperative* setting (where agents share the reward function), i.e., $r_{1,h} = \dots = r_{n,h} \forall h \in [H]$, we have that assuming, for all agents that $\phi_1 = \phi_2 = \dots = \phi_n$ satisfies the modeling requirement.

Remark 8 (Regret Bound). Theorem 2 claims in conjunction with Proposition 2 that **MultOVI** obtains Bayes regret of $\tilde{O}(\theta(G) \cdot \sqrt{T})$ even with limited communication. Note that for complete G , $\theta(G) = 1$ and the dependence on T matches that of MDP algorithms exactly (e.g., Jin et al. (2020)), demonstrating that our analysis is tight. Additionally, we see that this algorithm can easily be applied to an MDP by simply selecting $p_{\mathbf{r}}$ to be a point mass at the appropriate $\mathbf{v}$, with no increase in regret. Thirdly, we see that **MultOVI** can be simulated on a single agent with n objectives, where $S = 1$ and G is complete, which provides, to the best of our knowledge, the first no-regret algorithm for multi-objective reinforcement learning (Mossalam et al., 2016).

Remark 9 (Lower Bounds). For tabular multi-agent reinforcement learning, in the Appendix, we demonstrate that a collection of $\theta(G)$ episodic MDPs (one corresponding to each clique $C \in \mathcal{C}$, with state-action spaces $\mathcal{S}_C$ and $\mathcal{A}_C$) can be constructed such that the cumulative Bayes regret incurred by any algorithm is $\Omega(H\sqrt{dT})$ for each C , and therefore the total regret in the Markov game is $\Omega(\theta(G)H\sqrt{dT})$ where $n_{\min}$ is the size of the smallest clique of G , demonstrating that the $\theta(G)$ term is unavoidable in general. Further, the utilization of “Bernstein-type” confidence bonuses can shave off an additional factor of $\sqrt{H}$ in our regret (see discussion in Jin et al. (2020)). Regarding the optimal dependence on communication, we conjecture that our bound is almost-optimal, as similar lower bounds have been demonstrated for distributed exploration in multi-armed bandits (Hillel et al., 2013).

Remark 10 (Modeling influence and unknown dynamics). For arbitrary influence graphs G , the misspecification ε incurred by using a k -clique covering $\mathcal{C}$ of G (Assumption 2) can be unknown in general, and may be unique for each $\mathcal{C}$. In this setting, we conjecture that a corraling-type algorithm (Pacchiano et al., 2020; Agarwal et al., 2017) that adaptively selects the best clique covering $\mathcal{C}$ can provide regret close to our algorithm without knowing the misspecification $\varepsilon(k)$.

Remark 11 (Communication complexity). We can control the communication budget by adjusting the threshold parameter S . Note that when $S = 1$, communication will occur each round, as the threshold will be satisfied trivially by the rank-1 update to the Gram matrix. If the horizon T is known in advance, one can set $S = (1 + n_{\max}T/d)^{1/D}$ for some independent constant $D > 1$, to ensure that the total rounds of communication is a fixed constant $\theta(G)(dD + 1)H$, which provides us a group regret of $\tilde{\mathcal{O}}(\theta(G) \cdot n^{\frac{1}{2D}} \cdot T^{\frac{1}{2} + \frac{1}{2D}})$. A balance can be obtained by setting $S = C'$ for some absolute constant C' , leading to a total $\mathcal{O}(\theta(G) \cdot \log(n_{\max}T))$ rounds with $\tilde{\mathcal{O}}(\theta(G)\sqrt{T})$ regret.

C. Algorithm Design: Extended

The first step in our approach is to compute a k -clique covering $\mathcal{C}$ of the influence graph G . Recall that by Remark 6 that this can be done in polynomial time with a 1.25 approximation of $\mathcal{C}^*$. Since the game is clique-dominant (Assumption 2), we can learn k decentralized policies $\pi_1, \dots, \pi_k$, one corresponding to each clique of agents in $\mathcal{C}$ without incurring too much approximation error. Now, to motivate the design, we first observe that Assumption 3 implies that for each clique $C \in \mathcal{C}$, there exist a set of weights such that the scalarized Q -values for any parameter $\mathbf{v}_C$ are *almost* linear projections of the overall clique features $\Phi_C(\cdot)$, where the total error is no larger than $2H\varepsilon(k)$.

Lemma 2 (Almost linear weights in Markov Games). *Under Assumption 3 for graph G with k cliques ordered from $1, \dots, k$, we have, for any fixed decentralized policy $\pi = \{\pi_1, \dots, \pi_k\}$ and $\mathbf{v} = \{\mathbf{v}_1, \dots, \mathbf{v}_k\} \in \Upsilon$, there exist weights $\{\mathbf{w}_{\mathbf{v},h}^{\pi_\tau}\}_{h=1,\tau=1}^{H,k}$ such that $\left| Q_{\mathbf{v},h}^\pi(\mathbf{x}, \mathbf{a}) - \sum_{\tau=1}^k \mathbf{v}_\tau^\top \Phi_\tau(\mathbf{x}_\tau, \mathbf{a}_\tau)^\top \mathbf{w}_{\mathbf{v},h}^{\pi_\tau} \right| \leq 2H\varepsilon(k) \forall (\mathbf{x}, \mathbf{a}, h)$, where $\|\mathbf{w}_{\mathbf{v},h}^{\pi_k}\|_2 \leq 2H\sqrt{d}$.*

Armed with this observation, we design a policy using *vector-valued* linear least-squares regression, as the optimal policy is only at most $2H\varepsilon$ away from the best least-squares fit. In a nutshell, our approach can be summed up in two steps: (a) first, we approximate the Pareto frontier Π^* with the set of policies Π_Υ^* recoverable by scalarization (see Remark 2), (b) next, we empirically approximate Π_Υ^* with a collection of T policies (one for each episode), such that the *Bayes Regret* is minimized (Proposition 2). In each episode $t \in [T]$, we sample a scalarization parameter $\mathbf{v}_t \sim p_\Upsilon$, and run k *vector-valued* decentralized linear least-squares regressions to approximate the optimal policy $\pi_{\mathbf{v}_t}^*$ with k policies $\pi_1(t), \dots, \pi_k(t)$ such that the resulting Q -values overestimate $Q_{\mathbf{v}_t,h}^*$ with high probability. Then, each agent in clique τ selects the corresponding greedy action with respect to $\pi_\tau(t)$. This approach is carried out via *value iteration with optimism*, as described below.

We describe the policy for any clique $C \in \mathcal{C}$ of size n_C . For any scalarization $\mathbf{v}(t) \in \mathbb{R}^n$, the n_C values corresponding to agents in C is denoted by $\mathbf{v}_C(t)$. Now, consider the MDP $\widehat{\text{MDP}}_C$ formed by scalarizing the Markov game corresponding to the approximate rewards $\tilde{\mathbf{r}}_h^C$ and transition dynamics $\tilde{\mathbb{P}}_h^C$ with the parameter $\mathbf{v}_{C,t}$ (i.e., the reward function in $\widehat{\text{MDP}}_C$ is given by $\mathbf{v}_C(t)^\top \mathbf{r}_h^C$, transition by $\tilde{\mathbb{P}}_h$ and state-action spaces as $\mathcal{S}_C$ and $\mathcal{A}_C$ respectively). For each clique C , we will use value iteration to recover the optimal policy for this $\widehat{\text{MDP}}_C$ (let us call it $\tilde{\pi}_C^*(t)$). The algorithm is a distributed variant of least-squares value iteration with UCB exploration. Following Proposition 1, the key idea is to make sure that each agent in C acts according to the joint policy that is aiming to mimic $\tilde{\pi}_C^*(t)$. Therefore, we must ensure that the local estimate for the *joint* policy obtained by any agent must be identical, such that the *joint* action is in accordance with $\tilde{\pi}_C^*(t)$. To achieve this we will obtain the approximated (scalar) Q -values for $\tilde{\pi}_C^*(t)$ by recursively applying the Bellman equation and solving the resulting equations via a *vector-valued* regression. Since the policy variables are designed to be identical each agent in C at all times, we describe the procedure for any agent within C .

For any episode t , let us assume that the last round of synchronization between agents in C occurred at time s_t^C . Each

Algorithm 1 MultOVI: Decentralized Learning in Low-Rank Cooperative Markov Games

```

1: Input:  $T, \Phi, H, S$ , sequence  $\beta_h = \{(\beta_h^t)_t\}$ .
2: Initialize:  $\Lambda_C^h(t) = \lambda \mathbf{I}_d, \delta \Lambda_C^h(t) = \mathbf{0}, \mathcal{U}_\nu^h, \mathcal{W}_\nu^h = \emptyset$  for each  $\nu \in G$ , clique cover  $\hat{\mathcal{C}}$  of  $G$ .
3: for episode  $t = 1, 2, \dots, T$  do
4:   Sample  $\mathbf{v}_t \sim p_{\mathbf{r}}$  using public randomness.
5:   for clique  $C \in \hat{\mathcal{C}}$  do
6:     for agent  $\nu \in C$  do
7:       Set  $V_C^{h+1}(t)(\cdot) \leftarrow 0$ .
8:       for step  $h = H, \dots, 1$  do
9:         Compute  $Q_C^h(t)(\cdot, \cdot)$  using vector-valued least-squares regression on  $\mathcal{U}_\nu^h$ .
10:        Set  $V_C^{h+1}(t)(\cdot) \leftarrow \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(t)(\cdot, \mathbf{a})$ .
11:      end for
12:      for step  $h = 1, \dots, H$  do
13:        Take action  $a_\nu^h(t) \leftarrow [\arg \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(t)(\mathbf{x}_C^h(t), \mathbf{a})]_\nu$ .
14:        Observe  $r_\nu^h(t), \tilde{\mathbf{x}}_\nu^{h+1}$ .
15:        Update  $\delta \Lambda_C^h(t) \leftarrow \delta \Lambda_C^h(t-1) + \Phi_C(\mathbf{z}_C^h(t)) \Phi_C(\mathbf{z}_C^h(t))^\top$ .
16:        Update  $\mathcal{W}_\nu^h \leftarrow \mathcal{W}_\nu^h \cup (\nu, \mathbf{x}_\nu^h(t), r_\nu^h(t))$ .
17:        if  $\log \frac{\det(\Lambda_C^h(t) + \delta \Lambda_C^h(t) + \lambda \mathbf{I})}{\det(\Lambda_C^h(t) + \lambda \mathbf{I})} > S$  then
18:          SYNCHRONIZE  $\leftarrow$  TRUE.
19:        end if
20:      end for
21:    end for
22:    if SYNCHRONIZE then
23:      Assign arbitrary agent in  $C$  as the SERVER AGENT.
24:      for step  $h = H, \dots, 1$  do
25:         $\forall$  AGENTS Send  $\mathcal{W}_\nu^h \rightarrow$  SERVER AGENT.
26:        SERVER AGENT Aggregate  $\mathcal{W}^h \leftarrow \bigcup_{\nu \in C} \mathcal{W}_\nu^h$ .
27:        SERVER AGENT Communicate  $\mathcal{W}^h$  to each agent.
28:         $\forall$  AGENTS Set  $\delta \Lambda_C^h(t+1) \leftarrow 0, \mathcal{W}_\nu^h \leftarrow \emptyset$ .
29:         $\forall$  AGENTS Set  $\Lambda_C^h(t+1) \leftarrow \Lambda_C^h(t) + \sum_{(\mathbf{x}, \mathbf{a}) \in \mathcal{W}^h} \Phi_C(\mathbf{x}, \mathbf{a}) \Phi_C(\mathbf{x}, \mathbf{a})^\top$ .
30:         $\forall$  AGENTS Set  $\mathcal{U}_\nu^h \leftarrow \mathcal{U}_\nu^h \cup \mathcal{W}^h$ 
31:      end for
32:    end if
33:  end for
34: end for

