
Optimal Horizon-Free Reward-Free Exploration for Linear Mixture MDPs

Junkai Zhang¹ Weitong Zhang¹ Quanquan Gu¹

Abstract

We study reward-free reinforcement learning (RL) with linear function approximation, where the agent works in two phases: (1) in the exploration phase, the agent interacts with the environment but cannot access the reward; and (2) in the planning phase, the agent is given a reward function and is expected to find a near-optimal policy based on samples collected in the exploration phase. The sample complexities of existing reward-free algorithms have a polynomial dependence on the planning horizon, which makes them intractable for long planning horizon RL problems. In this paper, we propose a new reward-free algorithm for learning linear mixture Markov decision processes (MDPs), where the transition probability can be parameterized as a linear combination of known feature mappings. At the core of our algorithm is uncertainty-weighted value-targeted regression with exploration-driven pseudo-reward and a high-order moment estimator for the aleatoric and epistemic uncertainties. When the total reward is bounded by 1, we show that our algorithm only needs to explore $\tilde{O}(d^2\varepsilon^{-2})$ episodes to find an ε -optimal policy, where d is the dimension of the feature mapping. The sample complexity of our algorithm only has a polylogarithmic dependence on the planning horizon and therefore is “horizon-free”. In addition, we provide an $\Omega(d^2\varepsilon^{-2})$ sample complexity lower bound, which matches the sample complexity of our algorithm up to logarithmic factors, suggesting that our algorithm is optimal.

1. Introduction

In Reinforcement Learning (RL), the agent sequentially interacts with the environment by executing policies and re-

ceiving observations, including states and rewards. The goal of the agent is to maximize the total reward. To achieve this goal, the agent needs to explore the environment and exploit the collected information to find the optimal policy. The exploration has long been considered as a central challenge for RL, for which the agent needs to strategically visit states to learn transition dynamics and the value of different states. RL algorithms are often designed to exploit the transition and the reward information to achieve efficient exploration.

Unfortunately, in many real-world RL problems, reward functions are manually designed to incentive the agent to learn specific tasks, and they may change over time (Achiam et al., 2017; Tessler et al., 2018; Miryoosefi et al., 2019). To avoid learning the transition dynamics repeatedly, Jin et al. (2020a) proposed a new RL paradigm, *Reward-free Exploration*, which separates exploration and planning into two different phases. In the exploration phase, the agent cannot access the real reward function. It can only learn the transition dynamics based on the collected episodes. While in the planning phase, the agent can no longer interact with the environment but has access to the reward function. The goal is to find the optimal policy based on the reward function and previous exploration. A series of work (Jin et al., 2020a; Kaufmann et al., 2021b; Ménard et al., 2021; Zhang et al., 2020) have achieved the optimal sample complexity of $\tilde{O}(H^2S^2A\varepsilon^{-2})$, where H is the planning horizon, S is the number of states, and A is the number of actions.

The sample complexity for learning tabular MDPs shows that learning becomes intractable when the sizes of the state and action spaces increase without further structural assumptions. Linear function approximation is a classical approach to deal with this challenge, which approximates the transition dynamic or the value function by linear functions on compact feature mappings. To this end, Wang et al. (2020b); Zanette et al. (2020c) studied reward-free RL in linear MDPs (Yang & Wang, 2019; Jin et al., 2020b). Zhang et al. (2021a) studied reward-free exploration for linear mixture MDPs (Ayoub et al., 2020; Zhou et al., 2021c), where the transition probability is a linear combination of feature mappings. The subsequent work Chen et al. (2021) has achieved near optimal sample complexity $\tilde{O}(H^3d(H+d)\varepsilon^{-2})$. However, this sample complexity depends on the planning horizon, which will blow up for long planning horizon problems. Therefore, a natural question arises:

¹Department of Computer Science, University of California, Los Angeles, California, USA. Correspondence to: Quanquan Gu <qgu@cs.ucla.edu>.

Can we design horizon-free, minimax optimal reward-free RL algorithms with linear function approximation?

In this paper, we answer this question affirmatively for linear mixture MDPs. In detail, our contributions are highlighted as follows.

- We propose an algorithm for reward-free exploration in the linear mixture MDP setting. The algorithm guides the agent to collect samples using a well-designed exploration-driven pseudo-reward function. With a novel analysis based on high-order moment estimation that precisely controls the aleatoric and epistemic uncertainties, our algorithm can achieve an $\tilde{O}(d^2\varepsilon^{-2})$ sample complexity. This complexity only has polylogarithmic dependence on H .
- We show that any reward-free algorithm needs to explore at least $\Omega(d^2\varepsilon^{-2})$ episodes, to achieve an ε -optimal policy for any reward function, by constructing a special class of linear mixture MDPs. This lower bound matches the upper bound of our algorithm up to logarithmic factors, which indicates that our algorithm is optimal.
- When rescaling the reward to satisfy $\sum_{h=1}^H r_h(s_h, a_h) \leq H$, our algorithm achieves an $\tilde{O}(d^2H^2\varepsilon^{-2})$ sample complexity. This improves the sample complexity in the previous work Chen et al. (2021), which requires the $d > H$ condition to match the lower bound.

Notation We use the lowercase letter to denote scalars and lower and uppercase boldface letters to denote vectors and matrices, respectively. We denote by $[n]$ the set $\{1, \dots, n\}$, and by $\overline{[n]}$ the set $\{0, \dots, n-1\}$. For a vector $\mathbf{x}$ and a positive semi-definite matrix Σ , we denote by $\|\mathbf{x}\|_2$ the vector's Euclidean norm and define $\|\mathbf{x}\|_\Sigma = \sqrt{\mathbf{x}^\top \Sigma \mathbf{x}}$. For two positive sequences $\{a_n\}$ and $\{b_n\}$ with $n = 1, 2, \dots$, we write $a_n = O(b_n)$ if there exists an absolute constant $C > 0$ such that $a_n \leq Cb_n$ holds for all $n \geq 1$, write $a_n = \Omega(b_n)$ if there exists an absolute constant $C > 0$ such that $a_n \geq Cb_n$ holds for all $n \geq 1$, and write $a_n = o(b_n)$ if $a_n/b_n \rightarrow 0$ as $n \rightarrow \infty$. We use $\tilde{O}(\cdot)$ and $\tilde{\Omega}(\cdot)$ to further hide the polylogarithmic factors.

2. Related Work

RL with Linear Function Approximation In recent years, a series of works have been devoted to the study of RL with linear function approximation (Jiang et al., 2017; Dann et al., 2018; Yang & Wang, 2019; Wang et al., 2019; Du et al., 2019; Sun et al., 2019; Jin et al., 2020b; Zanette et al., 2020a;b; Yang & Wang, 2020; Modi et al., 2020; Ayoub et al., 2020; Jia et al., 2020; Cai et al., 2020; Weisz et al.,

2021; Zhou et al., 2021d;a; He et al., 2022; Agarwal et al., 2022). Our work belongs to the linear mixture MDP setting (Yang & Wang, 2019; Modi et al., 2020; Ayoub et al., 2020; Jia et al., 2020; Zhou et al., 2021a;d), where the transition kernel can be parameterized as a linear combination of some basic transition probability functions. Zhou et al. (2021a) firstly achieved minimax regret $\tilde{O}(dH\sqrt{T})$ in linear mixture MDPs by proposing a Bernstein-type concentration inequality for self-normalized martingales. Another kind of popular linearly parameterized MDP is linear MDP (Wang et al., 2019; Du et al., 2019; Yang & Wang, 2020; Jin et al., 2020b; Zanette et al., 2020a; Wang et al., 2020c; He et al., 2021), which assumes both transition probability and reward function are linear functions of known feature mappings on state-action pairs. Under this setting, Jin et al. (2020b) firstly proposed statistically and computationally efficient algorithm LSVI-UCB and achieved $\tilde{O}(\sqrt{d^3H^3T})$ regret bound. Recent works (He et al., 2022) further achieved nearly minimax optimal regret $\tilde{O}(d\sqrt{H^3K})$ by proposing computationally efficient algorithm LSVI-UCB++. Its concurrent work (Agarwal et al., 2022) achieves similar result under assumption $\sum_{h=1}^H r_h(s_h, a_h) \leq 1$ with regret upper bound of $\tilde{O}(d\sqrt{HT} + d^6H^5)$.

Reward-free RL Unlike standard RL settings, exploration and planning in reward-free RL are separated into two different phases. Jin et al. (2020a) achieved $\tilde{O}(H^5S^2A/\varepsilon^2)$ sample complexity in tabular MDPs by executing exploratory policy visiting states with probability proportional to its maximum visitation probability under any possible policy. Subsequent works (Kaufmann et al., 2021b; Ménard et al., 2021) proposed algorithms RF-UCRL and RF-Express to gradually improve the result to $\tilde{O}(H^3S^2A\varepsilon^{-2})$. The optimal sample complexity bound $\tilde{O}(H^2S^2A\varepsilon^{-2})$ was achieved by algorithm SSTP proposed in Zhang et al. (2020), which matched the lower bound provided in Jin et al. (2020a) up to logarithmic factors.

