
Beating Adversarial Low-Rank MDPs with Unknown Transition and Bandit Feedback

Haolin Liu*

University of Virginia
srs8rh@virginia.edu

Zakaria Mhammedi*

Google Research
mhammedi@google.com

Chen-Yu Wei*

University of Virginia
chenyu.wei@virginia.edu

Julian Zimmert*

Google Research
zimmert@google.com

Abstract

We consider regret minimization in low-rank MDPs with fixed transition and adversarial losses. Previous work has investigated this problem under either full-information loss feedback with unknown transitions (Zhao et al., 2024), or bandit loss feedback with known transition (Foster et al., 2022). First, we improve the $\text{poly}(d, A, H)T^{5/6}$ regret bound of Zhao et al. (2024) to $\text{poly}(d, A, H)T^{2/3}$ for the full-information unknown transition setting, where d is the rank of the transitions, A is the number of actions, H is the horizon length, and T is the number of episodes. Next, we initiate the study on the setting with bandit loss feedback and unknown transitions. Assuming that the loss has a linear structure, we propose both model-based and model-free algorithms achieving $\text{poly}(d, A, H)T^{2/3}$ regret, though they are computationally inefficient. We also propose oracle-efficient model-free algorithms with $\text{poly}(d, A, H)T^{4/5}$ regret. We show that the linear structure is necessary for the bandit case—without structure on the reward function, the regret has to scale polynomially with the number of states. This is contrary to the full-information case (Zhao et al., 2024), where the regret can be independent of the number of states even for unstructured reward function.

1 Introduction

We study online reinforcement learning (RL) in low-rank Markov Decision Processes (MDPs). Low-rank MDPs is a class of MDPs where the transition probability can be decomposed as an inner product between two low-dimensional features, i.e., $P(x' | x, a) = \phi^*(x, a)^\top \mu^*(x')$, where $P(x' | x, a)$ is the probability of transitioning to state x' when the learner takes action a on state x , and ϕ^*, μ^* are two feature mappings. The ground truth features ϕ^* and μ^* are unknown to the learner. This setting has recently caught theoretical attention due to its simplicity and expressiveness (Agarwal et al., 2020; Uehara et al., 2021; Zhang et al., 2022a; Cheng et al., 2023; Modi et al., 2024; Zhang et al., 2022b; Mhammedi et al., 2024a; Huang et al., 2023). In particular, since the learner does not know the features, it is necessary for the learner to perform *feature learning* (or *representation learning*) to approximate them. This allows low-rank MDPs to model the additional difficulty not present in traditional linear function approximation schemes where the features are given, such as in linear MDPs (Jin et al., 2020b) and in linear mixture MDPs (Ayoub et al., 2020). Since feature learning is an indispensable part of modern deep RL pipelines, low-rank MDP is a model that is closer to practice than traditional linear function approximation.

*The authors are listed in alphabetical order.

Table 1: Comparison of adversarial low-rank MDP algorithms. $\mathcal{O}$ here hides factors of order $\text{poly}(d, |\mathcal{A}|, \log T, \log |\Phi| |\Upsilon|)$. † Algorithm 5 assumes access to $\phi(x, a)$ for any $(\phi, x, a) \in \Phi \times \mathcal{X} \times \mathcal{A}$, while other algorithms only require access to $\phi(x, a)$ for any $(\phi, a) \in \Phi \times \mathcal{A}$ on *visited* x .

Feedback	Algorithm	Algorithm type	Regret	Efficiency	Loss
Full-info	Zhao et al. (2024)	Model-based	$\mathcal{O}(T^{5/6})$	Oracle-efficient	Arbitrary
	Algorithm 1	Model-based	$\mathcal{O}(T^{2/3})$	Oracle-efficient	Arbitrary
	Lower Bound		$\Omega(\sqrt{ \mathcal{A} T})$		Arbitrary
Bandit	Algorithm 2	Model-based	$\mathcal{O}(T^{2/3})$	Inefficient	Linear loss Unknown loss feature
	Algorithm 5 †	Model-free	$\mathcal{O}(T^{2/3})$	Inefficient	Linear loss Unknown loss feature
	Algorithm 3 (for oblivious adversary)	Model-free	$\mathcal{O}(T^{4/5})$	Oracle-efficient	Linear loss Unknown loss feature
	Algorithm 4 (for adaptive adversary)	Model-free	$\mathcal{O}(T^{4/5})$	Oracle-efficient	Linear loss Known loss feature
	Lower Bound		$\Omega(\sqrt{ \mathcal{X} } \mathcal{A} T)$		Arbitrary

Most prior theoretical work on low-rank MDPs focuses on reward-free learning; this is a setting where instead of focusing on a particular reward function, the goal is to learn a model for the transitions (or, in the model-free setting, a small set of policies with good state cover), that enables policy optimization for *any* downstream reward functions. While this is a reasonable setup in some cases, in other applications, the learner can only obtain loss information from interactions with the environment, and only observes the loss on the state-actions that have been visited (i.e., bandit feedback). This introduces additional challenges to the learner.

Furthermore, in many online learning scenarios, the loss function may change over time, reflecting the non-stationary nature of the environment or task switches (Padakandla et al., 2020). This could be modeled by the *adversarial* MDP setting, where the loss function changes arbitrarily from one episode to the next, and the changes might even depend on the behavior of the learner. This setting is also extensively studied, but mostly restricted to tabular MDPs (Rosenberg and Mansour, 2019; Jin et al., 2020a; Shani et al., 2020; Luo et al., 2021) or traditional linear function approximation schemes (Cai et al., 2020; Luo et al., 2021; He et al., 2022; Zhao et al., 2022; Sherman et al., 2023b; Dai et al., 2023; Liu et al., 2023). The work by Zhao et al. (2024) initiated the study on adversarial MDPs in low-rank MDPs, but their work is restricted to full-information loss feedback.

When feature learning, bandit feedback, and adversarial losses are combined, the problem becomes highly challenging, and to the best of our knowledge there are no provably efficient algorithms to tackle this setting. In this work, we provide the first result for this combination. We hope that our result would bring new ideas to RL in practice, where all three elements are usually present simultaneously. We give several main results, targeting at either tighter regret (i.e., the performance gap between the optimal policy and the learner) or computational efficiency, as summarized in Table 1. Below we give a brief introduction for each of them. A more thorough related work review is in Appendix A.

- **$T^{2/3}$ -regret algorithm under full-information feedback (Algorithm 1).** This setting is studied by the only prior work in adversarial low-rank MDPs (Zhao et al., 2024), and we greatly improve their $T^{5/6}$ regret bound to $T^{2/3}$. Our algorithm begins with a model-based initial exploration phase to estimate the transition. It then performs policy optimization where the critic is the Q value induced by the estimated transition and the full information loss.
- **$T^{2/3}$ -regret model-based/model-free inefficient algorithm under bandit feedback (Algorithm 2, Algorithm 5).** Algorithm 2 starts with a model-based initial exploration phase to learn an estimated transition, and then runs exponential weights over policy space for regret minimization in the second phase. To tackle bandit feedback, we construct a novel loss estimator that leverages the structure of low-rank MDP to perform accurate off-policy evaluation. Algorithm 5 starts with a

different exploration phase, where it calls VoX (Mhammedi et al., 2023) to learn a policy cover; VoX is a model-free, reward-free exploration algorithm. After this initial exploratory phase, the algorithm also applies exponential weights and utilizes the same loss estimator as in Algorithm 2. However, due to its model-free nature, certain components of the estimator cannot be directly accessed and must be derived through specific optimizations.

- $T^{4/5}$ -**regret model-free oracle-efficient algorithm under bandit feedback (Algorithm 3, Algorithm 4)**. Algorithm 3 also starts with the model-free exploration algorithm VoX (Mhammedi et al., 2023) to learn a policy cover. After that, the algorithm operates in epochs; during epoch k , the algorithm commits to a fixed mixture of policies. This mixture consists of certain exploratory policies (based on the policy cover from the initial phase) and a policy computed using an online learning algorithm based on estimated Q -functions from previous epochs (these serve as loss functions). Algorithm 4 deals with the much more challenging setting of an *adaptive* adversary with bandit feedback. Here, we make the additional assumption that the loss feature, which may be different from the feature of the low-rank decomposition, is given. The algorithm is similar to Algorithm 3 with key differences outlined in Section 4.3.

2 Preliminaries

We study the episodic online reinforcement learning setting with horizon H . We consider an MDP $\mathcal{M} = (\mathcal{X}, \mathcal{A}, P_{1:H}^*)$, where $\mathcal{X}$ represents a countable (possibly infinite) state space², $\mathcal{A}$ is a finite action space, and $P_h^* : \mathcal{X} \times \mathcal{A} \rightarrow \Delta(\mathcal{X})$ denotes the transition kernel from layer h to $h+1$. We assume that the initial state $x_1 \in \mathcal{X}$ is fixed for simplicity without loss of generality. For any policy $\pi : \mathcal{X} \rightarrow \Delta(\mathcal{A})$ and arbitrary set of transition kernels $\{P_h\}_{h \in [H]}$, we let $\mathbb{P}^{P, \pi}$ denote the law over $(x_1, a_1, \dots, x_H, a_H)$ induced by the process of setting $x_1 = x_1$, sampling $a_1 \sim \pi_1(\cdot | x_1)$, then for $h = 2, \dots, H$, $x_h \sim P_{h-1}(\cdot | x_{h-1}, a_{h-1})$ and $a_h \sim \pi_h(\cdot | x_h)$. We let $\mathbb{E}^{P, \pi}$ denote the corresponding expectations. Further, we let $d_h^{P, \pi}(x) := \mathbb{P}^{P, \pi}[x_h = x]$ denote the *occupancy* of $x \in \mathcal{X}$. We also let $d_h^{P, \pi}(x, a) := \mathbb{P}^{P, \pi}[x_h = x, a_h = a]$. Further, we let $\mathbb{E}^\pi = \mathbb{E}^{P^*, \pi}$, $\mathbb{P}^\pi = \mathbb{P}^{P^*, \pi}$, and $d_h^\pi = d_h^{P^*, \pi}$. We use $\pi \circ_h \pi'$ to denote a policy that follows $\pi_k(\cdot | \cdot)$ for $k < h$ and $\pi'_k(\cdot | \cdot)$ for $k \geq h$. Similarly, $\pi \circ_h \pi' \circ_{h'} \pi''$ denotes a policy that follows π_k for $k < h$, π'_k for $h \leq k < h'$ and π''_k for $k > h'$.

We consider a learner interacting with the MDP $\mathcal{M}$ for T episodes with adversarial loss functions. Before the game starts, an oblivious adversary chooses the loss functions for all episodes $(\ell_{1:H}^t : \mathcal{X} \times \mathcal{A} \rightarrow [0, 1])_{t=1}^T$. For each episode $t \in [T]$, the learner starts at state $x_1^t = x_1$, then for each step $h \in [H]$ within episode t , the learner observes state $x_h^t \in \mathcal{X}_h$, chooses an action $a_h^t \in \mathcal{A}$, then suffers loss $\ell_h^t(x_h^t, a_h^t)$. The state x_{h+1}^t at the next step is drawn from transition $P_h^*(\cdot | x_h^t, a_h^t)$. We consider *bandit feedback* setting where the learner could only observe the losses $\ell_1^t(x_1^t, a_1^t), \dots, \ell_H^t(x_H^t, a_H^t)$ at the visited state-action pairs.

We let $\Pi := \{\pi : \mathcal{X} \rightarrow \Delta(\mathcal{A})\}$ denote the set of Markovian policies. For policy $\pi \in \Pi$, loss ℓ and transition kernels $P_{1:H}$, we denote by $Q_h^{P, \pi}(\cdot, \cdot; \ell)$ the *state-action* value function (a.k.a. Q -function) at step $h \in [H]$ with respect to the transitions $P_{1:H}$ and loss ℓ ; that is

$$Q_h^{P, \pi}(x, a; \ell) := \mathbb{E}^{P, \pi} \left[\sum_{s=h}^H \ell_s(x_s^t, a_s^t) \mid x_h = x, a_h = a \right], \quad (1)$$

for all $(x, a) \in \mathcal{X} \times \mathcal{A}$. We let $V_h^{P, \pi}(x; \ell) := \max_{a \in \mathcal{A}} Q_h^{P, \pi}(x, a; \ell)$ be the corresponding *state* value function at layer h . Further, we write $Q_h^\pi(\cdot, \cdot; \ell) := Q_h^{P^*, \pi}(\cdot, \cdot; \ell)$ and $V_h^\pi(\cdot; \ell) := V^{P^*, \pi}(\cdot; \ell)$.

For all of our algorithms except for Algorithm 4, we aim to construct (possibly randomized) policies $\{\pi^t\}_{t \in [T]}$ that ensure a sublinear *pseudo-regret* with respect to the best-fixed policy; that is,

$$\text{Reg}_T := \min_{\pi \in \Pi} \text{Reg}_T(\pi) \quad \text{where} \quad \text{Reg}_T(\pi) := \mathbb{E} \left[\sum_{t=1}^T V_1^{\pi^t}(x_1; \ell^t) - \sum_{t=1}^T V_1^\pi(x_1; \ell^t) \right]. \quad (2)$$

For Algorithm 4, we bound the standard *regret*

$$\overline{\text{Reg}}_T := \sum_{t=1}^T V_1^{\pi^t}(x_1; \ell^t) - \min_{\pi \in \Pi} \sum_{t=1}^T V_1^\pi(x_1; \ell^t) \quad (3)$$

²We assume that $\mathcal{X}$ is countable only to simplify the presentation. Our results can easily be extended to a continuous state space with an appropriate measure-theoretic treatment (see e.g. Mhammedi et al. (2023)).

with high probability. This allows it to handle adaptive adversary.

Throughout, we will assume that the MDP $\mathcal{M}$ is low-rank with *unknown* feature maps ϕ_h^* and μ_h^* .

Assumption 2.1 (Low-Rank MDP). *There exist (unknown) features maps $\phi_{1:H}^* : \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^d$ and $\mu_{1:H}^* : \mathcal{X} \rightarrow \mathbb{R}^d$, such that for all $h \in [H - 1]$ and $(x, a, x') \in \mathcal{X} \times \mathcal{A} \times \mathcal{X}$:*

$$\mathbb{P}[\mathbf{x}_{h+1} = x' \mid \mathbf{x}_h = x, \mathbf{a}_h = a] = \phi_h^*(x, a)^\top \mu_{h+1}^*(x'). \quad (4)$$

Furthermore, for all $h \in [H]$, the feature maps μ_h^* and ϕ_h^* are such that $\sup_{(x,a) \in \mathcal{X} \times \mathcal{A}} \|\phi_h^*(x, a)\| \leq 1$ and $\|\sum_{x \in \mathcal{X}} g(x) \cdot \mu_h^*(x)\| \leq \sqrt{d}$, for all $g : \mathcal{X} \rightarrow [0, 1]$.

Loss function under bandit feedback. For bandit feedback setting, we make the following additional linear assumption on the losses; in the sequel, we will argue that this is necessary to avoid a sample complexity scaling with the number of states. This linear loss assumption also appears in Ren et al. (2022); Zhang et al. (2022a) for stochastic low-rank MDPs. Note that for the full-information feedback setting, such an assumption is not required.

Assumption 2.2 (Loss Representation). *For any $t \in [T]$ and layer h , there is a vector $g_h^t \in \mathbb{B}_d(1)$ such that the loss $\ell_h^t(x, a)$ at round t satisfies:*

$$\forall (x, a) \in \mathcal{X} \times \mathcal{A}, \quad \ell_h^t(x, a) = \phi_h^*(x, a)^\top g_h^t. \quad (5)$$

We note that there is no loss of generality in assuming that the losses are expressed using the same features $\phi_{1:H}^*$ as the low-rank structure in (4). This is because if the losses have different features, we can simply combine these features with the low-rank features, and redefine ϕ^* accordingly. For the bulk of our results (and as stated in the prequel), we will assume that the losses $\{\ell_h^t(\cdot, \cdot)\}_{h \in [H], t \in [T]}$ (or equivalently $\{g_h^t\}_{h \in [H], t \in [T]}$ under Assumption 2.2) are chosen by an adversary before the start of the game (i.e. oblivious adversary). In Section 4.3, we will present a model-free, oracle-efficient algorithm for an adaptive adversary.

Function approximation. So far, Assumption 2.1 and Assumption 2.2 are in line with assumptions made in the linear MDP setting (Jin et al., 2020b). However, unlike in linear MDPs, we do not assume that the feature maps $\phi_{1:H}^*$ are known. To facilitate representation learning and ultimately a sublinear regret, we need to make *realizability* assumptions. In particular, in the model-free setting, we assume we have a function class Φ that contains the true features $\phi_{1:H}^*$. In the model-based setting, we additionally assume access to a function class Υ that contains the feature maps $\mu_{1:H}^*$ ³. We will formalize these assumptions in their corresponding sections in the sequel.

Other notation. For $\psi : \Pi \rightarrow \mathbb{R}^d$, we define $\text{John}(\psi, \Pi)$ as a distribution $\mu \in \Delta(\Pi)$ such that $\|\psi(\pi)\|_{G^{-1}}^2 \leq d$ for all $\pi \in \Pi$, where $G = \sum_{\pi \in \Pi} \mu(\pi) \cdot \psi(\pi) \psi(\pi)^\top$. This is the standard John's exploration or G -optimal design, which always exists.

3 Model-based Algorithms for Adversarial Low-rank MDPs

In this section, we discuss adversarial low-rank MDPs under model-based assumption. The model-based assumption is formalized in the Assumption 3.1 below. This assumption is standard which also appears in prior works on model-based learning in low-rank MDPs (Agarwal et al., 2020; Uehara et al., 2021; Zhang et al., 2022a; Cheng et al., 2023; Zhao et al., 2024).

Assumption 3.1 (Model-based assumption). *The learner has access to two model spaces Φ and Υ such that $\phi^* \in \Phi$ and $\mu^* \in \Upsilon$. Moreover, for any $\phi \in \Phi$, $\mu \in \Upsilon$, and $h \in [2..H]$, we have $\sup_{(x,a) \in \mathcal{X} \times \mathcal{A}} \|\phi_{h-1}(x, a)\| \leq 1$, $\sum_{x' \in \mathcal{X}} \phi_{h-1}(x, a)^\top \mu_h(x') = 1$ and $\|\sum_{x \in \mathcal{X}} g(x) \cdot \mu_h(x)\| \leq \sqrt{d}$, for all $g : \mathcal{X} \rightarrow [0, 1]$.*

3.1 Adversarial Low-rank MDPs with Full Information

We first discuss learning adversarial low-rank MDPs with full information and model-based assumption. This setting aligns with Zhao et al. (2024), and our Algorithm 1 successfully improves their regret from $T^{5/6}$ to $T^{2/3}$.

³The setting where we assume access to function classes that realize both $\phi_{1:H}^*$ and $\mu_{1:H}^*$ is called *model-based* because it allows one to model the transition probabilities, thanks to the low-rank structure in (4).

Algorithm 1 Model-Based Algorithm for Full-Information Feedback

- 1: Let $\eta = \frac{1}{H\sqrt{T}}$, $\epsilon = (Hd^2|\mathcal{A}|(d^2 + |\mathcal{A}|))^{\frac{1}{3}}T^{-\frac{1}{3}}$, and $T_0 = \tilde{O}(\epsilon^{-2}H^3d^2|\mathcal{A}|(d^2 + |\mathcal{A}|))$. Let π^1 be a uniform policy.
- 2: Run (Cheng et al., 2023, Algorithm 1) for T_0 episodes and get outputs $\hat{\phi} \in \Phi$, $\hat{\mu} \in \Upsilon$.
- 3: Define transitions $\hat{P}_{1:H-1}$ as

$$\hat{P}_h(x' | x, a) = \hat{\phi}_h(x, a)^\top \hat{\mu}_{h+1}(x'), \quad \forall (x, a, x') \in \mathcal{X} \times \mathcal{A} \times \mathcal{X}.$$

- 4: **for** $t = T_0 + 1, T_0 + 2, \dots, T$ **do**
- 5: Execute policy π^t and observe trajectory $(\mathbf{x}_{1:H}, \mathbf{a}_{1:H})$ and full information loss ℓ^t .
- 6: Update policy for all $h \in [H]$:

$$\pi_h^{t+1}(a | x) \propto \exp\left(-\eta \sum_{i=1}^t \widehat{Q}_h^i(x, a)\right) \text{ where } \widehat{Q}_h^t(x, a) = Q_h^{\hat{P}, \pi^t}(x, a; \ell^t).$$

- 7: **end for**
-

As argued in Zhao et al. (2024), the challenge of learning adversarial low-rank MDPs lies in the need for balancing exploration and exploitation both in representation learning and policy optimization over adversarial losses. To tackle this doubled exploration and exploitation challenge, the algorithm of Zhao et al. (2024) performs *simultaneous* representation learning and policy optimization. With a closer look at their analysis, we find that there is a drawback of this approach: because their algorithm handles the two tasks at the same time, it spends less exploration for representation learning in the early phase of the algorithm. This results in larger error in the estimated Q-values (i.e., critic) fed to policy optimization, and worsens the overall regret.

To address this issue, we design Algorithm 1 as a simple two-phase algorithm that separates representation learning and policy optimization. In the first phase, following Cheng et al. (2023), we perform optimal reward-free exploration for low-rank MDPs to estimate the transition. The resulted estimator, $\hat{P}$, is able to accurately approximate the true transition and give accurate Q-value estimators for any policy. The more accurate Q-estimator allows for more effective policy optimization in the second phase. Theorem 3.1 shows the guarantee of Algorithm 1 where $\tilde{O}$ hides logarithmic factors of $d, H, T, |\Phi||\Upsilon|$.

Theorem 3.1. *Algorithm 1 ensures $\text{Reg}_T \leq \tilde{O}\left(H^3(d^2 + |\mathcal{A}|)T^{\frac{2}{3}}\right)$.*

The proof for Theorem 3.1 is given in Appendix B. Zhao et al. (2024) also constructs a lower bound $\Omega\left(H\sqrt{d|\mathcal{A}|T}\right)$ for this settings. Thus, the $\text{poly}(|\mathcal{A}|)$ -dependence is unavoidable.

3.2 Model-Based, Computationally Inefficient Algorithm for Bandit Feedback

In this section, following Assumption 3.1, we introduce the first (model-based) algorithm (Algorithm 2) for adversarial low-rank MDPs with bandit feedback and sublinear regret. Compared with linear MDPs, the key challenge for more general low-rank MDPs is to construct a proper loss estimator. For linear MDP, since the feature is known, the loss estimator closely resembles that of linear bandits. However, low-rank MDPs lack such structural simplicity, making standard loss estimators invalid. To overcome this challenge, we propose a new loss estimator that works for any loss function based on off-policy evaluation and the low-rank structure of transition. In this section, $\tilde{O}$ hides logarithmic factors of $d, H, T, |\Phi||\Upsilon|$.

In Algorithm 2, we first conduct an initial representation learning phase to establish accurate transition estimator $\hat{P}$ and its corresponding features $\hat{\phi}$ and $\hat{\mu}$ based on reward-free exploration algorithms in Cheng et al. (2023). Then, in the second phase, we use exponential weights to maintain a distribution over the policy space Π' where we mix a uniform policy with Π to enhance exploration. At every round t , a behavior policy π^t is chosen from the current policy distribution, and we use the data collected by π^t to estimate the value for every $\pi \in \Pi'$. The success of such off-policy evaluation is based on the following observations of low-rank MDP. Using the low-rank transition structure, for $h \geq 2$, we have

$$\forall \pi \in \Pi, \quad d_h^\pi(x) = \phi_{h-1}^*(\pi)^\top \mu_h^*(x), \quad \text{where} \quad \phi_{h-1}^*(\pi) := \mathbb{E}^\pi[\phi_{h-1}^*(\mathbf{x}_{h-1}, \mathbf{a}_{h-1})]. \quad (6)$$

Algorithm 2 Model-Based Algorithm for Bandit Feedback

Input: A policy class Π .