```

agent within the clique C obtains an *identical* sequence of value functions $\{Q_C^h(t)\}_{h \in [H]}$ by iteratively performing linear least-squares ridge regression from the history available from the previous s_t^C episodes by first learning a vector Q -function $\hat{\mathbf{Q}}_C^h(t)$ over $\mathbb{R}^{n_C}$, which is scalarized by $\mathbf{v}_C(t)$ to obtain the Q -values as $Q_C^h(t) = \mathbf{v}_C(t)^\top \hat{\mathbf{Q}}_C^h(t)$. Each agent m first sets $\hat{\mathbf{Q}}_C^{H+1}(t)$ to be a zero vector in $\mathbb{R}^{n_C}$, and for any $h \in [H]$, solves the following sequence of regressions to obtain Q -values. For each $h = H, \dots, 1$, for each agent computes,

$$\begin{aligned}
 V_C^{h+1}(t)(\mathbf{x}) &\leftarrow \arg \max_{\mathbf{a} \in \mathcal{A}} [\mathbf{v}_C(t)^\top (\mathbf{Q}_C^{h+1}(t)^\top \Phi_C(\mathbf{x}, \mathbf{a}))] \quad \forall \mathbf{x} \in \mathcal{S}_C, \\
 \hat{\mathbf{Q}}_C^h(t) &\leftarrow \arg \min_{\mathbf{w} \in \mathbb{R}^d} \left[\sum_{\tau \in [s_t^C]} \|\mathbf{y}_C^h(\tau) - \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))^\top \mathbf{w}\|_2^2 + \lambda \|\mathbf{w}\|_2^2 \right], \\
 Q_C^h(t)(\mathbf{x}, \mathbf{a}) &\leftarrow \mathbf{v}_C(t)^\top \left(\hat{\mathbf{Q}}_C^h(t)^\top \Phi_C(\mathbf{x}, \mathbf{a}) \right) + \beta_C^h(t) \cdot \|\Phi_C(\mathbf{x}, \mathbf{a})^\top \Lambda_C^h(t)^{-1} \Phi_C(\mathbf{x}, \mathbf{a})\|_2.
 \end{aligned}$$

Where the last equation holds for any $(\mathbf{x}, \mathbf{a}) \in (\mathcal{S}_C \times \mathcal{A}_C)$ and the targets $\mathbf{y}_C^h(\tau) = \mathbf{r}_C^h(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) + \mathbf{1}_{n_C} \cdot V_C^{h+1}(t)(\mathbf{x}_C^{h+1}(\tau))$, $\mathbf{1}_{n_C}$ denotes the all-ones vector in $\mathbb{R}^{n_C}$, $\beta_C^h(t)$ is selected such that with high probability the estimated Q -values overestimate the require Q -values, and $\Lambda_C^h(t)$ is described subsequently. Once all of these quantities are computed, each agent $\nu \in C$ selects the action $a_\nu^h(t) = [\arg \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(t)(\mathbf{x}_C^h(t), \mathbf{a})]_\nu$ for each $h \in [H]$. Hence, the joint clique action $\mathbf{a}_C^h(t) = \{a_\nu^h(t)\}_{\nu=1}^{n_C} = \arg \max_{\mathbf{a} \in \mathcal{A}_C} Q_C^h(\mathbf{x}_C^h(t), \mathbf{a})$. Observe that while the computation of the policy is decentralized, the policies executed for all agents $\nu \in C$ coincide at all times by the modeling assumption and the periodic synchronizations between agents. We now present the closed form of $\hat{\mathbf{Q}}_C^h(t)$. Consider the contraction

$\mathbf{z}_C^h(\tau) = (\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))$ and the map $\widehat{\Phi}_C^h(t) : \mathbb{R}^d \rightarrow \mathbb{R}^{tn_C}$ such that for any $\boldsymbol{\theta} \in \mathbb{R}^d$,

$$\Phi_C^h(t)\boldsymbol{\theta} \triangleq [(\Phi_C(\mathbf{z}_C^h(1))^\top \boldsymbol{\theta})^\top, \dots, (\Phi_C(\mathbf{z}_C^h(t))^\top \boldsymbol{\theta})^\top]^\top.$$

Now, consider $\Lambda_C^h(t) = \Phi_C^h(s_t^C)^\top \Phi_C^h(s_t^C) + \lambda \mathbf{I}_d \in \mathbb{R}^{d \times d}$, and $\mathbf{U}_C^h(t) = \sum_{\tau=1}^{s_t^C} \Phi_C(\mathbf{z}_C^h(\tau)) \mathbf{y}_C^h(\tau)$. Then, we have by a multi-task concentration (see Appendix B of Chowdhury and Gopalan (2020)),

$$\widehat{\mathbf{Q}}_C^h(t)(\mathbf{x}, \mathbf{a}) = \Phi_C(\mathbf{x}, \mathbf{a})^\top \Lambda_C^h(t)^{-1} \mathbf{U}_C^h(t).$$

The algorithm is presented in Algorithm 1. The algorithm is essentially learning k multi-agent policies by solving a vector-valued regression, one for each clique in the covering $\mathcal{C}$, such that the group of agents in each clique can learn the approximate clique-based MG (ref. Assumption 2). Since these approximate games themselves are bounded close to the true Markov game (by clique-dominance), this ensures that the agents incur low regret. We next present an analysis of communication cost.

Communication. Note that within a clique, the common state $\mathbf{x}_C$ is visible to all agents (Assumption 1), and hence the agents only require communication of rewards within a clique. To limit rounds in which communication occurs, we consider a synchronization criterion that is triggered whenever any agent in the clique explores a sufficiently novel part of the environment. Specifically, whenever $\det(\Lambda_C^h(t)) \geq S \cdot \det(\Lambda_C^h(s_t^C))$, for any $h \in [H]$ where S is a fixed constant determined in advance, the agents synchronize their rewards within their corresponding clique C . The synchronization can be done in $\mathcal{O}(n)$ messages by designating one agent per clique as the SERVER to aggregate messages.

Lemma 3 (Communication complexity). *Let the clique covering number be $\theta(G)$ and let $n_{\max} \leq n$ denote the size of the largest clique of G . If we use threshold $S > 1$, then the total number of episodes with communication $\gamma \leq dH \cdot \theta(G) \cdot \log_S(1 + \frac{T n_{\max}}{d}) + \theta(G) \cdot H$. When $S \leq 1$, $\gamma = T$.*

D. Full Proofs

D.1. Proof of Proposition 1

We restate the Proposition for clarity.

Proposition 1. For the scalarized value function, the Bellman optimality conditions are given as, for all $h \in [H]$, $\mathbf{x} \in \mathcal{S}$, $\mathbf{a} \in \mathcal{A}$ for any fixed $\mathbf{v} \in \Upsilon$,

$$Q_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}) = \mathbf{v}^\top \mathbf{r}_h(\mathbf{x}, \mathbf{a}) + \mathbb{P}_h V_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}), V_{\mathbf{v},h}^*(\mathbf{x}) = \max_{\mathbf{a} \in \mathcal{A}} Q_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}), \text{ and } V_{\mathbf{v},H+1}^*(\mathbf{x}) = 0.$$

Proof. We prove the above result by reducing the scalarized MMDP to an equivalent MDP. Observe that for any fixed $\mathbf{v} \in \Upsilon$, the (vector-valued) rewards can be scalarized to a scalar reward. For any step $h \in [H]$, for any fixed $\mathbf{v} \in \Upsilon$, consider the MDP with state space $\mathcal{S} = \mathcal{S}_1 \times \dots \times \mathcal{S}_n$, action space $\mathcal{A} = \mathcal{A}_1 \times \dots \times \mathcal{A}_n$ and reward function r'_h such that for all $(\mathbf{x}, \mathbf{a}) \in \mathcal{S} \times \mathcal{A}$, $r'_h(\mathbf{x}, \mathbf{a}) = \mathbf{v}^\top \mathbf{r}_h(\mathbf{x}, \mathbf{a})$. Therefore $r'_h(\mathbf{x}, \mathbf{a}) \in [0, 1]$ (since $\mathbf{r}_h$ lies on the n -dimensional simplex). Therefore, if the group of agents cooperate to optimize the scalarized reward (for any fixed scalarization parameter), the optimal (joint) policy coincides with the optimal policy for the aforementioned MDP defined over the *joint* state and action spaces. The optimal policy for the scalarized MDP is given by the greedy policy with respect to the following parameters:

$$Q_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}) = r'_h(\mathbf{x}, \mathbf{a}) + \mathbb{P}_h V_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}), V_{\mathbf{v},h}^*(\mathbf{x}) = \max_{\mathbf{a} \in \mathcal{A}} Q_{\mathbf{v},h}^*(\mathbf{x}, \mathbf{a}), \text{ and } V_{\mathbf{v},H+1}^*(\mathbf{x}) = 0. \quad (3)$$

Replacing the reward function with the vector-valued reward in terms of $\mathbf{v}$ provides us the result. $\square$

D.2. Proof of Theorem 1

We first restate the Theorem for clarity.

Theorem 1. For any Markov game with finite $\mathcal{A}$ and H , $\Pi_{\Upsilon}^* \subseteq \Pi^*$. If Π^* is convex, $\Pi_{\Upsilon}^* = \Pi^*$.

Proof. First, we prove the forward direction, i.e., that $\Pi_{\Upsilon}^* \subseteq \Pi^*$. The proof proceeds by contradiction. Assume that $\pi_{\mathbf{v}}^*$ does not lie in the Pareto frontier, then there exists a policy $\pi' \in \Pi$ such that $\mathbf{V}_1^{\pi'}(\mathbf{x}) \succeq \mathbf{V}_1^{\pi_{\mathbf{v}}^*}(\mathbf{x})$ for all $\mathbf{x} \in \mathcal{S}$ and $\pi \neq \pi_{\mathbf{v}}^*$.