Recent years have witnessed a trend of reward-free exploration in RL with linear function approximation (Wang et al., 2020b; Zanette et al., 2020c; Zhang et al., 2021a; Chen et al., 2021; Huang et al., 2022; Wagenmaker et al., 2022). The near minimax optimal sample complexity of reward-free exploration in linear mixture MDP was achieved by Chen et al. (2021) when $d > H$ using the well-designed exploration-driven pseudo reward function. On the other hand, in the linear MDP setting, Wang et al. (2020b) proposed the first efficient algorithm, which only required $\tilde{O}(H^6d^3\varepsilon^{-2})$ sample complexity. The subsequent works, Chen et al. (2021) and Wagenmaker et al. (2022), gave sample complexity of $\tilde{O}(H^4d^3\varepsilon^{-2})$ and $\tilde{O}(H^5d^2\varepsilon^{-2})$, which are the sharpest for H and d , respectively. Some significant works are summarized in Table 1.

Optimal Horizon-Free Reward-Free Exploration for Linear Mixture MDPs

Setting	Algorithm	Rewards Scale	Time Homo.	Sample Complexity
Tabular MDP	Jin et al. (2020a)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^5 S^2 A \varepsilon^{-2})$
	RF-UCRL (Kaufmann et al., 2021a)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^4 S^2 A \varepsilon^{-2})$
	RF-Express (Ménard et al., 2021)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^3 S^2 A \varepsilon^{-2})$
	SSTP (Zhang et al., 2020)	$\sum_{h=1}^H r_h(s_h, a_h) \leq 1$	✓	$\tilde{O}(S^2 A \varepsilon^{-2})$
	Lower bound (Jin et al., 2020a)	$r_h(s_h, a_h) \in [0, 1]$	×	$\Omega(H^2 S^2 A \varepsilon^{-2})$
	Lower bound (Zhang et al., 2020)	$\sum_{h=1}^H r_h(s_h, a_h) \leq 1$	✓	$\Omega(S^2 A \varepsilon^{-2})$
Linear MDP	Wang et al. (2020b)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^6 d^3 \varepsilon^{-2})$
	FRANCIS (Zanette et al., 2020c)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^5 d^3 \varepsilon^{-2})$
	RFLIN (Wagenmaker et al., 2022)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^5 d^2 \varepsilon^{-2})$
	UCRL-RFE+ (Zhang et al., 2021a)	$r_h(s_h, a_h) \in [0, 1]$	✓	$\tilde{O}(H^4 d(H+d) \varepsilon^{-2})$
Linear Mixture MDP	Chen et al. (2021)	$r_h(s_h, a_h) \in [0, 1]$	×	$\tilde{O}(H^3 d(H+d) \varepsilon^{-2})$
	Our work (Cor. 5.2)	$\sum_{h=1}^H r_h(s_h, a_h) \leq 1$	✓	$\tilde{O}(d^2 \varepsilon^{-2})$
	Our work (Cor. 5.4)	$\sum_{h=1}^H r_h(s_h, a_h) \leq H$	✓	$\tilde{O}(H^2 d^2 \varepsilon^{-2})$
	Lower bound (Thm. 5.6)	$\sum_{h=1}^H r_h(s_h, a_h) \leq 1$	✓	$\Omega(d^2 \varepsilon^{-2})$
	Lower bound (Cor. 5.8)	$r_h(s_h, a_h) \in [0, 1]$	✓	$\Omega(H^2 d^2 \varepsilon^{-2})$

Table 1. Comparison of episodic reward-free RL algorithms. Column **Time Homo.** means if the algorithm is time-homogeneous (✓) or not (×), rows with light blue background indicates our results

Horizon-free RL The long planning horizon has long been viewed as RL’s main challenge. However, a series of works has shown that RL is no more difficult than contextual bandits by removing the influence of the total reward scale. In tabular MDPs, the algorithm proposed in Wang et al. (2020a) firstly achieved polylogarithmic H dependency sample complexity bound $\tilde{O}(S^5 A^4 \varepsilon^{-2})$ by carefully reusing samples and avoid unnecessary sampling. Zhang et al. (2021b) further proposed an improved algorithm MVP to achieve near-optimal regret bound $\tilde{O}(\sqrt{SAK} + S^2 A)$ based on new Bernstein-type bonus. Similar polylogarithmic H dependency bounds had been established by Ren et al. (2021) for linear MDP with anchor points, Tarbouriech et al. (2021) for the stochastic shortest path. Li et al. (2022) achieved the surprising H independent sample complexity bound $O((SA)^{O(S)} \varepsilon^{-5})$ by building a connection between discounted MDPs and episodic MDPs and a novel perturbation analysis in MDPs. The algorithm proposed by Zhang et al. (2022) further improved the sample complexity to $O(S^9 A^3 \varepsilon^{-2} \text{polylog}(S, A, \varepsilon^{-1}))$ only depending on state and action spaces size polynomially by exploiting the power of stationary policy. Thanks to the linear function approximation, Zhou & Gu (2022) firstly achieve horizon-free regret bound $\tilde{O}(d\sqrt{K} + d^2)$ independent of state and ac-

tion spaces size. However, all the above works are limited to standard RL settings. In the paradigm of reward-free exploration, the only horizon-free result was achieved by Zhang et al. (2021a) with sample complexity bound of $\tilde{O}(S^2 A \varepsilon^{-2})$, where the polynomial dependency on S and A is still unacceptable when the states and actions spaces are large. Our algorithm HF-UCRL-RFE++ establishes the first horizon-free sample complexity bound independent of state and action spaces size in reward-free exploration.

3. Preliminaries

We consider an episodic finite horizon Markov Decision Process (MDP) $\mathcal{M} = (\mathcal{S}, \mathcal{A}, H, \{r_h\}_{h=1}^H, \mathbb{P})$, where $\mathcal{S}$ is the countable (and maybe infinite) state space, $\mathcal{A}$ is the action space, H is the length of the episode, $r_h : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$ is the deterministic reward function, and $\mathbb{P} : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$ is the time-homogeneous transition probability.

Based on the current state $s \in \mathcal{S}$ and the time step $h \in [H]$, a policy π chooses the action $a \in \mathcal{A}$ to be executed by the agent. Formally, we denote a policy as $\pi = \{\pi_h\}_{h=1}^H$, where $\pi_h : \mathcal{S} \rightarrow \mathcal{A}$ is a function which maps a state s to an action a . For any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and $h \in [H]$, we define the

action-value function $Q_h^\pi(s, a)$ and value function $V_h^\pi(s)$ as follows:

$$Q_h^\pi(s, a) = \mathbb{E} \left[\sum_{h'=h}^H r(s_{h'}, a_{h'}) \middle| s_h = s, a_h = a, \right. \\ \left. s_{h'} \sim \mathbb{P}(\cdot | s_{h'-1}, a_{h'-1}), a_{h'} = \pi_{h'}(s_{h'}) \right], \\ V_h^\pi(s) = Q_h^\pi(s, \pi_h(s)).$$

We define the optimal value function $V_h^*(\cdot)$ and optimal action-value function $Q_h^*(\cdot, \cdot)$ as $V_h^*(\cdot) = \sup_\pi V_h^\pi(\cdot)$ and $Q_h^*(\cdot, \cdot) = \sup_\pi Q_h^\pi(\cdot, \cdot)$, respectively. For any function $V : \mathcal{S} \rightarrow [0, 1]$, we introduce the following short-hands to denote the conditional expectation and variance of V :

$$[\mathbb{P}V](s, a) = \mathbb{E}_{s' \sim \mathbb{P}(\cdot | s, a)} V(s'), \\ [\mathbb{V}V](s, a) = [\mathbb{P}V^2](s, a) - [\mathbb{P}V](s, a)^2$$

For any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and $h \in [H]$, the Bellman equation and Bellman optimality equation can be defined recursively as

$$Q_h^\pi(s, a) = r(s, a) + [\mathbb{P}V_{h+1}^\pi](s, a), \\ Q_h^*(s, a) = r(s, a) + [\mathbb{P}V_{h+1}^*](s, a)$$

In this paper, we make the structural assumption that the transition dynamic has a linear structure, which has been considered in prior works as below:

Definition 3.1 (Linear Mixture MDPs, Jia et al. 2020; Ayoub et al. 2020; Zhou et al. 2021c). The unknown transition probability $\mathbb{P}$ is a linear combination of d signed basis measures $\phi_i(s'|s, a)$, i.e., $\mathbb{P}(s'|s, a) = \sum_{i=1}^d \phi_i(s'|s, a) \theta_i^*$. Meanwhile, for any $V : \mathcal{S} \rightarrow [0, 1]$, $i \in [d]$, $(s, a) \in \mathcal{S} \times \mathcal{A}$, the summation $\sum_{s' \in \mathcal{S}} \phi_i(s'|s, a) V(s')$ can be calculated efficiently within $\mathcal{O}$ time. For simplicity, let $\phi = [\phi_1, \dots, \phi_d]^\top$, $\theta^* = [\theta_1^*, \dots, \theta_d^*]^\top$ and $\phi_V(s, a) = \sum_{s' \in \mathcal{S}} \phi(s'|s, a) V(s')$. W.l.o.g., we assume $\|\theta^*\|_2 \leq B$, $\|\phi_V(s, a)\|_2 \leq 1$ for all $V : \mathcal{S} \rightarrow [0, 1]$ and $(s, a) \in \mathcal{S} \times \mathcal{A}$.

We further assume that the accumulated reward of an episode for any trajectory is upper bounded by 1, which ensures that the only factors affecting the final statistical complexity are difficulties brought by exploration and long planning horizon rather than the scale of the total reward.

Assumption 3.2. (Bounded total reward) For any trajectory $\{s_h, a_h\}_{h=1}^H$, we have $0 \leq \sum_{h=1}^H r_h(s_h, a_h) \leq 1$.

We denote the set of reward functions satisfying Assumption 3.2 by $\mathcal{R}$.