- 1: Set $\epsilon = T^{-\frac{1}{3}}$, $\gamma = T^{-\frac{1}{3}}$, $\beta = T^{-\frac{1}{3}}$, $\eta = (4Hd|\mathcal{A}|)^{-1}T^{-\frac{2}{3}}$, and $T_0 = \tilde{O}(\epsilon^{-2}H^3d^2|\mathcal{A}|(d^2 + |\mathcal{A}|))$.
- 2: Run (Cheng et al., 2023, Algorithm 1) for T_0 episodes and get outputs $\hat{\phi} \in \Phi$, $\hat{\mu} \in \Upsilon$.
- 3: Define transitions $\hat{P}_{1:H-1}$ as

$$\hat{P}_h(x' | x, a) = \hat{\phi}_h(x, a)^\top \hat{\mu}_{h+1}(x'), \quad \forall (x, a, x') \in \mathcal{X} \times \mathcal{A} \times \mathcal{X}.$$

- 4: For all $h \in [H-1]$, define $\hat{\phi}_h(\pi) = \sum_{(x,a) \in \mathcal{X} \times \mathcal{A}} \hat{d}_h^\pi(x, a) \cdot \hat{\phi}_h(x, a)$, where $\hat{d}_h^\pi := d_h^{\hat{P}, \pi}$.
- 5: Define the policy space $\Pi' = \{\pi' : \exists \pi \in \Pi, \pi'_h(\cdot | x) = (1 - \beta)\pi_h(\cdot | x) + \beta/|\mathcal{A}|, \forall x, h\}$.
- 6: **for** $t = T_0 + 1, T_0 + 2, \dots, T$ **do**
- 7: Define $p^t(\pi) \propto \exp(-\eta \sum_{i=1}^{t-1} (\hat{\ell}^i(\pi) - b^i(\pi)))$, for all $\pi \in \Pi'$.
- 8: Let $\rho^t(\pi) = (1 - \gamma)p^t(\pi) + \frac{\gamma}{H-1} \sum_{h=1}^{H-1} J_h$, where $J_h = \text{John}(\hat{\phi}_h(\cdot), \Pi')$. // John as in §2
- 9: Execute policy $\pi^t \sim \rho^t$ and observe trajectory $(\mathbf{x}_{1:H}^t, \mathbf{a}_{1:H}^t)$ and losses $\ell_h^t = \ell_h^t(\mathbf{x}_h^t, \mathbf{a}_h^t)$.
- 10: Define $\Sigma_h^t = \sum_{\pi \in \Pi'} \rho^t(\pi) \cdot \hat{\phi}_h(\pi) \hat{\phi}_h(\pi)^\top$, $b^t(\pi) = \sqrt{d}H\epsilon \cdot \sum_{h=1}^{H-1} \|\hat{\phi}_h(\pi)\|_{(\Sigma_h^t)^{-1}}$, and

$$\hat{\ell}^t(\pi) = \frac{\pi_1(\mathbf{a}_1^t | \mathbf{x}_1^t)}{\pi_1^t(\mathbf{a}_1^t | \mathbf{x}_1^t)} \ell_1^t + \sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \frac{\pi_h(\mathbf{a}_h^t | \mathbf{x}_h^t)}{\pi_h^t(\mathbf{a}_h^t | \mathbf{x}_h^t)} \ell_h^t.$$

11: **end for**

Thus, using the definition of the V -function from Section 2, we have for any loss function ℓ and π :

$$V_1^\pi(x_1; \ell) - \mathbb{E}[\ell_1(\mathbf{x}_1, \mathbf{a}_1)] = \sum_{h=2}^H \sum_{(x,a) \in \mathcal{X} \times \mathcal{A}} \phi_{h-1}^*(\pi)^\top \mu_h^*(x) \pi_h(a | x) \cdot \ell_h(x, a). \quad (7)$$

Letting $\Lambda_h^t := (\mathbb{E}_{\pi^t \sim p^t} [\phi_h^*(\pi^t) \phi_h^*(\pi^t)^\top])^{-1}$ and ignoring the loss term $\mathbb{E}[\ell_1(\mathbf{x}_1, \mathbf{a}_1)]$ from the first step (this term can easily be treated as in a bandit setting with $H = 1$), we have for all $\pi \in \Pi'$:

$$\begin{aligned} V_1^\pi(x_1; \ell) &= \mathbb{E}_{\pi^t \sim p^t} \left[\sum_{h=2}^H \sum_{(x,a) \in \mathcal{X} \times \mathcal{A}} \phi_{h-1}^*(\pi)^\top \Lambda_{h-1}^t \phi_{h-1}^*(\pi^t) \phi_{h-1}^*(\pi^t)^\top \mu_h^*(x) \pi_h(a | x) \cdot \ell_h(x, a) \right], \\ &= \mathbb{E}_{\pi^t \sim p^t} \left[\sum_{h=2}^H \sum_{x,a} \phi_{h-1}^*(\pi)^\top \Lambda_{h-1}^t \phi_{h-1}^*(\pi^t) \cdot d_h^{\pi^t}(x) \pi_h^t(a | x) \frac{\pi_h(a | x)}{\pi_h^t(a | x)} \cdot \ell_h(x, a) \right], \\ &= \mathbb{E}_{\pi^t \sim p^t} \mathbb{E}^{\pi^t} \left[\sum_{h=2}^H \phi_{h-1}^*(\pi)^\top \Lambda_{h-1}^t \phi_{h-1}^*(\pi^t) \cdot \frac{\pi_h(\mathbf{a}_h | \mathbf{x}_h)}{\pi_h^t(\mathbf{a}_h | \mathbf{x}_h)} \cdot \ell_h(\mathbf{x}_h, \mathbf{a}_h) \right]. \end{aligned}$$

Thus, for all ℓ and π , $\sum_{h=2}^H \phi_{h-1}^*(\pi)^\top \Lambda_{h-1}^t \phi_{h-1}^*(\pi^t) \cdot \pi_h(\mathbf{a}_h | \mathbf{x}_h) \pi_h^t(\mathbf{a}_h | \mathbf{x}_h)^{-1} \cdot \ell_h(\mathbf{x}_h, \mathbf{a}_h)$ for $\pi^t \sim p^t$ and $(\mathbf{x}_h, \mathbf{a}_h) \sim d_h^{\pi^t}$ is an unbiased estimator of $V_1^\pi(x_1, \ell)$. However, $\phi_{h-1}^*(\pi)$ is not accessible because both the true feature ϕ^* and occupancy d^π for the true transition are unknown. Thus, our estimator incorporates the learned feature $\hat{\phi}$ and the occupancy of $\hat{P}$ instead as shown in Line 4 and Line 10. Utilizing estimated features and transition could introduce additional bias but the initial representation learning already ensures such bias is small enough to tackle. We compensate for the bias by incorporating an exploration bonus $b^t(\pi)$ in exponential weights. To further encourage exploration, we additionally perform John's exploration together with exponential weights when selecting behavior policies. The main guarantee of Algorithm 2 is given in Theorem 3.2.

Theorem 3.2. Algorithm 2 achieves $\text{Reg}_T(\pi) \leq \tilde{O}(d^2 H^3 |\mathcal{A}| (d^2 + |\mathcal{A}|) T^{2/3} \log |\Pi|)$ for any $\pi \in \Pi$.

Note that the guarantee in Theorem 3.2 only holds for policy $\pi \in \Pi$. To ensure our regret bound is meaningful, at least a near-optimal policy should be contained in the given policy set Π . In general, the size of such a policy set would grow exponentially with the number of states (e.g. covering of all Markovian policies), making the regret have polynomial dependence on the number of states. In Theorem 3.3, we show that even for low-rank MDPs, if the loss function lacks structure, the regret cannot avoid polynomial dependence on the number of states. The detailed construction for this lower-bound is given in Appendix H.

Theorem 3.3. *There exists a low-rank MDP with $|\mathcal{X}|$ states, $|\mathcal{A}|$ actions and sufficiently large T with unstructured losses such that any agent suffers at least regret of $\Omega(\sqrt{|\mathcal{X}||\mathcal{A}|T})$.*

Theorem 3.3 shows that under bandit feedback, in general, we could not gain too much from low-rank transition structure compared with tabular MDPs. This contrasts with the $\Omega(\sqrt{|\mathcal{A}|T})$ lower bound in the full information settings (Zhao et al. (2024)). To get rid of any dependence on the number of states, we additionally introduce Assumption 2.2 to impose linear structure on the loss function. Unlike linear MDPs that require the loss feature to be known, our algorithm can even handle linear loss with unknown feature, since our Algorithm 2 never explicitly uses the loss feature. The linear structure is only used to control the size of the candidate policy class in the analysis (i.e., making $\log |\Pi|$ irrelevant to the number of states). Specifically, when both loss and transition are linear, the Q-function is also linear, making it sufficient to consider the following linear policy space:

$$\Pi_{\text{lin}} = \left\{ \pi : \mathcal{X} \rightarrow \mathcal{A} \mid \pi_h(a \mid x) = \mathbb{I}\{a = \underset{a \in \mathcal{A}}{\operatorname{argmin}} \phi_h(x, a)^\top \theta_h\}, h \in [H], \|\theta_h\|_2 \leq \sqrt{dHT}, \phi \in \Phi \right\}.$$

The $\frac{1}{T}$ -cover of Π_{lin} only have size $|\Phi| \cdot T^{\mathcal{O}(d)}$ following standard arguments (e.g, Exercise 27.6 of Lattimore and Szepesvári (2020)) and if we feed it into Algorithm 2, our regret could avoid dependence on the size of state space as shown in Corollary 3.1.

Corollary 3.1. *If the loss function satisfies Assumption 2.2, applying Algorithm 2 with Π as the $\frac{1}{T}$ -cover of Π_{lin} ensures $\text{Reg}_T \leq \tilde{\mathcal{O}}(d^3 H^3 |\mathcal{A}|(d^2 + |\mathcal{A}|)T^{2/3})$.*

4 Model-free Algorithms for Adversarial Low-rank MDPs

In this section, we consider the model-free setting, where we only assume access to a feature class Φ that contains the true feature map ϕ^* .

Assumption 4.1 (Model-free realizability). *The learner has access to a function class Φ such that*

$$\phi^* \in \Phi \quad \text{and} \quad \sup_{\phi \in \Phi} \sup_{(x,a) \in \mathcal{X} \times \mathcal{A}} \|\phi(x, a)\| \leq 1. \quad (8)$$

This is a standard assumption in the context of model-free RL (Modi et al., 2024; Zhang et al., 2022b; Mhammedi et al., 2024a). We note that having access to the function class Φ alone (instead of both Φ and Υ as in Assumption 3.1) is not sufficient to model the transition probabilities (unlike in the model-based case). This makes the model-free setting much more challenging; in fact, until the recent work by Mhammedi et al. (2023) there were no model-free, oracle-efficient algorithms for this setting that do not require any additional structural assumptions on the MDP.

4.1 Model-free, Inefficient Algorithm for Bandit Feedback

Our first algorithm follows the same structure as Algorithm 2 but incorporates a model-free initial exploration phase introduced by Mhammedi et al. (2023). Unlike the model-based exploration phase, which directly provides an estimated transition, the model-free exploration phase outputs a policy cover. This policy cover can be combined with the optimization in Algorithm 1 of Liu et al. (2023) to solve the expected feature $\hat{\phi}$, which is then used in the loss estimator in Line 10 of Algorithm 2. The algorithm, summarized in Algorithm 5, is inefficient but achieves $T^{2/3}$ regret. More details and proofs can be found in Appendix D.

4.2 Model-free, Oracle Efficient Algorithm for Bandit Feedback (Oblivious Adversary)

We now describe the key component of our efficient model-free algorithm (Algorithm 3).

Exploration phase and policy cover. Similar to algorithms in the previous section, Algorithm 3 begins with a reward-free exploration phase; Line 2 of Algorithm 3. However, unlike in the previous sections where the role of this exploration phase was to learn a model for the transition probabilities, here the goal is to compute a, so called, *policy cover* which is a small set of policies that can be used to effectively explore the state space.

Definition 4.1 (Approximate policy cover). *For $\alpha, \varepsilon \in (0, 1]$ and $h \in [H]$, a subset $\Psi \subseteq \Pi$ is an (α, ε) -policy cover for layer h if*

$$\max_{\pi \in \Psi} d_h^\pi(x) \geq \alpha \cdot \max_{\pi' \in \Pi} d_h^{\pi'}(x), \quad \text{for all } x \in \mathcal{X} \text{ such that } \max_{\pi' \in \Pi} d_h^{\pi'}(x) \geq \varepsilon \cdot \|\mu_h^*(x)\|. \quad (10)$$

Algorithm 3 Oracle Efficient Algorithm for Adversarial Low-Rank MDPs (Oblivious Adversary).

Input: Number of rounds T , feature class Φ , confidence parameter $\delta \in (0, 1)$.

- 1: Set $\varepsilon \leftarrow T^{-1/3}$, $N_{\text{reg}} \leftarrow T^{2/3}$, $\nu \leftarrow N_{\text{reg}}^{-1/2}$, and $T_0 \leftarrow \varepsilon^{-2} Ad^{13} H^6 \log(\Phi/\delta)$.
- 2: Get $\Psi_{1:H}^{\text{cov}} \leftarrow \text{VoX}(\Phi, \varepsilon, \delta)$. // Compute policy cover with VoX (Mhammedi et al., 2023).
- 3: **for** $k = 1, \dots, (T - T_0)/N_{\text{reg}}$ **do**
- 4: Define $\widehat{\pi}_h^{(k)}(a | x) \propto \exp\left(-\eta \sum_{s < k} \widehat{Q}_h^{(s)}(x, a)\right)$ for $h \in [H]$.
- 5: **for** $t = T_0 + (k - 1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}$ **do**
- 6: Sample variables $\zeta^t \sim \text{Ber}(\nu)$, $\mathbf{h}^t \sim \text{unif}([H])$, and $\pi^t \sim \text{unif}(\Psi_{\mathbf{h}^t}^{\text{cov}})$.
- 7: Set $\widehat{\pi}^t = \mathbb{I}\{\zeta^t = 0\} \cdot \widehat{\pi}^{(k)} + \mathbb{I}\{\zeta^t = 1\} \cdot \pi^t \circ_{\mathbf{h}^t} \pi_{\text{unif}} \circ_{\mathbf{h}^t+1} \widehat{\pi}^{(k)}$.
- 8: Execute $\widehat{\pi}^t$, and observe trajectory $(\mathbf{x}_1^t, \mathbf{a}_1^t, \dots, \mathbf{x}_H^t, \mathbf{a}_H^t)$.
- 9: For $h \in [H]$, observe loss $\ell_h^t := \ell_h^t(\mathbf{x}_h^t, \mathbf{a}_h^t)$.
- 10: **end for**
- 11: For $h \in [H]$ and $\mathcal{I}^{(k)} = \{T_0 + (k - 1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}\}$, compute $(\hat{\phi}_h^{(k)}, \hat{\theta}_h^{(k)})$:

$$(\hat{\phi}_h^{(k)}, \hat{\theta}_h^{(k)}) \leftarrow \underset{(\phi, \theta) \in \Phi \times \mathbb{B}_d(H\sqrt{d})}{\text{argmin}} \sum_{t \in \mathcal{I}^{(k)}} \left(\phi_h(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \theta - \sum_{s=h}^H \ell_s^t \right)^2 \cdot \mathbb{I}\{\zeta^t = 0 \text{ or } \mathbf{h}^t \leq h\}. \quad (9)$$

- 12: Set $\widehat{Q}_h^{(k)}(x, a) = \hat{\phi}_h^{(k)}(x, a)^\top \hat{\theta}_h^{(k)}$, for all $(x, a) \in \mathcal{X} \times \mathcal{A}$.
 - 13: **end for**
-

In Line 2, Algorithm 3 calls VoX (Mhammedi et al., 2023), a reward-free and model-free exploration algorithm to compute $(1/(8Ad), \varepsilon)$ -policy covers $\Psi_1^{\text{cov}}, \dots, \Psi_H^{\text{cov}}$ for layers $1, \dots, H$, respectively, with $|\Psi_h| = d$ for all $h \in [H]$. This call to VoX requires $O(1/\varepsilon^2)$ episodes; see the guarantee of VoX in Lemma G.1. After this initial phase, the algorithm operates in epochs, each consisting of $N_{\text{reg}} \in \mathbb{N}$ episodes, where in each epoch $k \in [K]$, the algorithm commits to executing policies sampled from a fixed policy distribution $\rho^{(k)} \in \Delta(\Pi)$ with support on the policy covers $\Psi_{1:H}^{\text{cov}}$ and a policy $\widehat{\pi}^{(k)}$ specified by an online learning algorithm. Next, we describe in more detail how $\rho^{(k)}$ is constructed and motivate the elements of its construction starting with the online learning policies $\{\widehat{\pi}^{(k)}\}_{k \in [K]}$.

Online learning policies. Given estimates $\{\widehat{Q}_{1:H}^{(s)}\}_{s < k}$ of the average Q -functions

$$\left\{ \frac{1}{N_{\text{reg}}} \sum_{t \text{ in epoch } s} Q_{1:H}^{\widehat{\pi}^{(s)}}(\cdot, \cdot; \ell^t) \right\}_{s < k} \quad (11)$$

from the previous epoch (we will describe how these estimates are computed in the sequel), Algorithm 3 computes policy $\widehat{\pi}^{(k)}$ for epoch k according to

$$\widehat{\pi}_h^{(k)}(a | x) \propto \exp\left(-\eta \sum_{s < k} \widehat{Q}_h^{(s)}(x, a)\right), \quad (12)$$

for all $h \in [H]$. Given a state $x \in \mathcal{X}$, such exponential weight update ensures a sublinear regret with respect to the sequence of loss functions given by $\{\pi(\cdot | x) \mapsto \sum_{h \in [H]} \widehat{Q}_h^{(k)}(x, \pi_h(\cdot | x))\}_{k \in [K]}$. Thanks to the performance difference lemma, and as shown in Luo et al. (2021), a sublinear regret with respect to these "surrogate" loss functions translates into a sublinear regret in the low-rank MDP game we are interested in, granted that $\{\widehat{Q}_{1:H}^{(k)}\}_{k \in [K]}$ are good estimates of the average Q -functions (Luo et al., 2021). In line with previous analyses, we require the Q -function estimates to ensure that the following bias term

$$\mathbb{E}^\pi \left[\max_{a \in \mathcal{A}} \left(\frac{1}{N_{\text{reg}}} \sum_{t \text{ in epoch } k} Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right)^2 \right] \quad (13)$$

is small for all $h \in [H]$, $k \in [K]$, and $\pi \in \Pi$.

Q-function estimates. Thanks to the low-rank MDP structure and the linear loss representation assumption (Assumption 2.2), the average Q -functions in (11) are linear in the feature maps ϕ^* . Thus, using the function class Φ in Assumption 2.1 we can estimate these average Q -functions by regressing the sum of losses $\sum_{s=h}^H \ell_s^t$ onto $(\mathbf{x}_h^t, \mathbf{a}_h^t)$ for t in the k th epoch (as in (9)). However, naïvely doing this using only trajectories generated by $\widehat{\pi}^{(k)}$ would only ensure that the bias term in (13) is small for $\pi = \widehat{\pi}^{(k)}$. To ensure that it is small for all possible policies π 's, we need to estimate the Q -functions on the trajectories of policies that are guaranteed to have good state coverage; this is where we use the policy cover from the initial phase.

Mixture of policies. At episode t in each epoch $k \in [K]$, we execute policy $\widehat{\pi}^t$ sampled from $\rho^{(k)}$, where $\rho^{(k)}$ is the distribution of the random policy:

$$\mathbb{I}\{\zeta^t = 0\} \cdot \widehat{\pi}^{(k)} + \mathbb{I}\{\zeta^t = 1\} \cdot \pi^t \circ_{\mathbf{h}^t} \pi_{\text{unif}} \circ_{\mathbf{h}^{t+1}} \widehat{\pi}^{(k)}, \quad (14)$$

with $\zeta^t \sim \text{Ber}(\nu)$, $\mathbf{h}^t \sim \text{unif}([H])$, and $\pi^t \sim \text{unif}(\Psi_{\mathbf{h}^t}^{\text{cov}})$. In words, at the start of each episode of any epoch k , we execute $\widehat{\pi}^{(k)}$ (see (12)) with probability $1 - \nu$; and with probability ν , we execute a policy in $\Psi_{1:H}^{\text{cov}}$ selected uniformly at random. As explained in the previous paragraph, this ensures a small bias for all choices of π in (13) thanks the policy cover property of $\Psi_{1:H}^{\text{cov}}$. We now state the guarantee of Algorithm 3.

Theorem 4.1. *Let $\delta \in (0, 1)$ be given and suppose Assumption 2.1 and Assumption 2.2 hold. This, for $T = \text{poly}(A, d, H, \log(|\Phi|/\delta))$ sufficiently large, Algorithm 3 guarantees $\text{Reg}_T \leq \text{poly}(A, d, H, \log(|\Phi|/\delta)) \cdot T^{4/5}$ regret against an oblivious adversary.*

The proof is in Appendix E. Note that the T -dependence in this regret even outperforms that of the previous best bound by Zhao et al. (2024) (see Table 1). Compared to their algorithm, Algorithm 3 is model-free and only requires bandit feedback. This makes the result in Theorem 4.1 rather surprising.

4.3 Model-free, Oracle Efficient Algorithm (Adaptive Adversary)

In this section, we present a variant of Algorithm 3 (Algorithm 4) that guarantees a sublinear regret against an adaptive adversary. Given the difficulty of this setting, we make the additional assumption that the algorithm has access to the loss feature ϕ^{loss} , which may be different than the low-rank MDP feature ϕ^* (unlike in Assumption 2.2).

Assumption 4.2 (Loss Representation). *There is a (known) feature map ϕ^{loss} satisfying $\sup_{h \in [H], (x, a) \in \mathcal{X} \times \mathcal{A}} \|\phi_h^{\text{loss}}(x, a)\| \leq 1$ and such that for any round $h \in [H]$, $t \in [T]$, and history $\mathcal{H}^{t-1} = (x_{1:H}^{1:t-1}, a_{1:H}^{1:t-1})$, the loss function at round t satisfies*

$$\forall (x, a) \in \mathcal{X} \times \mathcal{A}, \quad \ell_h(x, a; \mathcal{H}^{t-1}) = \phi_h^{\text{loss}}(x, a)^\top g_h^t, \quad (15)$$

for some $g_h^t \in \mathbb{B}_d(1)$.

Note that Assumption 4.2 asserts that the loss at round t depends only on the history $\mathcal{H}^{t-1}$ and the current state action pair. Before moving forward, we introduce some additional notation we will use throughout this section.

Additional notation. For any two feature maps $\phi, \psi : \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^d$, we denote by $[\phi, \psi] : \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^{2d}$ the vertical concatenation of the two feature maps. For any $h \in [H]$, $t \in [T]$, policy $\pi \in \Pi$, and history $\mathcal{H}^{t-1} = (x_{1:H}^{1:t-1}, a_{1:H}^{1:t-1})$, we denote by $Q_h^\pi(\cdot, \cdot; \mathcal{H}^{t-1})$ the Q -function at layer h corresponding to rollout policy π ; that is,

$$Q_h^\pi(x, a; \mathcal{H}^{t-1}) := \mathbb{E}^\pi \left[\sum_{s=h}^H \ell_s(\mathbf{x}_s, \mathbf{a}_s; \mathcal{H}^{t-1}) \mid \mathbf{x}_h = x, \mathbf{a}_h = a \right]. \quad (16)$$

Finally, we let $V_h^\pi(x; \mathcal{H}^{t-1}) := \max_{a \in \mathcal{A}} Q_h^\pi(x, a; \mathcal{H}^{t-1})$ denote the corresponding V -function.

Algorithm 4 is similar to Algorithm 3 with the following key differences; after computing a policy cover, the algorithm calls RepLearn (a representation learning algorithm initially introduced by Modi et al. (2024) and subsequently refined by Mhammedi et al. (2023)) to compute a feature map ϕ^{rep} . Then, for every $h \in [H]$, the algorithm computes a *spanner*; a set of policies $\Psi_h^{\text{span}} = \{\pi_{h,1}, \dots, \pi_{h,2d}\}$ that act as an approximate spanner for the set $\{\mathbb{E}^\pi[\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h), \phi_h^{\text{loss}}(\mathbf{x}_h, \mathbf{a}_h)] : \pi \in \Pi\} \subseteq \mathbb{R}^{2d}$, where we use $[\cdot, \cdot]$ to denote the vertical concatenation of vectors. These spanner

policies are then used as the exploratory policies after the initial phase; that is, at episode t in each epoch $k \in [K]$, we execute policy $\pi^{(k)}$ sampled from $\rho^{(k)}$, where $\rho^{(k)}$ is set to be the distribution of the random policy: $\mathbb{I}\{\zeta^t = 0\} \cdot \widehat{\pi}^{(k)} + \mathbb{I}\{\zeta^t = 1\} \cdot \pi^t \circ_{h^{t+1}} \widehat{\pi}^{(k)}$, with $\zeta^t \sim \text{Ber}(\nu)$, $h^t \sim \text{unif}([H])$, and $\pi^t \sim \text{unif}(\Psi_{h^t}^{\text{span}})$. Here, the main difference to Algorithm 3 (see also (47)) is that we use $\pi^t \sim \text{unif}(\Psi_{h^t}^{\text{span}})$ instead of $\pi^t \sim \text{unif}(\Psi_{h^t}^{\text{cov}})$. We require these spanner policies instead of policies in the policy cover, as an adaptive adversary's history-dependent losses prevent standard least squares regression due to the lack of permutation invariance of state-action pairs across episodes within an epoch. Estimating the Q-functions is thus more complex, and we approach it in expectation over roll-ins using policies in Ψ^{span} , the “in-expectation” estimation task is in a sense easier.

We now state the guarantee of Algorithm 4.

Theorem 4.2. *Let $\delta \in (0, 1)$ be given and suppose that Assumption 2.1 and Assumption 4.2 hold. Then, for $T = \text{poly}(A, d, H, \log(|\Phi|/\delta))$ sufficiently large, Algorithm 4 guarantees with probability at least $1 - \delta$,*

$$\sum_{t \in [T]} V_h^{\pi^t}(x_1; \mathcal{H}^{t-1}) - \min_{\pi \in \Pi} \sum_{t \in [T]} V_h^{\pi}(x_1; \mathcal{H}^{t-1}) \leq \text{poly}(A, d, H, \log(|\Phi|/\delta)) \cdot T^{4/5}, \quad (17)$$

where π^t is the policy that Algorithm 4 executes at episode $t \in [T]$.

Algorithm 4 Oracle Efficient Algorithm for Adversarial Low-Rank MDPs (Adaptive Adversary).

Input: Number of rounds T , feature class Φ , loss feature ϕ^{loss} , confidence parameter $\delta \in (0, 1)$.

- 1: Set $\varepsilon \leftarrow T^{-1/3}$, $N_{\text{reg}} \leftarrow T^{2/3}$, $\nu \leftarrow N_{\text{reg}}^{-1/4}$, and $\alpha \leftarrow (8Ad)^{-1}$, $T_{\text{cov}} \leftarrow \varepsilon^{-2} Ad^{13} H^6 \log(\Phi/\delta)$.
 - 2: Set $T_{\text{rep}} \leftarrow \alpha^{-1} \varepsilon^{-2} AH \log(|\Phi|/\delta)$, $T_{\text{span}} \leftarrow \alpha^{-2} \varepsilon^{-2} A \log(dH|\Phi|\varepsilon^{-1}\delta^{-1})$,
 - 3: Define $\mathcal{F}_h = \{(x, a) \mapsto \max_{\phi \in \Phi} \bar{\phi}_h(x, a)^\top \bar{\theta} \mid \bar{\phi}_h = [\phi_h^{\text{loss}}, \phi_h], \phi \in \Phi, \bar{\theta} \in \mathbb{B}_{2d}(1)\}$, $\forall h \in [H]$.
 - 4: Get $\Psi_{1:H}^{\text{cov}} \leftarrow \text{VoX}(\Phi, \varepsilon, \delta/4)$.
 - 5: Get $\phi_h^{\text{rep}} \leftarrow \text{RepLearn}(h, \mathcal{F}_{h+1}, \Phi, \text{unif}(\Psi_h^{\text{cov}}), T_{\text{rep}})$, for all $h \in [H - 1]$. // RepLearn as in Mhammedi et al. (2023)
 - 6: For all $h \in [H]$, set $\bar{\phi}_h^{\text{rep}} \leftarrow [\phi_h^{\text{loss}}, \phi_h^{\text{rep}}] \in \mathbb{R}^{2d}$.
 - 7: For $h \in [H]$, set $\Psi_h^{\text{span}} \leftarrow \text{Spanner}(h, \Phi, \Psi_{1:h}^{\text{cov}}, \bar{\phi}_h^{\text{rep}}, T_{\text{span}})$. // Algorithm 7
 - 8: Set $T_0 \leftarrow T_{\text{cov}} + T_{\text{rep}} + T_{\text{span}}$.
 - 9: **for** $k = 1, \dots, (T - T_0)/N_{\text{reg}}$ **do**
 - 10: Define $\widehat{\pi}_h^{(k)}(a \mid x) \propto \exp\left(-\eta \sum_{s < k} \widehat{Q}_h^{(s)}(x, a)\right)$ for $h \in [H]$.
 - 11: **for** $t = T_0 + (k - 1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}$ **do**
 - 12: Define the random variables $\zeta^t \sim \text{Ber}(\nu)$, $h^t \sim \text{unif}([H])$, and $\pi^t \sim \text{unif}(\Psi_{h^t}^{\text{span}})$.
 - 13: Set $\widehat{\pi}^t = \mathbb{I}\{\zeta^t = 0\} \cdot \widehat{\pi}^{(k)} + \mathbb{I}\{\zeta^t = 1\} \cdot \pi^t \circ_{h^{t+1}} \widehat{\pi}^{(k)}$.
 - 14: Execute $\widehat{\pi}^t$, and observe trajectory $(x_1^t, a_1^t, \dots, x_H^t, a_H^t)$.
 - 15: For $h \in [H]$, observe loss $\ell_h^t := \ell_h(x_h^t, a_h^t; \mathcal{H}^{t-1})$, where $\mathcal{H}^{t-1} := (x_{1:H}^{1:t-1}, a_{1:H}^{1:t-1})$.
 - 16: **end for**
 - 17: For $h \in [H]$ and $\mathcal{I}^{(k)} = \{T_0 + (k - 1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}\}$, compute $\hat{\theta}_h^{(k)}$ such that
$$\hat{\theta}_h^{(k)} \leftarrow \underset{\theta \in \mathbb{B}_{2d}(4Hd^2)}{\text{argmin}} \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{h^t = h, \pi^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(x_h^t, a_h^t)^\top \theta - \sum_{s=h}^H \ell_s^t \right) \right| \quad (18)$$
 - 18: Set $\widehat{Q}_h^{(k)}(x, a) = \bar{\phi}_h^{\text{rep}}(x, a)^\top \hat{\theta}_h^{(k)}$, for all $(x, a) \in \mathcal{X} \times \mathcal{A}$.
 - 19: **end for**
-

5 Conclusion

In this paper, we focus on learning low-rank MDPs with unknown transitions and adversarial losses. For the full-information setting, we improve upon previous regret bounds. More importantly, we initiate the study of the challenging bandit feedback setting, developing various algorithms that achieve sublinear regret under different assumptions. However, the optimal $\sqrt{T}$ regret remains out of reach due to the limitations of our two-phase design. An interesting direction for future work is to perform on-the-fly representation learning to adapt to adversarial losses and achieve optimal regret.