Consider the final step H . Then, for any state $x \in \mathcal{S}$, we have that if $\mathbf{V}_H^{\pi'}(x) \succeq \mathbf{V}_H^{\pi_v^*}(x)$, then,

$$r_H(x, \pi'(x)) \succeq r_H(x, \pi_v^*(x)) \implies s_v r_H(x, \pi'(x)) \geq s_v r_H(x, \pi_v^*(x)). \quad (4)$$

However, this is only true with equality if $\pi'(x) = \pi_v^*(x)$ for all $x \in \mathcal{S}$, as for any $x \in \mathcal{S}$, $\pi_{v,H}^*(x) = \arg \max[s_v r_H(x, a)] \geq s_v r_H(x, a')$ for any other $a' \in \mathcal{A}$. Therefore, we have that $\pi_H'(x) = \pi_{v,H}^*(x)$ for each $x \in \mathcal{S}$, and that $\mathbf{V}_H^{\pi'}(x) = \mathbf{V}_H^{\pi_v^*}(x)$. This implies that $\mathbb{P}_H \mathbf{V}_H^{\pi'}(x, a) = \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, a)$ for all $x \in \mathcal{S}$ and $a \in \mathcal{A}$. Now, if $\mathbf{V}_{H-1}^{\pi'}(x) \succeq \mathbf{V}_{H-1}^{\pi_v^*}(x)$, then we have that,

$$\begin{aligned} & r_{H-1}(x, \pi'_{H-1}(x)) + \mathbb{E}_{x' \sim \mathbb{P}_H(\cdot|x, \pi'_{H-1}(x))} [\mathbf{V}_H^{\pi'}(x')] \\ & \succeq r_{H-1}(x, \pi_v^*(x)) + \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, \pi_v^*(x)) \\ & \implies r_{H-1}(x, \pi'_{H-1}(x)) + \mathbb{E}_{x' \sim \mathbb{P}_H(\cdot|x, \pi'_{H-1}(x))} [\mathbf{V}_H^{\pi_v^*}(x')] \\ & \succeq r_{H-1}(x, \pi_v^*(x)) + \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, \pi_v^*(x)) \\ & \implies s_v \left(r_{H-1}(x, \pi'_{H-1}(x)) + \mathbb{E}_{x' \sim \mathbb{P}_H(\cdot|x, \pi'_{H-1}(x))} [\mathbf{V}_H^{\pi_v^*}(x')] \right) \\ & \geq s_v \left(r_{H-1}(x, \pi_v^*(x)) + \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, \pi_v^*(x)) \right) \\ & \implies s_v r_{H-1}(x, \pi'_{H-1}(x)) + \mathbb{E}_{x' \sim \mathbb{P}_H(\cdot|x, \pi'_{H-1}(x))} [\mathbf{V}_H^{\pi_v^*}(x')] \\ & \geq s_v r_{H-1}(x, \pi_v^*(x)) + \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, \pi_v^*(x)). \end{aligned}$$

This is true only if $\pi'_{H-1}(x) = \pi_{v,H}^*(x)$ for each $x \in \mathcal{S}$, as $\pi_{v,H}^*$ is the greedy policy with respect to $s_v r_{H-1}(x, a) + \mathbb{P}_H \mathbf{V}_H^{\pi_v^*}(x, a)$. Continuing this argument inductively for $h = H-2, H-3, \dots, 1$ we obtain that $\mathbf{V}_1^{\pi'}(x) \succeq \mathbf{V}_1^{\pi_v^*}(x)$ for each $x \in \mathcal{S}$ only if $\pi' = \pi_v^*$. This is a contradiction as we assumed that $\pi' \neq \pi_v^*$, and hence π_v^* lies in Π^* .

We now prove the other direction for convex Π^* , i.e., that if Π^* is convex, then $\Pi^* \subseteq \Pi_{\Upsilon}^*$ for $\Upsilon = \Delta^n$. This proof proceeds by contradiction as well. Let us assume that there exists a policy π in Π^* that is not present in Π_{Υ}^* . Therefore, there does not exist any $v \in \Delta^n$ such that π maximizes the value function of the scalarized MDP. Alternatively stated, for each $v \in \Upsilon$, there exists another policy $\pi_v^* \in \Pi_{\Upsilon}^*$ such that $\pi_v^* \neq \pi$ and it maximizes the scalarized value function $\mathbf{V}_{v,1}^*$. Now, observe that since $\pi \in \Pi^*$, it must be that for all $\pi' \in \Pi$, $\mathbf{V}_1^{\pi} \succeq \mathbf{V}_1^{\pi'}$. Additionally, since Π^* is convex and the scalarization function $v^\top(\cdot)$ is linear, each scalarization function $s_v(\cdot)$ for $v \in \Delta^n$ is convex over Π^* . Therefore, each policy that maximizes the scalarized value function corresponding to any v is a global optimum in Π^* .

Now, consider the scalarization v_* where $[v_*]_i = \frac{[\mathbf{V}_1^{\pi}]_i}{\|\mathbf{V}_1^{\pi}\|_1} \in \Delta^n$. Now, by our assumption, there must exist an alternative policy $\pi' \neq \pi$ in Π_{Υ}^* , such that (by the convexity of scalarization), $v_*^\top(\mathbf{V}_1^{\pi'} - \mathbf{V}_1^{\pi}) \leq 0$. This implies that $[\mathbf{V}_1^{\pi}]_i^2 \leq [\mathbf{V}_1^{\pi'}]_i [\mathbf{V}_1^{\pi}]_i \implies [\mathbf{V}_1^{\pi'}]_i \geq [\mathbf{V}_1^{\pi}]_i \implies \mathbf{V}_1^{\pi'} \succeq \mathbf{V}_1^{\pi}$. This is a contradiction as π is Pareto-optimal, and hence $\pi \in \Pi_{\Upsilon}^*$. $\square$

D.3. Proof of Proposition 2

We first restate Proposition 2 for clarity.

Proposition 2. For any scalarization s that is Lipschitz and bounded Υ , we have that $\mathfrak{R}_B(T) \leq \frac{1}{T} \mathfrak{R}_C(T) + o(1)$.

Proof. We will follow the approach in (Paria et al., 2020) (Appendix B.3). Recall that Υ is a bounded subset of $\mathbb{R}^n$. Now, we have that since $s_v(\cdot) = v^\top(\cdot)$, we have that s_v is Lipschitz with constant n with respect to the ℓ_1 -norm, i.e., for any $y \in \mathbb{R}^n$,

$$|s_v(y) - s_{v'}(y)| \leq n \|v - v'\|_1. \quad (5)$$

Now, consider the Wasserstein distance conditioned on the history $\mathcal{H}$ between the sampling distribution p_{Υ} on Υ and the empirical distribution $\hat{p}_{\Upsilon}$ corresponding to $\{v_t\}_{t=1}^T$,

$$W_1(p_{\Upsilon}, \hat{p}_{\Upsilon}) = \inf_q \{ \mathbb{E}_q \|X - Y\|_1, q(X) = p_{\Upsilon}, q(Y) = \hat{p}_{\Upsilon} \}, \quad (6)$$

where q is a joint distribution on the RVs X, Y with marginal distributions equal to $p_{\mathbf{r}}$ and $\hat{p}_{\mathbf{r}}$. We therefore have for some randomly drawn samples $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_T$ and for any arbitrary sequence of (joint) policies $\hat{\Pi}_T = \{\pi_1, \dots, \pi_T\}$, for any state $\mathbf{x} \in \mathcal{S}$,

$$\frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} \left[V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}) - \mathbb{E}_{\mathbf{v} \in \mathbf{r}} \left[\max_{\pi \in \hat{\Pi}_T} V_{\mathbf{v},1}^{\pi}(\mathbf{x}) \right] \right] \quad (7)$$

$$\leq \frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} \left[V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}) - \mathbb{E}_{\mathbf{v} \in \mathbf{r}} \left[\max_{\pi \in \hat{\Pi}_T} V_{\mathbf{v},1}^{\pi}(\mathbf{x}) \right] \right] \quad (8)$$

$$\leq \mathbb{E}_{q(X,Y)} \left[\max_{\mathbf{x} \in \mathcal{S}} \left[\max_{\pi \in \hat{\Pi}_T} V_{X,1}^{\pi}(\mathbf{x}) - \max_{\pi \in \hat{\Pi}_T} V_{Y,1}^{\pi}(\mathbf{x}) \right] \right] \quad (9)$$

$$\leq n \cdot \mathbb{E}_{q(X,Y)} [\|X - Y\|_1]. \quad (10)$$

Taking an expectation with respect to $\mathcal{H} = \{\mathbf{v}_1, \dots, \mathbf{v}_T\}$, we have,

$$\mathfrak{R}_B(T) - \frac{1}{T} \mathfrak{R}_C(T) \quad (11)$$

$$= \mathbb{E}_{\mathbf{v} \in \mathbf{r}} \left[\max_{\mathbf{x} \in \mathcal{S}} \left[V_{\mathbf{v},1}^{\star}(\mathbf{x}) - \max_{\pi \in \hat{\Pi}_T} V_{\mathbf{v},1}^{\pi}(\mathbf{x}) \right] \right] - \mathbb{E}_{\mathcal{H}} \left[\frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} [V_{\mathbf{v}_t,1}^{\star}(\mathbf{x}) - V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x})] \right] \quad (12)$$

$$= \mathbb{E}_{\mathcal{H}} \left[\frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} \left[V_{\mathbf{v}_t,1}^{\star}(\mathbf{x}) - \max_{\pi \in \hat{\Pi}_T} V_{\mathbf{v}_t,1}^{\pi}(\mathbf{x}) \right] \right] - \mathbb{E}_{\mathcal{H}} \left[\frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} [V_{\mathbf{v}_t,1}^{\star}(\mathbf{x}) - V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x})] \right] \quad (13)$$

$$\leq \mathbb{E}_{\mathcal{H}} \left[\frac{1}{T} \sum_{t=1}^T \max_{\mathbf{x} \in \mathcal{S}} \left[V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}) - \mathbb{E}_{\mathbf{v} \in \mathbf{r}} \left[\max_{\pi \in \hat{\Pi}_T} V_{\mathbf{v},1}^{\pi}(\mathbf{x}) \right] \right] \right] \quad (14)$$

$$\leq n \cdot \mathbb{E}_{q(X,Y)} [\|X - Y\|_1]. \quad (15)$$

The penultimate inequality follows from $\max$ being a contraction mapping in bounded domains, and the final inequality follows from the previous analysis. To complete the proof, we first take an infimum over q and observe that the subsequent RHS converges at a rate of $\tilde{O}(T^{-\frac{1}{n}})$ under mild regulatory conditions, as shown by [Canas and Rosasco \(2012\)](#). $\square$

D.4. Proof of Lemma 2

Follows from Lemma 11.

D.5. Proof of Theorem 2

The proof for this result is to essentially solve the approximate scalarized MDP for each clique. We first present a vector-valued concentration result.

Lemma 4. *Select any clique C in a clique covering $\hat{\mathcal{C}}$ such that $|C| = M$. For any $m \in [M], h \in [H]$ and $t \in [T]$, let k_t denote the episode after which the last local synchronization has taken place, and $\mathbf{S}_C^h(t)$ and $\mathbf{\Lambda}_C^h(t)$ be defined as follows.*

$$\begin{aligned} \mathbf{S}_C^h(t) &= \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \left[\mathbf{v}_{\mathbf{v},h+1}^t(\mathbf{x}_C^{h+1}(\tau)) - (\tilde{\mathbb{P}}_h^C \mathbf{v}_{\mathbf{v},h+1}^t)(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \right], \\ \mathbf{\Lambda}_C^h(t) &= \lambda \mathbf{I}_d + (\Phi_C^h(k_t))^{\top} (\Phi_C^h(k_t)). \end{aligned}$$

Where $\mathbf{v}_{\mathbf{v},h+1}^t(\mathbf{x}) = \mathbf{1}_M \cdot V_{\mathbf{v},h+1}^t(\mathbf{x}) \forall \mathbf{x} \in \mathcal{S}_C$, $\mathbf{1}_M$ denotes the all-ones vector in $\mathbb{R}^M$, and C_{β} is the constant such that $\beta_C^h(t) = C_{\beta} \cdot dH \sqrt{\log(TM H)}$. Then, there exists a constant B such that with probability at least $1 - \delta$,

$$\sup_{\mathbf{v}_C \in \mathbf{r}_C} \|\mathbf{S}_C^h(t)\|_{(\mathbf{\Lambda}_C^h(t))^{-1}} \leq B \cdot dH \sqrt{2 \log \left(\frac{(C_{\beta} + 2)dMTH}{\delta} \right)}.$$

Proof. The proof is done in two steps. The first step is to bound the deviations in $\mathbf{S}$ for any fixed function V by a martingale concentration. The second step is to bound the resulting concentration over all functions V by a covering argument. Finally, we select appropriate constants to provide the form of the result required.