Reward-free RL In reward-free RL, the real reward function is accessible only after the agent finishes the interactions with the environment. Specifically, the algorithm can be separated into two phases: (i) *Exploration phase*: the algorithm can't access the reward function but collect K episodes of samples by interacting with the environment. (ii) *Planning phase*: The algorithm is given reward function $\{r_h\}_{h=1}^H$ and is expected to find the optimal policy. (ε, δ) -learn and sample complexity of the algorithm is formally defined as follows.

Definition 3.3. (ε, δ) -learnability). Given an MDP transition kernel set $\mathcal{P}$, reward function set $\mathcal{R}$ and a initial state distribution μ , we say a reward-free algorithm can (ε, δ) -learn the problem $(\mathcal{P}, \mathcal{R})$ with sample complexity $K(\varepsilon, \delta)$, if for any transition kernel $P \in \mathcal{P}$, after receiving $K(\varepsilon, \delta)$ episodes in the exploration phase, for any reward function $r \in \mathcal{R}$, the algorithm returns a policy π in planning phase, such that with probability at least $1 - \delta$, $V_1^*(s_1; r) - V_1^\pi(s_1; r) \leq \varepsilon$.

4. Algorithms

In this section, we propose our reward-free exploration algorithm HF-UCRL-RFE++. This algorithm consists of two phases. In the exploration phase, it builds an estimator θ for the linear mixture MDP transition kernel parameter θ^* based on the sampled episodes. At a high level, the estimation follows the *value-targeted regression* (VTR) framework proposed by Jia et al. (2020). The VTR is basically a ridge regression with value functions as responses and feature mappings as predictors. However, value functions have no estimates since the reward function is not accessible. Therefore, the value functions and reward functions are replaced by well-designed exploration-driven pseudo-value functions and pseudo-reward functions. To achieve a better estimation, we further apply the *high-order moment estimation* (HOME) technique proposed by Zhou & Gu (2022). Then, during the planning phase, the algorithm uses the estimator θ acquired in the exploration phase to find the optimal policy π for the given reward functions. Our algorithm is described in Algorithm 1.

Exploration-driven Pseudo Value Function As mentioned above, in the paradigm of reward-free exploration, we have to construct the pseudo-reward function to guide the agent in taking actions in the absence of the real reward function. As we adopt in this work, the most natural idea is to construct the pseudo-reward function related to uncertainty, which urges the agent to collect information about the most uncertain states and actions. Two approaches follow this idea: one is constructing the pseudo reward function directly measuring and maximizing the uncertainty of each stage, and the other is constructing the pseudo reward function maximizing the overall uncertainty along trajectories.

Algorithm 1 HF-UCRL-RFE++

Input: Confidence radius $\{\beta_k\}$, regularization parameter λ , number of the high-order estimator M

- 1: **Phase I: Exploration Phase**
- 2: Initialize $\hat{\Sigma}_{1,1,m} \leftarrow \lambda \mathbf{I}$, $\tilde{\Sigma}_{1,1,m} \leftarrow \lambda \mathbf{I}$, $\hat{\mathbf{b}}_{1,1,m} \leftarrow \mathbf{0}$, $\tilde{\mathbf{b}}_{1,1,m} \leftarrow \mathbf{0}$ for all $m \in \overline{[M]}$, $\mathcal{U}_1 = \{\theta | \theta \in \mathbb{R}^d\}$.
- 3: Set $\hat{\theta}_{1,m} \leftarrow \hat{\Sigma}_{1,1,m}^{-1} \hat{\mathbf{b}}_{1,1,m}$, $\tilde{\theta}_{1,m} \leftarrow \tilde{\Sigma}_{1,1,m}^{-1} \tilde{\mathbf{b}}_{1,1,m}$ for all $m \in \overline{[M]}$.
- 4: **for** $k = 1, 2, \dots, K$ **do**
- 5: Calculate $\pi_k, \theta_k, r_k = \operatorname{argmax}_{\pi, \theta \in \mathcal{U}_k, r \in R} \hat{V}_{k,1}(s_1; \theta, \pi, r)$, where $\hat{V}_{k,1}$ is defined in (4.2). Denote $\{\tilde{V}_{k,h}(\cdot)\}_{h=1}^H = \{V_h(\cdot; \theta_k, \pi_k, r_k)\}_{h=1}^H$.
- 6: Receive initial state $s_1^k = s_1$.
- 7: **for** $h = 1, 2, \dots, H$ **do**
- 8: Execute $a_h^k = \pi_k^k(s_h^k)$, receive $s_{h+1}^k \sim \mathbb{P}(\cdot | s_h^k, a_h^k)$.
- 9: For $m \in \overline{[M]}$, denote $\hat{\phi}_{k,h,m} = \phi_{\hat{V}_{k,h+1}^{2m}}(s_h^k, a_h^k)$, $\tilde{\phi}_{k,h,m} = \phi_{\tilde{V}_{k,h+1}^{2m}}(s_h^k, a_h^k)$
- 10: Set $\{\hat{\sigma}_{k,h,m}\} \leftarrow \text{HOME}_{\text{Alg. 2}}(\{\hat{\phi}_{k,h,m}, \hat{\theta}_{k,m}, \hat{\Sigma}_{k,h,m}, \hat{\dot{\Sigma}}_{k,m}\}, \beta_k, \alpha, \gamma)$
- 11: Set $\{\tilde{\sigma}_{k,h,m}\} \leftarrow \text{HOME}_{\text{Alg. 2}}(\{\tilde{\phi}_{k,h,m}, \tilde{\theta}_{k,m}, \tilde{\Sigma}_{k,h,m}, \tilde{\dot{\Sigma}}_{k,m}\}, \beta_k, \alpha, \gamma)$.
- 12: Set $\tilde{\Sigma}_{k,h+1,m} \leftarrow \tilde{\Sigma}_{k,h,m} + \tilde{\phi}_{k,h,m} \tilde{\phi}_{k,h,m}^\top \tilde{\sigma}_{k,h,m}^{-2}$ for $m \in \overline{[M]}$
- 13: Set $\hat{\Sigma}_{k,h+1,m} \leftarrow \hat{\Sigma}_{k,h,m} + \hat{\phi}_{k,h,m} \hat{\phi}_{k,h,m}^\top \hat{\sigma}_{k,h,m}^{-2}$ for $m \in \overline{[M]}$
- 14: Set $\tilde{\mathbf{b}}_{k,h+1,m} \leftarrow \tilde{\mathbf{b}}_{k,h,m} + \tilde{\phi}_{k,h,m} \tilde{V}_{k,h+1}^{2m}(s_{h+1}^k) \tilde{\sigma}_{k,h,m}^{-2}$ for $m \in \overline{[M]}$
- 15: Set $\hat{\mathbf{b}}_{k,h+1,m} \leftarrow \hat{\mathbf{b}}_{k,h,m} + \hat{\phi}_{k,h,m} \hat{V}_{k,h+1}^{2m}(s_{h+1}^k) \hat{\sigma}_{k,h,m}^{-2}$ for $m \in \overline{[M]}$
- 16: **end for**
- 17: $\tilde{\Sigma}_{k+1,m} \leftarrow \tilde{\Sigma}_{k,H+1,m}$, $\hat{\Sigma}_{k+1,m} \leftarrow \hat{\Sigma}_{k,H+1,m}$
- 18: Set $\tilde{\Sigma}_{k+1,1,m} \leftarrow \tilde{\Sigma}_{k,H+1,m}$, $\tilde{\mathbf{b}}_{k+1,1,m} \leftarrow \tilde{\mathbf{b}}_{k,H+1,m}$, $\tilde{\theta}_{k+1,m} = \tilde{\Sigma}_{k+1,1,m}^{-1} \tilde{\mathbf{b}}_{k+1,1,m}$.
- 19: Set $\hat{\Sigma}_{k+1,1,m} \leftarrow \hat{\Sigma}_{k,H+1,m}$, $\hat{\mathbf{b}}_{k+1,1,m} \leftarrow \hat{\mathbf{b}}_{k,H+1,m}$, $\hat{\theta}_{k+1,m} = \hat{\Sigma}_{k+1,1,m}^{-1} \hat{\mathbf{b}}_{k+1,1,m}$.
- 20: Update the confidence set $\mathcal{U}_k$ to $\mathcal{U}_{k+1}$ by adding constraints (4.5), (4.6).
- 21: **end for**
- 22: **Phase II: Planning Phase**
- 23: Receive reward function r .
- 24: $\hat{\pi}_r = \operatorname{argmax}_{\pi} V_1(\cdot; \theta_K, \pi, r)$.
- 25: Return policy $\hat{\pi}_r$.

Zhang et al. (2021d) took the first approach, constructing the pseudo-reward function in the form of

$$r_h^k(s, a) = \min \left\{ 1, \frac{2\beta}{H} \sqrt{\max_{V \in \mathcal{S} \rightarrow [0, H-h]} \|\phi_V(s, a)\|_{\Sigma_{1,k}^{-1}}} \right\},$$

and the pseudo-value function to be the argument of the maxima for the above uncertainty measure. Under this construction, the suboptimality in the planning phase can be bounded by the accumulation of uncertainty. This approach is straightforward but has the following two drawbacks. Firstly, without the truncation for accumulation of uncertainty, the upper bound of overall suboptimality in the planning phase will be in the scale of $O(H)$, which is meaningless since the value function lies in the interval of $[0, 1]$ under our assumption. Second, since VTR utilizes value functions' variance information for θ estimation, it requires a Bellman-equation-type equality between two consecutive stages h and $h+1$. However, the first approach does not satisfy this requirement, preventing us from acquiring a more accurate estimate.