References

- Abbeel, P. and Ng, A. Y. (2005). Exploration and apprenticeship learning in reinforcement learning. In *Proceedings of the 22nd international conference on Machine learning*, pages 1–8.
- Agarwal, A., Kakade, S., Krishnamurthy, A., and Sun, W. (2020). Flambe: Structural complexity and representation learning of low rank mdps. *Advances in neural information processing systems*, 33:20095–20107.
- Ayoub, A., Jia, Z., Szepesvari, C., Wang, M., and Yang, L. (2020). Model-based reinforcement learning with value-targeted regression. In *International Conference on Machine Learning*, pages 463–474. PMLR.
- Cai, Q., Yang, Z., Jin, C., and Wang, Z. (2020). Provably efficient exploration in policy optimization. In *International Conference on Machine Learning*, pages 1283–1294. PMLR.
- Chen, F., Mei, S., and Bai, Y. (2022a). Unified algorithms for rl with decision-estimation coefficients: No-regret, pac, and reward-free learning. *arXiv preprint arXiv:2209.11745*.
- Chen, J., Modi, A., Krishnamurthy, A., Jiang, N., and Agarwal, A. (2022b). On the statistical efficiency of reward-free exploration in non-linear rl. *Advances in Neural Information Processing Systems*, 35:20960–20973.
- Chen, L. and Luo, H. (2021). Finding the stochastic shortest path with low regret: The adversarial cost and unknown transition case. In *International Conference on Machine Learning*, pages 1651–1660. PMLR.
- Cheng, Y., Huang, R., Yang, J., and Liang, Y. (2023). Improved sample complexity for reward-free reinforcement learning under low-rank mdps. *arXiv preprint arXiv:2303.10859*.
- Dai, Y., Luo, H., and Chen, L. (2022). Follow-the-perturbed-leader for adversarial markov decision processes with bandit feedback. *Advances in Neural Information Processing Systems*, 35:11437–11449.
- Dai, Y., Luo, H., Wei, C.-Y., and Zimmert, J. (2023). Refined regret for adversarial mdps with linear function approximation. In *International Conference on Machine Learning*, pages 6726–6759. PMLR.
- Dann, C., Wei, C.-Y., and Zimmert, J. (2023). Best of both worlds policy optimization. In *International Conference on Machine Learning*.
- Du, S., Kakade, S., Lee, J., Lovett, S., Mahajan, G., Sun, W., and Wang, R. (2021). Bilinear classes: A structural framework for provable generalization in rl. In *International Conference on Machine Learning*, pages 2826–2836. PMLR.
- Foster, D. J., Kakade, S. M., Qian, J., and Rakhlin, A. (2021). The statistical complexity of interactive decision making. *arXiv preprint arXiv:2112.13487*.
- Foster, D. J., Rakhlin, A., Sekhari, A., and Sridharan, K. (2022). On the complexity of adversarial decision making. *Advances in Neural Information Processing Systems*, 35:35404–35417.
- He, J., Zhou, D., and Gu, Q. (2022). Near-optimal policy optimization algorithms for learning adversarial linear mixture mdps. In *International Conference on Artificial Intelligence and Statistics*, pages 4259–4280. PMLR.
- Huang, A., Chen, J., and Jiang, N. (2023). Reinforcement learning in low-rank mdps with density features. In *International Conference on Machine Learning*, pages 13710–13752. PMLR.
- Jiang, N., Krishnamurthy, A., Agarwal, A., Langford, J., and Schapire, R. E. (2017). Contextual decision processes with low bellman rank are pac-learnable. In *International Conference on Machine Learning*, pages 1704–1713. PMLR.
- Jin, C., Jin, T., Luo, H., Sra, S., and Yu, T. (2020a). Learning adversarial markov decision processes with bandit feedback and unknown transition. In *International Conference on Machine Learning*, pages 4860–4869. PMLR.

- Jin, C., Liu, Q., and Miryoosefi, S. (2021a). Bellman eluder dimension: New rich classes of rl problems, and sample-efficient algorithms. *Advances in neural information processing systems*, 34:13406–13418.
- Jin, C., Yang, Z., Wang, Z., and Jordan, M. I. (2020b). Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pages 2137–2143. PMLR.
- Jin, T., Huang, L., and Luo, H. (2021b). The best of both worlds: stochastic and adversarial episodic mdps with unknown transition. *Advances in Neural Information Processing Systems*, 34:20491–20502.
- Kakade, S. and Langford, J. (2002). Approximately optimal approximate reinforcement learning. In *Proceedings of the Nineteenth International Conference on Machine Learning*, pages 267–274.
- Kong, F., Zhang, X., Wang, B., and Li, S. (2023). Improved regret bounds for linear adversarial mdps via linear optimization. *arXiv preprint arXiv:2302.06834*.
- Lattimore, T. and Szepesvári, C. (2020). *Bandit algorithms*. Cambridge University Press.
- Lee, C.-W., Luo, H., Wei, C.-Y., and Zhang, M. (2020). Bias no more: high-probability data-dependent regret bounds for adversarial bandits and mdps. In *Advances in Neural Information Processing Systems*.
- Liu, H., Wei, C.-Y., and Zimmert, J. (2023). Towards optimal regret in adversarial linear mdps with bandit feedback. *arXiv preprint arXiv:2310.11550*.
- Luo, H., Wei, C.-Y., and Lee, C.-W. (2021). Policy optimization in adversarial mdps: Improved exploration via dilated bonuses. *Advances in Neural Information Processing Systems*, 34:22931–22942.
- Mhammedi, Z., Block, A., Foster, D. J., and Rakhlin, A. (2023). Efficient model-free exploration in low-rank mdps. *arXiv preprint arXiv:2307.03997*.
- Mhammedi, Z., Block, A., Foster, D. J., and Rakhlin, A. (2024a). Efficient model-free exploration in low-rank mdps. *Advances in Neural Information Processing Systems*, 36.
- Mhammedi, Z., Foster, D. J., and Rakhlin, A. (2024b). The power of resets in online reinforcement learning. *arXiv preprint arXiv:2404.15417*.
- Modi, A., Chen, J., Krishnamurthy, A., Jiang, N., and Agarwal, A. (2024). Model-free representation learning and exploration in low-rank mdps. *Journal of Machine Learning Research*, 25(6):1–76.
- Padakandla, S., KJ, P., and Bhatnagar, S. (2020). Reinforcement learning algorithm for non-stationary environments. *Applied Intelligence*, 50(11):3590–3606.
- Ren, T., Zhang, T., Szepesvári, C., and Dai, B. (2022). A free lunch from the noise: Provable and practical exploration for representation learning. In *Uncertainty in Artificial Intelligence*, pages 1686–1696. PMLR.
- Rosenberg, A. and Mansour, Y. (2019). Online stochastic shortest path with bandit feedback and unknown transition function. *Advances in Neural Information Processing Systems*, 32.
- Shani, L., Efroni, Y., Rosenberg, A., and Mannor, S. (2020). Optimistic policy optimization with bandit feedback. In *International Conference on Machine Learning*, pages 8604–8613. PMLR.
- Sherman, U., Cohen, A., Koren, T., and Mansour, Y. (2023a). Rate-optimal policy optimization for linear markov decision processes. *arXiv preprint arXiv:2308.14642*.
- Sherman, U., Koren, T., and Mansour, Y. (2023b). Improved regret for efficient online reinforcement learning with linear function approximation. In *International Conference on Machine Learning*.
- Uehara, M., Zhang, X., and Sun, W. (2021). Representation learning for online and offline rl in low-rank mdps. *arXiv preprint arXiv:2110.04652*.

- Xie, T., Foster, D. J., Bai, Y., Jiang, N., and Kakade, S. M. (2022). The role of coverage in online reinforcement learning. *arXiv preprint arXiv:2210.04157*.
- Zhang, T., Ren, T., Yang, M., Gonzalez, J., Schuurmans, D., and Dai, B. (2022a). Making linear mdps practical via contrastive representation learning. In *International Conference on Machine Learning*, pages 26447–26466. PMLR.
- Zhang, X., Song, Y., Uehara, M., Wang, M., Agarwal, A., and Sun, W. (2022b). Efficient reinforcement learning in block mdps: A model-free representation learning approach. In *International Conference on Machine Learning*, pages 26517–26547. PMLR.
- Zhao, C., Yang, R., Wang, B., and Li, S. (2022). Learning adversarial linear mixture markov decision processes with bandit feedback and unknown transition. In *The Eleventh International Conference on Learning Representations*.
- Zhao, C., Yang, R., Wang, B., Zhang, X., and Li, S. (2024). Learning adversarial low-rank markov decision processes with unknown transition and full-information feedback. *Advances in Neural Information Processing Systems*, 36.
- Zhong, H., Xiong, W., Zheng, S., Wang, L., Wang, Z., Yang, Z., and Zhang, T. (2022). Gec: A unified framework for interactive decision making in mdp, pomdp, and beyond. *arXiv preprint arXiv:2211.01962*.

Appendices

A	Related Work	15
B	Proof of Theorem 3.1 (Model-Based, Full Information)	16
C	Proof of Theorem 3.2 (Model-Based, Bandit Feedback)	17
D	More Details of Inefficient Model-Free Algorithm in Section 4.1	23
	D.1 Algorithm Description	23
	D.2 Analysis of Occupancy Estimation from Algorithm 6	24
	D.3 Regret Analysis	27
E	Proof of Theorem 4.1 (Model-Free, Banfit Feedback)	32
	E.1 Bounding the Regret Term	32
	E.2 Bounding the Bias Term	33
	E.3 Bounding the Regression Error	34
	E.4 Putting It All Together	36
F	Proof of Theorem 4.2 (Model-Free, Bandit Feedback, Adaptive Adversary)	37
	F.1 Bounding the Regret Term	37
	F.2 Bounding the Bias Term	38
	F.3 Bound the Regression Error	39
	F.4 Putting It All Together	40
	F.5 Spanner Guarantee	40
	F.6 Representation + Spanner	42
	F.7 Martingal Concentration	42
G	Policy Cover and Representation Learning Algorithms	44
	G.1 Policy Cover	44
	G.2 Representation Learning	44
H	Lower Bound for Bandit feedback with Unstructured Losses	46
I	Helper Results	46
	I.1 Martingale Concentration and Regression Results	47
	I.2 Online Learning	48
	I.3 Reinforcement Learning	48

A Related Work

Learning low-rank MDPs in the stochastic setting. In the absence of adversarial losses, several general learning frameworks have been developed for super classes of low-rank MDPs which offer tight sample complexity for either reward-based (Jiang et al., 2017; Jin et al., 2021a; Du et al., 2021; Foster et al., 2021; Zhong et al., 2022) or reward-free (Chen et al., 2022b,a; Xie et al., 2022) settings. However, these algorithms require solving non-convex optimization problems on non-convex version spaces, making them computationally inefficient. Oracle-efficient algorithms for low-rank MDPs are first obtained by Agarwal et al. (2020) using a model-based approach, and the sample complexity bound has been largely improved in subsequent works (Uehara et al., 2021; Zhang et al., 2022a; Cheng et al., 2023). The model-based approach, however, necessitates the function class to accurately model the transition, which is a strong requirement. To relax it, Modi et al. (2024); Zhang et al. (2022b) developed oracle-efficient model-free algorithms, but both of them require additional assumptions on the MDP structure. Recently, Mhammedi et al. (2024a) proposed a satisfactory model-free algorithm that removes all these assumptions. Our work leverages their techniques to tackle the more challenging adversarial setting.

Learning adversarial MDPs. Learning adversarial tabular MDPs under bandit feedback and unknown transition has been extensively studied (Rosenberg and Mansour, 2019; Jin et al., 2020a; Lee et al., 2020; Jin et al., 2021b; Shani et al., 2020; Chen and Luo, 2021; Luo et al., 2021; Dai et al., 2022; Dann et al., 2023). This line of work has demonstrated not only $\sqrt{T}$ regret bounds but also several data-dependent bounds.

For adversarial MDPs with a large state space which necessitates the use of function approximation, if the transition is known, Foster et al. (2022) shows that adversarial setting is as easy as the stochastic setting even under general function approximation. For full-information loss feedback with unknown transition, $\sqrt{T}$ bound is derived for both linear mixture MDPs (Cai et al., 2020; He et al., 2022) and linear MDPs (Sherman et al., 2023a). For more challenging low-rank MDPs with unknown features, the best result only achieves $T^{5/6}$ regret (Zhao et al., 2024).

For function approximation with bandit feedback and unknown transition, Zhao et al. (2022) provides $\sqrt{T}$ bound for linear mixture MDPs, but their regret has polynomial dependence on the size of the state space due to the lack of structure on the loss function. For linear MDPs, a series of recent work has made significant progress in improving the regret bound (Luo et al., 2021; Dai et al., 2023; Sherman et al., 2023b; Kong et al., 2023; Liu et al., 2023). The state-of-the-art result by Liu et al. (2023) gives an inefficient algorithm with $\sqrt{T}$ regret and an efficient algorithm with $T^{3/4}$ regret. These regret bounds for linear MDPs do not depend on the state space size because of the linear loss assumption. We show in Appendix H that cross-state structure on the losses is necessary for low-rank MDPs with bandit feedback to achieve regret bound that do not scale with the number of states.

B Proof of Theorem 3.1 (Model-Based, Full Information)

Theorem B.1 (Theorem 3 in Cheng et al. (2023)). *With probability $1 - \delta$, for any policy π and layer h , Algorithm 1 in Cheng et al. (2023) outputs transition $\widehat{P}_{1:H}$ and features $\hat{\phi}_h, \hat{\mu}_h$ such that $\widehat{P}_h(x' | x, a) = \hat{\phi}_h(x, a)^\top \hat{\mu}_{h+1}(x')$ and*

$$\mathbb{E}^\pi \left[\left\| \widehat{P}_h(\cdot | \mathbf{x}_h, \mathbf{a}_h) - P_h^*(\cdot | \mathbf{x}_h, \mathbf{a}_h) \right\|_1 \right] \leq \epsilon,$$

if the number of collected trajectories is at least $\mathcal{O}\left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|)\right)$.

Define $\widehat{V}_1^\pi(x_1; \ell)$ as the value function of policy π under transition $\{\widehat{P}_h\}_{h=1}^H$ and loss ℓ . We have

$$\begin{aligned} \text{Reg}_T(\pi^*) &= \sum_{t=1}^T V_1^{\pi^*}(x_1; \ell^t) - \sum_{t=1}^T V_1^{\pi^*}(x_1; \ell^t) \\ &= \underbrace{\sum_{t=1}^T \left(V_1^{\pi^*}(x_1; \ell^t) - \widehat{V}_1^{\pi^*}(x_1; \ell^t) \right)}_{\text{Bias1}} + \underbrace{\sum_{t=1}^T \left(\widehat{V}_1^{\pi^*}(x_1; \ell^t) - V_1^{\pi^*}(x_1; \ell^t) \right)}_{\text{Bias2}} \\ &\quad + \underbrace{\sum_{t=1}^T \left(\widehat{V}_1^{\pi^*}(x_1; \ell^t) - \widehat{V}_1^{\pi^*}(x_1; \ell^t) \right)}_{\text{FTRL}} + \mathcal{O}\left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|)\right) \end{aligned}$$

Bounding the bias term. By Theorem B.1 and Lemma I.6, we have

$$\text{Bias1} + \text{Bias2} \leq 2H^2 \epsilon T.$$

Bounding the FTRL term. Since $\eta = \frac{1}{H\sqrt{T}}$, we have

$$\begin{aligned} \text{FTRL} &\leq \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}^{\widehat{P}, \pi^*} \left[\left\langle \widehat{Q}_h^t(\mathbf{x}_h, \cdot), \pi_h^t(\cdot | \mathbf{x}_h) - \pi_h^*(\cdot | \mathbf{x}_h) \right\rangle \right] && (\text{Lemma I.7}) \\ &\leq \frac{H \log |\mathcal{A}|}{\eta} + \eta \sum_{h=1}^H \sum_{t=1}^T \mathbb{E}^{\widehat{P}, \pi^* \circ_h \pi^t} \left[\left(\widehat{Q}_h^t(\mathbf{x}_h, \mathbf{a}_h) \right)^2 \right], && (\text{Lemma I.5}) \\ &\leq \frac{H \log |\mathcal{A}|}{\eta} + 2H^3 \eta T && (\widehat{Q}_h^t(\mathbf{x}_h, \mathbf{a}_h) \leq H) \\ &= \mathcal{O}\left(H^2 \sqrt{T} \log |\mathcal{A}|\right). \end{aligned}$$

Thus, by setting $\epsilon = (Hd^2|\mathcal{A}|(d^2 + |\mathcal{A}|))^{\frac{1}{3}} T^{-\frac{1}{3}}$, we have

$$\begin{aligned} \text{Reg}_T &\leq \mathcal{O}\left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|) + 2H^2 \epsilon T + H^2 \sqrt{T} \log |\mathcal{A}|\right) \\ &\leq \mathcal{O}\left(H^3 (d^2 + |\mathcal{A}|) T^{\frac{2}{3}} \log(|\mathcal{A}| + dH|\Phi||\Upsilon|T)\right). \end{aligned}$$

C Proof of Theorem 3.2 (Model-Based, Bandit Feedback)

Lemma C.1. With $\widehat{P}_{1:H}$ as in Theorem B.1, we have for any $h \in [H]$ and any policy π ,

$$\sum_{x \in \mathcal{X}} |\hat{d}_h^\pi(x) - d_h^\pi(x)| \leq \sum_{i=1}^{h-1} \mathbb{E}^\pi [\|\widehat{P}_i(\cdot | \mathbf{x}_i, \mathbf{a}_i) - P_i(\cdot | \mathbf{x}_i, \mathbf{a}_i)\|_1] \leq (h-1) \cdot \epsilon.$$

where $\hat{d}_h^\pi := d_h^{\widehat{P}, \pi}$.

Proof. We prove the claim by induction. When $h = 1$, given that the $\|d_1^\pi - \hat{d}_1^\pi\|_1 = 0$. Assume

$$\sum_{x \in \mathcal{X}} |\hat{d}_h^\pi(x) - d_h^\pi(x)| \leq \sum_{i=1}^{h-1} \mathbb{E}^\pi [\|\widehat{P}_i(\cdot | \mathbf{x}_i, \mathbf{a}_i) - P_i(\cdot | \mathbf{x}_i, \mathbf{a}_i)\|_1].$$

We have

$$\begin{aligned} & \sum_{x \in \mathcal{X}_{h+1}} |\hat{d}_{h+1}^\pi(x) - d_{h+1}^\pi(x)| \\ &= \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \sum_{x' \in \mathcal{X}_{h+1}} |\hat{d}_h^\pi(x) \pi_h(a | x) \cdot \widehat{P}_h(x' | x, a) - d_h^\pi(x) \pi_h(a | x) \cdot P_h^*(x' | x, a)|, \\ &\leq \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \sum_{x' \in \mathcal{X}_{h+1}} |\hat{d}_h^\pi(x) \pi_h(a | x) \cdot \widehat{P}_h(x' | x, a) - d_h^\pi(x) \pi_h(a | x) \cdot \widehat{P}_h(x' | x, a)| \\ &\quad + \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \sum_{x' \in \mathcal{X}_{h+1}} |d_h^\pi(x) \pi_h(a | x) \cdot \widehat{P}_h(x' | x, a) - d_h^\pi(x) \pi_h(a | x) \cdot P_h^*(x' | x, a)|, \\ &\leq \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \sum_{x' \in \mathcal{X}_{h+1}} \widehat{P}_h(x' | x, a) \pi_h(a | x) \cdot |\hat{d}_h^\pi(x) - d_h^\pi(x)| \\ &\quad + \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \sum_{x' \in \mathcal{X}_{h+1}} d_h^\pi(x) \pi_h(a | x) \cdot |\widehat{P}_h(x' | x, a) - P_h^*(x' | x, a)|, \\ &\leq \sum_{x \in \mathcal{X}} |\hat{d}_h^\pi(x) - d_h^\pi(x)| + \mathbb{E}^\pi [\|\widehat{P}_h(\cdot | \mathbf{x}_h, \mathbf{a}_h) - P_h(\cdot | \mathbf{x}_h, \mathbf{a}_h)\|_1], \\ &\leq \sum_{i=1}^h \mathbb{E}^\pi [\|\widehat{P}_i(\cdot | \mathbf{x}_i, \mathbf{a}_i) - P_i(\cdot | \mathbf{x}_i, \mathbf{a}_i)\|_1], \end{aligned}$$

where the last step follows by the induction hypothesis. The second inequality of Lemma C.1 directly comes from Theorem B.1. $\square$

Our candidate policy space Π' has a β mixture of the random policy. For any deterministic policy π_0^* , define policy π^* such that for any state $x \in \mathcal{X}$, we have $\pi^*(\cdot | x) = (1 - \beta)\pi_0^*(\cdot | x) + \frac{\beta}{|\mathcal{A}|}$. We have $\pi^* \in \Pi'$. Define $\widehat{V}_1^\pi(x_1; \ell)$ as the value function of policy π under transition $\{\widehat{P}_h\}_{h=1}^H$ and loss ℓ .

For any policy π_0^* , we have

$$\text{Reg}_T(\pi_0^*) \tag{19}$$

$$\begin{aligned} &= \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} \rho^t(\pi) V_1^\pi(\mathbf{x}_1; \ell^t) - V_1^{\pi_0^*}(\mathbf{x}_1; \ell^t) \right] + \mathcal{O} \left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|) \right), \\ &= \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) V_1^\pi(\mathbf{x}_1; \ell^t) - V_1^{\pi^*}(\mathbf{x}_1; \ell^t) \right] + \mathcal{O} \left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|) \right) \tag{20} \\ &\quad + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} (\rho^t(\pi) - p^t(\pi)) V_1^\pi(\mathbf{x}_1; \ell^t) \right]}_{\text{Error1}} + \underbrace{\mathbb{E} \left[\sum_{t=1}^T V_1^{\pi^*}(\mathbf{x}_1; \ell^t) - V_1^{\pi_0^*}(\mathbf{x}_1; \ell^t) \right]}_{\text{Error2}}, \\ &= \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) (V_1^\pi(\mathbf{x}_1; \ell^t) - \widehat{V}_1^\pi(\mathbf{x}_1; \ell^t)) \right]}_{\text{Bias1}} + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \widehat{V}_1^{\pi^*}(\mathbf{x}_1; \ell^t) - V_1^{\pi^*}(\mathbf{x}_1; \ell^t) \right]}_{\text{Bias2}} \end{aligned}$$

$$\begin{aligned}
& + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) \widehat{V}_1^{\pi}(\mathbf{x}_1; \ell^t) - \widehat{V}_1^{\pi^*}(\mathbf{x}_1; \ell^t) \right]}_{\text{EXP}} + \mathbf{Error1} + \mathbf{Error2} \\
& + \mathcal{O} \left(\frac{H^3 d^2 |\mathcal{A}| (d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2 (TdH|\Phi||\Upsilon|) \right). \tag{21}
\end{aligned}$$

Recall that $J = \text{John}\{\hat{\phi}_h(\pi)\}_{\pi \in \Pi', h \in [H]} \in \Delta(\Pi' \times H)$, and we have $\rho^t(\pi) = (1 - \gamma)p^t(\pi) + \gamma \sum_{h=1}^H J(\pi, h)$ where $p^t(\pi)$ is defined in Line 7 of Algorithm 2.

Lemma C.2. *We have*

$$\mathbf{Error1} + \mathbf{Error2} \leq H\gamma T + 2H^2\beta T.$$

Proof.

$$\begin{aligned}
\mathbf{Error1} &= \gamma \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} \left(\sum_{h=1}^H J(\pi, h) - p^t(\pi) \right) \cdot V_1^{\pi}(\mathbf{x}_1; \ell^t) \right] \leq H\gamma T. \\
\mathbf{Error2} &= \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} |\pi_0^*(a | \mathbf{x}_h) - \pi^*(a | \mathbf{x}_h)| \cdot Q_h^{\pi_0^*}(\mathbf{x}_h, a; \ell^t) \right] \leq 2H^2\beta T,
\end{aligned}$$

where the last step follows by Lemma I.7. $\square$

Lemma C.3. *We have*

$$\mathbf{Bias1} + \mathbf{Bias2} \leq H^2 T \epsilon.$$

Proof. This is a direct result combining Theorem B.1 and Lemma I.6. $\square$

For $(x_{1:H}, a_{1:H}) \in \mathcal{X}^H \times \mathcal{A}^H$ and $\pi \in \Pi'$ where Π' defined in Line 5 in Algorithm 2 is the mix of a given policy class Π and a uniform policy, and is also the policy class we play with. Recall that the loss estimator

$$\begin{aligned}
\hat{\ell}^t(\pi; \pi^t, x_{1:H}, a_{1:H}) &:= \frac{\pi_1(a_1 | x_1)}{\pi_1^t(a_1 | x_1)} \ell_1^t(x_1, a_1) \\
&+ \sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \frac{\pi_h(a_h | x_h)}{\pi_h^t(a_h | x_h)} \ell_h^t(x_h, a_h). \tag{22}
\end{aligned}$$

defined in Line 10 of Algorithm 2.

Lemma C.4. *For any episode $t \in [T]$, for any policy π we have*

$$\begin{aligned}
\widehat{V}_1^{\pi}(x_1; \ell^t) &\leq \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] + \sqrt{d}H\epsilon \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}, \\
\widehat{V}_1^{\pi}(x_1; \ell^t) &\geq \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - \sqrt{d}H\epsilon \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}.
\end{aligned}$$

Proof. First, from the definition in Line 4 of Algorithm 2, for any $x \in \mathcal{X}$ with $h \geq 2$, we have

$$\hat{d}_h^{\pi}(x) = \sum_{x', a'} \hat{d}_{h-1}^{\pi}(x', a') \cdot \widehat{P}_{h-1}(x | x', a') = \hat{\phi}_{h-1}(\pi)^\top \hat{\mu}_h(x)$$

where $\hat{\phi}_{h-1}(\pi) = \sum_{x', a'} \hat{d}_{h-1}^{\pi}(x', a') \hat{\phi}_{h-1}(x', a')$.