Step 1. Recall that $\mathbf{S}_C^h(t) = \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [V_{\mathbf{v},h+1}^t(\mathbf{x}_C^{h+1}(\tau)) - (\tilde{\mathbb{P}}_h^C V_{\mathbf{v},h+1}^t)(\mathbf{z}_C^h(\tau))]$, where $\mathbf{v}_{\mathbf{v},h+1}^t$ is the vector with each entry being $V_{\mathbf{v},h+1}^t$. We have that $V_{\mathbf{v},h+1}^t(\mathbf{x}_C^{h+1}(\tau)) - (\tilde{\mathbb{P}}_h^C V_{\mathbf{v},h+1}^t)(\mathbf{z}_C^h(\tau)) = \mathbf{v}_{\mathbf{v},h+1}^t - \tilde{\mathbb{P}}_h^C \mathbf{v}_{\mathbf{v},h+1}^t$. Consider the following distance metric dist_{Υ_C} ,

$$\text{dist}_{\Upsilon_C}(\mathbf{v}, \mathbf{v}') = \sup_{\mathbf{x} \in \mathcal{S}_C, \mathbf{v} \in \Upsilon_C} \|\mathbf{v}(\mathbf{x}) - \mathbf{v}'(\mathbf{x})\|_1. \quad (16)$$

Let $\mathcal{V}_{\Upsilon_C}$ be the family of all vector-valued UCB value functions that can be produced by the algorithm on clique C , and now let $\mathcal{N}_\epsilon$ be an ϵ -covering of $\mathcal{V}_{\Upsilon_C}$ under dist_{Υ_C} , i.e., for every $\mathbf{v} \in \mathcal{V}_{\Upsilon_C}$, there exists $\mathbf{v}' \in \mathcal{N}_\epsilon$ such that $\text{dist}_{\Upsilon_C}(\mathbf{v}, \mathbf{v}') \leq \epsilon$. Now, here again, we adopt a similar strategy as the independent case. To bound the RHS, we decompose $\mathbf{S}_C^h(t)$ in terms of the covering described earlier. We know that since $\mathcal{N}_\epsilon$ is an ϵ -covering of $\mathcal{V}_{\Upsilon_C}$, there exists a $\mathbf{v}' \in \mathcal{N}_\epsilon$ and $\Delta = \mathbf{v}_{\mathbf{v},h+1}^t - \mathbf{v}'$ such that,

$$\mathbf{S}_C^h(t) = \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [\mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) - \tilde{\mathbb{P}}_h^C \mathbf{v}'(\mathbf{z}_C^h(\tau))] + \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [\Delta(\mathbf{x}_C^{h+1}(\tau)) - \tilde{\mathbb{P}}_h^C \Delta(\mathbf{z}_C^h(\tau))].$$

Now, observe that by the definition of the covering, we have that $\|\Delta\|_1 \leq \epsilon$. Therefore, we have that $\|\Delta(\mathbf{x})\|_{(\Lambda_C^h(t))^{-1}} \leq \epsilon/\sqrt{\lambda}$, and $\|\tilde{\mathbb{P}}_h^C \Delta(\mathbf{z})\|_{(\Lambda_C^h(t))^{-1}} \leq \epsilon/\sqrt{\lambda}$ for all $\mathbf{z} \in \mathcal{Z}, \mathbf{x} \in \mathcal{S}, h \in [H]$. Therefore, since $\|\Phi_C(\mathbf{z})\|_2 \leq \sqrt{M}$,

$$\|\mathbf{S}_C^h(t)\|_{(\Lambda_C^h(t))^{-1}}^2 \leq 2 \left\| \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [\mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) - \tilde{\mathbb{P}}_h^C \mathbf{v}'(\mathbf{z}_C^h(\tau))] \right\|_{(\Lambda_C^h(t))^{-1}} + \left\| \sum_{\tau=1}^{k_t} \Phi_C^T \bar{\epsilon}_\tau \right\|_{(\Lambda_C^h(t))^{-1}} + \frac{8Mt^2\epsilon^2}{\lambda}.$$

Here $\bar{\epsilon}_\tau$ denotes the maximum misspecification incurred from observing $\tilde{\mathbb{P}}_h^C$ instead of the true $\mathbb{P}_h$. By a standard argument from misspecified bandits (see, e.g., (Ghosh et al., 2017)), using the fact that Φ_C has maximum norm $\sqrt{M}$ and the misspecification is bounded by $2\epsilon(k)$, we can bound the second term by $\epsilon(k) \cdot \sqrt{dtM \log \left(\frac{\det(\Lambda_C^h(t))}{\lambda \mathbf{I}_d} \right)}$. To bound the first term on the RHS, we consider the substitution $\epsilon_{\tau,h}^t = \mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) - \tilde{\mathbb{P}}_h^C \mathbf{v}'(\mathbf{z}_C^h(\tau))$. To bound the first term on the RHS, we consider the filtration $\{\mathcal{F}_\tau\}_{\tau=0}^\infty$ where $\mathcal{F}_0$ is empty, and $\mathcal{F}_\tau = \sigma \left(\bigcup_{i \leq \tau} (\mathbf{x}_{h+1}^i, \Phi_C(\mathbf{z}_h^i)) \right)$, and σ denotes the σ -algebra generated by a finite set. Then, we have that,

$$\begin{aligned} & \left\| \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [\mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) - \tilde{\mathbb{P}}_h^C \mathbf{v}'(\mathbf{z}_C^h(\tau))] \right\|_{(\Lambda_C^h(t))^{-1}} \\ &= \left\| \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) [\mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) - \mathbb{E}[\mathbf{v}'(\mathbf{x}_C^{h+1}(\tau)) | \mathcal{F}_{\tau-1}]] \right\|_{(\Lambda_C^h(t))^{-1}} = \left\| \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) \epsilon_{\tau,h}^t \right\|_{(\Lambda_C^h(t))^{-1}}. \end{aligned}$$

Note that for each $\epsilon_{\tau,h}^t$, each entry is bounded by $2H$, and therefore we have that the vector $\epsilon_{\tau,h}^t$ is H -sub-Gaussian. Then, applying Lemma 14, we have that,

$$\left\| \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{z}_C^h(\tau)) \epsilon_{\tau,h}^t \right\|_{(\Lambda_C^h(t))^{-1}} \leq H^2 \log \left(\frac{\det(\Lambda_h^t)}{\det(\lambda \mathbf{I}_d) \delta^2} \right) \leq H^2 \log \left(\frac{\det(\bar{\Lambda}_h^t)}{\det(\lambda \mathbf{I}_d) \delta^2} \right). \quad (17)$$

Replacing this result for each $\mathbf{v} \in \mathcal{N}_\epsilon$, we have by a union bound over each $t \in [T], h \in [H]$, we have with probability at least $1 - \delta$, simultaneously for each $t \in [T], h \in [H]$,

$$\sup_{\mathbf{v}_t \in \Upsilon, \mathbf{v} \in \mathcal{V}_\Upsilon} \|\mathbf{S}_C^h(t)\|_{(\Lambda_C^h(t))^{-1}} \leq 2H \sqrt{\log \left(\frac{\det(\bar{\Lambda}_h^t)}{\det(\lambda \mathbf{I}_d)} \right) + \log \left(\frac{HT|\mathcal{N}_\epsilon|}{\delta} \right) + \frac{2Mt^2\epsilon^2}{H^2\lambda}} \quad (18)$$

$$\leq 2H \sqrt{d \log \left(\frac{Mt + \lambda}{\lambda} \right) + \log \left(\frac{|\mathcal{N}_\epsilon|}{\delta} \right) + \log(HT) + \frac{2Mt^2\epsilon^2}{H^2\lambda}}. \quad (19)$$

The last step follows once again by first noticing that $\|\Phi_C(\cdot)\| \leq \sqrt{M}$ and then applying an AM-GM inequality, and then using the determinant-trace inequality.

Step 2. Here $\mathcal{N}_\epsilon$ is an ϵ -covering of the function class $\mathcal{V}_{\Upsilon_C}$ for any $h \in [H], m \in [M]$ or $t \in [T]$ under the distance function $\text{dist}_{\Upsilon_C}(\mathbf{v}, \mathbf{v}') = \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{v} \in \Upsilon} \|\mathbf{v}(\mathbf{x}) - \mathbf{v}'(\mathbf{x}')\|_1$. To bound this quantity by the appropriate covering number, we first observe that for any $V \in \mathcal{V}_{\Upsilon_C}$, we have that the policy weights are bounded as $2HM\sqrt{dT/\lambda}$ (Lemma 12). Therefore, by Lemma 10 we have for any constant B such that $\beta_h^t \leq B$,

$$\log(\mathcal{N}_\epsilon) \leq d \cdot \log \left(1 + \frac{8HM^3}{\epsilon} \sqrt{\frac{dT}{\lambda}} \right) + d^2 \log \left(1 + \frac{8Md^{1/2}B^2}{\lambda\epsilon^2} \right). \quad (20)$$

Recall that we select the hyperparameters $\lambda = 1$ and $\beta = \mathcal{O}(dH\sqrt{\log(TM H)})$, and to balance the terms in $\bar{\beta}_C^h(t)$ we select $\epsilon = \epsilon^* = dH/\sqrt{MT^2}$. Finally, we obtain that for some absolute constant C_β , by replacing the above values,

$$\log(\mathcal{N}_\epsilon) \leq d \cdot \log \left(1 + \frac{8M^{7/2}T^{3/2}}{d^{1/2}} \right) + d^2 \log \left(1 + 8C_\beta d^{1/2}MT^2 \log(TM H) \right). \quad (21)$$

Therefore, for some absolute constant C' independent of M, T, H, d and C_β , we have,

$$\log|\mathcal{N}_\epsilon| \leq C'd^2 \log(C_\beta \cdot dMT \log(TM H)). \quad (22)$$

Replacing this result in the result from Step 1, we have that with probability at least $1 - \delta'/2$ for all $t \in [T], h \in [H]$ simultaneously,

$$\begin{aligned} \|\mathbf{S}_C^h(t)\|_{(\Lambda_C^h(t))^{-1}}^2 &\leq 2H \left((d + 2 + \varepsilon(k)dMT) \log \frac{MT + \lambda}{\lambda} + 2 \log \left(\frac{1}{\delta'} \right) \right. \\ &\quad \left. + C'd^2 \log(C_\beta \cdot dMT \log(TM H)) + 2 + 4 \log(TH) \right). \end{aligned}$$

This implies that there exists an absolute constant B independent of M, T, H, d and C_β , such that, with probability at least $1 - \delta'/2$ for all $t \in [T], h \in [H], \mathbf{v}_C \in \Upsilon_C$ simultaneously,

$$\|\mathbf{S}_C^h(t)\|_{(\Lambda_C^h(t))^{-1}} \leq B \cdot (dH + \varepsilon(k)H\sqrt{dMT}) \sqrt{2 \log \left(\frac{(C_\beta + 2)dMTH}{\delta'} \right)}. \quad (23)$$

□

Next, we present the key result for cooperative value iteration, which demonstrates that for any agent the estimated Q -values have bounded error for any policy π .