To address the above issues, we follow the design of pseudo value function proposed in Chen et al. (2021). In particular, we are constructing the pseudo-reward function aiming to maximize the overall uncertainty along trajectories. We view the uncertainty of states and actions as a function of (pseudo) reward function r , policy π , and transition kernel parameter θ defined as follows

$$u_{k,h}(s, a; \theta, \pi, r) = \min \left\{ 1, \beta \left\| \phi_{V_h(\cdot; \theta, \pi, r)}(s, a) \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right\}, \quad (4.1)$$

where $V_h(\cdot; \theta, \pi, r)$ is the the value function of policy π for linear mixture MDP with transition kernel parameter θ and the reward function r , and the overall uncertainty along the trajectory is the truncated sum of each step uncertainty defined as

$$\bar{V}_{k,h}(s; \theta, \pi, r) = \min \left\{ 1, u_{k,h}(s, \pi(s); \theta, \pi, r) + \phi_{\tilde{V}_{k,h+1}^\top(\cdot; \theta, \pi, r)}(s, \pi(s)) \theta^* \right\}.$$

However, the definition of $\bar{V}_{k,h}(s; \theta, \pi, r)$ involves θ^* , which is unknown to the agent. Hence, we construct the optimistic estimation of $\bar{V}_{k,h}(s; \theta, \pi, r)$ as $\hat{V}_{k,h}(s; \theta, \pi, r)$ defined as

$$\begin{aligned} \hat{V}_{k,h}(s; \theta, \pi, r) = & \min \left\{ 1, u_{k,h}(s, \pi(s); \theta, \pi, r) \right. \\ & + 2\beta \left\| \phi_{\hat{V}_{k,h+1}(\cdot; \theta, \pi, r)}(s, \pi(s)) \right\|_{\hat{\Sigma}_{k,0}^{-1}} \\ & \left. + \phi_{\hat{V}_{k,h+1}(\cdot; \theta, \pi, r)}^\top(s, \pi(s)) \theta \right\}. \end{aligned} \quad (4.2)$$

Notable, the definitions of $u_{k,h}$ and $\hat{V}_{k,h}$ involve the covariance matrices $\hat{\Sigma}_{k,0}$ and $\hat{\Sigma}_{k,0}$, which are computed at the end of the preceding episode at Line 17 of Algorithm 1. In the following content, when there is no confusion, we may write $\hat{V}_{k,h}(\cdot) = \hat{V}_{k,h}(\cdot; \theta_k, \pi_k, r_k)$, $u_{k,h}(\cdot, \cdot) = u_{k,h}(\cdot, \cdot; \theta_k, \pi_k, r_k)$. In order to collect more information, the agent is expected to transit through the trajectory with the largest uncertainty $\hat{V}_{k,h}$. It is notable that $\hat{V}_{k,h}$ is a function of (pseudo) reward function r_k , policy π_k , and transition kernel parameter θ_k . Thus, at the beginning of each episode, we set r_k , π_k , and θ_k to be arguments of the maxima, as presented in Line 5 in Algorithm 1. Through this process, we acquire the pseudo value function r_k , which is essential for reward-free exploration. Afterward, the algorithm collects samples along trajectories induced by policy π_k and improves the estimation of θ_k in Line 6 to Line 21. In this stage, Algorithm 1 encounters two series of functions in the form of Bellman equations; one is the sum of pseudo rewards r , $\tilde{V}_{k,h}(\cdot) = V_h(\cdot; \theta_k, \pi_k, r_k)$, which we refer as pseudo value function, and one is the uncertainty along the trajectory, $\hat{V}_{k,h}$. These two series of functions are both eligible for refined VTR and thus help estimate θ , as we will explain in the following.

High-order Moment Estimation The key technique used in our algorithm consists of two series of high-order estimations for the transition kernel parameter θ . The algorithm for high-order moment estimation is stated in Algorithm 2. In the exploration phase, the agent learns the environment with the help of two series of value functions $\tilde{V}_{k,h}$ and $\hat{V}_{k,h}$. They serve to characterize different aspects of the model, one for pseudo values and one for trajectory uncertainty. And thus, they rely on different estimations of transition kernel parameter θ . Two independent series of higher-order moment estimations are necessary for achieving accurate estimation. In the Algorithm 1, both estimations of θ are the solutions to the weighted regression problem in the following form:

$$\begin{aligned} \operatorname{argmin}_{\theta} \left(\lambda \|\theta\|_2^2 \right. \\ \left. + \sum_{j=1}^{k-1} \sum_{h=1}^H (\phi_{j,h,0}^\top \theta - V_{j,h}(s_{h+1}^j))^2 / \bar{\sigma}_{j,h,0}^2 \right), \end{aligned} \quad (4.3)$$

where the regression weight $\bar{\sigma}_{j,h,0}$ is set as Equation (4.4).

$$\begin{aligned} \bar{\sigma}_{k,h,0}^2 \leftarrow & \max \left\{ \gamma^2 \|\phi_{k,h,0}\|_{\Sigma_{k,h,0}^{-1}}^2, \right. \\ & \left. [\bar{\mathbb{V}}_{k,0} V_{k,h+1}](s_h^k, a_h^k) + E_{k,h,0}, \alpha^2 \right\}. \end{aligned} \quad (4.4)$$

$\bar{\sigma}_{j,h,0}$ can be considered as an combination of *aleatoric uncertainty* and *epistemic uncertainty* (Kendall & Gal, 2017; Mai et al., 2022). The first term $\gamma^2 \|\phi_{k,h,0}\|_{\Sigma_{k,h,0}^{-1}}^2$ in (4.4) is the *epistemic uncertainty* caused by limited available data. And the second term in Equation (4.4) is supposed to be the *aleatoric uncertainty* $\mathbb{V}_{k,0} V_{k,h+1}$ characterizing the inherent non-determinism of the transition kernel, which is irreducible. Here the $\mathbb{V}_{k,m} V_{k,h+1}$ is the variance of $V_{k,h+1}$ to 2^m defined as $[\mathbb{P}V_{k,h+1}^{2^{m+1}}](s_h^k, a_h^k) - [\mathbb{P}V_{k,h+1}^{2^m}](s_h^k, a_h^k)^2$. Then, $[\mathbb{V}_{k,0} V_{k,h+1}](s_h^k, a_h^k)$ is further replaced with its estimate $[\bar{\mathbb{V}}_{k,0} V_{k,h+1}](s_h^k, a_h^k)$ plus its error bound $E_{k,h,0}$ since real variance $[\mathbb{V}_{k,0} V_{k,h+1}](s_h^k, a_h^k)$ is unknown to the agent. Because $[\mathbb{V}_{k,0} V_{k,h+1}](s_h^k, a_h^k)$ is a quadratic function of the real transition kernel parameter θ^* , its estimate can be achieved as

$$\begin{aligned} [\bar{\mathbb{V}}_{k,0} V_{k,h+1}](s_h^k, a_h^k) = & \left[\left\langle \phi_{k,h,1}, \theta_{k,1} \right\rangle \right]_{[0,1]} \\ & - \left[\left\langle \hat{\phi}_{k,h,0}, \theta_{k,0} \right\rangle \right]_{[0,1]}^2, \end{aligned}$$

where $\theta_{k,1}$ is again the solution to the weighted regression problem similar to (4.4) with predictors $\phi_{k,h,1} = \phi_{V_{k,h+1}^2}(s_h^k, a_h^k)$, responses $V_{k,h+1}^2(s_{h+1}^k)$ and weight $\bar{\sigma}_{k,h,1}$. Following the above idea, the value of weight $\bar{\sigma}_{k,h,1}$ further relies on $\theta_{k,2}$, which is the solution to a weighted regression problem involving another weight $\bar{\sigma}_{k,h,2}$. The algorithm carried out this process recursively until $\bar{\sigma}_{k,h,M-1}$, where its second term is replaced by the trivial upper bound of aleatoric uncertainty.

Applying HOME to the reward-free setting brings additional difficulties in controlling the error of our estimate for the model, as the error introduced by using the pseudo reward function instead of the real reward function and the error introduced by estimating the true transition kernel must be controlled separately. To address this problem, we carefully estimate variables indicating different kinds of error into two series of HOME in Line 10 and Line 11. Since the separation of variables deeply exploits the inner structure of the problem, the two series of HOME can be merged in the end to achieve a unified control for both kinds of error.

Previous work Chen et al. (2021) implemented the weighted value regression in a more crude way. The weights are

Algorithm 2 High-order moment estimator (HOME)

Input: Features $\{\phi_{k,h,m}\}_{m \in [M]}$, vector estimators $\{\theta_{k,m}\}_{m \in [M]}$, covariance matrix $\{\Sigma_{k,h,m}, \dot{\Sigma}_{k,m}\}_{m \in [M]}$, confidence radius β_k, α, γ

- 1: **for** $m = 0, \dots, M - 2$ **do**
- 2: Set $[\bar{\nabla}_{k,m} V_{k,h+1}^{2m}](s_h^k, a_h^k) \leftarrow [\phi_{k,h,m+1}^\top \theta_{k,m+1}]_{[0,1]} - [\phi_{k,h,m}^\top \theta_{k,m}]_{[0,1]}^2$
- 3: Set $E_{k,h,m} \leftarrow [2\beta_k \|\phi_{k,h,m}\| \dot{\Sigma}_{k,m}^{-1}]_{[0,1]} + [\beta_k \|\phi_{k,h,m+1}\| \dot{\Sigma}_{k,m+1}^{-1}]_{[0,1]}$
- 4: Set $\bar{\sigma}_{k,h,m}^2 \leftarrow \max \left\{ \gamma^2 \|\phi_{k,h,m}\|_{\Sigma_{k,h,m}^{-1}}, [\bar{\nabla}_{k,m} V_{k,h+1}^{2m}](s_h^k, a_h^k) + E_{k,h,m}, \alpha^2 \right\}$
- 5: **end for**
- 6: Set $\bar{\sigma}_{k,h,M-1}^2 \leftarrow \max \left\{ \gamma^2 \|\phi_{k,h,M-1}\|_{\Sigma_{k,h,M-1}^{-1}}, 1, \alpha^2 \right\}$

Output: $\{\bar{\sigma}_{k,h,m}\}_{m \in [M]}$

constructed only on aleatoric uncertainty, totally ignoring epistemic uncertainty. In addition, they use the same instead of different transition kernel parameters to calculate different order moments of the value function and stop target value regression at second order moment, which increased avoidable error. As a result, Chen et al. (2021) can only replace factor Hd with factor $H + d$ when trying to improve the dependency on d in the upper bound. In contrast, our work further improves factor $H + d$ to factor H through the well-designed target value regression, as we can see in Corollary 5.4.