We now prove the first inequality:

$$\widehat{V}_1^{\pi}(x_1; \ell^t)$$

$$\begin{aligned}
&= \sum_{h=1}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{d}_h^\pi(x_h) \pi_h(a_h | x_h) \cdot \ell_h^t(x_h, a_h), \\
&= \underbrace{\sum_{a_1 \in \mathcal{A}} \pi_1(a_1 | x_1) \cdot \ell_1^t(x_1, a_1)}_{\text{First}} + \underbrace{\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top \hat{\mu}_h(x_h) \pi_h(a_h | x_h) \cdot \ell_h^t(x_h, a_h)}_{\text{Remain}} \quad (23)
\end{aligned}$$

Through importance sampling, we have

$$\text{First} = \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\frac{\pi_1(a_1 | x_1)}{\pi_1^t(a_1 | x_1)} \cdot \ell_1^t(x_1, a_1) \right].$$

We now bound the remaining term in (42) Since $\Sigma_{h-1}^t = \mathbb{E}_{\pi^t \sim \rho^t} [\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top]$, we have

Remain

$$\begin{aligned}
&= \sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \mathbb{E}_{\pi^t \sim \rho^t} [\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top] \hat{\mu}_h(x_h) \pi_h(a_h | x_h) \ell_h^t(x_h, a_h), \\
&= \mathbb{E}_{\pi^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \underbrace{\hat{\phi}_{h-1}(\pi^t)^\top \hat{\mu}_h(x_h)}_{\hat{d}_h^{\pi^t}(x_h)} \pi_h(a_h | x_h) \ell_h^t(x_h, a_h) \right], \\
&= \mathbb{E}_{\pi^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \hat{d}_h^{\pi^t}(x_h) \pi_h(a_h | x_h) \ell_h^t(x_h, a_h) \right], \\
&= \mathbb{E}_{\pi^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) d_h^{\pi^t}(x_h) \pi_h(a_h | x_h) \ell_h^t(x_h, a_h) \right] \\
&\quad + \mathbb{E}_{\pi^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \left(\hat{d}_h^{\pi^t}(x_h) - d_h^{\pi^t}(x_h) \right) \pi_h(a_h | x_h) \ell_h^t(x_h, a_h) \right], \quad (24)
\end{aligned}$$

$$\begin{aligned}
&\leq \mathbb{E}_{\pi^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) d_h^{\pi^t}(x_h, a_h) \frac{\pi_h(a_h | x_h)}{\pi_h^t(a_h | x_h)} \ell_h^t(x_h, a_h) \right] \\
&\quad + \sum_{h=2}^H \left\| \hat{\phi}_{h-1}(\pi) \right\|_{(\Sigma_{h-1}^t)^{-1}} \mathbb{E}_{\pi^t \sim \rho^t} \left[\left\| \hat{\phi}_{h-1}(\pi^t) \right\|_{(\Sigma_{h-1}^t)^{-1}} \sum_{x_h \in \mathcal{X}} \left| \hat{d}_h^{\pi^t}(x_h) - d_h^{\pi^t}(x_h) \right| \right], \quad (25)
\end{aligned}$$

$$\begin{aligned}
&\leq \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \frac{\pi_h(a_h | x_h)}{\pi_h^t(a_h | x_h)} \ell_h^t(x_h, a_h) \right] \\
&\quad + \sqrt{d} H \epsilon \sum_{h=2}^H \left\| \hat{\phi}_{h-1}(\pi) \right\|_{(\Sigma_{h-1}^t)^{-1}}, \quad (26)
\end{aligned}$$

where (25) follows by Cauchy-Schwarz, and the last inequality uses Lemma C.1. Adding up **First** and **Remain**, and using the definition of $\hat{\ell}$ in (22) implies the first inequality of the lemma.

The second inequality follows the same procedure except for applying Cauchy-Schwarz inequality in the opposite direction in Eq. (24). $\square$

Lemma C.5. If $\eta \leq \left(\frac{2|A|Hd}{\beta\gamma} + dH^2 \frac{\epsilon}{\sqrt{\gamma}} \right)^{-1}$, then

$$\mathbf{EXP} \leq \frac{\log |\Pi|}{\eta} + \frac{6d\eta H^2 |A|T}{\beta} + 4\eta d^2 H^4 \epsilon^2 T + 4dH^2 \epsilon T,$$

where **EXP** is as in (21).

Proof. Recall in Line 10 of Algorithm 2, $b^t(\pi) := \sqrt{d} H \epsilon \Sigma_{h=2}^H \left\| \hat{\phi}_{h-1}(\pi) \right\|_{(\Sigma_{h-1}^t)^{-1}}$, for all $\pi \in \Pi'$. By Lemma C.4, we have

$$\mathbf{EXP} = \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) \widehat{V}_1^\pi(x_1; \ell^t) - \widehat{V}_1^{\pi^*}(x_1; \ell^t) \right]$$

$$\begin{aligned}
&\leq \sum_{t=1}^T \sum_{\pi} p^t(\pi) \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - \sum_{t=1}^T \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi^*; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] \\
&\quad + \sqrt{d}H\epsilon \sum_{t=1}^T \sum_{\pi} p^t(\pi) \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}} + \sqrt{d}H\epsilon \sum_{t=1}^T \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi^*)\|_{(\Sigma_{h-1}^t)^{-1}}, \\
&= \sum_{t=1}^T \sum_{\pi} p^t(\pi) \cdot \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi) \right) \\
&\quad - \sum_{t=1}^T \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi^*; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi^*) \right) \\
&\quad + \sum_{t=1}^T \sum_{\pi} p^t(\pi) b^t(\pi) - \sum_{t=1}^T b^t(\pi^*) + \sqrt{d}H\epsilon \sum_{t=1}^T \sum_{\pi} p^t(\pi) \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}} \\
&\quad + \sqrt{d}H\epsilon \sum_{t=1}^T \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi^*)\|_{(\Sigma_{h-1}^t)^{-1}} \\
&= \sum_{t=1}^T \sum_{\pi} p^t(\pi) \cdot \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi) \right) \\
&\quad - \sum_{t=1}^T \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi^*; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi^*) \right) \\
&\quad + 4\sqrt{d}H\epsilon \sum_{t=1}^T \sum_{\pi} p^t(\pi) \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}, \quad (\text{since } p^t(\pi) \leq 2\rho^t(\pi)) \tag{27} \\
&\leq \mathbf{FTRL} + 4dH^2\epsilon T,
\end{aligned}$$

where

$$\begin{aligned}
\mathbf{FTRL} &:= \sum_{t=1}^T \sum_{\pi} p^t(\pi) \cdot \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi) \right) \\
&\quad - \sum_{t=1}^T \left(\mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi^*; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - b^t(\pi^*) \right). \tag{28}
\end{aligned}$$

We now bound the FTRL term. Since $\rho^t(\pi) = (1 - \gamma)p^t(\pi) + \gamma \sum_{h=1}^H J(\pi, h)$, define $\Sigma_J = \sum_{\pi \in \Pi} \sum_{h=1}^H J(\pi, h) \hat{\phi}_h(\pi) \hat{\phi}_h(\pi)^\top$. Note that for any $t \in [T]$ and $h \in [H]$, we have $\Sigma_h^t \geq \gamma \Sigma_J$.

By the triangle inequality, we have for any π, π^t and $(x_{1:H}, a_{1:H})$,

$$\begin{aligned}
|\hat{\ell}^t(\pi; \pi^t, x_{1:H}, a_{1:H})| &\leq \left| \frac{\pi_1(\mathbf{a}_1 | \mathbf{x}_1)}{\pi_1^t(\mathbf{a}_1 | \mathbf{x}_1)} \cdot \ell_1^t(\mathbf{x}_1, \mathbf{a}_1) \right| \\
&\quad + \sum_{h=2}^H \left| \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \frac{\pi_h(a_h | x_h)}{\pi_h^t(a_h | x_h)} \ell_h^t(x_h, a_h) \right|, \\
&\leq \frac{|\mathcal{A}|}{\beta} + \frac{|\mathcal{A}|}{\beta} \sum_{h=2}^H \left| \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \right|, \\
&\hspace{15em} (\beta\text{-mixture of uniform policy}) \\
&\leq \frac{|\mathcal{A}|}{\beta} + \frac{|\mathcal{A}|}{\beta\gamma} \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{\Sigma_J^{-1}} \|\hat{\phi}_{h-1}(\pi^t)\|_{\Sigma_J^{-1}}, \\
&\leq \frac{|\mathcal{A}|}{\beta} + \frac{|\mathcal{A}|Hd}{\beta\gamma}, \\
&\leq \frac{2|\mathcal{A}|Hd}{\beta\gamma},
\end{aligned}$$

and

$$|b^t(\pi)| = \left| \sqrt{d}H\epsilon \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}} \right| \leq dH^2 \frac{\epsilon}{\sqrt{\gamma}}.$$

To ensure $\eta |\hat{\ell}^t(\pi; \pi^t, x_{1:H}, a_{1:H}) - b^t(\pi)| \leq 1$, it suffices to set $\eta \leq \left(\frac{2|\mathcal{A}|Hd}{\beta\gamma} + dH^2 \frac{\epsilon}{\sqrt{\gamma}} \right)^{-1}$. Under this constraint, from Lemma I.5, we have

$$\begin{aligned} \mathbf{FTRL} &\leq \frac{\log |\Pi|}{\eta} + \underbrace{2\eta \sum_{t=1}^T \sum_{\pi \in \Pi'} p^t(\pi) \cdot \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})^2]}_{\text{Stability-1}} \\ &\quad + \underbrace{\mathbb{E} \left[2\eta \sum_{t=1}^T \sum_{\pi \in \Pi'} p^t(\pi) b^t(\pi)^2 \right]}_{\text{Stability-2}} \end{aligned}$$

For any $t \in [T]$, we have

$$\begin{aligned} &\sum_{\pi \in \Pi'} p^t(\pi) \cdot \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})^2] \\ &\leq H \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\frac{\pi_1(\mathbf{a}_1 | \mathbf{x}_1)^2}{\pi_1^t(\mathbf{a}_1 | \mathbf{x}_1)^2} \ell_1^t(\mathbf{x}_1, \mathbf{a}_1)^2 \right] \\ &\quad + H \sum_{h=2}^H \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\sum_{\pi \in \Pi'} p^t(\pi) \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi) \frac{\pi_h(\mathbf{a}_h | \mathbf{x}_h)^2}{\pi_h^t(\mathbf{a}_h | \mathbf{x}_h)^2} \ell_h^t(\mathbf{x}_h, \mathbf{a}_h)^2 \right], \\ &\leq \frac{H|\mathcal{A}|}{\beta} + 2H \sum_{h=2}^H \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\text{Tr} \left(\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top (\Sigma_{h-1}^t)^{-1} \right) \sum_{a \in \mathcal{A}} \frac{\pi_h(a | \mathbf{x}_h)^2}{\pi_h^t(a | \mathbf{x}_h)} \ell_h^t(\mathbf{x}_h, a)^2 \right], \\ &\quad (29) \\ &\leq \frac{H|\mathcal{A}|}{\beta} + \frac{2H|\mathcal{A}|}{\beta} \sum_{h=2}^H \mathbb{E}_{\pi^t \sim \rho^t} \left[\text{Tr} \left(\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top (\Sigma_{h-1}^t)^{-1} \right) \right], \\ &\leq \frac{3dH^2|\mathcal{A}|}{\beta}. \end{aligned}$$

where (29) follows by the fact that $p^t(\pi) \leq \frac{1}{1-\gamma} \rho^t(\pi)$ and $\frac{1}{1-\gamma} \leq 2$. Thus,

$$\mathbf{Stability-1} = 2\eta \sum_{\pi \in \Pi'} p^t(\pi) \cdot \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})^2] \leq \frac{6d\eta H^2|\mathcal{A}|T}{\beta}.$$

Moreover,

$$\begin{aligned} \mathbf{Stability-2} &= 2\eta \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi \in \Pi'} p^t(\pi) b^t(\pi)^2 \right], \\ &= 2\eta dH^3 \epsilon^2 \mathbb{E} \left[\sum_{t=1}^T \sum_{h=2}^H \sum_{\pi \in \Pi'} p^t(\pi) \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}^2 \right], \\ &\leq 4\eta dH^3 \epsilon^2 \mathbb{E} \left[\sum_{t=1}^T \sum_{h=1}^{H-1} \sum_{\pi \in \Pi'} \rho^t(\pi) \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}^2 \right], \\ &\leq 4\eta d^2 H^4 \epsilon^2 T. \end{aligned}$$

□

Combing Lemma C.2, Lemma C.3 and Lemma C.5, if $\eta \leq \left(\frac{2|\mathcal{A}|Hd}{\beta\gamma} + dH^2 \frac{\epsilon}{\sqrt{\gamma}} \right)^{-1}$, we have

$$\text{Reg}_T(\pi_0^*) \leq H\gamma T + 2H^2\beta T + H^2T\epsilon + \frac{\log |\Pi|}{\eta} + \frac{6d\eta H^2|\mathcal{A}|T}{\beta}$$

$$+ 4\eta d^2 H^4 \epsilon^2 T + 4dH^2 \epsilon T + \mathcal{O}\left(\frac{H^3 d^2 |\mathcal{A}|(d^2 + |\mathcal{A}|)}{\epsilon^2} \log^2(TdH|\Phi||\Upsilon|)\right).$$

By setting $\epsilon = T^{-\frac{1}{3}}$, $\gamma = T^{-\frac{1}{3}}$, $\beta = T^{-\frac{1}{3}}$, $\eta = \frac{1}{4Hd|\mathcal{A}|} T^{-\frac{2}{3}}$, we have for any $\pi_0^* \in \Pi$,

$$\mathbf{Reg}_T(\pi_0^*) \leq \mathcal{O}\left(d^2 H^3 |\mathcal{A}|(d^2 + |\mathcal{A}|) T^{\frac{2}{3}} \log |\Pi| \log^2(TdH|\Phi||\Upsilon|)\right).$$

D More Details of Inefficient Model-Free Algorithm in Section 4.1

D.1 Algorithm Description

In this section, we give a more detailed introduction of the algorithm mentioned in Section 4.1. This algorithm is model-free and achieves $T^{\frac{2}{3}}$ regret, but it is computationally inefficient. We consider the low-rank MDPs with linear losses that satisfies Assumption 2.1 and Assumption 2.2.

Let $\mathcal{C}(S, \epsilon')$ be ϵ' -net of space S . We define necessary policy and function classes in Definition D.1.

Definition D.1. We define linear policy class and its discretization as

$$\Pi_{\text{lin}} = \left\{ \pi : \mathcal{X} \rightarrow \Delta(\mathcal{A}) \mid \pi_h(a \mid x) = \mathbb{I}\{a = \underset{a \in \mathcal{A}}{\operatorname{argmin}} \phi_h(x, a)^\top \theta_h\}, h \in [H], \theta_h \in \mathbb{B}_d(\sqrt{d}HT), \phi \in \Phi \right\}.$$

$$\Pi_{\text{lin}}^{\text{cov}}(\epsilon') = \left\{ \pi : \mathcal{X} \rightarrow \Delta(\mathcal{A}) \mid \pi_h(a \mid x) = \mathbb{I}\{a = \underset{a \in \mathcal{A}}{\operatorname{argmin}} \phi_h(x, a)^\top \theta_h\}, h \in [H], \theta_h \in \mathcal{C}(\mathbb{B}_d(\sqrt{d}HT), \epsilon'), \phi \in \Phi \right\}.$$

Define corresponding function class as follows

$$\mathcal{F}^\pi = \left\{ f : \mathcal{X} \rightarrow [-1, 1] \mid f(x) = \sum_a \pi(a \mid x) \phi(x, a)^\top \theta, \text{ for } \theta \in \mathbb{B}^d(\sqrt{d}) \text{ and } \phi \in \Phi \right\}$$

$$\mathcal{F} = \left\{ f : \mathcal{X} \rightarrow [-1, 1] \mid f \in \bigcup_{\pi \in \Pi_{\text{lin}}} \mathcal{F}^\pi \right\}.$$

Our main algorithm is given in Algorithm 5, which shares the same structure as Algorithm 2, but with a different initial phase to learn expected feature estimator $\hat{\phi}_h(\pi) = \sum_{(x,a) \in \mathcal{X} \times \mathcal{A}} \hat{d}_h^\pi(x) \pi(a \mid x) \phi_h(x, a)$ to approximate $\phi_h^*(\pi) = \sum_{(x,a) \in \mathcal{X} \times \mathcal{A}} d_h^\pi(x) \pi(a \mid x) \phi_h(x, a)$ for every $h \in [H]$. In Algorithm 2, under Assumption 3.1, it is feasible to use established model-based approach to learn an accurate estimated transition $\hat{P}$ together with its feature $\hat{\phi}$. The occupancy estimator $\hat{d}_{1:H}^\pi$ is induced by $\hat{P}$ which also enjoy small errors. However, when we move to model-free settings with Assumption 4.1, there is no existing approach that could guarantee a good estimation for $\hat{d}_h^\pi(x)$ and $\hat{\phi}$.

To tackle this challenge, we first call VoX (Mhammedi et al., 2023) to construct a policy cover $\Psi_{1:H}^{\text{cov}}$, and then play every policy in $\Psi_{1:H}^{\text{cov}}$ for n episodes to collect data. Subsequently, these data are fed into Algorithm 6 to jointly solve estimated occupancy $\hat{d}_{1:H}^\pi$ and feature $\hat{\phi}_{1:H}$. Algorithm 6 is similar to (Liu et al., 2023, Algorithm 1), which is used to estimate occupancy on the fly for linear MDP. In Algorithm 6, given a target policy π , we jointly solve $\hat{\phi}_{1:H} \in \Phi$, $\hat{d}_{1:H}^\pi \in [0, 1]^{|\mathcal{X}|}$, and $(\hat{\xi}_{1:H,f})_{f \in \mathcal{F}^\pi} \subset \mathbb{B}^d(\sqrt{d})$ that satisfies four constrains, where $\hat{\xi}_{h,f}$ is the estimation of $\xi_{h,f}^* := \sum_{x' \in \mathcal{X}} \mu_{h+1}^*(x') f(x')$. The first constraint Eq. (30) ensures the estimated occupancy $\hat{d}_{1:H}^\pi$ are valid distributions. The second constraint Eq. (31) enforces the estimated values to follow the dynamic programming relationship between the occupancy of layer h and layer $h+1$, which helps to control the propagation of estimation errors across layers through the bias of $\hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f}$. The third constrain Eq. (32) and fourth constrain Eq. (33) are then used to bound the estimated bias of $\hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f}$ by utilizing the data collected from policies in policy covers $\Psi_{1:H}^{\text{cov}}$. Note that applying Eq. (33) requires access to the whole state space, which is an additional assumption not needed in previous algorithms. The gurantee of Algorithm 5 is given in Theorem D.1, where the $\tilde{O}$ hides the logarithmic dependence on $d, H, |\mathcal{A}|, T$.

Theorem D.1. Algorithm 5 ensures $\text{Reg}_T \leq \tilde{O}\left(d^8 H^6 |\mathcal{A}| T^{\frac{2}{3}} \log(|\Phi|)\right)$.

Algorithm 5 Model-Free Algorithm for Bandit Feedback

Input: Policy class $\Pi = \Pi_{\text{lin}}^{\text{cov}}(\frac{1}{T})$.

- 1: Set $\epsilon = 18^{-1} d^{\frac{5}{2}} T^{-\frac{1}{3}}, \gamma = T^{-\frac{1}{3}}, \beta = T^{-\frac{1}{3}}, \eta = (4Hd|\mathcal{A}|)^{-1} T^{-\frac{2}{3}}, n = 11250d^{\frac{5}{2}} |\mathcal{A}| T^{\frac{2}{3}} \log \frac{3dnHT|\Phi|}{\delta}$ and $T_0 = \tilde{O}(\epsilon^{-2} |\mathcal{A}| d^{13} H^6 \log(|\Phi|/\delta))$.
- 2: Get $\Psi_{1:H}^{\text{cov}} \leftarrow \text{VoX}(\Phi, \epsilon, \delta)$ using T_0 episodes.
- 3: For every policy $\pi' \in \Psi_{1:H}^{\text{cov}}$, play it for n episodes and get the data set $(\mathcal{D}_h^{\pi'})_{h \in [H]}$ where $\mathcal{D}_h^{\pi'}$ consists of tuples (x, a, x') such that $(x, a) \sim d_h^{\pi'}$ and $x' \sim P^*(\cdot | x, a)$.
- 4: Define the policy space $\Pi' = \{\pi' : \exists \pi \in \Pi, \pi'_h(\cdot | x) = (1 - \beta)\pi_h(\cdot | x) + \beta/|\mathcal{A}|, \forall x, h\}$.
- 5: Get $\hat{\phi}_h(\cdot) \leftarrow \text{EOM-PC}(\Pi', (\mathcal{D}_h^{\pi'})_{h \in [H], \pi' \in \Psi_{1:H}^{\text{cov}}})$ from Algorithm 6.
- 6: **for** $t = T_0 + 1, T_0 + 2, \dots, T$ **do**
- 7: Define $p^t(\pi) \propto \exp(-\eta \sum_{i=1}^{t-1} (\hat{\ell}^i(\pi) - b^i(\pi)))$, for all $\pi \in \Pi'$.
- 8: Let $\rho^t(\pi) = (1 - \gamma)p^t(\pi) + \frac{\gamma}{H-1} \sum_{h=1}^{H-1} J_h$, where $J_h = \text{John}(\hat{\phi}_h(\cdot), \Pi')$. // John as in §2
- 9: Execute policy $\pi^t \sim \rho^t$ and observe trajectory $(\mathbf{x}_{1:H}^t, \mathbf{a}_{1:H}^t)$ and losses $\ell_h^t = \ell_h^t(\mathbf{x}_h^t, \mathbf{a}_h^t)$.
- 10: Define $\Sigma_h^t = \sum_{\pi \in \Pi'} \rho^t(\pi) \cdot \hat{\phi}_h(\pi) \hat{\phi}_h(\pi)^\top$, $b^t(\pi) = d^{\frac{11}{2}} HT^{-\frac{1}{3}} \cdot \sum_{h=1}^{H-1} \|\hat{\phi}_h(\pi)\|(\Sigma_h^t)^{-1}$, and

$$\hat{\ell}^t(\pi) = \frac{\pi_1(\mathbf{a}_1^t | \mathbf{x}_1^t)}{\pi_1^t(\mathbf{a}_1^t | \mathbf{x}_1^t)} \ell_1^t + \sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\pi^t) \frac{\pi_h(\mathbf{a}_h^t | \mathbf{x}_h^t)}{\pi_h^t(\mathbf{a}_h^t | \mathbf{x}_h^t)} \ell_h^t.$$

11: **end for**

Algorithm 6 EOM-PC $(\Pi, (\mathcal{D}_h^{\pi'})_{h \in [H], \pi' \in \Psi_{1:H}^{\text{cov}}})$ (Estimate Occupancy Measure with Policy Cover)

Input: The policy class Π , datasets $(\mathcal{D}_h^{\pi'})_{h \in [H]}$ for every $\pi' \in \Psi_{1:H}^{\text{cov}}$

Jointly find $\hat{\phi}_h \in \Phi$, $(\hat{d}_h^\pi)_{\pi \in \Pi} \in [0, 1]^{|\mathcal{X}|}$, and $(\hat{\xi}_{h,f})_{f \in \mathcal{F}} \subset \mathbb{B}^d(\sqrt{d})$ for any $h \in [H]$ such that for all $\pi \in \Pi$,

$$\sum_{x \in \mathcal{X}} \hat{d}_h^\pi(x) = 1, \quad \forall h \in [H] \quad (30)$$

$$\sum_{x' \in \mathcal{X}} \hat{d}_{h+1}^\pi(x') f(x') = \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} \hat{d}_h^\pi(x) \pi(a|x) \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f}, \quad \forall f \in \mathcal{F}, h \in [H] \quad (31)$$

$$\begin{aligned} \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f})^2 - \min_{(\phi, \xi) \in \Phi \times \mathbb{B}^d(\sqrt{d})} \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi(x, a)^\top \xi)^2 \\ \leq 132d^{\frac{3}{2}} \log(3dnHT|\Phi|/\delta), \quad \forall \pi' \in \Psi_{1:H}^{\text{cov}}, f \in \mathcal{F}, h \in [H] \end{aligned} \quad (32)$$

$$\max_{x, a} |\hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f}| \leq 1 \quad \forall f \in \mathcal{F}, h \in [H] \quad (33)$$

Output: $\hat{\phi}_h : \Pi \rightarrow \mathbb{R}^d$, $\hat{\phi}_h(\pi) = \sum_{(x, a) \in \mathcal{X} \times \mathcal{A}} \hat{d}_h^\pi(x) \pi(a|x) \hat{\phi}_h(x, a)$, $\forall h \in [H]$.

D.2 Analysis of Occupancy Estimation from Algorithm 6

Lemma D.1. With probability $1 - \delta$, $\phi_{1:H}^*$, $(d_{1:H}^\pi)_{\pi \in \Pi'}$, and $\xi_{h,f}^* := \sum_{x' \in \mathcal{X}} \mu_{h+1}^*(x') f(x')$, $\forall f \in \mathcal{F}, \forall h \in [H]$ is a solution to Algorithm 6.