Lemma 5. Fix a clique $C \in \hat{\mathcal{C}}$ such that $|C| = M$. For each C , there exists an absolute constant c_β such that for $\beta_C^h(t) = c_\beta \cdot (dH + \varepsilon(k) \cdot \sqrt{dM}) \sqrt{\log(2dM H t / \delta')}$ for any policy π , there exists a constant C'_β such that for each $x \in \mathcal{S}, a \in \mathcal{A}$ we have for all $m \in C, t \in [T], h \in [H]$ simultaneously, with probability at least $1 - \delta'/2$,

$$\begin{aligned} |\langle \Phi_C(\mathbf{x}_C, \mathbf{a}_C), \mathbf{w}_{m,h}^t - \mathbf{w}_h^\pi \rangle| &\leq \mathbb{P}_h(V_{m,h+1}^t - V_{m,h+1}^\pi)(\mathbf{x}, \mathbf{a}) + 4H\varepsilon(k) \\ &\quad + C'_\beta \cdot dH \cdot \|\Phi_C(\mathbf{z}_C)\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)}. \end{aligned}$$

Proof. By the Bellman equation and Assumptions 1, 2, 3, we have that for any policy π , and $\mathbf{v}_C \in \Upsilon_C$, there exist weights $\mathbf{w}_{\mathbf{v}_C, h}^\pi$ such that, for all $\mathbf{z} = \{\mathbf{z}_C, \bar{\mathbf{z}}_C\} \in \mathcal{Z} = \mathcal{S} \times \mathcal{A}$,

$$\mathbf{v}_C^\top \Phi_C(\mathbf{z}_C)^\top \mathbf{w}_{\mathbf{v}_C, h}^\pi = \mathbf{v}_C^\top \tilde{\mathbf{r}}_h^C(\mathbf{z}_C) + \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi(\mathbf{z}) = \mathbf{v}_C^\top \left(\tilde{\mathbf{r}}_h^C(\mathbf{z}) + \mathbf{1}_M \cdot \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi(\mathbf{z}) \right).$$

We have,

$$\mathbf{w}_{\mathbf{v}_C, h}^t - \mathbf{w}_{\mathbf{v}_C, h}^\pi \quad (24)$$

$$= (\Lambda_C^h(t))^{-1} \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [r_h(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) + \mathbf{1}_M \cdot V_{\mathbf{v}_C, h+1}^t(\mathbf{x}_\tau)]] - \mathbf{w}_{\mathbf{v}_C, h}^\pi \quad (25)$$

$$= (\Lambda_C^h(t))^{-1} \left\{ -\lambda \mathbf{w}_{\mathbf{v}_C, h}^\pi + \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [\mathbf{1}_M \cdot (V_{\mathbf{v}_C, h+1}^t(\mathbf{x}'_\tau) - \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)))] \right\}. \quad (26)$$

$$\begin{aligned} \mathbf{w}_{\mathbf{v}_C, h}^t - \mathbf{w}_{\mathbf{v}_C, h}^\pi &= \underbrace{-\lambda (\Lambda_C^h(t))^{-1} \mathbf{w}_{\mathbf{v}_C, h}^\pi}_{\mathbf{v}_1} \\ &\quad + \underbrace{(\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [\mathbf{1}_M \cdot (V_{\mathbf{v}_C, h+1}^t(\mathbf{x}'_\tau) - \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^t(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)))] \right\}}_{\mathbf{v}_2} \\ &\quad + \underbrace{(\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [\mathbf{1}_M \cdot (\tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^t - \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi)(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))]] \right\}}_{\mathbf{v}_3} + \\ &\quad + \underbrace{(\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [\mathbf{1}_M \cdot (\mathbb{P}_h V_{\mathbf{v}_C, h+1}^t - \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^t)(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))]] \right\}}_{\mathbf{v}_4} \\ &\quad + \underbrace{(\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} [\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [r_{[C]}(\mathbf{x}_h(t), \mathbf{a}_h(t)) - \tilde{r}_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))]] \right\}}_{\mathbf{v}_5}. \quad (27) \end{aligned}$$

Now, we know that for any $\mathbf{z} \in \mathcal{Z}$ for any policy π ,

$$\begin{aligned} &\|\langle \Phi_C(\mathbf{z}), \mathbf{v}_1 \rangle\|_2 \\ &\leq \lambda \|\langle \Phi_C(\mathbf{z}), (\Lambda_C^h(t))^{-1} \mathbf{w}_{\mathbf{v}_C, h}^\pi \rangle\|_2 \leq \lambda \cdot \|\mathbf{w}_{\mathbf{v}_C, h}^\pi\| \|\Phi_C(\mathbf{z})\|_{(\Lambda_C^h(t))^{-1}} \leq 2HM\lambda\sqrt{d} \cdot \|\Phi_C(\mathbf{z})\|_{(\Lambda_C^h(t))^{-1}} \end{aligned}$$

Here the last inequality follows from Lemma 11. For the second term, we have by Lemma 4 that there exists an absolute constant C_β , independent of M, T, H, d such that, with probability at least $1 - \delta'/2$ for all $t \in [T], h \in [H], \mathbf{v}_C \in \Upsilon_C$ simultaneously,

$$\|\langle \Phi_C(\mathbf{z}), \mathbf{v}_2 \rangle\|_2 \leq \|\Phi_C(\mathbf{z})\|_{(\Lambda_C^h(t))^{-1}} \cdot C_\beta \cdot dH \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)}. \quad (28)$$

For the third term, note that,

$$\langle \Phi_C(\mathbf{x}, \mathbf{a}), \mathbf{v}_3 \rangle \quad (29)$$

$$= \left\langle \Phi_C(\mathbf{z}), (\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) [\mathbf{1}_M \cdot (\mathbb{P}_h V_{\mathbf{v}_C, h+1}^t - \mathbb{P}_h V_{\mathbf{v}_C, h+1}^\pi)(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))] \right\} \right\rangle \quad (30)$$

$$= \left\langle \Phi_C(\mathbf{z}), (\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))^\top \int (V_{\mathbf{v}_C, h+1}^t - V_{\mathbf{v}_C, h+1}^\pi)(\mathbf{x}') d\mu_h(\mathbf{x}') \right\} \right\rangle \quad (31)$$

$$= \left\langle \Phi_C(\mathbf{z}), (\Lambda_C^h(t))^{-1} \left\{ \sum_{\tau=1}^{k_t} \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau))^\top \int (V_{\mathbf{v}_C, h+1}^t - V_{\mathbf{v}_C, h+1}^\pi)(\mathbf{x}') d\mu_h(\mathbf{x}') \right\} \right\rangle \quad (32)$$

$$= \left\langle \Phi_C(z), \int (V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x') d\mu_h(x') \right\rangle - \lambda \left\langle \Phi_C(z), (\Lambda_h^t)^{-1} \int (V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x') d\mu_h(x') \right\rangle \quad (33)$$

$$= \int (V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x') \langle \Phi_C(z), \mu_h(x') \rangle - \lambda \left\langle \Phi_C(z), (\Lambda_h^t)^{-1} \int (V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x') d\mu_h(x') \right\rangle \quad (34)$$

$$= \mathbf{1}_M \cdot \left(\tilde{\mathbb{P}}_h^C(V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x, a) \right) - \lambda \left\langle \Phi_C(z), (\Lambda_h^t)^{-1} \int (V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x') d\mu_h(x') \right\rangle \quad (35)$$

$$\leq \mathbf{1}_M \cdot \left(\tilde{\mathbb{P}}_h^C(V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x, a) + 2H\sqrt{d\lambda} \|\Phi_C(x, a)\|_{(\Lambda_C^h(t))^{-1}} \right). \quad (36)$$

For the last two terms, we can bound them by a similar argument of misspecification as Lemma 4. We can bound both terms by $\mathbf{1}_M \cdot \left(\varepsilon(k) \cdot H\sqrt{dMT} \|\Phi_C(x, a)\|_{(\Lambda_C^h(t))^{-1}} \right)$. Putting it all together, we have that since $\langle \Phi_C(x, a), w_{v_C, h}^t - w_{v_C, h}^\pi \rangle = \langle \Phi_C(x, a), v_1 + v_2 + v_3 + v_4 + v_5 \rangle$, there exists an absolute constant C_β independent of M, T, H, d , such that, with probability at least $1 - \delta'/2$ for all $t \in [T], h \in [H], v_C \in \Upsilon_C$ simultaneously,

$$\begin{aligned} |\langle v_C^\top \Phi_C(x, a), w_{v_C, h}^t - w_{v_C, h}^\pi \rangle| &\leq v_C^\top \mathbf{1}_M \cdot \left(\mathbb{P}_h(V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x, a) \right) + 4H\varepsilon(k) \\ &\quad + \|\Phi_C(x, a)\|_{(\Lambda_C^h(t))^{-1}} \|v_C\|_2 \left(2H\sqrt{d\lambda} + C_\beta \cdot dH \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)} + 2HM\lambda\sqrt{d} + 2H\varepsilon(k)\sqrt{dMT} \right) \end{aligned} \quad (37)$$

Since $\lambda \leq 1$ and $\|v_C\|_2 \leq 1$, there exists a constant C'_β that we have the following for any $(x, a) \in \mathcal{S} \times \mathcal{A}$ with probability $1 - \delta'/2$ simultaneously for all $h \in [H], v_C \in \Upsilon_C, t \in [T]$,

$$\begin{aligned} |\langle v_C^\top \Phi_C(x, a), w_{v_C, h}^t - w_{v_C, h}^\pi \rangle| &\leq \mathbb{P}_h(V_{v_C, h+1}^t - V_{v_C, h+1}^\pi)(x, a) + 4H\varepsilon(k) \\ &\quad + C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \|\Phi(z)\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)}. \end{aligned}$$

□

Lemma 6 (UCB in the Multiagent Setting). *For each $C \in \hat{\mathcal{C}}$, with probability at least $1 - \delta'/2$, we have that for all $(x_C, a_C, h, t, v_C) \in \mathcal{S}_C \times \mathcal{A}_C \times [H] \times [T] \times \Upsilon_C$,*

$$Q_{v, h}^t(x_C, a_C) \geq Q_{v, h}^*(x_C, a_C) - 4H(H+1-h)\varepsilon(k).$$

Proof. We prove this result by induction. First, for the last step H , note that the statement holds as $Q_{v_C, H}^t(x_C, a_C) \geq Q_{v_C, H}^*(x_C, a_C) - 4H\varepsilon(k)$ for all v_C . Recall that the value function at step $H+1$ is zero. Therefore, by Lemma 5, we have that, for any $v_C \in \Upsilon_C$,

$$\begin{aligned} |\langle v_C^\top \Phi_C(x_C, a_C), w_{v_C, H}^t - Q_{v_C, H}^*(x_C, a_C) \rangle| \\ \leq C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \|\Phi(z_C)\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)} + 4H\varepsilon(k). \end{aligned}$$

We have $Q_{v_C, H}^*(x_C, a_C) \leq \langle v_C^\top \Phi_C(x_C, a_C), w_{v_C, H}^t \rangle + C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \|\Phi(z)\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)} = Q_{v_C, H}^t$. Now, for the inductive case, we have by Lemma 5 for any $h \in [H], v_C \in \Upsilon_C$,

$$\begin{aligned} |\langle v_C^\top \Phi_C(x_C, a_C), w_{v_C, h}^t - w_{v_C, h}^* \rangle - (\mathbb{P}_h V_{v_C, h+1}^*(x_C, a_C) - \mathbb{P}_h V_{v_C, h+1}^t(x_C, a_C))| \\ \leq C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \|\Phi(z)\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)}. \end{aligned}$$

By the inductive assumption we have $Q_{v_C, h+1}^t(x_C, a_C) \geq Q_{v_C, h+1}^*(x_C, a_C)$ implying $\mathbb{P}_h V_{v_C, h+1}^*(x_C, a_C) - \mathbb{P}_h V_{v_C, h+1}^t(x_C, a_C) \geq 0$. Substituting the appropriate Q value formulations we have,

$$Q_{\mathbf{v}_C, h}^* \leq \langle \mathbf{v}_C^\top \Phi_C(\mathbf{x}_C, \mathbf{a}_C), \mathbf{w}_{\mathbf{v}_C, h}^t \rangle + 4H\varepsilon(k) \\ + C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \|\Phi(\mathbf{z})\|_{(\Lambda_C^h(t))^{-1}} \cdot \sqrt{2 \log \left(\frac{dMTH}{\delta'} \right)} = Q_{\mathbf{v}_C, h}^t(\mathbf{x}_C, \mathbf{a}_C).$$