High Confidence Set At the end of each episode, we add the following constraints into $\mathcal{U}_k$ to update the high confidence set in Line 20 of Algorithm 1.

$$\|\theta - \hat{\theta}_{k,m}\|_{\dot{\Sigma}_{k,m}} \leq \beta_k, m \in [M], \quad (4.5)$$

$$\|\theta - \tilde{\theta}_{k,m}\|_{\dot{\Sigma}_{k,m}} \leq \beta_k, m \in [M], \quad (4.6)$$

High confidence set $\mathcal{U}_k$ ensures that the estimate θ_k lies in a neighborhood of real transition kernel parameter θ^* . Here the algorithm adds $2M$ inequalities to constraints in each episode. These inequalities guarantee that estimations of the variance of $\hat{V}_{k,h}$ and $\tilde{V}_{k,h}$ up to M -th order are near the real values.

Planning Phase After finishing the exploration, the agent enters the planning phase and receives the real reward function. Depending on optimal Bellman equations, the agent is able to obtain the optimal policy backward from state H to state 1 by dynamic programming based on real reward function r and transition kernel parameter estimate θ_K . And then, the algorithm outputs the optimal policy.

Computational Complexity of HF-UCRL-RFE++ Similar with Chen et al. (2021), we assume that the optimization over θ , π , and r in Line 5 of Algorithm 1

can be accomplished with an oracle which is obvious to be called for K times. At each episode k and each stage h , **HF-UCRL-RFE++** computes $\{\hat{\phi}_{k,h,m}\}_{m \in [M]}$, $\{\tilde{\phi}_{k,h,m}\}_{m \in [M]}$, $\{\hat{\sigma}_{k,h,m}\}_{m \in [M]}$, $\{\tilde{\sigma}_{k,h,m}\}_{m \in [M]}$, and updates $\{\hat{\Sigma}_{k,h,m}\}_{m \in [M]}$, $\{\tilde{\Sigma}_{k,h,m}\}_{m \in [M]}$. The computation of $\{\hat{\phi}_{k,h,m}\}_{m \in [M]}$ and $\{\tilde{\phi}_{k,h,m}\}_{m \in [M]}$ require $O(\mathcal{O}M)$ times. According to Algorithm 2, calculating $\{\hat{\sigma}_{k,h,m}\}_{m \in [M]}$ and $\{\tilde{\sigma}_{k,h,m}\}_{m \in [M]}$ require $O(Md^2)$ time since the computation of the inner-product an inversion of matrix and a vector needs $O(d^2)$. The updates of $\{\hat{\Sigma}_{k,h,m}\}_{m \in [M]}$ and $\{\tilde{\Sigma}_{k,h,m}\}_{m \in [M]}$ further require $O(Md^2)$ time. Lastly, determining the optimal policy during the planning phase takes $O(H(SAd + \mathcal{O}))$ time. Therefore, the total time complexity of **HF-UCRL-RFE++** is $O(KH(\mathcal{O}M + Md^2) + HSA d)$.

5. Main Results

5.1. Upper Bounds

We first provide the suboptimality upper bound of our algorithm HF-UCRL-RFE++.

Theorem 5.1. For Algorithm 1, set $M = \log(7KH)/\log(2)$, $\alpha = H^{-1/2}$, $\gamma = d^{-1/4}$, $\lambda = d/B^2$, $\{\beta_k\}_{k \geq 1}$ as

$$\beta_k = 12\sqrt{d\eta\tau} + 30\tau/\gamma^2 + \sqrt{\lambda}B,$$

and denote $\beta = \beta_K$, where $\eta = \log(1 + kH/(\alpha^2 d \lambda))$ and $\tau = \log(32(\log(\gamma^2/\alpha) + 1)k^2 H^2/\delta)$. Then, for any $0 < \delta < 1$, we have with probability at least $1 - \delta$, after collecting K episodes of samples, algorithm 1 returns a policy $\hat{\pi}_r$ satisfying the following sub-optimality bound,

$$V_1^*(s_1; r) - V_1(s_1; \theta^*, \hat{\pi}_r, r) = \tilde{O}\left(\frac{d^2}{K} + \frac{d}{\sqrt{K}}\right).$$

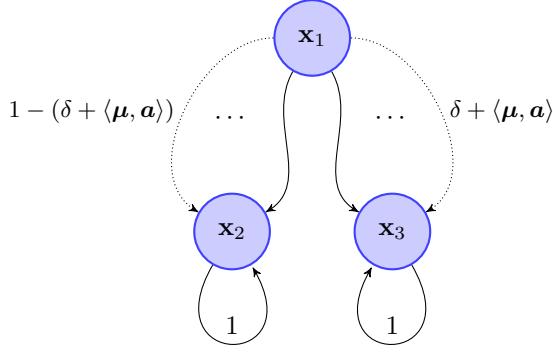


Figure 1. The transition kernel of the hard-to-learn linear mixture MDPs.

The next corollary specifies the sample complexity of our algorithm.

Corollary 5.2. Under the same conditions as in Theorem 5.1, Algorithm 1 has sample complexity of $m(\varepsilon, \delta) = \tilde{O}(d^2\varepsilon^{-2})$ to output an ε -optimal policy in the planning phase. The exact expression of sample complexity is delayed to Appendix in Lemma A.11.

Remark 5.3. To the best of our knowledge, Corollary 5.2 provides the first horizon-free sample complexity upper bound independent of state space size S and action space size A for reward-free exploration. This result shows that long-horizon planning does not add extra difficulty to reward-free exploration.

Corollary 5.4. When rescaling the assumption $\sum_{h=1}^H r_h(s_h, a_h) \leq 1$ to $\sum_{h=1}^H r_h(s_h, a_h) \leq H$, under the same conditions as in Theorem 5.1, Algorithm 1 has sample complexity of $m(\varepsilon, \delta) = O(H^2 d^2 \varepsilon^{-2})$ to output an ε -optimal policy in the planning phase.

Remark 5.5. The assumption $\sum_{h=1}^H r_h(s_h, a_h) \leq H$ covers the vanilla assumption $r_h(s_h, a_h) \in [0, 1]$. Therefore, compared with Chen et al. (2021), our analysis does not require the $d > H$ assumption and achieves the same sample complexity bound up to logarithmic factors except for the trivial $\tilde{O}(H)$ difference between time-homogeneous and time-inhomogeneous models with a milder assumption. This improvement can be attributed to the refined value target regression technique, high-order moment estimation (HOME), adopted in our approach. We provide a detailed analysis of this improvement in the “High-order Moment Estimation” part in the Section 4.

5.2. Lower Bounds

The following results provide lower bounds of the sample complexity and suggest that our algorithm is minimax optimal. We will consider the *hard-to-learn linear*

mixture MDPs constructed in Zhou & Gu (2022). The state space is $\mathcal{S} = \{x_1, x_2, x_3\}$ and the action space is $\mathcal{A} = \{\mathbf{a}\} = \{-1, 1\}^{d-1}$. The reward function satisfies $r(x_1, \cdot) = r(x_2, \cdot) = 0$, and $r(x_3, \cdot) = \frac{1}{H}$. The transition probability is defined to be $\mathbb{P}(x_2 | x_1, \mathbf{a}) = 1 - (\delta + \langle \boldsymbol{\mu}, \mathbf{a} \rangle)$ and $\mathbb{P}(x_3 | x_1, \mathbf{a}) = \delta + \langle \boldsymbol{\mu}, \mathbf{a} \rangle$, where $\delta = 1/6$ and $\boldsymbol{\mu} \in \{-\Delta, \Delta\}^{d-1}$ with $\Delta = \sqrt{\delta/K}/(4\sqrt{2})$.

Theorem 5.6. Suppose $B > 1$. Then for any algorithm ALG_{Free} solving reward-free linear mixture MDP problems satisfying assumption 3.2, there exist a linear mixture MDP $\mathcal{M}$ such that ALG_{Free} needs to collect at least $\Omega(d^2\varepsilon^{-2})$ episodes of samples to output an ε -optimal policy with probability at least $1 - \delta$. This lower bound matches the sample complexity upper bound provided in Corollary 5.2, which shows our upper bound is optimal.

Remark 5.7. The lower bound is similar to the lower bound provided in Chen et al. (2021). The first difference is that we rescale the non-zero reward in hard-to-learn cases from 1 to $\frac{1}{H}$ in order to satisfy Assumption 3.2. The second difference is that we consider the time-homogeneous model instead of the time-inhomogeneous one in theirs. By these changes, our lower bound for reward-free exploration provided in Theorem 5.6 removes the unnecessary polynomial dependency on episode length H introduced by the scale of total reward.