Proof. Since for any policy π and any $h \in [H]$, $\sum_{x \in \mathcal{X}} d_h^\pi(x) = 1$, Eq. (30) holds. For any policy π , any $f \in \mathcal{F}^\pi$ and any $h \in [H]$, we have

$$\begin{aligned} \sum_{x' \in \mathcal{X}} d_{h+1}^\pi(x') f(x') &= \sum_{x' \in \mathcal{X}} \sum_{x \in \mathcal{X}} \sum_{a \in \mathcal{A}} d_h^\pi(x) \pi(a|x) P_h^*(x' | x, a) f(x') \\ &= \sum_{x \in \mathcal{X}_h} \sum_{a \in \mathcal{A}} d_h^\pi(x) \pi(a|x) \phi_h^*(x, a)^\top \sum_{x' \in \mathcal{X}_{h+1}} \mu_{h+1}^*(x') f(x') \\ &= \sum_{x \in \mathcal{X}_h} \sum_{a \in \mathcal{A}} d_h^\pi(x) \pi(a|x) \xi_{h,f}^*. \end{aligned}$$

Thus, Eq. (31) holds. From Exercise 27.6 of Lattimore and Szepesvári (2020), the ϵ -net of $\mathbb{B}_d(R)$ is $\left(\frac{3R}{\epsilon}\right)^d$. Thus, $|\Pi'| = |\Pi_{\text{lin}}^{\text{cov}}(\frac{1}{T})| = |\Phi| (3\sqrt{d}HT^2)^d$. We also have $|\mathcal{C}(\mathbb{B}_d(\sqrt{d}), \frac{1}{T})| = (3\sqrt{d}T)^d$ and for any policy π , $|\mathcal{C}(\mathcal{F}^\pi, \frac{1}{T})| = |\Phi| (3\sqrt{d}T)^d$. To consider all possible instances, define

$$\mathcal{N}_T := |\Pi'| |\mathcal{C}(\mathbb{B}_d(\sqrt{d}), \epsilon)| |\mathcal{C}(\mathcal{F}^\pi, \epsilon)| |\Psi_{1:H}^{\text{cov}}| |\Phi| H \leq dH^2 |\Phi|^3 (3\sqrt{d}HT^2)^{3d}$$

Thus, by union bound, with probability of $1 - \delta$, for every $\pi \in \Pi'$, every $\pi' \in \Psi_{1:H}^{\text{cov}}$, every $f \in \mathcal{C}(\mathcal{F}^\pi, \epsilon)$, every $\xi \in \mathcal{C}(\mathbb{B}_d(\sqrt{d}), \epsilon)$, every $\phi \in \Phi$ and every $h \in [H]$, we have

$$\begin{aligned} &\sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h,f}^*)^2 - \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi(x, a)^\top \xi)^2 \\ &= -2 \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h,f}^*) (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi) \\ &\quad - \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi)^2 \\ &= -2 \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \mathbb{E}_{x' \sim P^*(\cdot|x,a)}[f(x')]) (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi) \\ &\quad - \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi)^2 \\ &\leq 8 \sqrt{\sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi)^2 \log \frac{|\mathcal{N}_T|}{\delta}} + 4\sqrt{d} \log \frac{|\mathcal{N}_T|}{\delta} \\ &\quad - \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h,f}^* - \phi(x, a)^\top \xi)^2 \quad \text{(Freedman's Inequality)} \\ &\leq 20\sqrt{d} \log \frac{|\mathcal{N}_T|}{\delta} \quad \text{(AM-GM)} \\ &\leq 120d^{\frac{3}{2}} \log \frac{3dHT|\Phi|}{\delta} \end{aligned}$$

Bounding the distance through $\frac{1}{T}$ -net, we have with probability of $1 - \delta$, for every $\pi \in \Pi'$, every $\pi' \in \Psi_{1:H}^{\text{cov}}$, every $f \in \mathcal{F}^\pi$, every $\xi \in \mathbb{B}_d(\sqrt{d})$, every $\phi \in \Phi$ and every $h \in [H]$,

$$\begin{aligned} \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h,f}^*)^2 - \sum_{x,a,x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi(x, a)^\top \xi)^2 &\leq 120d^{\frac{3}{2}} \log \frac{3dHT|\Phi|}{\delta\epsilon} + \frac{12\sqrt{d}n}{T} \\ &\leq 132d^{\frac{3}{2}} \log \frac{3dnHT|\Phi|}{\delta} \end{aligned}$$

Thus, Eq. (32) also holds. Finally, for all x, a , we have

$$|\phi_h^*(x, a)^\top \xi_{h,f}^*| = \left| \sum_{x' \in \mathcal{X}} \phi_h^*(x, a)^\top \mu_{h+1}^*(x') f(x') \right| = \left| \sum_{x' \in \mathcal{X}} P^*(x' | x, a) f(x') \right| \leq 1.$$

Thus, ϕ_h^* , d_h^π , and $\xi_{h,f}^* := \sum_{x' \in \mathcal{X}} \mu_{h+1}^*(x') f(x')$, $\forall f \in \mathcal{F}^\pi, h \in [H]$ satisfy all Eq. (30) – Eq. (33) and is a solution to Algorithm 6. $\square$

Lemma D.2. With probability $1 - \delta$, for all $f \in \mathcal{F}$, any solution $\hat{d}^\pi$ from Algorithm 6 for any $\pi \in \Pi'$ satisfies

$$\left| \sum_x (\hat{d}_h^\pi(x) - d_h^\pi(x)) f(x) \right| \leq d^5 HT^{-\frac{1}{3}}$$

Proof.

For every solution $\hat{\phi}_{1:H}, \hat{\xi}_{1:H, f \in \mathcal{F}^\pi}, (\hat{d}_{1:H}^\pi)_{\pi \in \Pi'}$ of Algorithm 6, we have

$$\begin{aligned} & \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 - \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h, f}^*)^2 \\ &= \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 - \min_{(\phi, \xi)} \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi(x, a)^\top \xi)^2 \\ & \quad + \underbrace{\min_{(\phi, \xi)} \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi(x, a)^\top \xi)^2 - \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h, f}^*)^2}_{\leq 0} \\ &\leq 132d^{\frac{3}{2}} \log \frac{3dnHT|\Phi|}{\delta} \end{aligned} \quad (34)$$

where the last step comes from the constrain Eq. (32). On the other hand

$$\begin{aligned} & 2 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 - 2 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h, f}^*)^2 \\ &= 4 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \phi_h^*(x, a)^\top \xi_{h, f}^*) (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f}) \\ & \quad + 2 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 \\ &= 4 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \mathbb{E}_{x' \sim P^*(\cdot|x, a)}[f(x')]) (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f}) \\ & \quad + 2 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 \end{aligned} \quad (35)$$

Combing Eq. (34) and Eq. (35), we have

$$\begin{aligned} & \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 \\ &\leq -4 \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (f(x') - \mathbb{E}_{x' \sim P^*(\cdot|x, a)}[f(x')]) (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f}) \\ & \quad - \sum_{x, a, x' \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 \\ & \quad + 264d^{\frac{3}{2}} \log \frac{3dnHT|\Phi|}{\delta} \\ &\leq 288d^{\frac{3}{2}} \log \frac{3dnHT|\Phi|}{\delta} \end{aligned} \quad (\text{Lemma I.2 with } \lambda = \frac{1}{8})$$

Since for every data tuple $(x, a, x') \in \mathcal{D}_h^{\pi'}$, $(x, a) \sim d_h^{\pi'}$ independently, by Lemma I.3, for every $\pi' \in \Psi_{1:H}^{\text{cov}}$ we have

$$\mathbb{E}^{\pi'} \left[(\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 \right] \leq \frac{2}{n} \sum_{x, a \in \mathcal{D}_h^{\pi'}} (\phi_h^*(x, a)^\top \xi_{h, f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h, f})^2 + \frac{24d \log \left(\frac{6\sqrt{d}|\Phi|}{\delta} \right)}{n}$$

$$\leq \frac{600d^{\frac{3}{2}}}{n} \log \frac{3dnHT|\Phi|}{\delta}$$

This implies there exists a representation $\hat{\phi}_{1:H}(x, a)$ such that for every $f \in \mathcal{F}$ and any $h \in [H]$, there exists $\hat{\xi}_{h,f}$ such that

$$\max_{\pi' \in \Psi_{1:H}^{\text{cov}}} \mathbb{E}^{\pi'} \left[\left(\phi_h^*(x, a)^\top \xi_{h,f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} \right)^2 \right] \leq \frac{600d^{\frac{3}{2}}}{n} \log \frac{3dnHT|\Phi|}{\delta}. \quad (36)$$

Eq. (36) matches Eq. (99) with a different error bound on the right hand. From Lemma G.1, we have Ψ_h^{cov} is a $(\frac{1}{8Ad}, \varepsilon)$ -policy cover for layer h , following the rest of the proof in Lemma G.2, for every π and every $f \in \mathcal{F}, h \in [H]$, we have

$$\left| \mathbb{E}^\pi \left[\phi_h^*(x, a)^\top \xi_{h,f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} \right] \right| \leq \frac{75d^{\frac{5}{4}}}{\sqrt{n}} \sqrt{|\mathcal{A}| \log \frac{3dnHT|\Phi|}{\delta}} + 9d^{\frac{5}{2}}\epsilon. \quad (37)$$

The $d^{\frac{5}{2}}$ in the second term of Eq. (37) improves the $d^{\frac{7}{2}}$ term in Eq. (98) because $\hat{\xi}_{h,f} \in \mathbb{B}_d(\sqrt{d})$ rather than $\mathbb{B}_d(d^{\frac{3}{2}})$. Putting the choice $n = 11250d^{\frac{5}{2}}|\mathcal{A}|T^{\frac{2}{3}} \log \frac{3dnHT|\Phi|}{\delta}$ and $\epsilon = \frac{d^{\frac{5}{2}}}{18}T^{-\frac{1}{3}}$ into Eq. (37), we have for every π and every $f \in \mathcal{F}, h \in [H]$, we have

$$\left| \mathbb{E}^\pi \left[\phi_h^*(x, a)^\top \xi_{h,f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} \right] \right| \leq d^5 T^{-\frac{1}{3}} \quad (38)$$

Utilizing above results, for every π , every $f \in \mathcal{F}$ and any $h \in [H - 1]$, we have

$$\begin{aligned} & \left| \sum_x (\hat{d}_{h+1}^\pi(x) - d_{h+1}^\pi(x)) f(x) \right| \\ &= \left| \sum_{x,a} \hat{d}_h^\pi(x) \pi_h(a|x) \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} - \sum_{x,a} d_h^\pi(x) \pi(a|x) \phi_h^*(x, a)^\top \xi_{h,f}^* \right| \quad (\text{Eq. (31)}) \\ &\leq \left| \sum_{x,a} d_h^\pi(x) \pi_h(a|x) \left(\phi_h^*(x, a)^\top \xi_{h,f}^* - \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} \right) \right| \\ &\quad + \left| \sum_x (\hat{d}_h^\pi(x) - d_h^\pi(x)) \underbrace{\sum_a \pi_h(a|x) \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f}}_{\in \mathcal{F}} \right| \\ &\leq d^5 T^{-\frac{1}{3}} + \left| \sum_x (\hat{d}_h^\pi(x) - d_h^\pi(x)) f'(x) \right| \quad (\text{Eq. (38)}) \end{aligned}$$

□

where in the last step, we define $f'(x) = \sum_a \pi_h(a|x) \hat{\phi}_h(x, a)^\top \hat{\xi}_{h,f} \in \mathcal{F}$. This allow us to use recursion to finish the proof.

D.3 Regret Analysis

We begin by proving Lemma D.3, showing that the policy class $\Pi_{\text{lin}}^{\text{cov}}(\epsilon')$ suffices to approximate all policies with small error.

Lemma D.3. *For any policy π , there exists a policy $\pi' \in \Pi_{\text{lin}}^{\text{cov}}(\epsilon')$ such that*

$$\sum_{t=1}^T V_1^{\pi'}(x_1; \ell^t) - \sum_{t=1}^T V_1^\pi(x_1; \ell^t) \leq H\epsilon'$$

Proof. Let $\theta_h^{t,\pi} \in \mathbb{B}_d(H\sqrt{d})$ be such that

$$\forall (x, a) \in \mathcal{X} \times \mathcal{A}, \quad Q_h^\pi(x, a; \ell^t) = \phi_h^*(x, a)^\top \theta_h^{\pi,t}.$$

Such a $\theta_h^{\pi,t}$ is guaranteed to exist by the low-rank MDP structure (Assumption 2.1) and Assumption 2.2. For every $h \in [H]$, since $\sum_{t=1}^T \theta_h^{\pi,t} \in \mathbb{B}_d(\sqrt{d}HT)$, we define $\theta'_h \in \mathcal{C}(\mathbb{B}_d(\sqrt{d}HT), \epsilon')$ such that $\|\theta'_h - \sum_{t=1}^T \theta_h^{\pi,t}\|_2 \leq \epsilon'$, and let $\pi_h(a|x) = \mathbb{I}\{a = \operatorname{argmin}_{a \in \mathcal{A}} \phi_h^*(x, a)^\top \theta'_h\}$ for every $h \in [H]$. We have $\pi' \in \Pi_{\text{lin}}^{\text{cov}}(\epsilon')$. From Lemma I.7, we have

$$\begin{aligned} & \sum_{t=1}^T V_1^{\pi'}(x_1; \ell^t) - \sum_{t=1}^T V_1^\pi(x_1; \ell^t) \\ &= \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{x \sim d_h^{\pi'}} \left[\sum_{a \in \mathcal{A}} (\pi'_h(a|x) - \pi_h(a|x)) Q_h^\pi(x, a; \ell^t) \right] \\ &= \sum_{h=1}^H \mathbb{E}_{x \sim d_h^{\pi'}} \left[\sum_{a \in \mathcal{A}} (\pi'_h(a|x) - \pi_h(a|x)) \phi_h^*(x, a)^\top \sum_{t=1}^T \theta_h^{\pi,t} \right] \\ &= \underbrace{\sum_{h=1}^H \mathbb{E}_{x \sim d_h^{\pi'}} \left[\sum_{a \in \mathcal{A}} (\pi'_h(a|x) - \pi_h(a|x)) \phi_h^*(x, a)^\top \theta'_h \right]}_{\leq 0} + \sum_{h=1}^H \mathbb{E}_{x \sim d_h^{\pi'}} \left[\sum_{a \in \mathcal{A}} (\pi'_h(a|x) - \pi_h(a|x)) \phi_h^*(x, a)^\top \left(\sum_{t=1}^T \theta_h^{\pi,t} - \theta'_h \right) \right] \\ &\leq H\epsilon' \end{aligned}$$

where the last inequality comes from the fact that $\pi_h(a|x) = \mathbb{I}\{a = \operatorname{argmin}_{a \in \mathcal{A}} \phi_h^*(x, a)^\top \theta'_h\}$ and $\|\theta'_h - \sum_{t=1}^T \theta_h^{\pi,t}\|_2 \leq \epsilon'$ for every $h \in [H]$. $\square$

From Lemma D.3, for any policy π_1^* , there exists a policy $\pi_0^* \in \Pi_{\text{lin}}^{\text{cov}}(\frac{1}{T})$ such that

$$\begin{aligned} \text{Reg}_T(\pi_1^*) &= \mathbb{E} \left[\sum_{t=1}^T V_1^{\pi^t}(x_1; \ell^t) - \sum_{t=1}^T V_1^{\pi_1^*}(x_1; \ell^t) \right] \\ &\leq \mathbb{E} \left[\sum_{t=1}^T V_1^{\pi^t}(x_1; \ell^t) - \sum_{t=1}^T V_1^{\pi_0^*}(x_1; \ell^t) \right] + 1 \end{aligned} \quad (39)$$

Our candidate policy space Π' has a β mixture of the random policy. For any policy $\pi_0^* \in \Pi_{\text{lin}}^{\text{cov}}(\frac{1}{T})$, define policy π^* such that for any state $x \in \mathcal{X}$, we have $\pi^*(\cdot|x) = (1-\beta)\pi_0^*(\cdot|x) + \frac{\beta}{|\mathcal{A}|}$. We have $\pi^* \in \Pi'$. Define

$$\widehat{V}_1^\pi(x_1; \ell) = \sum_{h=1}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{d}_h^\pi(x_h) \pi_h(a_h|x_h) \cdot \ell_h^t(x_h, a_h)$$

Utilizing Eq. (39) and the decomposition in Eq. (21), together with the fact that $|\Psi_h^{\text{cov}}| \leq d$ from Lemma G.1, we have for any policy π_1^* ,

$$\begin{aligned} & \text{Reg}_T(\pi_1^*) \\ &\leq \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) (V_1^\pi(x_1; \ell^t) - \widehat{V}_1^\pi(x_1; \ell^t)) \right]}_{\text{Bias1}} + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \widehat{V}_1^{\pi^*}(x_1; \ell^t) - V_1^{\pi^*}(x_1; \ell^t) \right]}_{\text{Bias2}} \\ &\quad + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} p^t(\pi) \widehat{V}_1^\pi(x_1; \ell^t) - \widehat{V}_1^{\pi^*}(x_1; \ell^t) \right]}_{\text{EXP}} + \underbrace{\mathbb{E} \left[\sum_{t=1}^T \sum_{\pi} (\rho^t(\pi) - p^t(\pi)) V_1^\pi(x_1; \ell^t) \right]}_{\text{Error1}} \\ &\quad + \underbrace{\mathbb{E} \left[\sum_{t=1}^T V_1^{\pi^*}(x_1; \ell^t) - V_1^{\pi_0^*}(x_1; \ell^t) \right]}_{\text{Error2}} + \widetilde{O}(\epsilon^{-2} |\mathcal{A}| d^{13} H^6 \log(|\Phi|/\delta)) + dHn \end{aligned} \quad (40)$$

Following Lemma C.2, we have

$$\text{Error1} + \text{Error2} \leq H\gamma T + 2H^2\beta T \leq 3H^2T^{\frac{2}{3}}. \quad (41)$$

Lemma D.4.

$$\mathbf{Bias1} + \mathbf{Bias2} \leq 2d^5 H^2 T^{\frac{2}{3}}$$

Proof. For any policy π and any $t \in [T]$,

$$\begin{aligned} |V_1^\pi(\mathbf{x}_1; \ell^t) - \widehat{V}_1^\pi(\mathbf{x}_1; \ell^t)| &= \sum_{h=1}^H \sum_{x_h \in \mathcal{X}} |d_h^\pi(x_h) - \hat{d}_h^\pi(x_h)| \sum_{a_h \in \mathcal{A}} \pi_h(a_h | x_h) \cdot \ell_h^t(x_h, a_h) \\ &\leq d^5 H^2 T^{-\frac{1}{3}} \end{aligned} \quad (\text{Lemma D.2})$$

Thus,

$$\mathbf{Bias1} \leq d^5 H^2 T^{\frac{2}{3}} \quad \text{and} \quad \mathbf{Bias2} \leq d^5 H^2 T^{\frac{2}{3}}$$

□

We now prove a modle-free counterpart of Lemma C.4 in Lemma D.5.

Lemma D.5. For any episode $t \in [T]$, for any policy π we have

$$\begin{aligned} \widehat{V}_1^\pi(x_1; \ell^t) &\leq \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] + d^{\frac{11}{2}} H T^{-\frac{1}{3}} \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}, \\ \widehat{V}_1^\pi(x_1; \ell^t) &\geq \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} [\hat{\ell}^t(\pi; \pi^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H})] - d^{\frac{11}{2}} H T^{-\frac{1}{3}} \sum_{h=2}^H \|\hat{\phi}_{h-1}(\pi)\|_{(\Sigma_{h-1}^t)^{-1}}. \end{aligned}$$

Proof.

We now prove the first inequality:

$$\begin{aligned} \widehat{V}_1^\pi(x_1; \ell^t) &= \sum_{h=1}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{d}_h^\pi(x_h) \pi_h(a_h | x_h) \cdot \ell_h^t(x_h, a_h), \\ &= \underbrace{\sum_{a_1 \in \mathcal{A}} \pi_1(a_1 | x_1) \cdot \ell_1^t(x_1, a_1)}_{\text{First}} + \underbrace{\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{d}_h^\pi(x_h) \pi_h(a_h | x_h) \phi_h(x_h, a_h)^\top g_h^t}_{\text{Remain}} \end{aligned} \quad (42)$$

Through importance sampling, we have

$$\mathbf{First} = \mathbb{E}_{\pi^t \sim \rho^t} \mathbb{E}^{\pi^t} \left[\frac{\pi_1(\mathbf{a}_1 | \mathbf{x}_1)}{\pi_1^t(\mathbf{a}_1 | \mathbf{x}_1)} \cdot \ell_1^t(\mathbf{x}_1, \mathbf{a}_1) \right].$$

We now bound the remaining term in (42) Since $\Sigma_{h-1}^t = \mathbb{E}_{\pi^t \sim \rho^t} [\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top]$, we have

Remain

$$\begin{aligned} &= \sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \hat{d}_h^\pi(x_h) \underbrace{\sum_{a_h \in \mathcal{A}} \pi_h(a_h | x_h) \phi_h(x_h, a_h)^\top g_h^t}_{:= f(x_h)}, \\ &= \sum_{h=2}^H \sum_{x_{h-1} \in \mathcal{X}_{h-1}} \sum_{a_{h-1}} \hat{d}_{h-1}^\pi(x) \pi_{h-1}(a_{h-1} | x_{h-1}) \hat{\phi}(x_{h-1}, a_{h-1})^\top \hat{\xi}_{h-1, f} \quad (\text{by Eq. (31)}) \\ &= \sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top \hat{\xi}_{h-1, f} \\ &= \sum_{h=2}^H \hat{\phi}_{h-1}(\pi)^\top (\Sigma_{h-1}^t)^{-1} \mathbb{E}_{\pi^t \sim \rho^t} [\hat{\phi}_{h-1}(\pi^t) \hat{\phi}_{h-1}(\pi^t)^\top] \hat{\xi}_{h-1, f} \end{aligned}$$

$$\begin{aligned}
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \hat{\phi}_{h-1}(\boldsymbol{\pi}^t)^\top \hat{\xi}_{h-1,f} \right] \\
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \sum_{x_{h-1} \in \mathcal{X}_{h-1}} \sum_{a_{h-1}} \hat{d}_{h-1}^{\boldsymbol{\pi}^t}(x) \pi_{h-1}(a_{h-1}|x_{h-1}) \hat{\phi}(x_{h-1}, a_{h-1})^\top \hat{\xi}_{h-1,f} \right] \\
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \sum_{x_h \in \mathcal{X}_h} \hat{d}_h^{\boldsymbol{\pi}^t}(x_h) f(x_h) \right] \\
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \sum_{x_h \in \mathcal{X}_h} \sum_{a_h} \hat{d}_h^{\boldsymbol{\pi}^t}(x_h) \pi_h(a_h|x_h) \phi_h(x_h, a_h)^\top g_h^t \right] \\
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \sum_{x_h \in \mathcal{X}_h} \sum_{a_h} \hat{d}_h^{\boldsymbol{\pi}^t}(x_h) \pi_h(a_h|x_h) \phi_h(x_h, a_h)^\top g_h^t \right] \\
&\quad + \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \sum_{x_h \in \mathcal{X}_h} \sum_{a_h} \left(\hat{d}_h^{\boldsymbol{\pi}^t}(x_h) - d_h^{\boldsymbol{\pi}^t}(x_h) \right) \pi_h(a_h|x_h) \phi_h(x_h, a_h)^\top g_h^t \right] \\
&= \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) d_h^{\boldsymbol{\pi}^t}(x_h) \pi_h(a_h|x_h) \ell_h^t(x_h, a_h) \right] \\
&\quad + d^5 HT^{-\frac{1}{3}} \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \left| \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \right| \right], \quad (\text{by Lemma D.2}) \\
&\leq \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\sum_{h=2}^H \sum_{x_h \in \mathcal{X}} \sum_{a_h \in \mathcal{A}} \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) d_h^{\boldsymbol{\pi}^t}(x_h, a_h) \frac{\pi_h(a_h|x_h)}{\pi_h^t(a_h|x_h)} \ell_h^t(x_h, a_h) \right] \\
&\quad + d^5 HT^{-\frac{1}{3}} \sum_{h=2}^H \left\| \hat{\phi}_{h-1}(\boldsymbol{\pi}) \right\|_{(\Sigma_{h-1}^t)^{-1}} \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \left[\left\| \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \right\|_{(\Sigma_{h-1}^t)^{-1}} \right], \quad (43) \\
&\leq \mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \mathbb{E}^{\boldsymbol{\pi}^t} \left[\sum_{h=2}^H \hat{\phi}_{h-1}(\boldsymbol{\pi})^\top (\Sigma_{h-1}^t)^{-1} \hat{\phi}_{h-1}(\boldsymbol{\pi}^t) \frac{\pi_h(\mathbf{a}_h^t | \mathbf{x}_h^t)}{\pi_h^t(\mathbf{a}_h^t | \mathbf{x}_h^t)} \ell_h^t(\mathbf{x}_h^t, \mathbf{a}_h^t) \right] \\
&\quad + d^{\frac{11}{2}} HT^{-\frac{1}{3}} \sum_{h=2}^H \left\| \hat{\phi}_{h-1}(\boldsymbol{\pi}) \right\|_{(\Sigma_{h-1}^t)^{-1}},
\end{aligned}$$

Adding up **First** and **Remain**, and using the definition of $\hat{\ell}$ in (22) implies the first inequality of the lemma.

The second inequality follows the same procedure except for applying Cauchy-Schwarz inequality in the opposite direction in Eq. (43). $\square$

Given Lemma D.5, we could follow the same procedure in Lemma C.5 except replacing the factor of bonus $b^t(\boldsymbol{\pi})$ from $\sqrt{d}H\epsilon$ to $d^{\frac{11}{2}}HT^{-\frac{1}{3}}$. This leads to the following changes. Firstly, by a similar argument as Eq. (27), we have

$$\mathbf{EXP} = \mathbf{FTRL} + 4d^6 H^2 T^{\frac{2}{3}}$$

where

$$\begin{aligned}
\mathbf{FTRL} &:= \sum_{t=1}^T \sum_{\boldsymbol{\pi}} p^t(\boldsymbol{\pi}) \cdot \left(\mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \mathbb{E}^{\boldsymbol{\pi}^t} \left[\hat{\ell}^t(\boldsymbol{\pi}; \boldsymbol{\pi}^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H}) \right] - b^t(\boldsymbol{\pi}) \right) \\
&\quad - \sum_{t=1}^T \left(\mathbb{E}_{\boldsymbol{\pi}^t \sim \rho^t} \mathbb{E}^{\boldsymbol{\pi}^t} \left[\hat{\ell}^t(\boldsymbol{\pi}^*; \boldsymbol{\pi}^t, \mathbf{x}_{1:H}, \mathbf{a}_{1:H}) \right] - b^t(\boldsymbol{\pi}^*) \right).
\end{aligned}$$

Secondly, now we have

$$|b^t(\boldsymbol{\pi})| = \left| d^{\frac{11}{2}} HT^{-\frac{1}{3}} \sum_{h=2}^H \left\| \hat{\phi}_{h-1}(\boldsymbol{\pi}) \right\|_{(\Sigma_{h-1}^t)^{-1}} \right| \leq d^6 H^2 \frac{T^{-\frac{1}{3}}}{\sqrt{\gamma}} = d^6 H^2 T^{-\frac{1}{6}}. \quad (\gamma = T^{-\frac{1}{3}})$$

To ensure $\eta \left| \hat{\ell}^t(\pi; \pi^t, x_{1:H}, a_{1:H}) - b^t(\pi) \right| \leq 1$, it suffices to set $\eta \leq \left(\frac{2|\mathcal{A}|Hd}{\beta\gamma} + d^6 H^2 T^{-\frac{1}{6}} \right)^{-1} = \left(2|\mathcal{A}|HdT^{\frac{2}{3}} + d^6 H^2 T^{-\frac{1}{6}} \right)^{-1}$ from $\beta = \gamma = T^{-\frac{1}{3}}$. Thus our choice $\eta = (4Hd|\mathcal{A}|)^{-1} T^{-\frac{2}{3}}$ satisfies the condition if we assume $T \geq d^6 H^2$. Moreover, now we have

$$\begin{aligned}
\textbf{Stability-2} &= 2\eta \mathbb{E} \left[\sum_{t=1}^T \sum_{\pi \in \Pi'} p^t(\pi) b^t(\pi)^2 \right], \\
&= 2\eta d^{11} H^3 T^{-\frac{2}{3}} \mathbb{E} \left[\sum_{t=1}^T \sum_{h=2}^H \sum_{\pi \in \Pi'} p^t(\pi) \left\| \hat{\phi}_{h-1}(\pi) \right\|_{(\Sigma_{h-1}^t)^{-1}}^2 \right], \\
&\leq 4\eta d^{11} H^3 T^{-\frac{2}{3}} \mathbb{E} \left[\sum_{t=1}^T \sum_{h=1}^{H-1} \sum_{\pi \in \Pi'} \rho^t(\pi) \left\| \hat{\phi}_{h-1}(\pi) \right\|_{(\Sigma_{h-1}^t)^{-1}}^2 \right], \\
&\leq 4\eta d^{11} H^4 T^{\frac{1}{3}} \\
&\leq d^{10} H^3 T^{-\frac{1}{3}} \\
&\leq d^4 H T^{\frac{2}{3}}. \tag*{(Assume } T \geq d^6 H^2 \text{)}
\end{aligned}$$

Putting these two changes back to the proof into Lemma C.5, given $\eta = (4Hd|\mathcal{A}|)^{-1} T^{-\frac{2}{3}}$ we have

$$\begin{aligned}
\textbf{EXP} &= \frac{\log |\Pi_{\text{lin}}^{\text{cov}}(\frac{1}{T})|}{\eta} + \frac{6d\eta H^2 |\mathcal{A}| T}{\beta} + 4\eta d^2 H^4 \epsilon^2 T + d^4 H T^{\frac{2}{3}} + 4d^6 H^2 T^{\frac{2}{3}} \\
&= \tilde{\mathcal{O}} \left(d^6 |\mathcal{A}| H^2 T^{\frac{2}{3}} \log(|\Phi|) \right) \tag{44}
\end{aligned}$$

Putting Eq. (41), Lemma D.4, Eq. (44) into Eq. (40) together with $\epsilon = 18^{-1} d^{\frac{5}{2}} T^{-\frac{1}{3}}$, we have

$$\text{Reg}_T \leq \tilde{\mathcal{O}} \left(d^8 H^6 |\mathcal{A}| T^{\frac{2}{3}} \log(|\Phi|) \right)$$

E Proof of Theorem 4.1 (Model-Free, Banfit Feedback)

We start by introducing some notation. We let $\mathcal{I}^{(k)}$ denote the rounds in the k -th epoch:

$$\mathcal{I}^{(k)} := \{T_0 + (k-1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}\}, \quad (45)$$

where T_0 is as in Line 1 of Algorithm 3. Throughout the analysis, we condition on the event

$$\mathcal{E} := \mathcal{E}^{\text{cov}}, \quad (46)$$

where $\mathcal{E}^{\text{cov}}$ is as in Lemma G.1. Further, for any $k \in [K]$, let $\rho^{(k)}$ be the distribution of the random policy:

$$\mathbb{I}\{\zeta = 0\} \cdot \widehat{\pi}^{(k)} + \mathbb{I}\{\zeta = 1\} \cdot \pi \circ_{\mathbf{h}} \pi_{\text{unif}} \circ_{\mathbf{h}+1} \widehat{\pi}^{(k)}, \quad (47)$$

with $\zeta \sim \text{Ber}(\nu)$, $\mathbf{h} \sim \text{unif}([H])$, and $\pi \sim \text{unif}(\Psi_{\mathbf{h}}^{\text{cov}})$.