□

Lemma 7 (Recursive Relation in Multiagent MDP Settings). *Fix a clique $C \in \widehat{\mathcal{C}}$ of size M . For any $\mathbf{v}_C \in \Upsilon_C$, let $\delta_{\mathbf{v}_C, h}^t = V_{\mathbf{v}_C, h}^t(\mathbf{x}_C^h(t)) - V_{\mathbf{v}_C, h}^{\pi_t}(\mathbf{x}_C^h(t))$, and $\xi_{\mathbf{v}_C, h+1}^t = \mathbb{E} \left[\delta_{\mathbf{v}_C, h}^t | \mathbf{x}_C^h(t), \mathbf{a}_C^h(t) \right] - \delta_{\mathbf{v}_C, h}^t$. Then, with probability at least $1 - \alpha$, for all $(t, h) \in [T] \times [H]$ simultaneously,*

$$\delta_{\mathbf{v}_C, h}^t \leq \delta_{\mathbf{v}_C, h+1}^t + \xi_{\mathbf{v}_C, h+1}^t + 4H\varepsilon(k) \\ + 2 \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\Lambda_C^h(t))^{-1}} \cdot C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \sqrt{2 \log \left(\frac{dMTH}{\alpha} \right)}.$$

Proof. By Lemma 5, we have that for any $(\mathbf{x}_C, \mathbf{a}_C, h, \mathbf{v}_C, t) \in \mathcal{S}_C \times \mathcal{A}_C \times [H] \times \Upsilon_C \times [T]$ with probability at least $1 - \alpha/2$,

$$Q_{\mathbf{v}_C, h}^t(\mathbf{x}_C, \mathbf{a}_C) - Q_{\mathbf{v}_C, h}^{\pi_t}(\mathbf{x}_C, \mathbf{a}_C) \leq \mathbb{P}_h(V_{\mathbf{v}_C, h+1}^t - V_{\mathbf{v}_C, h}^{\pi_t})(\mathbf{x}_C, \mathbf{a}_C) + 4H\varepsilon(k) \\ + 2 \|\Phi_C(\mathbf{x}_C, \mathbf{a}_C)\|_{(\Lambda_C^h(t))^{-1}} \cdot C'_\beta \cdot \left(dH + \varepsilon(k)H\sqrt{dMT} \right) \cdot \sqrt{2 \log \left(\frac{dMTH}{\alpha} \right)}.$$

Replacing the definition of $\delta_{\mathbf{v}_C, h}^t$ and $V_{\mathbf{v}_C, h}^{\pi_t}$ finishes the proof. □

Lemma 8. *For each clique $C \in \widehat{\mathcal{C}}$ and each $\xi_{\mathbf{v}_C, h}^t$ as defined earlier and any $\delta \in (0, 1)$, we have that with probability at least $1 - \delta/2$,*

$$\sum_{t=1}^T \sum_{h=1}^H \sum_{C \in \widehat{\mathcal{C}}} \xi_{\mathbf{v}_C, h}^t \leq \sqrt{2H^3 T |\widehat{\mathcal{C}}| \log \left(\frac{2}{\alpha} \right)}. \quad (38)$$

Proof. Observe that following the reasoning in Theorem 3.1 of Jin et al. (2020), we can see that $\{\xi_{\mathbf{v}_C, h}^t\}_{h, t, C}$ is a martingale difference sequence (computation within each clique at any instant is independent of the current state of other cliques). Furthermore, since $|\xi_{\mathbf{v}_C, h}^t| \leq H$ regardless of $\mathbf{v}_C$, which allows us to apply Azuma-Hoeffding inequality. We have, for any $t > 0$,

$$\mathbb{P} \left(\sum_{t=1}^T \sum_{h=1}^H \sum_{C \in \widehat{\mathcal{C}}} \xi_{\mathbf{v}_C, h}^t > t \right) \leq \exp \left(-\frac{t^2}{2T|\widehat{\mathcal{C}}|H^2} \right).$$

Rearranging provides us the final result. □

We are now ready to prove Theorem 2. We first restate the Theorem for completeness.

Theorem 2. Algorithm 1 when run on a game with n agents satisfying Assumptions 1, 2, 3 with error ε_* , approximate clique covering $\widehat{\mathcal{C}}$, and $\kappa \cdot dH \cdot \theta(G)$ rounds of communication for some $\kappa > 1$, $\beta_C^h(t) = \mathcal{O}(H\sqrt{d \log(ntH)} + \varepsilon_* \sqrt{dT}) \forall C \in \widehat{\mathcal{C}}$ obtains, with probability at least $1 - \alpha$, regret:

$$\mathfrak{R}_C(T) = \widetilde{\mathcal{O}} \left(\theta(G) \cdot d^{\frac{3}{2}} H^2 \cdot \max \left\{ 1, (2T \cdot n_{\max})^{\frac{2}{\kappa}} \right\} \left(\sqrt{T \log \left(\frac{1}{\alpha} \right)} + 2T \cdot \varepsilon_* \right) \right).$$

Where $\theta(G)$ denotes the clique covering number of G , and $n_{\max}$ is the size of the largest clique in $\widehat{\mathcal{C}}$.

Proof. We have by the definition of cumulative regret:

$$\mathfrak{R}_C(T) = \sum_{t=1}^T \mathbb{E}_{\mathbf{v}_t \sim \Upsilon} \left[\max_{\mathbf{x}_1^t \in \mathcal{S}} [V_{\mathbf{v}_t,1}^*(\mathbf{x}_1^t) - V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}_1^t)] \right] = \mathbb{E}_{\mathbf{v}_t \sim \Upsilon} \left[\sum_{t=1}^T \max_{\mathbf{x}_1^t \in \mathcal{S}} [V_{\mathbf{v}_t,1}^*(\mathbf{x}_1^t) - V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}_1^t)] \right]. \quad (39)$$

Our analysis focuses only on the term inside the expectation, which we will bound via terms that are independent of $\mathbf{v}_1, \dots, \mathbf{v}_T$, bounding $\mathfrak{R}_C$. We bound the cumulative regret incurred by each clique, summing over which gives us the cumulative regret.

$$\sum_{t=1}^T \max_{\mathbf{x}_1^t \in \mathcal{S}} [V_{\mathbf{v}_t,1}^*(\mathbf{x}_1^t) - V_{\mathbf{v}_t,1}^{\pi_t}(\mathbf{x}_1^t)] \leq \sum_{C \in \hat{\mathcal{C}}} \left(\sum_{t=1}^T \max_{\mathbf{x}_C \in \mathcal{S}_C} [V_{\mathbf{v}_t,C,1}^*(\mathbf{x}_C) - V_{\mathbf{v}_t,C,1}^{\pi_t,C}(\mathbf{x}_C)] \right).$$

We can bound the clique-wise regret for any $C \in \hat{\mathcal{C}}$ of size M as follows.

$$\begin{aligned} \sum_{t=1}^T \max_{\mathbf{x}_C \in \mathcal{S}_C} [V_{\mathbf{v}_t,C,1}^*(\mathbf{x}_C) - V_{\mathbf{v}_t,C,1}^{\pi_t,C}(\mathbf{x}_C)] &\leq \sum_{t=1}^T \max_{\mathbf{x}_C \in \mathcal{S}_C} \delta_{\mathbf{v}_t,C,1}^t + 4HT\varepsilon(k) \\ &\leq \sum_{t,h}^{T,H} \xi_{\mathbf{v}_t,C,h}^t + 2C'_\beta \cdot (dH + \varepsilon(k)H\sqrt{dMT}) \cdot \sqrt{2 \log \left(\frac{dMTH}{\alpha} \right)} \left(\sum_{t,h}^{T,H} \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\bar{\Lambda}_C^h(t))^{-1}} \right) + 4H\varepsilon(k). \end{aligned}$$

Where the last inequality holds with probability at least $1 - \alpha/2$, via Lemma 7 and Lemma 6. To bound the second summation, we can use the technique in Theorem 4 of Abbasi-Yadkori et al. (2011). Assume that the last time synchronization of rewards occurred was at instant k_T . We therefore have, by Lemma 12 of Abbasi-Yadkori et al. (2011), for any $h \in [H]$

$$\begin{aligned} \sum_{t=1}^T \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\Lambda_C^h(t))^{-1}} &\leq \frac{\det(\bar{\Lambda}_C^h(t))}{\det(\Lambda_C^h(t))} \sum_{t=1}^T \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\bar{\Lambda}_C^h(t))^{-1}} \\ &\leq \sqrt{S} \sum_{t=1}^T \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\bar{\Lambda}_C^h(t))^{-1}} \end{aligned}$$

Here $\bar{\Lambda}_C^h(t) = \sum_{t=1}^T \Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t)) \Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))^\top$ and the last inequality follows from the algorithms' synchronization condition. Replacing this result, we have that,

$$\begin{aligned} \sum_{t=1}^T \sum_{h=1}^H \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\Lambda_C^h(t))^{-1}} &\leq 2 \sum_{h=1}^H \left(\sqrt{S} \sum_{t=1}^T \|\Phi_C(\mathbf{x}_C^h(t), \mathbf{a}_C^h(t))\|_{(\bar{\Lambda}_C^h(t))^{-1}} \right) \\ &\leq 2H \sqrt{ST \cdot d \log \frac{MT + \lambda}{\lambda}}. \end{aligned}$$

Where the last inequality is an application of Lemma 13 and using the fact that $\|\Phi_C(\cdot)\|_2 \leq \sqrt{M}$. Replacing this result, we have that with probability at least $1 - \alpha/2$, by a union bound over all cliques in $\hat{\mathcal{C}}$,

$$\begin{aligned} \sum_{C \in \hat{\mathcal{C}}} \left(\sum_{t=1}^T \max_{\mathbf{x}_C \in \mathcal{S}_C} [V_{\mathbf{v}_t,C,1}^*(\mathbf{x}_C) - V_{\mathbf{v}_t,C,1}^{\pi_t,C}(\mathbf{x}_C)] \right) \\ \leq \sum_{t,h,C}^{T,H,\hat{\mathcal{C}}} \xi_{\mathbf{v}_t,C,h}^t + 2C'_\beta \cdot H^2 \left(d + \varepsilon(k)\sqrt{dMT} \right) \cdot \sqrt{2ST \log \left(\frac{dMTH|\hat{\mathcal{C}}|}{\alpha} \right)} \cdot d \log \frac{MT + \lambda}{\lambda} + 4HT|\hat{\mathcal{C}}|\varepsilon(k). \end{aligned}$$

We can bound the second term via Lemma 8. Taking expectation of the RHS over $\mathbf{v}_1, \dots, \mathbf{v}_T$ and rewriting S in terms of κ via Lemma 3 gives us the final result (the $\tilde{\mathcal{O}}$ notation hides polylogarithmic factors). $\square$

D.6. Proof of Lemma 3

Proof. Let the total rounds of communication triggered by the threshold condition in any step $h \in [H]$ in any clique C of size M be given by $n_h(C)$. Then, we have, by the communication criterion,

$$S^{n_h(C)} < \frac{\det(\Lambda_C^h(t))}{\det(\lambda \mathbf{I}_d)} \leq (1 + MT/d)^d. \quad (40)$$

Where the last inequality follows from Lemma 13 and the fact that $\|\Phi\| \leq \sqrt{M} \leq \sqrt{n_{\max}}$. This gives us that $n_h(C) \leq d \log_S(1 + n_{\max}T/d) + 1$. Furthermore, by noticing that $\gamma \leq \sum_{C \in \hat{\mathcal{C}}} \sum_{h=1}^H n_h(C)$, and that $|\hat{\mathcal{C}}| \leq 1.25 \cdot \theta(G)$, we have the final result. $\square$