Corollary 5.8. Under the same conditions as Theorem 5.6 except replacing $\sum_{h=1}^H r_h(s_h, a_h) \leq 1$ with $r_h \in [0, 1]$, for any algorithm ALG_{Free} solving reward-free linear mixture MDP problems satisfying assumption 3.2, there exist a linear mixture MDP $\mathcal{M}$ such that ALG_{Free} needs to collect at least $\tilde{\Omega}(H^2 d^2 \varepsilon^{-2})$ episodes to output an ε -optimal policy with probability at least $1 - \delta$. This means our upper bound is optimal.

6. Proof Sketch of Theorem 5.1

We provide the proof sketch of Theorem 5.1 along with several key lemmas in big-O notation. The detail of these lemmas is restated in Appendix A. The following lemmas are conditioned on some good events.

Firstly, Lemma 6.1 controls the suboptimality gap between optimal value functions and our estimated value function in the planning phase with the uncertainty along trajectories.

Lemma 6.1. For any reward function r in the planning phase, the suboptimality gap of the outputted policy $\hat{\pi}_r$ can be bounded as

$$V_1^*(s_1; r) - V_1(s_1; \boldsymbol{\theta}^*, \hat{\pi}_r, r) \leq 4\hat{V}_{K,1}(s_1). \quad (6.1)$$

Then, the next lemma shows that the uncertainty along trajectories decreases with respect to episodes. This lemma is

intuitively right since the uncertainty should decrease with more information collected.

Lemma 6.2. For uncertainty along trajectories, we have

$$\widehat{V}_{K,1}(s_1; \theta_K, \pi_K, r_K) \leq \frac{1}{K} \left(\sum_{k=1}^K \widehat{V}_{k,1}(s_1; \theta_k, \pi_k, r_k) \right).$$

The last lemma upper bounds the sum of the uncertainty along trajectories.

Lemma 6.3. For any $0 < \delta < 1$, with probability at least $1 - 4M\delta$, we have

$$\sum_{k=1}^K \widehat{V}_{k,1}(s_1; \theta_k, \widehat{\pi}_k, r_k) = \widetilde{O}(d\sqrt{K} + d^2). \quad (6.2)$$

Equipped with the above lemmas, we are ready to prove Theorem 5.1.

Proof of Theorem 5.1. The suboptimality of the policy output in the planning phase can be bounded by the uncertainty along trajectories in the last episode of the exploration phase as the following equation according to Lemma 6.1.

$$V_1^*(s_1; r) - V_1(s_1; \theta^*, \widehat{\pi}_r, r) \leq 4\widehat{V}_{K,1}(s_1) \quad (6.3)$$

Since Lemma 6.2 indicates that the uncertainty is decreasing with the episodes, the uncertainty of the last episode can be further upper bounded by the sum of uncertainty in each episode by substituting (6.1) in Lemma 6.1 into the above inequality:

$$V_1^*(s_1; r) - V_1(s_1; \theta^*, \widehat{\pi}_r, r) \leq \frac{4}{K} \sum_{k=1}^K \widehat{V}_{k,1}(s_1). \quad (6.4)$$

At last, the sum of value functions can be upper bounded according to Lemma 6.3. Thus, plugging (6.2) into the (6.4) as follows concludes our proof.

$$V_1^*(s_1; r) - V_1(s_1; \theta^*, \widehat{\pi}_r, r) = \widetilde{O} \left(\frac{d^2}{K} + \frac{d}{\sqrt{K}} \right).$$

□

7. Conclusion

We study model-based reward-free exploration for learning the linear mixture MDPs. Our algorithm is guaranteed to have horizon-free sample complexity in the exploration phase to find a near-optimal policy in the planning phase for any given reward function. To our knowledge, these are the first horizon-free result for reward-free RL with function

approximation. We also give a sample complexity lower bound for reward-free exploration in linear mixture MDPs under our assumptions. The sample complexity upper bound of our algorithm matches the lower bound up to logarithmic factors.

Acknowledgements

We thank the anonymous reviewers for their helpful comments. WZ, JZ and QG are supported in part by the National Science Foundation CAREER Award 1906169 and research fund from UCLA-Amazon Science Hub. The views and conclusions contained in this paper are those of the authors and should not be interpreted as representing any funding agencies.

References

- Achiam, J., Held, D., Tamar, A., and Abbeel, P. Constrained policy optimization. In *International conference on machine learning*, pp. 22–31. PMLR, 2017.
- Agarwal, A., Jin, Y., and Zhang, T. Vo q 1: Towards optimal regret in model-free rl with nonlinear function approximation. *arXiv preprint arXiv:2212.06069*, 2022.
- Ayoub, A., Jia, Z., Szepesvari, C., Wang, M., and Yang, L. Model-based reinforcement learning with value-targeted regression. In *International Conference on Machine Learning*, pp. 463–474. PMLR, 2020.
- Cai, Q., Yang, Z., Jin, C., and Wang, Z. Provably efficient exploration in policy optimization. In *International Conference on Machine Learning*, pp. 1283–1294. PMLR, 2020.
- Chen, X., Hu, J., Yang, L., and Wang, L. Near-optimal reward-free exploration for linear mixture mdps with plug-in solver. In *International Conference on Learning Representations*, 2021.
- Dann, C., Jiang, N., Krishnamurthy, A., Agarwal, A., Langford, J., and Schapire, R. E. On oracle-efficient pac rl with rich observations. *Advances in neural information processing systems*, 31, 2018.
- Du, S. S., Kakade, S. M., Wang, R., and Yang, L. F. Is a good representation sufficient for sample efficient reinforcement learning? *arXiv preprint arXiv:1910.03016*, 2019.
- He, J., Zhou, D., and Gu, Q. Logarithmic regret for reinforcement learning with linear function approximation. In *International Conference on Machine Learning*, pp. 4171–4180. PMLR, 2021.

- He, J., Zhao, H., Zhou, D., and Gu, Q. Nearly minimax optimal reinforcement learning for linear markov decision processes. *arXiv preprint arXiv:2212.06132*, 2022.
- Huang, J., Chen, J., Zhao, L., Qin, T., Jiang, N., and Liu, T.-Y. Towards deployment-efficient reinforcement learning: Lower bound and optimality. *arXiv preprint arXiv:2202.06450*, 2022.
- Jia, Z., Yang, L., Szepesvari, C., and Wang, M. Model-based reinforcement learning with value-targeted regression. In *Learning for Dynamics and Control*, pp. 666–686. PMLR, 2020.
- Jiang, N., Krishnamurthy, A., Agarwal, A., Langford, J., and Schapire, R. E. Contextual decision processes with low bellman rank are pac-learnable. In *International Conference on Machine Learning*, pp. 1704–1713. PMLR, 2017.
- Jin, C., Krishnamurthy, A., Simchowitz, M., and Yu, T. Reward-free exploration for reinforcement learning. In *International Conference on Machine Learning*, pp. 4870–4879. PMLR, 2020a.
- Jin, C., Yang, Z., Wang, Z., and Jordan, M. I. Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pp. 2137–2143, 2020b.
- Kaufmann, E., Ménard, P., Domingues, O. D., Jonsson, A., Leurent, E., and Valko, M. Adaptive reward-free exploration. In *Algorithmic Learning Theory*, pp. 865–891. PMLR, 2021a.
- Kaufmann, E., Ménard, P., Domingues, O. D., Jonsson, A., Leurent, E., and Valko, M. Adaptive reward-free exploration. In *Algorithmic Learning Theory*, pp. 865–891. PMLR, 2021b.
- Kendall, A. and Gal, Y. What uncertainties do we need in bayesian deep learning for computer vision? *Advances in neural information processing systems*, 30, 2017.
- Li, Y., Wang, R., and Yang, L. F. Settling the horizon-dependence of sample complexity in reinforcement learning. In *2021 IEEE 62nd Annual Symposium on Foundations of Computer Science (FOCS)*, pp. 965–976. IEEE, 2022.
- Mai, V., Mani, K., and Paull, L. Sample efficient deep reinforcement learning via uncertainty estimation. *arXiv preprint arXiv:2201.01666*, 2022.
- Ménard, P., Domingues, O. D., Jonsson, A., Kaufmann, E., Leurent, E., and Valko, M. Fast active learning for pure exploration in reinforcement learning. In *International Conference on Machine Learning*, pp. 7599–7608. PMLR, 2021.
- Miryoosefi, S., Brantley, K., Daume III, H., Dudik, M., and Schapire, R. E. Reinforcement learning with convex constraints. *Advances in Neural Information Processing Systems*, 32, 2019.
- Modi, A., Jiang, N., Tewari, A., and Singh, S. Sample complexity of reinforcement learning using linearly combined model ensembles. In *International Conference on Artificial Intelligence and Statistics*, pp. 2010–2020, 2020.
- Ren, T., Li, J., Dai, B., Du, S. S., and Sanghavi, S. Nearly horizon-free offline reinforcement learning. *Advances in neural information processing systems*, 34:15621–15634, 2021.
- Sun, W., Jiang, N., Krishnamurthy, A., Agarwal, A., and Langford, J. Model-based rl in contextual decision processes: Pac bounds and exponential improvements over model-free approaches. In *Conference on Learning Theory*, pp. 2898–2933. PMLR, 2019.
- Tarbouriech, J., Zhou, R., Du, S. S., Pirota, M., Valko, M., and Lazaric, A. Stochastic shortest path: Minimax, parameter-free and towards horizon-free regret. *Advances in Neural Information Processing Systems*, 34: 6843–6855, 2021.
- Tessler, C., Mankowitz, D. J., and Mannor, S. Reward constrained policy optimization. *arXiv preprint arXiv:1805.11074*, 2018.
- Wagenmaker, A. J., Chen, Y., Simchowitz, M., Du, S., and Jamieson, K. Reward-free rl is no harder than reward-aware rl in linear markov decision processes. In *International Conference on Machine Learning*, pp. 22430–22456. PMLR, 2022.
- Wang, R., Du, S. S., Yang, L. F., and Kakade, S. M. Is long horizon reinforcement learning more difficult than short horizon reinforcement learning? *arXiv preprint arXiv:2005.00527*, 2020a.
- Wang, R., Du, S. S., Yang, L. F., and Salakhutdinov, R. On reward-free reinforcement learning with linear function approximation. *Advances in neural information processing systems*, 2020b.
- Wang, R., Salakhutdinov, R. R., and Yang, L. Reinforcement learning with general value function approximation: Provably efficient approach via bounded eluder dimension. *Advances in Neural Information Processing Systems*, 33:6123–6135, 2020c.
- Wang, Y., Wang, R., Du, S. S., and Krishnamurthy, A. Optimism in reinforcement learning with generalized linear function approximation. In *International Conference on Learning Representations*, 2019.