We start our analysis by applying the performance difference lemma.

Applying the performance difference lemma. For any $k \in [K]$, $t \in \mathcal{I}^{(k)}$, and $\rho^{(k)}$ as just defined, we have

$$\begin{aligned} & \mathbb{E}_{\pi \sim \rho^{(k)}} \mathbb{E} [V_1^\pi(x_1; \ell^t)] - V_1^{\pi^*}(x_1; \ell^t) \\ &= (1 - \nu) \cdot \left(V_1^{\widehat{\pi}^{(k)}}(x_1; \ell^t) - V_1^{\pi^*}(x_1; \ell^t) \right) \\ & \quad + \frac{\nu}{Hd} \sum_{h \in [H]} \sum_{\pi \in \Psi_h^{\text{cov}}} \left(V_1^{\pi \circ_{\mathbf{h}} \pi_{\text{unif}} \circ_{\mathbf{h}+1} \widehat{\pi}^{(k)}}(x_1; \ell^t) - V_1^{\pi^*}(x_1; \ell^t) \right), \\ &\leq (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\widehat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) \right] + H\nu, \\ &= (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\widehat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right] + H\nu \\ & \quad + (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \widehat{\pi}_h^{(k)}(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \\ & \quad + (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi_h^*(a | \mathbf{x}_h) \cdot \left(\widehat{Q}_h^{(k)}(\mathbf{x}_h, a) - Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) \right) \right]. \end{aligned}$$

Thus, by the triangle inequality, we have

$$\begin{aligned} & \sum_{t \in \mathcal{I}^{(k)}} \left(\mathbb{E}_{\pi \sim \rho^{(k)}} \mathbb{E} [V_1^\pi(x_1; \ell^t)] - V_1^{\pi^*}(x_1; \ell^t) \right) \\ &\leq (1 - \nu) \cdot \sum_{t \in \mathcal{I}^{(k)}} \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\widehat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right] + HN_{\text{reg}}\nu, \\ & \quad + 2(1 - \nu) \cdot \sum_{h=1}^H \max_{\pi' \in \Pi} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right|. \quad (48) \end{aligned}$$

We start by bound the first term in the right-hand side of (48) (the regret term).

E.1 Bounding the Regret Term

Fix $h \in [H]$ and define $\tilde{\pi}_h^*(\cdot | x) := (1 - \gamma/T) \cdot \pi_h^*(\cdot | x) + \gamma \pi_{\text{unif}}(\cdot | x)/T$ for all x , where γ is as in Algorithm 3. We have that $|\widehat{Q}_h^{(k)}(x, a)| = |\hat{\phi}_h^{(k)}(x, a)^\top \hat{\theta}_h^{(k)}| \leq H\sqrt{d}$ (since $\hat{\theta}_h^{(k)} \in \mathbb{B}_d(H\sqrt{d})$ and $\|\hat{\phi}_h^{(k)}(x, a)\| \leq 1$). By applying Lemma I.5, we have that for any $\eta \leq \frac{1}{H\sqrt{d}}$ and $x \in \mathcal{X}$:

$$\begin{aligned} & \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \left(\widehat{\pi}_h^{(k)}(a | x) - \pi_h^*(a | x) \right) \cdot \widehat{Q}_h^{(k)}(x, a) \\ &\leq \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \left(\widehat{\pi}_h^{(k)}(a | x) - \tilde{\pi}_h^*(a | x) \right) \cdot \widehat{Q}_h^{(k)}(x, a) + H\gamma, \end{aligned}$$

$$\begin{aligned}
&\leq \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + \eta \sum_{k=1}^K \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \widehat{\pi}_h^{(k)}(a | x) \cdot \widehat{Q}_h^{(k)}(x, a)^2 + H\gamma, \\
&\leq \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + H^2 d \eta T + H\gamma,
\end{aligned} \tag{49}$$

E.2 Bounding the Bias Term

Fix epoch $k \in [K]$, round $t \in \mathcal{I}^{(k)}$, layer $h \in [2..H]$, and $\pi' \in \Pi$. Further, let

$$\mathcal{X}_{h,\varepsilon} := \left\{ x \in \mathcal{X} : \max_{\pi \in \Pi} d_h^\pi(x) \geq \varepsilon \cdot \|\mu_h^*(x)\| \right\} \tag{50}$$

be the set of ε -reachable states. With this notation, we now bound the bias term in (48); we have

$$\begin{aligned}
&\left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right| \\
&\leq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \notin \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right| \\
&\quad + \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right|,
\end{aligned}$$

and so by Lemma I.1,

$$\begin{aligned}
&\leq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right| + 2N_{\text{reg}} H d^2 \varepsilon, \\
&= N_{\text{reg}} \cdot \left| \mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(\frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right| + 2N_{\text{reg}} H d^2 \varepsilon.
\end{aligned}$$

Thus, by letting

$$\overline{Q}_h^{\widehat{\pi}^{(k)}}(\cdot, \cdot) := \frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} Q_h^{\widehat{\pi}^{(k)}}(\cdot, \cdot; \ell^t), \tag{51}$$

and using Jensen's inequality (twice), we get

$$\begin{aligned}
&\left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a; \ell^t) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \right] \right| \\
&\leq N_{\text{reg}} \cdot \mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left| \overline{Q}_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right| \right] + 2N_{\text{reg}} H d^2 \varepsilon, \\
&\leq N_{\text{reg}} \cdot \sqrt{\mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(\overline{Q}_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right)^2 \right]} + 2N_{\text{reg}} H d^2 \varepsilon, \\
&\leq N_{\text{reg}} \cdot \sqrt{\mathbb{E}^{\pi^*} \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot \max_{a \in \mathcal{A}} \left(\overline{Q}_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right)^2 \right]} + 2N_{\text{reg}} H d^2 \varepsilon,
\end{aligned}$$

and now by the fact that Ψ_h^{cov} is an $(\frac{1}{8Ad}, \varepsilon)$ policy cover for layer h (see Lemma G.1 and Definition 4.1):

$$\begin{aligned}
&\leq N_{\text{reg}} \cdot \sqrt{8Ad \max_{\pi \in \Psi_h^{\text{cov}}} \mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot \max_{a \in \mathcal{A}} \left(\overline{Q}_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, a) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right)^2 \right]} + 2N_{\text{reg}} H d^2 \varepsilon, \\
&\leq N_{\text{reg}} \cdot \sqrt{8A^2 d \sum_{\pi \in \Psi_h^{\text{cov}}} \mathbb{E}^{\pi \circ h} \pi_{\text{unif}} \left[\left(\overline{Q}_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h, \mathbf{a}_h) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, \mathbf{a}_h) \right)^2 \right]} + 2N_{\text{reg}} H d^2 \varepsilon.
\end{aligned} \tag{52}$$

Next, we bound the regression error term in the right-hand side of (52).

E.3 Bounding the Regression Error

Lemma E.1. Let $\delta \in (0, 1)$ and $\nu \in (0, 1/4)$ be given. There is an event $\mathcal{E}^{\text{reg}}$ of probability at least $1 - 2\delta$ under which

$$\sum_{\pi \in \Psi_h} \mathbb{E}^{\pi \circ_h \pi_{\text{unif}}} \left[\left(\hat{\phi}_h^{(k)}(\mathbf{x}_h, \mathbf{a}_h)^\top \hat{\theta}_h^{(k)} - \bar{Q}_h^{\pi^{(k)}}(\mathbf{x}_h, \mathbf{a}_h) \right)^2 \right] \leq \varepsilon_{\text{reg}}^2 := \frac{40H^3 d \log(2|\mathcal{F}|/\delta)}{\nu N_{\text{reg}}}. \quad (53)$$

Proof. Fix $\delta \in (0, 1)$, $h \in [H]$ and $k \in [K]$, and let $(\mathbf{x}_h^t, \mathbf{a}_h^t, \zeta^t, \mathbf{h}^t)$ be as in Algorithm 3. With this, define

$$\mathbf{I}_h^t := \mathbb{I}\{\zeta^t = 0 \text{ or } \mathbf{h}^t \leq h\}, \quad (54)$$

and note that $(\mathbf{x}_h^t, \mathbf{a}_h^t, \mathbf{I}_h^t)_{t \in \mathcal{I}^{(k)}}$ are identically and independently distributed. Further, for $t \in \mathcal{I}^{(k)}$ and $\pi \in \Pi$, let $\theta_h^{t, \pi} \in \mathbb{B}_d(H\sqrt{d})$ be such that

$$\forall (x, a) \in \mathcal{X} \times \mathcal{A}, \quad Q_h^\pi(x, a; \ell^t) = \phi_h^*(x, a)^\top \theta_h^{\pi, t}.$$

Such a $\theta_h^{\pi, t}$ is guaranteed to exist by the low-rank MDP structure (Assumption 2.1) and Assumption 2.2. With this, note that for $\bar{Q}_h^{\pi^{(k)}}$ as in (51), we have

$$\bar{Q}_h^{\pi^{(k)}}(x, a) = \phi_h^*(x, a)^\top \theta_h^{(k)}, \quad \text{where} \quad \theta_h^{(k)} := \frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} \theta_h^{\pi^{(k)}, t}. \quad (55)$$

For the rest of this proof, we let $\mathcal{F}$ be the function class

$$\mathcal{F} := \{f : (x, a) \mapsto \phi_h(x, a)^\top \theta \mid \theta \in \mathcal{C} \cup \{\theta_h^{(k)}\}, \phi \in \Phi\},$$

where $\mathcal{C}$ is a minimal $(N_{\text{reg}})^{-1}$ -cover of $\mathbb{B}(H\sqrt{d})$ in $\|\cdot\|$ distance. Further, for $\tau \in \mathcal{I}^{(k)}$ we let

$$\begin{aligned} \mathbf{z}_h^\tau &:= \sum_{l=h}^H \ell_l^\tau(\mathbf{x}_l^\tau, \mathbf{a}_l^\tau); \\ \varepsilon_h^\tau &:= \sum_{l=h}^H \ell_l^\tau(\mathbf{x}_l^\tau, \mathbf{a}_l^\tau) - \frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} Q_h^{\pi^{(k)}}(\mathbf{x}_h^\tau, \mathbf{a}_h^\tau; \ell^t); \end{aligned} \quad (56)$$

and

$$\widehat{L}(f) := \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot (f(\mathbf{x}_h^t, \mathbf{a}_h^t) - \mathbf{z}_h^t)^2, \quad (57)$$

for $f \in \mathcal{F}$. Finally, let $f_*(x, a) := \phi_h^*(x, a)^\top \theta_h^{(k)}$, where $\theta_h^{(k)}$ is as in (55) and $\hat{f}(x, a) := \widehat{Q}_h^{(k)}(x, a)$. With this, note that f_* and $\hat{f}$ satisfy $f_*(\mathbf{x}_h^t, \mathbf{a}_h^t) = \mathbf{z}_h^t - \varepsilon_h^t$ and $\hat{f} \in \arg\min_{f \in \mathcal{F}} \widehat{L}(f)$.

Now, since $\hat{f} \in \arg\min_{f \in \mathcal{F}} \widehat{L}(f)$, we have

$$0 \geq \widehat{L}(\hat{f}) - \widehat{L}(f_*) = \nabla \widehat{L}(f_*)[\hat{f} - f_*] + \|\hat{f} - f_*\|^2, \quad (58)$$

where ∇ denotes directional derivative and

$$\|\hat{f} - f_*\|^2 := \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2.$$

Rearranging (58) and using that $f_*(\mathbf{x}_h^t, \mathbf{a}_h^t) = \mathbf{z}_h^t - \varepsilon_h^t$, we get that

$$\begin{aligned} \|\hat{f} - f_*\|^2 &\leq -\nabla \widehat{L}(f_*)[\hat{f} - f_*], \\ &= 2 \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot (\mathbf{z}_h^t - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))(\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t)), \\ &= 2 \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot \varepsilon_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t)). \end{aligned} \quad (59)$$

We now bound the right-hand side of (59). For any $h \in [H]$, $k \in [K]$ and any $\hat{f}, f_* \in \mathcal{F}$, we apply Lemma I.2 with

- $\mathfrak{F}^i = \sigma(\{(\mathbf{x}_h^j, \mathbf{a}_h^j, \boldsymbol{\varepsilon}_h^j, \zeta^j) : j \leq t_i\})$ where $t_i := (k-1) \cdot N_{\text{reg}} + i$;
- The random variable $\mathbf{w}^i$ set as the difference

$$\begin{aligned} \mathbf{w}^i &= \mathbb{I}\{\zeta^{t_i} = 0 \text{ or } \mathbf{h}^{t_i} \leq h\} \cdot \boldsymbol{\varepsilon}_h^{t_i} \cdot (\hat{f}(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i}) - f_\star(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i})) \\ &\quad - \mathbb{E}[\mathbb{I}\{\zeta^{t_i} = 0 \text{ or } \mathbf{h}^{t_i} \leq h\} \cdot \boldsymbol{\varepsilon}_h^{t_i} \cdot (\hat{f}(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i}) - f_\star(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i})) \mid \mathfrak{F}^{i-1}] \end{aligned}$$

where $t_i := (k-1) \cdot N_{\text{reg}} + i$;

- $n = N_{\text{reg}} = |\mathcal{I}^{(k)}|$.
- $R = 4H^2$; and
- $\lambda = 1/(16H^2)$;

to get that there is an event $\mathcal{E}$ of probability at least $1 - \delta$ under which

$$\begin{aligned} &\sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t)) \\ &\leq \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t))] \\ &\quad + \frac{1}{8H^2} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\boldsymbol{\varepsilon}_h^t)^2 \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] + 16H^2 \log(|\mathcal{F}|/\delta), \\ &\leq \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t))] \\ &\quad + \frac{1}{4} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] + 16H^2 \log(|\mathcal{F}|/\delta), \end{aligned} \quad (60)$$

where $\mathbb{E}_t[\cdot]$ is defined as $\mathbb{E}[\cdot \mid \mathfrak{F}^{t-1}]$ and the last step uses that $|\boldsymbol{\varepsilon}_h^t| \leq 2H$, for all $t \in \mathcal{I}^{(k)}$. For the rest of the proof, we condition on $\mathcal{E}$ and to simplify notation let

$$\Delta_h^t := \hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_\star(\mathbf{x}_h^t, \mathbf{a}_h^t). \quad (61)$$

Using the expression of $\boldsymbol{\varepsilon}_h^t$ in (56), we have

$$\begin{aligned} &\sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot \Delta_h^t] \\ &= \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t \left[\mathbf{I}_h^t \cdot \left(\sum_{l=h}^H \ell_l^t(\mathbf{x}_l^t, \mathbf{a}_l^t) - \frac{1}{N_{\text{reg}}} \sum_{\tau \in \mathcal{I}^{(k)}} Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^\tau) \right) \cdot \Delta_h^t \right]. \end{aligned} \quad (62)$$

Now, since $\mathbf{I}_h^t = \mathbb{I}\{\zeta^t = 0 \text{ or } \mathbf{h}^t \leq h\}$, we have

$$\mathbb{E}_t \left[\sum_{l=h}^H \ell_l^t(\mathbf{x}_l^t, \mathbf{a}_l^t) \mid \mathbf{x}_h^t, \mathbf{a}_h^t, \mathbf{I}_h^t = 1 \right] = Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^t).$$

Plugging this into (62) and using the law of total expectation, we have

$$\begin{aligned} &\sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot \Delta_h^t] \\ &= \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t \left[\mathbf{I}_h^t \cdot \left(Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^t) - \frac{1}{N_{\text{reg}}} \sum_{\tau \in \mathcal{I}^{(k)}} Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^\tau) \right) \cdot \Delta_h^t \right], \\ &= \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t \left[\mathbf{I}_h^t \cdot Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^t) \cdot \Delta_h^t \right] - \frac{1}{N_{\text{reg}}} \sum_{\tau \in \mathcal{I}^{(k)}} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t \left[\mathbf{I}_h^t \cdot Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^\tau) \cdot \Delta_h^t \right]. \end{aligned} \quad (63)$$

On the other hand, since $(\mathbf{x}_h^t, \mathbf{a}_h^t, \mathbf{I}_h^t)_{t \in \mathcal{I}^{(k)}}$ are i.i.d. and ℓ^t is chosen by an oblivious adversary, we have

$$\forall t \in \mathcal{I}^{(k)}, \quad \mathbb{E}_t \left[\mathbf{I}_h^t \cdot Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^t, \mathbf{a}_h^t; \ell^t) \cdot \Delta_h^t \right] = \frac{1}{N_{\text{reg}}} \sum_{\tau \in \mathcal{I}^{(k)}} \mathbb{E}_\tau \left[\mathbf{I}_h^\tau \cdot Q_h^{\widehat{\pi}^{(k)}}(\mathbf{x}_h^\tau, \mathbf{a}_h^\tau; \ell^\tau) \cdot \Delta_h^\tau \right].$$

Plugging this into (63) shows that

$$\sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot \boldsymbol{\varepsilon}_h^t \cdot \boldsymbol{\Delta}_h^t] = 0. \quad (64)$$

Combining this with (60) and (59), we get that

$$\begin{aligned} \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2 &\leq \frac{1}{4} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] \\ &\quad + 16H^2 \log(|\mathcal{F}|/\delta). \end{aligned} \quad (65)$$

Now, since $(\mathbf{x}_h^t, \mathbf{a}_h^t, \mathbf{I}_h^t)_{t \in \mathcal{I}^{(k)}}$ are i.i.d., we have by Lemma I.3 that there is an event $\mathcal{E}'$ of probability at least $1 - \delta$ under which we have

$$\begin{aligned} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] &\leq 2 \sum_{t \in \mathcal{I}^{(k)}} \mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2 \\ &\quad + 8H^2 \log(2|\mathcal{F}|/\delta). \end{aligned} \quad (66)$$

Combining this with (65) and rearranging, we get that under $\mathcal{E}^{\text{reg}} := \mathcal{E} \cap \mathcal{E}'$:

$$\frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] \leq \frac{40H^2 \log(2|\mathcal{F}|/\delta)}{N_{\text{reg}}}. \quad (67)$$

On the other hand, we have that for all $t \in \mathcal{I}^{(k)}$:

$$\begin{aligned} \mathbb{E}_t[\mathbf{I}_h^t \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2] &\geq \mathbb{E}_t[\mathbb{I}\{\boldsymbol{\zeta}^t = 1, \mathbf{h}^t = h\} \cdot (\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2], \\ &= \frac{\nu}{Hd} \sum_{\pi \in \Psi_h} \mathbb{E}^{\pi \circ \pi_{\text{unif}}}[(\hat{f}(\mathbf{x}_h^t, \mathbf{a}_h^t) - f_*(\mathbf{x}_h^t, \mathbf{a}_h^t))^2]. \end{aligned} \quad (68)$$

Plugging this into (67) and using the expressions of $\hat{f}$ and f_* , we get

$$\sum_{\pi \in \Psi_h} \mathbb{E}^{\pi \circ \pi_{\text{unif}}} \left[\left(\hat{\phi}_h^{(k)}(\mathbf{x}_h, \mathbf{a}_h)^\top \hat{\theta}_h^{(k)} - \bar{Q}_h^{\pi^{(k)}}(\mathbf{x}_h, \mathbf{a}_h) \right)^2 \right] \leq \frac{40H^3 d \log(2|\mathcal{F}|/\delta)}{\nu N_{\text{reg}}}.$$

□

E.4 Putting It All Together

By combining (48), (52), (53), and (49), we have that under the event $\mathcal{E} := \mathcal{E}^{\text{cov}} \cap \mathcal{E}^{\text{reg}}$ (where $\mathcal{E}^{\text{cov}}$ and $\mathcal{E}^{\text{reg}}$ are as in Lemma G.1 and Lemma E.1, respectively):

$$\begin{aligned} \text{Reg}_T &= HT_0 + \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \left(\mathbb{E}_{\pi \sim \rho^{(k)}} \mathbb{E}[V_1^\pi(x_1; \ell^t)] - V_1^{\pi^*}(x_1; \ell^t) \right), \\ &\leq HT_0 + HT\nu + T \cdot \sqrt{8A^2 d \cdot \varepsilon_{\text{reg}}^2} + \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + H^2 d \eta T + H\gamma. \end{aligned} \quad (69)$$

Thus, plugging in the expression of T_0 from Algorithm 3, and ignoring polynomial factors $d, A, H, \log(|\Phi| \varepsilon^{-1} \delta^{-1})$, we get that

$$\begin{aligned} \text{Reg}_T &< \frac{1}{\varepsilon^2} + T\varepsilon + T\nu + N_{\text{reg}} \cdot \sqrt{8A^2 d \cdot \varepsilon_{\text{reg}}^2} + 2N_{\text{reg}} H d \varepsilon + \frac{N_{\text{reg}}}{\eta} + \eta T, \\ &< T^{2/3} + \nu T + T \cdot \sqrt{8A^2 d \cdot \varepsilon_{\text{reg}}^2} + \sqrt{TN_{\text{reg}}}, \quad (\text{by setting } \varepsilon = T^{-1/3} \text{ and } \eta = (N_{\text{reg}}/T)^{1/2}) \\ &= T^{2/3} + \nu T + T \cdot \sqrt{\frac{1}{\nu N_{\text{reg}}}} + \sqrt{TN_{\text{reg}}}, \quad (\text{used the expression of } \varepsilon_{\text{reg}}^2 \text{ in (53)}) \\ &< T^{2/3} + \nu T + \sqrt{T} \cdot \left(\frac{T}{\nu} \right)^{1/4}, \quad (\text{by setting } N_{\text{reg}} = (T/\nu)^{1/2}) \\ &< T^{4/5}, \end{aligned} \quad (70)$$

where the last step follows by setting $\nu = T^{-1/5}$.

F Proof of Theorem 4.2 (Model-Free, Bandit Feedback, Adaptive Adversary)

We let $\mathcal{I}^{(k)}$ denote the rounds in the k th epoch:

$$\mathcal{I}^{(k)} := \{T_0 + (k-1) \cdot N_{\text{reg}} + 1, \dots, T_0 + k \cdot N_{\text{reg}}\}, \quad (71)$$

where T_0 is as in Line 8 of Algorithm 4. Throughout the analysis, we condition on the event

$$\mathcal{E} := \mathcal{E}^{\text{cov}} \cap \mathcal{E}^{\text{rep+span}} \cap \mathcal{E}^{\text{freed}}, \quad (72)$$

where $\mathcal{E}^{\text{cov}}$, $\mathcal{E}^{\text{rep+span}}$, and $\mathcal{E}^{\text{freed}}$ are as in Lemma G.1, Corollary F.1, and Lemma F.3, respectively.

We start our analysis by applying the performance difference lemma.

Applying the performance difference lemma. For any $k \in [K]$, $t \in \mathcal{I}^{(k)}$, and let $\rho^{(k)}$ be the distribution of the random policy:

$$\mathbb{I}\{\zeta^t = 0\} \cdot \hat{\pi}^{(k)} + \mathbb{I}\{\zeta^t = 1\} \cdot \pi^t \circ_{\mathbf{h}^{t+1}} \hat{\pi}^{(k)}, \quad (73)$$

with $\zeta^t \sim \text{Ber}(\nu)$, $\mathbf{h}^t \sim \text{unif}([H])$, and $\pi^t \sim \text{unif}(\Psi_{\mathbf{h}^t}^{\text{span}})$.

$$\begin{aligned} & \mathbb{E}_{\pi \sim \rho^{(k)}} \mathbb{E} [V_1^\pi(x_1; \mathcal{H}^{t-1})] - V_1^{\pi^*}(x_1; \mathcal{H}^{t-1}) \\ &= (1 - \nu) \cdot \left(V_1^{\hat{\pi}^{(k)}}(x_1; \mathcal{H}^{t-1}) - V_1^{\pi^*}(x_1; \mathcal{H}^{t-1}) \right) \\ & \quad + \frac{\nu}{Hd} \sum_{h \in [H]} \sum_{\pi \in \Psi_h^{\text{cov}}} \left(V_1^{\pi \circ_{\mathbf{h}^{t+1}} \hat{\pi}^{(k)}}(x_1; \mathcal{H}^{t-1}) - V_1^{\pi^*}(x_1; \mathcal{H}^{t-1}) \right), \\ &= (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\hat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot Q_{\hat{\pi}_h^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right] + HN_{\text{reg}}\nu, \\ &= (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\hat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \mid \mathcal{H}^{t-1} \right] + HN_{\text{reg}}\nu \\ & \quad + (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \hat{\pi}_h^{(k)}(a | \mathbf{x}_h) \cdot \left(Q_{\hat{\pi}_h^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \mid \mathcal{H}^{t-1} \right] \\ & \quad + (1 - \nu) \cdot \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi_h^*(a | \mathbf{x}_h) \cdot \left(\widehat{Q}_h^{(k)}(\mathbf{x}_h, a) - Q_{\hat{\pi}_h^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \right) \mid \mathcal{H}^{t-1} \right]. \end{aligned}$$

Thus, by the triangle inequality, we have

$$\begin{aligned} & \sum_{t \in \mathcal{I}^{(k)}} \left(\mathbb{E}_{\pi \sim \rho^{(k)}} \mathbb{E} [V_1^\pi(x_1; \mathcal{H}^{t-1})] - V_1^{\pi^*}(x_1; \mathcal{H}^{t-1}) \right) \\ & \leq (1 - \nu) \cdot \sum_{t \in \mathcal{I}^{(k)}} \sum_{h=1}^H \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \left(\hat{\pi}_h^{(k)}(a | \mathbf{x}_h) - \pi_h^*(a | \mathbf{x}_h) \right) \cdot \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \mid \mathcal{H}^{t-1} \right] + HN_{\text{reg}}\nu \\ & \quad + 2(1 - \nu) \cdot \sum_{h=1}^H \max_{\pi' \in \Pi} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | \mathbf{x}_h) \cdot \left(Q_{\hat{\pi}_h^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) - \widehat{Q}_h^{(k)}(\mathbf{x}_h, a) \right) \mid \mathcal{H}^{t-1} \right] \right|. \end{aligned} \quad (74)$$

We start by bounding the first term on the right-hand side of (74) (the regret term).