E. Auxiliary Results

E.1. Covering Number Bounds

Lemma 9 (Covering Number of the Euclidean Ball). *For any $\varepsilon > 0$, the ε -covering number of the Euclidean ball in $\mathbb{R}^d$ with radius $R > 0$ is less than $(1 + 2R/\varepsilon)^d$.*

Lemma 10 (Covering number for Markov game UCB-style functions). *Let $\mathcal{V}$ denote a class of functions mapping from $\mathcal{S}$ to $\mathbb{R}$ with the following parametric form*

$$\mathbf{v}_v(\cdot) = \mathbf{1}_M \cdot \min \left\{ \max_{\mathbf{a} \in \mathcal{A}} [\langle \mathbf{v}, \mathbf{v}(\cdot, \mathbf{a}) \rangle + \beta \|\Phi_C(\cdot, \mathbf{a})^\top \Lambda^{-1} \Phi_C(\cdot, \mathbf{a})\|], H \right\}, \mathbf{v}(\cdot, \mathbf{a}) = \mathbf{w}^\top \Phi_C(\cdot, \mathbf{a})$$

where the parameters $(\mathbf{w}, \beta, \Lambda)$ are such that $\mathbf{w} \in \mathbb{R}^d$, $\|\mathbf{w}\|_2 \leq L$, $\beta \in (0, B]$, $\|\Phi_C(\mathbf{x}, \mathbf{a})\| \leq \sqrt{M} \forall (\mathbf{x}, \mathbf{a}) \in \mathcal{S} \times \mathcal{A}$, and the minimum eigenvalue of Λ satisfies $\lambda_{\min}(\Lambda) \geq \lambda$. Let $\mathcal{N}_\varepsilon$ be the ε -covering number of $\mathcal{V}$ with respect to the distance $\text{dist}(\mathbf{v}, \mathbf{v}') = \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{v} \in \Upsilon} |\mathbf{v}_v(\mathbf{x}) - \mathbf{v}'_v(\mathbf{x})|$. Then,

$$\log(\mathcal{N}_\varepsilon) \leq d \cdot \log \left(1 + \frac{4LM^2}{\varepsilon} \right) + d^2 \log \left(1 + \frac{8Md^{1/2}B^2}{\lambda\varepsilon^2} \right).$$

Proof. We have that for two matrices $\mathbf{A}_1 = \beta^2 \Lambda_1^{-1}$, $\mathbf{A}_2 = \beta^2 \Lambda_2^{-1}$ and weight matrices $\mathbf{w}_1, \mathbf{w}_2 \in \mathbb{R}^d$,

$$\sup_{\mathbf{v} \in \Upsilon, \mathbf{x} \in \mathcal{S}} |\mathbf{v}_v(\mathbf{x}) - \mathbf{v}'_v(\mathbf{x})|_1 \quad (41)$$

$$= M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{v} \in \Upsilon} |\mathbf{v}^\top \mathbf{v}(\mathbf{x}) - \mathbf{v}^\top \mathbf{v}'(\mathbf{x})| \quad (42)$$

$$\leq M \cdot \sup_{\mathbf{x} \in \mathcal{S}} |\mathbf{v}(\mathbf{x}) - \mathbf{v}'(\mathbf{x})|_1 \quad (43)$$

$$\leq M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left| [\mathbf{w}_1^\top \Phi_C(\cdot, \mathbf{a}) + \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_1 \Phi_C(\cdot, \mathbf{a})\|_2] - [\mathbf{w}_2^\top \Phi_C(\cdot, \mathbf{a}) + \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_2 \Phi_C(\cdot, \mathbf{a})\|_2] \right|_1 \quad (44)$$

$$\leq M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left| (\mathbf{w}_1 - \mathbf{w}_2)^\top \Phi_C(\cdot, \mathbf{a}) + \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_1 \Phi_C(\cdot, \mathbf{a})\|_2 - \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_2 \Phi_C(\cdot, \mathbf{a})\|_2 \right|_1 \quad (45)$$

$$\leq M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left| (\mathbf{w}_1 - \mathbf{w}_2)^\top \Phi_C(\cdot, \mathbf{a}) \right|_1 + M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left| \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_1 \Phi_C(\cdot, \mathbf{a})\|_2 - \|\Phi_C(\cdot, \mathbf{a})^\top \mathbf{A}_2 \Phi_C(\cdot, \mathbf{a})\|_2 \right| \quad (46)$$

$$\leq M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left| (\mathbf{w}_1 - \mathbf{w}_2)^\top \Phi_C(\cdot, \mathbf{a}) \right|_1 + M \cdot \sup_{\mathbf{x} \in \mathcal{S}, \mathbf{a} \in \mathcal{A}} \left\| \Phi_C(\cdot, \mathbf{a})^\top (\mathbf{A}_1 - \mathbf{A}_2) \Phi_C(\cdot, \mathbf{a}) \right\|_2 \quad (47)$$

$$\leq M^{3/2} \cdot \sup_{\Phi: \|\Phi\| \leq \sqrt{M}} \left[\left\| (\mathbf{w}_1 - \mathbf{w}_2)^\top \Phi \right\|_2 \right] + M \cdot \sup_{\Phi: \|\Phi\| \leq \sqrt{M}} \left\| \Phi^\top (\mathbf{A}_1 - \mathbf{A}_2) \Phi \right\|_2 \quad (48)$$

$$\leq M^2 \cdot \|\mathbf{w}_1 - \mathbf{w}_2\|_2 + M^2 \|\mathbf{A}_1 - \mathbf{A}_2\|_2 \quad (49)$$

$$\leq M^2 \cdot \|\mathbf{w}_1 - \mathbf{w}_2\|_2 + M^2 \|\mathbf{A}_1 - \mathbf{A}_2\|_F \quad (50)$$

$$(51)$$

Now, let $\mathcal{C}_w$ be an $\varepsilon/(2M^2)$ cover of $\{\mathbf{w} \in \mathbb{R}^d \mid \|\mathbf{w}\|_2 \leq L\}$ with respect to the Frobenius-norm, and $\mathcal{C}_A$ be an $\varepsilon^2/4$ cover of $\{\mathbf{A} \in \mathbb{R}^{d \times d} \mid \|\mathbf{A}\|_F \leq (M^2 d)^{1/2} B^2 \lambda^{-1}\}$ with respect to the Frobenius norm. By Lemma 9 we have,

$$|\mathcal{C}_w| \leq (1 + 4LM^2/\varepsilon)^d, |\mathcal{C}_A| \leq (1 + 8(M^2 d)^{1/2} B^2 / (\lambda \varepsilon^2))^d. \quad (52)$$

Therefore, we can select, for any $\mathbf{v}_v(\cdot)$, corresponding weight $\mathbf{w} \in \mathcal{C}_w$, and matrix $\mathbf{A} \in \mathcal{C}_A$. Therefore, $\mathcal{N}_\varepsilon \leq |\mathcal{C}_A| \cdot |\mathcal{C}_w|$. This gives us,

$$\log(\mathcal{N}_\varepsilon) \leq d \cdot \log\left(1 + \frac{4LM^2}{\varepsilon}\right) + d^2 \log\left(1 + \frac{8Md^{1/2}B^2}{\lambda \varepsilon^2}\right). \quad (53)$$

□

Lemma 11 (Linearity of weights in Markov game). *In a game with n agents satisfying Assumptions 1, 2, 3, for any policy π , clique $C \in \hat{\mathcal{C}}$ of size M , and $\mathbf{v}_C \in \Upsilon_C$, there exists weights $\{\mathbf{w}_{\mathbf{v}_C, h}^\pi\}_{h \in [H]}$ such that $|Q_{\mathbf{v}_C, h}^\pi(\mathbf{x}_C, \mathbf{a}_C) - \mathbf{v}_C^\top \Phi_C(\mathbf{x}_C, \mathbf{a}_C)^\top \mathbf{w}_{\mathbf{v}_C, h}^\pi| \leq 2H\varepsilon(k)$ for all $(\mathbf{x}_C, \mathbf{a}_C, h) \in \mathcal{S}_C \times \mathcal{A}_C \times [H]$, where $\|\mathbf{w}_{\mathbf{v}_C, h}^\pi\|_2 \leq 2H\sqrt{d}$.*

Proof. By the Bellman equation and Proposition 1, we have that for any MDP corresponding to the scalarization parameter $\mathbf{v}_C \in \Upsilon_C$ and any policy π , state $\mathbf{x} \in \mathcal{S}_C$, joint action $\mathbf{a} \in \mathcal{A}_C$,

$$Q_{\mathbf{v}_C, h}^\pi(\mathbf{x}_C, \mathbf{a}_C) \quad (54)$$

$$\leq \mathbf{v}_C^\top \tilde{\mathbf{r}}_h^C(\mathbf{x}_C, \mathbf{a}_C) + \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}_C, \mathbf{a}_C) + 2H\varepsilon(k) \quad (55)$$

$$\leq \mathbf{v}_C^\top \left(\tilde{\mathbf{r}}_h^C(\mathbf{x}_C, \mathbf{a}_C) + \mathbf{1}_M \cdot \tilde{\mathbb{P}}_h^C V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}_C, \mathbf{a}_C) \right) + 2H\varepsilon(k) \quad (56)$$

$$\leq \mathbf{v}_C^\top \left(\Phi_C(\mathbf{x}_C, \mathbf{a}_C)^\top \begin{bmatrix} \theta_h \\ \mathbf{0}_d \end{bmatrix} + \int V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}'_C) \Phi_C(\mathbf{x}_C, \mathbf{a}_C)^\top \begin{bmatrix} \mathbf{0}_d \\ d\mu_h(\mathbf{x}'_C) \end{bmatrix} d\mathbf{x}'_C \right) + 2H\varepsilon(k) \quad (57)$$

$$\leq \mathbf{v}_C^\top \Phi_C(\mathbf{x}_C, \mathbf{a}_C)^\top \mathbf{w}_{\mathbf{v}_C, h}^\pi + 2H\varepsilon(k). \quad (58)$$

The first inequality follows from Assumption 2. Here $\mathbf{w}_{\mathbf{v}_C, h}^\pi = \left[\int V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}'_C) d\mu(\mathbf{x}'_C) \right] \theta_h$. Therefore, since $\|\theta_h\| \leq \sqrt{d}$ and $\|\int V_{\mathbf{v}_C, h+1}^\pi(\mathbf{x}'_C) d\mu(\mathbf{x}'_C)\| \leq H\sqrt{d}$, the result follows. □