- Weisz, G., Amortila, P., and Szepesvári, C. Exponential lower bounds for planning in mdps with linearly-realizable optimal action-value functions. In *Algorithmic Learning Theory*, pp. 1237–1264. PMLR, 2021.
- Yang, L. and Wang, M. Sample-optimal parametric q-learning using linearly additive features. In *International Conference on Machine Learning*, pp. 6995–7004. PMLR, 2019.
- Yang, L. and Wang, M. Reinforcement learning in feature space: Matrix bandit, kernels, and regret bound. In *International Conference on Machine Learning*, pp. 10746–10756. PMLR, 2020.
- Zanette, A., Brandfonbrener, D., Brunskill, E., Pirotta, M., and Lazaric, A. Frequentist regret bounds for randomized least-squares value iteration. In *International Conference on Artificial Intelligence and Statistics*, pp. 1954–1964, 2020a.
- Zanette, A., Lazaric, A., Kochenderfer, M., and Brunskill, E. Learning near optimal policies with low inherent bellman error. In *International Conference on Machine Learning*, pp. 10978–10989. PMLR, 2020b.
- Zanette, A., Lazaric, A., Kochenderfer, M. J., and Brunskill, E. Provably efficient reward-agnostic navigation with linear value iteration. *Advances in Neural Information Processing Systems*, 2020c.
- Zhang, W., Zhou, D., and Gu, Q. Reward-free model-based reinforcement learning with linear function approximation. *Advances in Neural Information Processing Systems*, 34, 2021a.
- Zhang, Z., Du, S. S., and Ji, X. Nearly minimax optimal reward-free reinforcement learning. *arXiv preprint arXiv:2010.05901*, 2020.
- Zhang, Z., Ji, X., and Du, S. Is reinforcement learning more difficult than bandits? a near-optimal algorithm escaping the curse of horizon. In *Conference on Learning Theory*, pp. 4528–4531. PMLR, 2021b.
- Zhang, Z., Yang, J., Ji, X., and Du, S. S. Improved variance-aware confidence sets for linear bandits and linear mixture mdp. *Advances in Neural Information Processing Systems*, 34, 2021c.
- Zhang, Z., Zhou, Y., and Ji, X. Model-free reinforcement learning: from clipped pseudo-regret to sample complexity. In *International Conference on Machine Learning*, pp. 12653–12662. PMLR, 2021d.
- Zhang, Z., Ji, X., and Du, S. Horizon-free reinforcement learning in polynomial time: the power of stationary policies. In *Conference on Learning Theory*, pp. 3858–3904. PMLR, 2022.
- Zhou, D. and Gu, Q. Computationally efficient horizon-free reinforcement learning for linear mixture mdps. *arXiv preprint arXiv:2205.11507*, 2022.
- Zhou, D., Gu, Q., and Szepesvari, C. Nearly minimax optimal reinforcement learning for linear mixture markov decision processes. In *Conference on Learning Theory*. PMLR, 2021a.
- Zhou, D., Gu, Q., and Szepesvari, C. Nearly minimax optimal reinforcement learning for linear mixture markov decision processes. In *Conference on Learning Theory*, pp. 4532–4576. PMLR, 2021b.
- Zhou, D., He, J., and Gu, Q. Provably efficient reinforcement learning for discounted mdps with feature mapping. In *International Conference on Machine Learning*. PMLR, 2021c.
- Zhou, D., He, J., and Gu, Q. Provably efficient reinforcement learning for discounted mdps with feature mapping. In *International Conference on Machine Learning*, pp. 12793–12802. PMLR, 2021d.

A. Omitted Proof in Section 6

In this section, we provide detailed proof for Theorem 5.1. For $k \in [K]$, $h \in [H]$, let $\mathcal{F}_{k,h}$ be the σ -algebra generated by the random variables representing the state-action pairs up to and including those that appear stage h of episode k . That is, $\mathcal{F}_{k,h}$ is generated by

$$\begin{array}{ccccc} s_1^1, a_1^1 & \cdots & s_h^1, a_h^1 & \cdots & s_H^1, a_H^1 \\ s_1^2, a_1^2 & \cdots & s_h^2, a_h^2 & \cdots & s_H^2, a_H^2 \\ \vdots & & \vdots & & \vdots \\ s_1^k, a_1^k & \cdots & s_h^k, a_h^k & & \end{array}$$

A.1. Notations

To establish clarity and facilitate understanding, we provide the Table 2 that outlines the notations which will be utilized throughout the proof.

Notation	Meaning
s_h, a_h	States and actions introduced by a general policy π (not specified).
s_h^k, a_h^k	States and actions introduced in the episode k by the policy π_k .
$u_{k,h}(s, a; \theta, \pi, r)$	The uncertainty of states and actions, defined in Equation (4.1).
$\theta_k, \pi_k, r_k = \{r_{k,h}\}_{h \in [H]}$	The transition kernel parameter, the exploration policy, and the pseudo reward function obtained via the optimization oracle in Line 5 of Algorithm 1.
$V_h(s; \theta, \pi, r)$	The value function of policy π in the linear mixture MDP with transition kernel parameter θ and reward function r .
$\hat{V}_{k,h}(s; \theta, \pi, r)$	The uncertainty along the trajectory, defined in Equation (4.2)
$\tilde{V}_{k,h}(s)$	The pseudo value function, equal to $V_h(s; \theta_k, \pi_k, r_k)$
$\hat{\theta}_{k,m}$	The estimated transition kernel parameter obtained by value regression targeting $\hat{V}_{k,h}^{2^m}$.
$\tilde{\theta}_{k,m}$	The estimated transition kernel parameter obtained by value regression targeting $\tilde{V}_{k,h}^{2^m}$.
θ^*	The ground-truth transition kernel parameter.
$\mathcal{U}_k$	The confidence set containing θ^* with high probability.
$\beta_k, (\beta = \beta_K)$	The radius of confidence set $\mathcal{U}_k$.
$\tilde{\Sigma}_{k,h,m}, \hat{\Sigma}_{k,h,m}$	The covariance matrix for $\tilde{V}_{k,h}^{2^m}$ and $\hat{V}_{k,h}^{2^m}$, respectively.
$\tilde{\Sigma}_{k,m}, \hat{\Sigma}_{k,m}$	Equal to $\tilde{\Sigma}_{k-1,H+1,m}$ and $\hat{\Sigma}_{k-1,H+1,m}$, respectively.
$\tilde{\sigma}_{k,h,m}, \hat{\sigma}_{k,h,m}$	The weights for regression problems targeting $\tilde{V}_{k,h}^{2^m}$ and $\hat{V}_{k,h}^{2^m}$ respectively, defined in Equation (4.4).
$\tilde{\phi}_{k,h,m}, \hat{\phi}_{k,h,m}$	Equal to $\phi_{\tilde{V}_{k,h+1}^{2^m}}(s_h^k, a_h^k)$ and $\phi_{\hat{V}_{k,h+1}^{2^m}}(s_h^k, a_h^k)$, respectively.
$\hat{\pi}_r$	The policy obtained in the planning phase.

Table 2. Important Notations

A.2. Proof of Lemma 6.1

Lemma A.1. For all $0 < \delta < 1$, suppose β_k is set as in Theorem 5.1, the following event happens with probability at least $1 - 2M\delta$

$$\left\| \hat{\theta}_{k,m} - \theta^* \right\|_{\hat{\Sigma}_{k,m}} \leq \beta_k \quad (\text{A.1})$$

$$\left\| \tilde{\theta}_{k,m} - \theta^* \right\|_{\tilde{\Sigma}_{k,m}} \leq \beta_k \quad (\text{A.2})$$

$$\left\| \theta_k - \theta^* \right\|_{\hat{\Sigma}_{k,0}} \leq 2\beta_k \quad (\text{A.3})$$

$$\|\boldsymbol{\theta}_k - \boldsymbol{\theta}^*\|_{\dot{\Sigma}_{k,0}} \leq 2\beta_k. \quad (\text{A.4})$$

In the following, we define the event that Lemma A.1 holds to be $\mathcal{E}_{A.1}$. And the following lemmas are conditioned on $\mathcal{E}_{A.1}$ by default. We define function W_h for certain sequence $\{R_h\}$ recursively as

$$W_h(\{R_h\}) = \min\{1, R_h + W_{h+1}(\{R_h\})\}.$$

In addition we denote the trajectory of first h steps as $\text{traj}_h := (s_1, a_1, \dots, s_{h-1}, a_{h-1}, s_h)$, and the trajectory sampled from $(\pi, \mathbb{P})$ conditioned on traj_h as $\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h$.