F.1 Bounding the Regret Term

Fix $h \in [H]$ and define $\tilde{\pi}_h^*(\cdot | x) := (1 - \gamma/T) \cdot \pi_h^*(\cdot | x) + \gamma \pi_{\text{unif}}(\cdot | x)/T$ for all x , where γ is as in Algorithm 4. We have that $|\widehat{Q}_h^{(k)}(x, a)| = |\bar{\phi}_h^{\text{rep}}(x, a)^\top \hat{\theta}_h^{(k)}| \leq 8Hd^2$ (since $\hat{\theta}_h^{(k)} \in \mathbb{B}_{2d}(4Hd^2)$ and $\|\bar{\phi}_h^{\text{rep}}(x, a)\| \leq \|\phi_h^{\text{loss}}(x, a)\| + \|\phi_h^{\text{rep}}(x, a)\| \leq 2$). By applying I.5, we have that for any $\eta \leq \frac{1}{8Hd^2}$ and $x \in \mathcal{X}$:

$$\begin{aligned} & \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \left(\hat{\pi}_h^{(k)}(a | x) - \pi_h^*(a | x) \right) \cdot \widehat{Q}_h^{(k)}(x, a) \\ & \leq \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \left(\hat{\pi}_h^{(k)}(a | x) - \tilde{\pi}_h^*(a | x) \right) \cdot \widehat{Q}_h^{(k)}(x, a) + H\gamma, \end{aligned}$$

$$\begin{aligned}
&\leq \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + \eta \sum_{k=1}^K \sum_{t \in \mathcal{I}^{(k)}} \sum_{a \in \mathcal{A}} \widehat{\pi}_h^{(k)}(a | x) \cdot \widehat{Q}_h^{(k)}(x, a)^2 + H\gamma, \\
&\leq \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + 64H^2 d^4 \eta T + H\gamma,
\end{aligned} \tag{75}$$

F.2 Bounding the Bias Term

To bound the bias term (second term in (74)), we make use of the following result.

Lemma F.1. *For $t \in [T]$, $h \in [H]$, and $\pi \in \Pi$, there exists $\bar{\theta}_h^{t, \pi} \in \mathbb{B}_{2d}(H\sqrt{d})$ such that for all $(x, a) \in \mathcal{X} \times \mathcal{A}$ and history $\mathcal{H}^{t-1} = (x_{1:t-1}^1, a_{1:t-1}^1)$,*

$$Q_h^\pi(x, a; \mathcal{H}^{t-1}) = \bar{\phi}_h^*(x, a)^\top \bar{\theta}_h^{t, \pi}, \quad \text{where } \bar{\phi}_h^* := [\phi^{\text{loss}}, \phi^*] \in \mathbb{R}^{2d}. \tag{76}$$

Proof. Fix $t \in [T]$, $h \in [H]$, and $\pi \in \Pi$. By the low-rank MDP structure and the normalizing assumption on μ_h^* in Assumption 2.1, there exists $w_{h+1}^{t, \pi} \in \mathbb{B}_d((H-s)\sqrt{d})$ such that for all $(x, a) \in \mathcal{X} \times \mathcal{A}$ and history $\mathcal{H}^{t-1} = (x_{1:t-1}^1, a_{1:t-1}^1)$,

$$\mathbb{E}^\pi \left[\sum_{s=h+1}^H \ell_s(x_s, a_s; \mathcal{H}^{t-1}) \mid x_h = x, a_h = a \right] = \phi_h^*(x, a)^\top w_{h+1}^{t, \pi}. \tag{77}$$

Now, with g_h^t as in Assumption 4.2, (15) and (77) implies that $\bar{\theta}_h^{t, \pi} := [g_h^t, w_{h+1}^{t, \pi}] \in \mathbb{R}^{2d}$ (where $[\cdot, \cdot]$ denotes the vertical stacking of vectors) satisfies the desired property. $\square$

Fix epoch $k \in [K]$, round $t \in \mathcal{I}^{(k)}$, layer $h \in [2..H]$, and $\pi' \in \Pi$. By Lemma F.1, there exists $\bar{\theta}_h^{t, \pi^{(k)}} \in \mathbb{B}_{2d}(Hd^{1/2})$ such that for all $(x, a) \in \mathcal{X} \times \mathcal{A}$,

$$Q_{h+1}^{\pi^{(k)}}(x, a; \mathcal{H}^{t-1}) = \bar{\phi}_{h+1}^*(x, a)^\top \bar{\theta}_{h+1}^{t, \pi^{(k)}}.$$

With this, let $w_{h+1}^{t, \pi^{(k)}} \in \mathbb{B}_d(3Hd^2)$ be as in Corollary F.1 with $f(x) := \frac{1}{Hd^{1/2}} \max_{a \in \mathcal{A}} \bar{\phi}_{h+1}^*(x, a)^\top \bar{\theta}_{h+1}^{t, \pi^{(k)}}$; note that this function belongs to the function class $\mathcal{F}_{h+1}$ in Algorithm 4. With this, we define

$$\bar{\vartheta}_h^{t, \pi^{(k)}} := [g_h^t, w_{h+1}^{t, \pi^{(k)}}] \in \mathbb{B}_{2d}(4Hd^2), \tag{78}$$

where $g_h^t \in \mathbb{B}_d(1)$ is such that $\ell_h(x, a; \mathcal{H}^{t-1}) = \phi_h^{\text{loss}}(x, a)^\top g_h^t$.

With this notation, we now bound the bias term in Eq. (74): we have

$$\begin{aligned}
&\left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | x_h) \cdot \left(Q_h^{\pi^{(k)}}(x_h, a; \mathcal{H}^{t-1}) - \widehat{Q}_h^{(k)}(x_h, a) \right) \mid \mathcal{H}^{t-1} \right] \right| \\
&\leq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | x_h) \cdot \left(\bar{\phi}_h^{\text{rep}}(x_h, a)^\top \bar{\vartheta}_h^{t, \pi^{(k)}} - \widehat{Q}_h^{(k)}(x_h, a) \right) \mid \mathcal{H}^{t-1} \right] \right| \\
&\quad + \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | x_h) \cdot \left(Q_h^{\pi^{(k)}}(x_h, a; \mathcal{H}^{t-1}) - \bar{\phi}_h^{\text{rep}}(x_h, a)^\top \bar{\vartheta}_h^{t, \pi^{(k)}} \right) \mid \mathcal{H}^{t-1} \right] \right|,
\end{aligned}$$

and by Corollary F.1 (in particular (92))

$$\leq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi^*} \left[\sum_{a \in \mathcal{A}} \pi'_h(a | x_h) \cdot \left(\bar{\phi}_h^{\text{rep}}(x_h, a)^\top \bar{\vartheta}_h^{t, \pi^{(k)}} - \bar{\phi}_h^{\text{rep}}(x_h, a)^\top \hat{\theta}_h^{(k)} \right) \mid \mathcal{H}^{t-1} \right] \right| + N_{\text{reg}} \cdot \varepsilon_{\text{rep}},$$

and by Corollary F.1 again (in particular (93)) and the triangle inequality

$$\begin{aligned}
&\leq 2 \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(x_h, a_h)^\top \bar{\vartheta}_h^{t, \pi^{(k)}} - \bar{\phi}_h^{\text{rep}}(x_h, a_h)^\top \hat{\theta}_h^{(k)} \mid \mathcal{H}^{t-1} \right] \right| + N_{\text{reg}} \cdot \varepsilon_{\text{span}} + N_{\text{reg}} \cdot \varepsilon_{\text{rep}}, \\
&\leq 2 \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[Q_h^{\pi^{(k)}}(x_h, a; \mathcal{H}^{t-1}) - \bar{\phi}_h^{\text{rep}}(x_h, a_h)^\top \hat{\theta}_h^{(k)} \mid \mathcal{H}^{t-1} \right] \right| + N_{\text{reg}} \cdot \varepsilon_{\text{span}} + N_{\text{reg}} \cdot \varepsilon_{\text{rep}}
\end{aligned}$$

$$\begin{aligned}
& + 2 \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\boldsymbol{\vartheta}}_h^{t, \pi^{(k)}} - Q_h^{\pi^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right] \right|, \quad (\text{triangle inequality}) \\
& \leq 2 \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[Q_h^{\pi^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) - \bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \hat{\theta}_h^{(k)} \mid \mathcal{H}^{t-1} \right] \right| + N_{\text{reg}} \cdot \varepsilon_{\text{span}} + 3dN_{\text{reg}} \cdot \varepsilon_{\text{rep}},
\end{aligned} \tag{79}$$

where the last inequality follows by Corollary F.1 (in particular (92)).

Next, we bound the estimation error term in the right-hand side of (79).

F.3 Bound the Regression Error

By Lemma F.3 (Freedman's inequality), we have that for all $\pi \in \Psi_h^{\text{span}}$ and $\bar{\theta} \in \mathbb{B}_{2d}(4Hd^2)$:

$$\begin{aligned}
N_{\text{reg}} \cdot \varepsilon_{\text{freed}} & \geq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \boldsymbol{\zeta}^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \right. \\
& \quad \left. - \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \boldsymbol{\zeta}^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \mid \mathcal{H}^{t-1} \right] \right|. \tag{80}
\end{aligned}$$

On the other hand, we have

$$\begin{aligned}
& \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \boldsymbol{\zeta}^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \mid \mathcal{H}^{t-1} \right] \\
& = \frac{\nu}{Hd} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^{\pi \circ_{h+1} \hat{\pi}^{(k)}} \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \mid \mathcal{H}^{t-1} \right], \\
& = \frac{\nu}{Hd} \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\theta} - Q_h^{\pi^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right].
\end{aligned} \tag{81}$$

Now, by Corollary F.1 (in particular (92)) and the triangle inequality, we have

$$\begin{aligned}
N_{\text{reg}} \cdot \varepsilon_{\text{rep}} & \geq \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\theta} - Q_h^{\pi^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right] \right. \\
& \quad \left. - \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\theta} - \bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\boldsymbol{\vartheta}}_h^{t, \pi^{(k)}} \mid \mathcal{H}^{t-1} \right] \right|, \\
& = \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \bar{\theta} - Q_h^{\pi^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right] \right. \\
& \quad \left. - \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \right] \left(N_{\text{reg}} \cdot \bar{\theta} - \sum_{t \in \mathcal{I}^{(k)}} \bar{\boldsymbol{\vartheta}}_h^{t, \pi^{(k)}} \right) \right|.
\end{aligned} \tag{82}$$

Thus, by combining (80), (81), and (82), we have

$$\begin{aligned}
& \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \boldsymbol{\zeta}^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \right| \\
& \leq dN_{\text{reg}} \cdot \varepsilon_{\text{freed}} + \frac{\nu N_{\text{reg}} \varepsilon_{\text{rep}}}{H} + \frac{\nu}{Hd} \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \right] \left(N_{\text{reg}} \cdot \bar{\theta} - \sum_{t \in \mathcal{I}^{(k)}} \bar{\boldsymbol{\vartheta}}_h^{t, \pi^{(k)}} \right) \right|.
\end{aligned}$$

Using that $\hat{\theta}_h^{(k)}$ is the minimizer over $\mathbb{B}_{2d}(4Hd^2)$ of the left-hand side, and the right-hand side evaluated to $dN_{\text{reg}} \cdot \varepsilon_{\text{freed}} + \frac{\nu N_{\text{reg}} \varepsilon_{\text{rep}}}{H}$ with $\bar{\theta} = \frac{1}{N_{\text{reg}}} \sum_{t \in \mathcal{I}^{(k)}} \bar{\boldsymbol{\vartheta}}_h^{t, \pi^{(k)}} \in \mathbb{B}_{2d}(4Hd^2)$, we have that

$$\sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \boldsymbol{\zeta}^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \hat{\theta}_h^{(k)} - \sum_{s=h}^H \ell_s^t \right) \right| \leq dN_{\text{reg}} \cdot \varepsilon_{\text{freed}} + \frac{\nu N_{\text{reg}} \varepsilon_{\text{rep}}}{H}.$$

Now, combining this with (80) and (81) with $\bar{\theta} = \hat{\theta}_k^{(k)}$, we get that

$$\begin{aligned}
& \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E}^\pi \left[\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top \hat{\theta}_h^{(k)} - Q_h^{\bar{\pi}^{(k)}}(\mathbf{x}_h, a; \mathcal{H}^{t-1}) \mid \mathcal{H}^{t-1} \right] \right| \\
& \leq \frac{Hd^2 N_{\text{reg}} \varepsilon_{\text{freed}}}{\nu} + \frac{Hd}{\nu} \sum_{\pi \in \Psi_h^{\text{span}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \pi^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \hat{\theta}_h^{(k)} - \sum_{s=h}^H \ell_s^t \right) \right|, \\
& \leq \frac{2Hd^2 N_{\text{reg}} \varepsilon_{\text{freed}}}{\nu} + dN_{\text{reg}} \varepsilon_{\text{rep}}. \tag{83}
\end{aligned}$$

F.4 Putting It All Together

By combining (74), (79), (83), and (75), we have that under the event $\mathcal{E}$ in (72):

$$\begin{aligned}
\text{Reg}_T &= HT_0 + \sum_{k \in [K]} \sum_{t \in \mathcal{I}^{(k)}} \left(\mathbb{E}_{\pi \sim \rho^{(k)}} [V_1^\pi(x_1; \mathcal{H}^{t-1})] - V_1^{\pi^*}(x_1; \mathcal{H}^{t-1}) \right), \\
&\leq HT_0 + HT\nu + T \cdot \varepsilon_{\text{span}} + 3dT \cdot \varepsilon_{\text{rep}} + \frac{2Hd^2 T \varepsilon_{\text{freed}}}{\nu} \\
&\quad + dT \varepsilon_{\text{rep}} + \frac{N_{\text{reg}} \log(T/\gamma)}{\eta} + 64H^2 d^4 \eta T + H\gamma. \tag{84}
\end{aligned}$$

Thus, plugging in the expression of T_0 , $\varepsilon_{\text{freed}}$, and $(\varepsilon_{\text{span}}, \varepsilon_{\text{rep}})$ from Algorithm 4, Lemma F.3, and Corollary F.1, respectively, and ignoring polynomial factors $d, A, H, \log(|\Phi| \varepsilon^{-1} \delta^{-1})$, we get that

$$\begin{aligned}
\text{Reg}_T &< \frac{1}{\varepsilon^2} + T\varepsilon + T\nu + \frac{T}{\nu} \sqrt{\frac{\nu}{N_{\text{reg}}}} + \frac{N_{\text{reg}}}{\eta} + \eta T, \\
&< T^{2/3} + TN_{\text{reg}}^{-1/3} + \sqrt{TN_{\text{reg}}}, \quad (\text{by setting } \varepsilon = T^{-1/3}, \eta = (N_{\text{reg}}/T)^{1/2}, \nu = N_{\text{reg}}^{-1/3}) \tag{85}
\end{aligned}$$

$$< T^{4/5} \quad (N_{\text{reg}} = T^{3/5}), \tag{86}$$

where the last step follows by setting $N_{\text{reg}} = T^{2/3}$.

F.5 Spanner Guarantee

Algorithm 7 Spanner: Computing an Approximate Spanner.

Require: Layer h , feature classes Φ , policy covers $\Psi_{1:H}$, feature map $\bar{\phi} : \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^{2d}$, # of episodes n .

- 1: Define $\mathcal{G} = \{g : (x, a) \mapsto \phi(x, a)^\top w \mid \phi \in \Phi, w \in \mathbb{B}_d(2\sqrt{d})\}$.
- 2: For $\theta \in \mathbb{R}^{2d}$ and $(x, a) \in \mathcal{X} \times \mathcal{A}$, define

$$r_t(x, a; \theta) := \begin{cases} \bar{\phi}_h(x, a)^\top \theta, & \text{for } t = h, \\ 0, & \text{otherwise.} \end{cases} \tag{87}$$

- 3: Set $\mathcal{G}_h = \{(x, a) \mapsto \bar{\phi}_h(x, a)^\top \theta : \theta \in \mathbb{B}_{2d}(1)\}$, and for $t \in [h-1]$, set $\mathcal{G}_t = \mathcal{G}$.
 - 4: For each $t \in [h]$, set $P_t = \text{unif}(\Psi_t)$.
 - 5: For $\theta \in \mathbb{R}^{2d}$, define $\text{LinOpt}(\theta) = \text{PSDP}(h, r_{1:h}(\cdot, \cdot; \theta), \mathcal{G}_{1:h}, P_{1:h}, n) \in \Pi$. // PSDP as in Mhammedi et al. (2023).
 - 6: For $\theta \in \mathbb{R}^{2d}$ and $\pi \in \Pi$, define $\text{LinEst}(\pi) = \text{EstVec}(h, \bar{\phi}_h, \pi, n)$. // EstVec as in Mhammedi et al. (2023).
 - 7: Set $\pi_{1:2d} = \text{RobustSpanner}(\text{LinOpt}(\cdot), \text{LinEst}(\cdot), 2, \sqrt{\frac{Ad^2 \log(ndH|\Phi|/\delta)}{\alpha n}})$. // RobustSpanner as in Mhammedi et al. (2023).
 - 8: **Return:** Policy cover $\{\pi_1, \dots, \pi_{2d}\}$.
-

Lemma F.2 (Spanner Guarantee). *Let $\varepsilon, \alpha, \delta \in (0, 1)$, $h \in [H]$, $n \geq 1$, and $\bar{\phi}_h : \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^{2d}$ be given. Suppose that Assumption 2.1 and Assumption 4.1 hold, and let $\Psi_{1:h}$ be such that for all $s \in [h]$, Ψ_s is an (α, ε) -policy cover for layer s with $|\Psi_s| = 2d$. Then, the output $\Psi_h^{\text{span}} =$*

$\text{Spanner}(h, \Phi, \Psi_{1:h}, \bar{\phi}_h, n)$ (Algorithm 7) is such that $|\Psi_h^{\text{span}}| = 2d$ and, with probability at least $1 - \delta$, for all $\pi' \in \Pi$, there exist $\{\beta_\pi \in [-2, 2] : \pi \in \Psi_h^{\text{span}}\}$ such that

$$\left\| \mathbb{E}^{\pi'}[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] - \sum_{\pi \in \Psi_h^{\text{span}}} \beta_\pi \cdot \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] \right\| \leq \varepsilon_{\text{span}}(n, \alpha, \delta), \quad (88)$$

where

$$\varepsilon_{\text{span}}(n, \alpha, \delta) := cH^2 d \sqrt{\frac{dA \cdot (d \log(2n\sqrt{d}H) + \log(n|\Phi|/\delta))}{\alpha n}} + H^2 d^{5/2} \varepsilon. \quad (89)$$

where $c > 0$ is a large enough absolute constant. Furthermore, the number of episodes $T_{\text{span}}(n)$ used by the call to Spanner is at most $\tilde{O}(H^2 d^2 n)$.

Proof. To derive the desired bound, we will use the generic guarantee of RobustSpanner from (Mhammedi et al., 2023, Proposition E.1). To invoke this result, we first need to derive guarantees for the optimization and estimation subroutines LinOpt and LinEst within the Spanner algorithm (Algorithm 7). In particular, we need to show that there is some $\varepsilon' \in (0, 1)$ such that (with high probability) for any $\bar{\theta} \in \mathbb{R}^{2d} \setminus \{0\}$ and $\pi \in \Pi$, the outputs $\hat{\pi}_{\bar{\theta}} := \text{LinOpt}(\bar{\theta}/\|\bar{\theta}\|)$ and $\hat{\phi}^\pi := \text{LinEst}(\pi)$ satisfy

$$\sup_{\pi \in \Pi} \bar{\theta}^\top \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] \leq \bar{\theta}^\top \mathbb{E}^{\hat{\pi}_{\bar{\theta}}}[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] + \varepsilon' \cdot \|\bar{\theta}\| \quad \text{and} \quad \|\hat{\phi}^\pi - \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)]\| \leq \varepsilon'. \quad (90)$$

With this, we can apply (Mhammedi et al., 2023, Proposition E.1) to get that the output

$$\pi_{1:2d} = \text{RobustSpanner}(\text{LinOpt}(\cdot), \text{LinEst}(\cdot), 2, \varepsilon)$$

for $\varepsilon \leq 2\varepsilon'$ is such that for all $\pi \in \Pi$, there exist $\beta_1, \dots, \beta_d \in [-2, 2]$ satisfying

$$\left\| \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] - \sum_{i=1}^d \beta_i \cdot \mathbb{E}^{\pi_i}[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] \right\| \leq 6d\varepsilon'. \quad (91)$$

Since LinOpt is based on PSDP as in Line 5 of Algorithm 7 and $\Psi_1, \dots, \Psi_h$ are (α, ε) -policy covers for layers 1 to h , respectively, (Mhammedi et al., 2023, Corollary H.1) implies that there is an event $\mathcal{E}^{\text{PSDP}}$ of probability at least $1 - \delta/2$ under which for any $\bar{\theta} \in \mathbb{R}^d \setminus \{0\}$, the output $\hat{\pi}_{\bar{\theta}} = \text{LinOpt}(\bar{\theta})$ satisfies

$$\begin{aligned} \sup_{\pi \in \Pi} \bar{\theta}^\top \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] &\leq \bar{\theta}^\top \mathbb{E}^{\hat{\pi}_{\bar{\theta}}}[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] \\ &+ \|\bar{\theta}\| \cdot \left(cH^2 d \sqrt{\frac{dA \cdot (d \log(2n\sqrt{d}H) + \log(n|\Phi|/\delta))}{\alpha n}} + H^2 d^{3/2} \varepsilon \right), \end{aligned}$$

for a large enough absolute constant $c > 0$. On the other hand, since LinEst is based on EstVec as in Line 6, (Mhammedi et al., 2023, Lemma G.3) implies that there is an event $\mathcal{E}^{\text{EstVec}}$ of probability at least $1 - \delta/2$ under which for all $\pi \in \Pi$, the output $\hat{\phi}^\pi := \text{LinEst}(\pi)$ satisfies

$$\|\hat{\phi}^\pi - \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)]\| \leq c \cdot \sqrt{\frac{\log(2/\delta)}{n}},$$

for a large enough absolute constant $c > 0$. Therefore, under $\mathcal{E}^{\text{PSDP}} \cap \mathcal{E}^{\text{EstVec}}$, LinOpt and LinEst satisfy (90) with

$$\varepsilon' := cH^2 d \sqrt{\frac{dA \cdot (d \log(2n\sqrt{d}H) + \log(n|\Phi|/\delta))}{\alpha n}} + H^2 d^{3/2} \varepsilon.$$

Therefore, by (Mhammedi et al., 2023, Proposition) and the fact that $d \sqrt{\frac{A \log(ndH|\Phi|/\delta)}{\alpha n}} \leq \varepsilon'$, the output

$$\pi_{1:2d} = \text{RobustSpanner} \left(\text{LinOpt}(\cdot), \text{LinEst}(\cdot), 2, d \sqrt{\frac{A \log(ndH|\Phi|/\delta)}{\alpha n}} \right)$$

is such that for all $\pi \in \Pi$, there exist $\beta_1, \dots, \beta_{2d} \in [-2, 2]$ satisfying

$$\left\| \mathbb{E}^\pi[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] - \sum_{i=1}^{2d} \beta_i \cdot \mathbb{E}^{\pi_i}[\bar{\phi}_h(\mathbf{x}_h, \mathbf{a}_h)] \right\| \leq \varepsilon_{\text{span}}(n, \alpha, \delta),$$

where $\varepsilon_{\text{span}}(n, \alpha, \delta)$ is as in (89).