Lemma 12 (Bound on Weights). *For any $C \in \hat{\mathcal{C}}$, $|C| = M$, $t \in [T]$, $h \in [H]$, $\mathbf{v} \in \Upsilon$, the weights $\mathbf{w}_{\mathbf{v}_C, h}^t$ satisfy*

$$\|\mathbf{w}_{\mathbf{v}_C, h}^t\|_2 \leq 2HM\sqrt{dt/\lambda}.$$

Proof. For any vector $\mathbf{v} \in \mathbb{R}^d \mid \|\mathbf{v}\| = 1$,

$$|\mathbf{v}^\top \mathbf{w}_{\mathbf{v}_C, h}^t| = \left| \mathbf{v}^\top (\Lambda_h^t)^{-1} \left(\sum_{\tau=1}^{k_t} \left[\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \left[\mathbf{r}_h(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) + \max_{\mathbf{a} \in \mathcal{A}} Q_{\mathbf{v}, h+1}(\mathbf{x}'_\tau, \mathbf{a}) \right] \right] \right) \right| \quad (59)$$

$$\leq \sqrt{k_t \cdot \sum_{\tau=1}^{k_t} \left(\mathbf{v}^\top (\Lambda_h^t)^{-1} \left[\Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \left[\mathbf{r}_h(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) + \max_{\mathbf{a} \in \mathcal{A}} Q_{\mathbf{v}, h+1}(\mathbf{x}'_\tau, \mathbf{a}) \right] \right] \right)^2 } \quad (60)$$

$$\leq HM \sqrt{k_t \cdot \sum_{\tau=1}^{k_t} \left\| \mathbf{v}^\top (\Lambda_h^t)^{-1} \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \right\|_2^2} \quad (61)$$

$$\leq 2HM \sqrt{k_t \cdot \sum_{\tau=1}^{k_t} \left\| \mathbf{v} \right\|_{(\Lambda_C^h(t))^{-1}}^2 \left\| \Phi_C(\mathbf{x}_C^h(\tau), \mathbf{a}_C^h(\tau)) \right\|_{(\Lambda_C^h(t))^{-1}}^2} \quad (62)$$

$$\leq 2HM \|\mathbf{v}\| \sqrt{dk_t/\lambda} \leq 2HM \sqrt{dt/\lambda}. \quad (63)$$

The penultimate inequality follows from Lemma 13 and the final inequality follows from the fact that $k_t \leq t$. The remainder of the proof follows from the fact that for any vector $\mathbf{w}$, $\|\mathbf{w}\| = \max_{\mathbf{v}: \|\mathbf{v}\|=1} |\mathbf{v}^\top \mathbf{w}|$. □

Lemma 13 (Elliptical Potential, Lemma 3 of Abbasi-Yadkori et al. (2011)). *Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ be vectors such that $\|\mathbf{x}\|_2 \leq L$. Then, for any positive definite matrix $\mathbf{U}_0 \in \mathbb{R}^{d \times d}$, define $\mathbf{U}_t := \mathbf{U}_{t-1} + \mathbf{x}_t \mathbf{x}_t^\top$ for all t . Then, for any $\nu > 1$,*

$$\sum_{t=1}^n \|\mathbf{x}_t\|_{\mathbf{U}_{t-1}^{-1}}^2 \leq 2d \log_\nu \left(\frac{\text{tr}(\mathbf{U}_0) + nL^2}{d \det^{1/d}(\mathbf{U}_0)} \right).$$

E.2. Multi-task concentration bound (Chowdhury and Gopalan, 2020)

We assume the multi-agent kernel Γ to be continuous relative to the operator norm on $\mathcal{L}(\mathbb{R}^n)$, the space of bounded linear operators from $\mathbb{R}^n$ to itself (for some $n > 0$). Then the RKHS $\mathcal{H}_\Gamma(\mathcal{X}^n)$ associated with the kernel Γ is a subspace of the space of continuous functions from $\mathcal{X}^n$ to $\mathbb{R}^n$, and hence, Γ is a Mercer kernel. Let μ be a measure on the (compact) set $\mathcal{X}^n$. Since Γ is a Mercer kernel on $\mathcal{X}$ and $\sup_{\mathbf{X} \in \mathcal{X}^n} \|\Gamma(\mathbf{X}, \mathbf{X})\| < \infty$, the RKHS $\mathcal{H}_\Gamma(\mathcal{X}^n)$ is a subspace of $L^2(\mathcal{X}^n, \mu; \mathbb{R}^n)$, the Banach space of measurable functions $g : \mathcal{X}^n \rightarrow \mathbb{R}^n$ such that $\int_{\mathcal{X}^n} \|g(\mathbf{X})\|^2 d\mu(\mathbf{X}) < \infty$, with norm $\|g\|_{L^2} = (\int_{\mathcal{X}^n} \|g(\mathbf{X})\|^2 d\mu(\mathbf{X}))^{1/2}$. Since $\Gamma(\mathbf{X}, \mathbf{X}) \in \mathcal{L}(\mathbb{R}^n)$ is a compact operator, by the Mercer theorem

We can therefore define a feature map $\Phi : \mathcal{X}^M \rightarrow \mathcal{L}(\mathbb{R}^n, \ell^2)$ of the multi-agent kernel Γ by

$$\Phi(\mathbf{X})^\top \mathbf{y} = (\sqrt{\nu_1} \psi_1(\mathbf{x}_1)^\top \mathbf{y}, \sqrt{\nu_2} \psi_2(\mathbf{x}_2)^\top \mathbf{y}, \dots, \sqrt{\nu_M} \psi_M(\mathbf{x}_M)^\top \mathbf{y}), \quad \forall \mathbf{X} \in \mathcal{X}^M, \mathbf{y} \in \mathbb{R}^m. \quad (64)$$

We then obtain $F(\mathbf{X}) = \Phi(\mathbf{X})^\top \Theta^*$ and $\Gamma(\mathbf{X}, \mathbf{X}') = \Phi(\mathbf{X})^\top \Phi(\mathbf{X}') \quad \forall \mathbf{X}, \mathbf{X}' \in \mathcal{X}^M$.

Define $\mathbf{S}_t = \sum_{\tau=1}^t \Phi(\mathbf{X}_\tau)^\top \varepsilon_\tau$, where $\varepsilon_1, \dots, \varepsilon_t$ are the noise vectors in $\mathbb{R}^M$. Now consider $\mathcal{F}_{t-1}$, the σ -algebra generated by the random variables $\{\mathbf{X}_\tau, \varepsilon_\tau\}_{\tau=1}^{t-1}$ and $\mathbf{X}_t$. We can see that $\mathbf{S}_t$ is $\mathcal{F}_t$ -measurable, and additionally, $\mathbb{E}[\mathbf{S}_t | \mathcal{F}_{t-1}] = \mathbf{S}_{t-1}$. Therefore, $\{\mathbf{S}_t\}_{t \geq 1}$ is a martingale with outputs in ℓ^2 space. Following (Chowdhury and Gopalan, 2020), consider now the map $\Phi_{\mathcal{X}_t} : \ell^2 \rightarrow \mathbb{R}^{Mt}$:

$$\Phi_{\mathcal{X}_t} \boldsymbol{\theta} = \left[(\Phi(\mathbf{X}_1)^\top \boldsymbol{\theta})^\top, (\Phi(\mathbf{X}_1)^\top \boldsymbol{\theta})^\top, \dots, (\Phi(\mathbf{X}_t)^\top \boldsymbol{\theta})^\top \right]^\top, \quad \forall \boldsymbol{\theta} \in \ell^2. \quad (65)$$

Additionally, denote $\mathbf{V}_t := \Phi_{\mathcal{X}_t}^\top \Phi_{\mathcal{X}_t}$ be a map from ℓ^2 to itself, with $\mathbf{I}$ being the identity operator in ℓ^2 . We have the following result from (Chowdhury and Gopalan, 2020) that provides us with a self-normalized martingale bound.

Lemma 14 (Lemma 3 of (Chowdhury and Gopalan, 2020)). *Let the noise vectors $\{\varepsilon_t\}_{t \geq 1}$ be σ -sub-Gaussian. Then, for any $\eta > 0$ and $\delta \in (0, 1]$, with probability at least $1 - \delta$, the following holds uniformly over all $t \geq 1$:*

$$\|\mathbf{S}_t\|_{(\mathbf{V}_t + \eta \mathbf{I})^{-1}} \leq \sigma \sqrt{2 \log(1/\delta) + \log \det(\mathbf{I} + \eta^{-1} \mathbf{V}_t)}.$$

Alternatively stated, we have again that with probability at least $1 - \delta$, the following holds uniformly over all $t \geq 1$:

$$\|\boldsymbol{\varepsilon}_t\|_{((\mathbf{K}_t + \eta \mathbf{I})^{-1} + \mathbf{I})^{-1}}^2 \leq 2\sigma^2 \log \left[\frac{\sqrt{\det(\mathbf{I}(1 + \eta) + \mathbf{K}_t)}}{\delta} \right].$$

F. Lower Bound

The central observation in this setting is that under the clique-dominance assumption (Assumption 2), it is impossible to obtain regret that avoids the $\theta(G)$ factor. Rather than provide a formal proof, we can provide a straightforward outline to obtain the guarantee. For any influence graph G , we can construct a minimal clique covering $\mathcal{C}_*$ and we can construct a unique Markov game for each clique in $\mathcal{C}_*$. For any clique C we construct a Markov game MG_C such that the reward functions for each agent in C are identical functions of the clique state-action (let us call it r_h^c for any agent c and step h), i.e., $r_h^c = r_h^C \forall c \in C$ and the marginal transition probability is only a function of the clique state-action as well. Now, observe that under this criterion, the scalarized reward is independent of the parameter $\mathbf{v}$ and is always r_h^C and hence one can find the regret in Π_Γ^* for the MG by simply choosing an arbitrary value of $\mathbf{v}$ and solving the scalarized MDP. Since we are considering the tabular setting, for any clique C , we can set $d_C = |\mathcal{S}_C| \times |\mathcal{A}_C|$. Note that for any MDP we have that the regret obeys $\Omega(H^2 \sqrt{\|S\| |\mathcal{A}| T})$, which gives us the regret within a clique as $\Omega(H^2 \sqrt{d_C T})$. Summing over the clique cover we obtain the lower bound for $\mathfrak{R}_C$.

G. Experiments

We run experiments on a basic cooperative multi-agent RL grid-world environment, `GridExplore` described as follows. In the first game, `GridExplore`, the agents are randomly placed in a grid of blank cells. Agents explore the grid by observing cells which are denoted as ‘explored’. Each agent obtains a reward for the number of cells they have explored. Each agent has the following actions $\{L, R, U, D, LU, LD, RU, RD\}$ and there are a total of $n = 8$ agents. The visibility of other agents is examined under 3 settings: (a) each agent can see all others (full), (b) each agent can only observe a random half of agents (partial), and (c) each agent can only observe their actions (self). The board is of size 10×10 and p_{Υ} is the uniform distribution over Δ_n . The game runs in episodes of length 200 and $T = 5 \times 10^6$.

For each agent c in clique C , the feature ϕ_c is given as the combined action of all visible agents (of dimensionality $8n'$) and ψ_C is the joint state of all agents in C (of dimensionality $100n'$) where $n' = 1$ in the self setting, $n' = \lfloor n/2 \rfloor + 1$ in the partial setting, $n' = n$ in the full setting. For our algorithm, we select $\varepsilon = 0.5 \times 10^{-6}$.

We present the average reward (over all agents) over the last 1000 episodes for 100 repeated trials in the table below. As baselines we consider a group of n individual Q-learning on the agents personal state space `Q-ind`, n individual `DQN` agents using a custom CNN with 3 hidden layers: 2 convolutional layers with filter size 4 and 32 filters each, and 1 fully-connected layer of dimensionality 256; and the `LSVI-UCB` algorithm proposed by Jin et al. (2020) on the same input space as ours.

Baseline	Full	Partial	Self
Q-ind	20.885 ± 2.833	16.294 ± 4.239	13.202 ± 4.887
DQN	35.932 ± 3.094	27.587 ± 5.059	18.478 ± 2.093
LSVI-ind	17.439 ± 2.192	9.847 ± 4.292	7.340 ± 3.778
MultOVI	31.294 ± 3.776	22.119 ± 5.882	15.395 ± 3.098

Table 1. Results on `GridExplore` environment.

We observe that our algorithm comfortably outperforms the individual baselines `Q-ind` and `LSVI` in all three settings, however, `DQN` outperforms our algorithm, presumably owing to better feature representations learnt from the deep neural networks. Future work may consider approaches to combine deep neural network based approaches with the multi-agent UCB algorithm as ours.