Bounding the number of episodes By (Mhammedi et al., 2023, Proposition E.1), RobustSpanner calls LinOpt and LinEst as most $\tilde{O}(d^2)$ times. Each call to LinOpt [resp. LinEst] requires H^2n episodes. This implies the desired bound on the number of iterations. $\square$

F.6 Representation + Spanner

Corollary F.1. Let $\varepsilon, \delta \in (0, 1)$, $\phi_{1:H}^{\text{rep}}$, and Ψ_h^{span} be as in Algorithm 4. Then, for all $h \in [H]$, $\phi_h^{\text{rep}} \in \Phi$ and $|\Psi_h^{\text{span}}| = 2d$ and there is an event $\mathcal{E}^{\text{rep+span}}$ of probability $1 - 3\delta/4$ under which for all $h \in [H]$:

- For $f \in \mathcal{F}_{h+1}$, with $\mathcal{F}_{h+1}$ as in (97), there exists $w_{h+1}^f \in \mathbb{B}_d(3d^{3/2})$ such that:

$$\forall \pi \in \Pi, \quad \left| \mathbb{E}^\pi \left[\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h] \right] \right| \leq \varepsilon_{\text{rep}} := 10d^{7/2}\varepsilon; \quad (92)$$

- For all $\pi' \in \Pi$, there exist $\{\beta_\pi \in [-2, 2] : \pi \in \Psi_h^{\text{span}}\}$ such that for $\bar{\phi}_h^{\text{rep}} := [\phi_h^{\text{loss}}, \phi_h^{\text{rep}}]$

$$\left\| \mathbb{E}^{\pi'} [\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)] - \sum_{\pi \in \Psi_h^{\text{span}}} \beta_\pi \cdot \mathbb{E}^\pi [\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)] \right\| \leq \varepsilon_{\text{span}} := 2H^2d^{5/2}\varepsilon. \quad (93)$$

Proof. From Algorithm 4, we have

$\phi_h^{\text{rep}} = \text{RepLearn}(h, \mathcal{F}_{h+1}, \Phi, \text{unif}(\Psi_h^{\text{cov}}), T_{\text{rep}})$ and $\Psi_h^{\text{span}} = \text{Spanner}(h, \Phi, \Psi_{1:h}^{\text{cov}}, \bar{\phi}_h^{\text{rep}}, T_{\text{span}})$, where

$$\Psi_{1:H}^{\text{cov}} = \text{VoX}(\Phi, \varepsilon, \delta/4), \quad T_{\text{rep}} := \frac{AH \log(|\Phi|/\delta)}{\alpha \varepsilon^2}, \quad T_{\text{span}} = \frac{A \log(dH|\Phi|\varepsilon^{-1}\delta^{-1})}{\alpha \varepsilon^2}, \quad (94)$$

and $\alpha := \frac{1}{8Ad}$. By Lemma G.1, there is an event $\mathcal{E}^{\text{cov}}$ of probability at least $1 - \delta/4$ under which, for all $h \in [H]$, Ψ_h^{cov} is an (α, ε) -policy cover for layer h with $|\Psi_h^{\text{cov}}| = d$. In what follows, we condition on $\mathcal{E}^{\text{cov}}$. By Lemma G.2 and Lemma F.2, there are events $\mathcal{E}^{\text{rep}}$ and $\mathcal{E}^{\text{span}}$ of probability at least $1 - \delta/4$ each such that under $\mathcal{E}^{\text{rep}} \cap \mathcal{E}^{\text{span}}$ (92) and (93) hold; this follows from (98) and (88) and the choices of T_{rep} and T_{span} in (94). Finally, by the union bound, we have $\mathbb{P}[\mathcal{E}^{\text{cov}} \cap \mathcal{E}^{\text{rep}} \cap \mathcal{E}^{\text{span}}] \geq 1 - \delta$ which completes the proof. $\square$

F.7 Martingale Concentration

Lemma F.3. Let $K, N_{\text{reg}}, \bar{\phi}_{1:H}^{\text{rep}}$, and $\mathcal{I}^{(k)}$ be as in Algorithm 4 for $k \in [K]$. There is an event $\mathcal{E}^{\text{freed}}$ of probability at least $1 - \delta/4$ under which for all $\bar{\theta} \in \mathbb{B}_{2d}(4Hd^2)$, $h \in [H]$, $k \in [K]$, and $\pi \in \Psi_h^{\text{span}}$:

$$\begin{aligned} & \frac{1}{N_{\text{reg}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \right. \\ & \left. - \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \mid \mathcal{H}^{t-1} \right] \right| \leq \varepsilon_{\text{freed}} := 4Hd^2 \sqrt{\frac{\nu \log(dKH N_{\text{reg}}/\delta)}{N_{\text{reg}}}}, \end{aligned}$$

where the random variables $\mathbf{h}^t, \zeta^t, \boldsymbol{\pi}^t$, and $\mathcal{H}^{t-1}$ are as in Algorithm 7.

Proof. Fix $\bar{\theta} \in \mathbb{B}_{2d}(4d^2)$, $h \in [H]$, $k \in [K]$, and $\pi \in \Psi_h^{\text{span}}$. We apply Lemma I.2 (Freedman's inequality) with

- $R = 4Hd^2$;
- $n = N_{\text{reg}}$;
- The random variable w^i set as the difference

$$\begin{aligned} w^i &:= \mathbb{I}\{\mathbf{h}^{t_i} = h, \boldsymbol{\pi}^{t_i} = \pi, \zeta^{t_i} = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i})^\top \bar{\theta} - \sum_{s=h}^H \ell_s^{t_i} \right) \\ &\quad - \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^{t_i} = h, \boldsymbol{\pi}^{t_i} = \pi, \zeta^{t_i} = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i})^\top \bar{\theta} - \sum_{s=h}^H \ell_s^{t_i} \right) \mid \mathcal{H}^{t-1} \right], \quad (95) \end{aligned}$$

where $t_i := (k-1) \cdot N_{\text{reg}} + i$;

- The filtration $\mathfrak{F}^i$ set as the σ -algebra $\sigma(\mathcal{H}^{t_i-1})$;
- The variance term V_n has the following upper bound

$$V_n = \sum_{i=1}^{N_{\text{reg}}} \mathbb{E}[(\mathbf{w}^i)^2 \mid \mathfrak{F}^{i-1}] \leq \sum_{i=1}^{N_{\text{reg}}} \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^{t_i} = h, \boldsymbol{\pi}^{t_i} = \pi, \zeta^{t_i} = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^{t_i}, \mathbf{a}_h^{t_i})^\top \bar{\theta} - \sum_{s=h}^H \ell_s^{t_i} \right)^2 \mid \mathfrak{F}^{i-1} \right] \\ \leq 8Hd^3 \nu N_{\text{reg}};$$

- $\lambda = H^{-1} \left(\frac{d^2 \nu N_{\text{reg}}}{\log(KHN_{\text{reg}}/\delta)} \right)^{-1/2}$;

to get that there is an event $\mathcal{E}_{h,k,\pi}^{\text{freed}}(\bar{\theta})$ of probability at least $1 - (N_{\text{reg}})^{-d} H^{-1} K^{-1} d^{-1} \delta/8$ under which

$$\frac{1}{N_{\text{reg}}} \left| \sum_{t \in \mathcal{I}^{(k)}} \mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \right. \\ \left. - \sum_{t \in \mathcal{I}^{(k)}} \mathbb{E} \left[\mathbb{I}\{\mathbf{h}^t = h, \boldsymbol{\pi}^t = \pi, \zeta^t = 1\} \cdot \left(\bar{\phi}_h^{\text{rep}}(\mathbf{x}_h^t, \mathbf{a}_h^t)^\top \bar{\theta} - \sum_{s=h}^H \ell_s^t \right) \mid \mathcal{H}^{t-1} \right] \right| \\ \leq 4Hd^2 \sqrt{\frac{\nu \log(dKH N_{\text{reg}}/\delta)}{N_{\text{reg}}}}. \quad (96)$$

Let $\mathcal{C}$ be a minimal $(dHN_{\text{reg}})^{-1}$ -cover of $\mathbb{B}_{2d}(4Hd^2)$ with respect to the $\|\cdot\|$ distance. Under the event

$$\mathcal{E}^{\text{freed}} := \bigcap_{h \in [H], k \in [K], \pi \in \Psi_h^{\text{span}}, \bar{\theta} \in \mathbb{B}_{2d}(4Hd^2)} \mathcal{E}_{h,k,\pi}^{\text{freed}}(\bar{\theta}),$$

Eq. (96) holds for all $h \in [H]$, $k \in [K]$, and $\bar{\theta} \in \mathbb{B}_{2d}(4Hd^2)$ up to an additive $O(1/N_{\text{reg}})$ error. By the union bound, we have $\mathbb{P}[\mathcal{E}^{\text{freed}}] \geq 1 - \delta/4$ which completes the proof. $\square$

G Policy Cover and Representation Learning Algorithms

In this section, we present guarantees for VoX, RepLearn, and RobustSpanner which we need in the analysis of our oracle efficient algorithm. The results are based on (Mhammedi et al., 2023).

G.1 Policy Cover

The following result is a restatement of (Mhammedi et al., 2023, Theorem 12).

Lemma G.1 (VoX Guarantee). *Let $\varepsilon, \delta \in (0, 1)$ be given. Suppose Assumption 2.1 and Assumption 4.1 hold. Then, there is an event $\mathcal{E}^{\text{cov}}$ of probability at least $1 - \delta$ under which the output $\Psi_{1:H}^{\text{cov}} = \text{VoX}(\Phi, \varepsilon, \delta)$ is such that for all $h \in [H]$:*

- Ψ_h^{cov} is a $(\frac{1}{8Ad}, \varepsilon)$ -policy cover for layer h ;
- $|\Psi_h^{\text{cov}}| \leq d$.

Furthermore, the number of episodes $T_{\text{cov}}(\varepsilon)$ used by the call to VoX is bounded by $\tilde{O}(Ad^{13}H^6 \log(\Phi/\delta))/\varepsilon^2$.

G.2 Representation Learning

Lemma G.2 (Representation Learning Guarantee). *Let $\varepsilon, \alpha, \delta \in (0, 1)$, $h \in [H - 1]$, and $n \geq 1$ be given and define the function class*

$$\mathcal{F}_{h+1} := \left\{ f : (x, a) \mapsto \max_{a \in \mathcal{A}} \bar{\phi}_{h+1}(x, a)^\top \bar{\theta} \mid \bar{\phi}_{h+1} = [\phi_{h+1}^{\text{loss}}, \phi_{h+1}], \phi \in \Phi, \bar{\theta} \in \mathbb{B}_{2d}(1) \right\}. \quad (97)$$

Further, let Ψ be an (α, ε) -policy cover for layer h with $|\Psi| = d$, and suppose Assumption 2.1 and Assumption 4.1 hold. Then, with probability at least $1 - \delta$, the output $\bar{\phi}_h^{\text{rep}} = \text{RepLearn}(h, \mathcal{F}_{h+1}, \Phi, \text{unif}(\Psi), n)$ is such that for all $f \in \mathcal{F}_{h+1}$ there exists $w_{h+1}^f \in \mathbb{B}_d(3d^{3/2})$ such that:

$$\forall \pi \in \Pi, \quad \left| \mathbb{E}^\pi \left[\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h] \right] \right| \leq c \cdot \sqrt{\frac{AHd^5 \log(|\Phi|/\delta)}{\alpha n}} + 9d^{7/2}\varepsilon, \quad (98)$$

where $c > 0$ is a large enough absolute constant. Furthermore, the number of episodes $T_{\text{rep}}(\varepsilon)$ used by the call to RepLearn is equal to n .

Proof. By (Mhammedi et al., 2023, Theorem F.1) and the assumption that $|\Psi| = d$, there is an event $\mathcal{E}$ of probability at least $1 - \delta/2$ under which ϕ_h^{rep} satisfies:

$$\sup_{f \in \mathcal{F}_{h+1}} \inf_{w \in \mathbb{B}_d(3d^{3/2})} \max_{\pi' \in \Psi} \mathbb{E}^{\pi'} \left[(\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h])^2 \right] \leq c \cdot \frac{Ad^5 \log(|\Phi|/\delta)}{n}, \quad (99)$$

where c is a large enough absolute constant. We use this to show (98). In what follows, we condition on $\mathcal{E}$. Fix $\pi \in \Pi$ and $f \in \mathcal{G}$ and let w_{h+1}^f be the vector $w \in \mathbb{B}_d(3d^{3/2})$ achieving the infimum in (99) for the given choice of f . Let $\mathcal{X}_{h,\varepsilon}$ be the set of ε -reachable states at layer h as defined in (50). With this this, we have for all $h \in [H]$,

$$\begin{aligned} & \left| \mathbb{E}^\pi \left[\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h] \right] \right| \\ & \leq \left| \mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \notin \mathcal{X}_{h,\varepsilon}\} \cdot (\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h]) \right] \right| \\ & \quad + \left| \mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot (\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h])^2 \right] \right|, \\ & \leq \left| \mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot (\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h]) \right] \right| + 9d^{7/2}\varepsilon, \end{aligned} \quad (100)$$

where the last inequality follows by Lemma I.1.

We now bound the first term on the right-hand side of (100). By Jensen's inequality, we have

$$\left| \mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot (\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h]) \right] \right|$$

$$\leq \sqrt{\mathbb{E}^\pi \left[\mathbb{I}\{\mathbf{x}_h \in \mathcal{X}_{h,\varepsilon}\} \cdot (\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h])^2 \right]},$$

and so using that Ψ is a (α, ε') -policy cover (see Definition 4.1), we have

$$\begin{aligned} &\leq \sqrt{\alpha^{-1} \cdot \max_{\pi' \in \Psi} \mathbb{E}^{\pi'} \left[(\phi_h^{\text{rep}}(\mathbf{x}_h, \mathbf{a}_h)^\top w_{h+1}^f - \mathbb{E}[f(\mathbf{x}_{h+1}) \mid \mathbf{x}_h, \mathbf{a}_h])^2 \right]}, \\ &\leq \sqrt{c \cdot \frac{A H d^5 \log(|\Phi|/\delta)}{\alpha n}}, \end{aligned}$$

where the last step follows by (99). Combining this with (100) yields (98). $\square$

H Lower Bound for Bandit feedback with Unstructured Losses

In full-information, one does not require any structure on the losses. We show that this is not the case for the bandit case via a lower bound depending polynomially on the number of states. This lower bound implies that the low-rank transition structure with unstructured losses does not give any significant improvements over the tabular setting.

Theorem H.1. *There exists a low-rank MDP with S states, A actions and sufficiently large time step T with unstructured losses such that any agent suffers at least regret of $\Omega(\sqrt{SAT})$.*

Proof. We assume $4S < \sqrt{T}$. The construction is an $H = 1$ (i.e. contextual bandit) MDP with uniform initial distribution over states. Each state is a copy of an A -armed bandit problem with Bernoulli losses with mean $\frac{1}{2}$, and one randomly chosen optimal arm with mean $\frac{1}{2} - \Delta$. Following the standard lower bound construction for bandits (Lattimore and Szepesvári, 2020), there exists $\Delta = \Theta(1/\sqrt{TAS})$ such that the regret of playing any individual bandit problem for $N \leq 2T/S$ rounds is lower bounded by $\Omega(N\Delta)$. Let denote $N(s)$ the number of time the agent receives the initial state s , then any agent suffers a regret lower bound in our MDP of

$$\Omega\left(\sum_{s=1}^S \mathbb{E}[\min\{N(s), 2T/S\}]\right).$$

$N(s)$ is the sum of T Bernoulli random variables with mean $1/S$. We have by Hoeffding's inequality

$$\mathbb{P}[N(s) > T/S + x] \leq \exp(-2x^2/T).$$

This allows to upper bound the tail

$$\begin{aligned} \mathbb{E}[N(s)\mathbb{I}(N(s) > 2T/S)] &\leq \int_{T/S}^{\infty} x(2x/T \exp(-2x^2/T)) dx \\ &\leq \int_{T/S}^{\infty} 4 \exp(x/\sqrt{T} - 2x^2/T) dx \quad (\tfrac{1}{2}(x/\sqrt{T})^2 < \exp(x/\sqrt{T})) \\ &\leq \int_{T/S}^{\infty} 4 \exp(-x/\sqrt{T}) dx \quad (x \geq T/S > 4\sqrt{T}) \\ &= 4\sqrt{T} \exp(-\sqrt{T}/S) \leq 4\sqrt{T} \exp(-4) \leq T/(2S). \end{aligned}$$

Hence $\mathbb{E}[\min\{N(s), 2T/S\}] = T/S - \mathbb{E}[N(s)\mathbb{I}(N(s) > 2T/S)] \geq T/(2S)$ and the regret in the MDP is lower bounded by $\Omega(\sqrt{TSA})$. $\square$

I Helper Results

Lemma I.1. *Let $\varepsilon, B > 0$ and $h \in [2..H]$ be given. For any function $f : \mathcal{X} \rightarrow [-B, B]$ and $\pi \in \Pi$, we have*

$$\mathbb{E}^{\pi}[\mathbb{I}\{\mathbf{x}_h \notin \mathcal{X}_{h,\varepsilon}\} \cdot f(\mathbf{x}_h)] \leq B\sqrt{d\varepsilon}, \quad (101)$$

where

$$\mathcal{X}_{h,\varepsilon} := \left\{x \in \mathcal{X} : \max_{\pi \in \Pi} d_h^{\pi}(x) \geq \varepsilon \cdot \|\mu_h^*(x)\|\right\}, \quad (102)$$

denotes the set of states that are ε -reachable at layer h .

Proof. Fix $f : \mathcal{X} \rightarrow [-B, B]$ and $\pi \in \Pi$. Using the definition of $\mathcal{X}_{h,\varepsilon}$ in (102), we have that $x \notin \mathcal{X}_{h,\varepsilon}$ only if $d_h^{\pi}(x) < \varepsilon \|\mu_h^*(x)\|$. Using this, we have

$$\begin{aligned} \mathbb{E}^{\pi}[\mathbb{I}\{\mathbf{x}_h \notin \mathcal{X}_{h,\varepsilon}\} \cdot f(\mathbf{x}_h)] &\leq \sum_{x \in \mathcal{X}} \mathbb{I}\{x \notin \mathcal{X}_{h,\varepsilon}\} \cdot d_h^{\pi}(x) \cdot f(x), \\ &\leq B\varepsilon \sum_{x \in \mathcal{X}} \|\mu_h^*(x)\|, \\ &\leq Bd^{3/2}\varepsilon, \end{aligned} \quad (103)$$

where the last step follows by the normalizing assumption on μ^* (see Assumption 2.1) and (Mhammedi et al., 2023, Lemma I.3). $\square$

I.1 Martingale Concentration and Regression Results

Lemma I.2. Let $R > 0$ be given and let $\mathbf{w}^1, \dots, \mathbf{w}^n$ be a sequence of real-valued random variables adapted to filtration $\mathfrak{F}^1, \dots, \mathfrak{F}^n$. Assume that for all $i \in [n]$, $\mathbf{w}^i \leq R$ and $\mathbb{E}[\mathbf{w}^i | \mathfrak{F}^{i-1}] = 0$. Define $S_n := \sum_{i=1}^n \mathbf{w}^i$ and $V_n := \sum_{i=1}^n \mathbb{E}[(\mathbf{w}^i)^2 | \mathfrak{F}^{i-1}]$. Then, for any $\delta \in (0, 1)$ and $\lambda \in [0, 1/R]$, with probability at least $1 - \delta$,

$$S_n \leq \lambda V_n + \ln(1/\delta)/\lambda. \quad (104)$$

We now state two helpful results from Mhammedi et al. (2024b) without a proof.

Lemma I.3. Let $B > 0$ and $n \in \mathbb{N}$ be given. abstract set. Further, let $\mathcal{Q} \subseteq \{g : \mathcal{X} \times \mathcal{A} \rightarrow [0, B]\}$ be a finite function class and $(\mathbf{x}^1, \mathbf{a}^1, \varepsilon^1), \dots, (\mathbf{x}^n, \mathbf{a}^n, \varepsilon^n)$ be a sequence of i.i.d. random variables in $\mathcal{X} \times \mathcal{A} \times \mathbb{R}$. Then, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$, we have

$$\forall g \in \mathcal{Q}, \quad \frac{1}{2} \|g\|^2 - 2B^2 \log(2|\mathcal{Q}|/\delta) \leq \|g\|_n^2 \leq 2\|g\|^2 + 2B^2 \log(2|\mathcal{Q}|/\delta), \quad (105)$$

where $\|g\|^2 := \sum_{i \in [n]} \mathbb{E}[g(\mathbf{x}^i, \mathbf{a}^i)^2 | \mathfrak{F}^{i-1}]$ and $\|g\|_n^2 := \sum_{i=1}^n g(\mathbf{x}^i, \mathbf{a}^i)^2$.

Lemma I.4 (Generic regression guarantee). Let $B > 0$, $n \in \mathbb{N}$, and $f_\star : \mathcal{X} \times \mathcal{A} \rightarrow [0, B]$ be given. Further, let $\mathcal{F} \subseteq \{f : \mathcal{X} \times \mathcal{A} \rightarrow [0, B]\}$ be a finite function class and $(\mathbf{x}^1, \mathbf{a}^1, \varepsilon^1), \dots, (\mathbf{x}^n, \mathbf{a}^n, \varepsilon^n)$ be a sequence of i.i.d. random variables in $\mathcal{X} \times \mathcal{A} \times \mathbb{R}$. Suppose that

- $f_\star \in \mathcal{F}$;
- $\mathbf{z}^i = f_\star(\mathbf{x}^i, \mathbf{a}^i) + \varepsilon^i + \mathbf{b}^i$, for all $i \in [n]$;
- $\mathbf{b}^1, \dots, \mathbf{b}^n \in \mathbb{R}$ (not necessarily i.i.d.);
- $\varepsilon^i \in [-B, B]$, for all $i \in [n]$; and
- $\mathbb{E}[\varepsilon^i | \mathbf{x}^i, \mathbf{a}^i] = 0$.

Then, for $\hat{f} \in \operatorname{argmin}_{f \in \mathcal{F}} \sum_{i=1}^n (f(\mathbf{x}^i, \mathbf{a}^i) - \mathbf{z}^i)^2$ and any $\delta \in (0, 1)$, with probability at least $1 - \delta/2$,

$$\|\hat{f} - f_\star\|_n^2 \leq 8B^2 \log(2|\mathcal{F}|/\delta) + 8 \sum_{i=1}^n (\mathbf{b}^i)^2, \quad (106)$$

where $\|\hat{f} - f_\star\|_n^2 := \sum_{i=1}^n (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i))^2$.

Proof. Fix $\delta \in (0, 1)$ and let $\widehat{L}_n(f) := \sum_{i=1}^n (f(\mathbf{x}^i, \mathbf{a}^i) - \mathbf{z}^i)^2$, for $f \in \mathcal{F}$, and note that since $\hat{f} \in \operatorname{argmin}_{f \in \mathcal{F}} \widehat{L}_n(f)$, we have

$$0 \geq \widehat{L}_n(\hat{f}) - \widehat{L}_n(f_\star) = \nabla \widehat{L}_n(f_\star)[\hat{f} - f_\star] + \|\hat{f} - f_\star\|_n^2, \quad (107)$$

where ∇ denotes directional derivative. Rearranging, we get that

$$\begin{aligned} \|\hat{f} - f_\star\|_n^2 &\leq -2\nabla \widehat{L}_n(f_\star)[\hat{f} - f_\star] - \|\hat{f} - f_\star\|_n^2, \\ &= 4 \sum_{i=1}^n (\mathbf{z}^i - f_\star(\mathbf{x}^i, \mathbf{a}^i))(\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - \|\hat{f} - f_\star\|_n^2, \\ &= 4 \sum_{i=1}^n (\varepsilon^i + \mathbf{b}^i)(\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - \|\hat{f} - f_\star\|_n^2, \\ &= 4 \sum_{i=1}^n \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - \|\hat{f} - f_\star\|_n^2 + 4 \sum_{i=1}^n \mathbf{b}^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)), \\ &\leq 4 \sum_{i=1}^n \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - \|\hat{f} - f_\star\|_n^2 + 4 \sum_{i=1}^n (\mathbf{b}^i)^2 + \frac{1}{2} \sum_{i=1}^n (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i))^2, \\ &= 4 \sum_{i=1}^n \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - \|\hat{f} - f_\star\|_n^2 + 4 \sum_{i=1}^n (\mathbf{b}^i)^2 + \frac{1}{2} \|\hat{f} - f_\star\|_n^2. \end{aligned} \quad (108)$$

Thus, rearranging, we get

$$\|\hat{f} - f_\star\|_n^2 \leq 8 \sum_{i=1}^n \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) - 2\|\hat{f} - f_\star\|_n^2 + 8 \sum_{i=1}^n (\mathbf{b}^i)^2. \quad (110)$$

We now bound the first term on the right-hand side of (110). For this, we apply Lemma I.2 with $\mathbf{w}^i = \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i))$, $R = B^2$, $\lambda = 1/(8B^2)$, and $\mathfrak{F}^i = \emptyset$, and use

1. the union bound over $f \in \mathcal{F}$; and
2. the facts that $\mathbb{E}[\varepsilon^i \mid \mathbf{x}^i, \mathbf{a}^i] = 0$,

to get that with probability at least $1 - \delta/2$,

$$\sum_{i=1}^n \varepsilon^i \cdot (\hat{f}(\mathbf{x}^i, \mathbf{a}^i) - f_\star(\mathbf{x}^i, \mathbf{a}^i)) \leq \frac{1}{4} \|\hat{f} - f_\star\|_n^2 + B^2 \log(2|\mathcal{F}|/\delta). \quad (111)$$

Combining this with (110), we get that with probability at least $1 - \delta/2$,

$$\|\hat{f} - f_\star\|_n^2 \leq 8B^2 \log(2|\mathcal{F}|/\delta) + 8 \sum_{i=1}^n (\mathbf{b}^i)^2. \quad (112)$$

This completes the proof. $\square$

I.2 Online Learning

The following is the standard guarantee of exponential weights (e.g. Lemma F.4 of Sherman et al. (2023b)).

Lemma I.5 (Exponential Weights). *Given a sequence of loss functions $\{g^t\}_{t=1}^T$ over a decision set Π , $\{p^t\}_{t=1}^T$ is a distribution sequence with $p^t \in \Delta(\Pi)$, $\forall t \in [T]$ such that*

$$p^{t+1}(\pi) \propto \exp\left(-\eta \sum_{t=1}^T g^t(\pi)\right).$$

If p^1 is a uniform distribution over $|\Pi|$ and $\eta g^t(\pi) \geq -1$ for all $t \in [T]$ and $\pi \in \Pi$. Then

$$\max_{p \in \Delta(\Pi)} \left\{ \sum_{t=1}^T \langle g^t, p^t - p \rangle \right\} \leq \frac{\log(|\Pi|)}{\eta} + \eta \sum_{t=1}^T \sum_{\pi \in \Pi} p^t(\pi) g^t(\pi)^2$$

I.3 Reinforcement Learning

The following is standard simulation lemma which is first proposed by Abbeel and Ng (2005).

Lemma I.6 (Simulation Lemma). *For two finite-horizon MDPs $\widehat{M} = \{\mathcal{X}, \mathcal{A}, \ell, \{\widehat{P}_h\}_{h=1}^H\}$ and $M = \{\mathcal{X}, \mathcal{A}, \ell, \{P_h\}_{h=1}^H\}$ with horizon H and $\|\ell\|_\infty \leq 1$. Let the corresponding value function be $\widehat{V}_h^\pi(x; \ell)$ and $V_h^\pi(x; \ell)$ for step $h \in [H]$. For any policy $\pi : \mathcal{X} \rightarrow \Delta(\mathcal{A})$, we have*

$$|\widehat{V}_1^\pi(x_1; \ell) - V_1^\pi(x_1; \ell)| \leq H \sum_{h=1}^H \mathbb{E}_{x, a \sim d_h^\pi} [\|\widehat{P}_h(\cdot \mid x, a) - P_h(\cdot \mid x, a)\|_1].$$

The following is the standard performance difference lemma which is first proposed by Kakade and Langford (2002).

Lemma I.7 (Performance Difference Lemma). *For a finite-horizon MDPs $M = \{\mathcal{X}, \mathcal{A}, \ell, \{P_h\}_{h=1}^H\}$ starting at x_1 , and two policies $\pi, \pi' : \mathcal{X} \rightarrow \Delta(\mathcal{A})$, we have*

$$V_1^{\pi'}(x_1; \ell) - V_1^\pi(x_1; \ell) = \sum_{h=1}^H \mathbb{E}_{x \sim d_h^\pi} \left[\sum_{a \in \mathcal{A}} (\pi'_h(a|x) - \pi_h(a|x)) Q_h^{\pi'}(x, a; \ell) \right]$$

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes] , [No] , or [NA] .
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]" , it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes] .

Justification: The main claims made in the abstract and introduction accurately reflect the paper's contributions and scope. The claims are validated by detailed proofs.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes] .

Justification: The paper discusses the limitations of the work in the discussion section.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.

- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#) .

Justification: The paper provides detailed assumptions and proofs.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[NA\]](#) .

Justification: This is a theoretical paper.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.

- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [NA] .

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA] .

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA] .

Justification: This paper does not include experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines?>

Answer: [Yes] .

Justification: The research conducted in the paper conforms, in every respect, with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: This is a theoretical work. There is no societal impact of the work performed.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA] .

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA] .

Justification: This paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA] .

Justification: This paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA] .

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA] .

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.

- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.