Lemma A.2. For any policy π and reward function $r \in R$, we have

$$V_1(s_1; \boldsymbol{\theta}_K, \pi, r) - V_1(s_1; \boldsymbol{\theta}^*, \pi, r) = \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{(\mathbb{P}_K - \mathbb{P})V_{h+1}(s_h; \boldsymbol{\theta}_K, \pi, r)\}) \quad (\text{A.5})$$

Lemma A.3. For any policy π and reward function $r \in R$, we have

$$\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{u_{k,h}(s_h, \pi(s_h); \boldsymbol{\theta}_K, \pi, r)\}) \leq \widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K).$$

Proof of Lemma 6.1. The proof follows the proof of Lemma 15 in Zhang et al. (2020). Firstly,

$$\begin{aligned} & V_1^*(s_1; r) - V_1(s_1; \boldsymbol{\theta}^*, \widehat{\pi}_r, r) \\ &= (V_1^*(s_1; r) - V_1(s_1; \boldsymbol{\theta}_K, \widehat{\pi}_r, r)) + (V_1(s_1; \boldsymbol{\theta}_K, \widehat{\pi}_r, r) - V_1(s_1; \boldsymbol{\theta}^*, \widehat{\pi}_r, r)) \\ &\leq (V_1^*(s_1; r) - V_1(s_1; \boldsymbol{\theta}_K, \pi_r^*, r)) + (V_1(s_1; \boldsymbol{\theta}_K, \widehat{\pi}_r, r) - V_1(s_1; \boldsymbol{\theta}^*, \widehat{\pi}_r, r)), \end{aligned} \quad (\text{A.6})$$

where π_r^* is the optimal policy for $(\boldsymbol{\theta}, r)$, and $\widehat{\pi}_r$ is the optimal policy for $(\boldsymbol{\theta}_K, r)$. Then for any policy $\pi \in \Pi$,

$$\begin{aligned} & |V_1(s_1; \boldsymbol{\theta}_K, \pi, r) - V_1(s_1; \boldsymbol{\theta}^*, \pi, r)| \\ &= |\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{(\mathbb{P}_K - \mathbb{P})V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r)\})| \\ &= |\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{(\boldsymbol{\theta}_K - \boldsymbol{\theta}^*)\phi_{V_{h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, a_h)\})| \\ &\leq \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1\left(\left\{\|\boldsymbol{\theta}_K - \boldsymbol{\theta}^*\|_{\dot{\Sigma}_{k,0}} \|\phi_{V_{h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, a_h)\|_{\dot{\Sigma}_{k,0}^{-1}}\right\}\right) \\ &\leq \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1\left(\left\{2\beta \|\phi_{V_{h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, a_h)\|_{\dot{\Sigma}_{k,0}^{-1}}\right\}\right) \\ &= 2\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{u_h(s_h, a_h; \boldsymbol{\theta}_K, \pi, r)\}) \\ &\leq 2\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K). \end{aligned} \quad (\text{A.7})$$

The first equality holds due to Lemma A.2, the second inequality holds due to Cauchy-Schwartz inequality, the third inequality holds due to Lemma A.1, and the last inequality holds due to Lemma A.3. Plugging (A.7) into (A.7), we obtain

$$\begin{aligned} V_1^*(s_1; r) - V_1(s_1; \boldsymbol{\theta}^*, \widehat{\pi}_r, r) &\leq 2\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K) + 2\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K) \\ &= 4\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K). \end{aligned}$$

□

A.3. Proof of Lemma 6.2

Proof of Lemma 6.2. The proof follows the proof of Lemma 14 in Chen et al. (2021). Firstly, we prove that $\widehat{V}_{k,1}(s; \boldsymbol{\theta}, \pi, r)$ is non-increasing w.r.t. k for any fixed $\boldsymbol{\theta}, \pi, r$ by induction in h . Suppose for any $k_1 \leq k_2$, $\widehat{V}_{k_1, h+1}(s; \boldsymbol{\theta}, \pi, r) \geq \widehat{V}_{k_2, h+1}(s; \boldsymbol{\theta}, \pi, r)$ for any s . By definition,

$$\begin{aligned} \widehat{V}_{k,h}(s; \boldsymbol{\theta}, \pi, r) &= \min \left\{ 1, u_{k,h}(s, a; \boldsymbol{\theta}, \pi, r) + 2\beta \left\| \phi_{\widehat{V}_{k,h+1}(\cdot; \boldsymbol{\theta}, \pi, r)}(s, \pi(s)) \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right. \\ &\quad \left. + \phi_{\widehat{V}_{k,h+1}(\cdot; \boldsymbol{\theta}, \pi, r)}^\top(s, \pi(s))\boldsymbol{\theta} \right\} \end{aligned}$$

$$u_{k,h}(s, a; \boldsymbol{\theta}, \pi, r) = \beta \left\| \phi_{V_h}(\cdot; \boldsymbol{\theta}, \pi, r)(s, a) \right\|_{\dot{\boldsymbol{\Sigma}}_{k,0}^{-1}}$$

Since $\dot{\boldsymbol{\Sigma}}_{k_1,0} \preceq \dot{\boldsymbol{\Sigma}}_{k_2,0}$ and $\dot{\boldsymbol{\Sigma}}_{k_1,0} \preceq \dot{\boldsymbol{\Sigma}}_{k_2,0}$, we have

$$\begin{aligned} u_{k_1,h}(s, a; \boldsymbol{\theta}, \pi, r) &\geq u_{k_2,h}(s, a; \boldsymbol{\theta}, \pi, r) \\ \left\| \phi_{\widehat{V}_{k_1,h+1}}(\cdot; \boldsymbol{\theta}, \pi, r)(s, \pi(s)) \right\|_{\dot{\boldsymbol{\Sigma}}_{k,0}^{-1}} &\geq \left\| \phi_{\widehat{V}_{k_2,h+1}}(\cdot; \boldsymbol{\theta}, \pi, r)(s, \pi(s)) \right\|_{\dot{\boldsymbol{\Sigma}}_{k,0}^{-1}} \\ \phi_{\widehat{V}_{k_1,h+1}}^\top(\cdot; \boldsymbol{\theta}, \pi, r)(s, \pi(s))\boldsymbol{\theta} &\geq \phi_{\widehat{V}_{k_2,h+1}}^\top(\cdot; \boldsymbol{\theta}, \pi, r)(s, \pi(s))\boldsymbol{\theta} \end{aligned}$$

Thus $\widehat{V}_{k_1,h}(s; \boldsymbol{\theta}, \pi, r) \geq \widehat{V}_{k_2,h}(s; \boldsymbol{\theta}, \pi, r)$ for any $k_1 \leq k_2$. Furthermore, since $\mathcal{U}_{k_2} \subset \mathcal{U}_{k_1}$, and $\boldsymbol{\theta}_k, \pi_k, r_k$ are argmax over $\mathcal{U}_k$, we have

$$\widehat{V}_{k_1,1}(s_1; \boldsymbol{\theta}_{k_1}, \pi_{k_1}, r_{k_1}) \geq \widehat{V}_{k_1,1}(s_1; \boldsymbol{\theta}_{k_2}, \pi_{k_2}, r_{k_2}) \geq \widehat{V}_{k_2,1}(s_1; \boldsymbol{\theta}_{k_2}, \pi_{k_2}, r_{k_2})$$

It follows that $\widehat{V}_{k,1}(s_1^k; \boldsymbol{\theta}_k, \pi_k, r_k)$ is non-increasing w.r.t. k . Thus,

$$K\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K) \leq \sum_{k=1}^K \widehat{V}_{k,1}(s_1; \boldsymbol{\theta}_k, \pi_k, r_k)$$

□

A.4. Proof of Lemma 6.3

Lemma A.4. Conditioned on the event $\mathcal{E}$, let $\widetilde{V}_{k,h}, \widehat{V}_{k,h}, \dot{\boldsymbol{\Sigma}}_{k,m}, \widehat{\boldsymbol{\Sigma}}_{k,m}, \widetilde{\boldsymbol{\phi}}_{k,h,m}, \widehat{\boldsymbol{\phi}}_{k,h,m}$ be defined in Algorithm 1, for any $k \in [K], h \in [H], m \in [\overline{M}]$, we have

$$\widehat{V}_{k,h}(s_h^k) - u_{k,h}(s_h^k, a_h^k) - \mathbb{P}\widehat{V}_{k,h+1}(s_h^k, a_h^k) \leq 4 \min \left\{ 1, \beta \left\| \widehat{\boldsymbol{\phi}}_{k,h,0} \right\|_{\dot{\boldsymbol{\Sigma}}_{k,0}^{-1}} \right\} \quad (\text{A.8})$$

$$\widetilde{V}_{k,h}(s_h^k) - r_{k,h}(s_h^k, a_h^k) - \mathbb{P}\widetilde{V}_{k,h+1} \leq 2 \min \left\{ 1, \beta \left\| \widetilde{\boldsymbol{\phi}}_{k,h,0} \right\|_{\dot{\boldsymbol{\Sigma}}_{k,0}^{-1}} \right\} \quad (\text{A.9})$$

In order to prove Lemma 6.3, we introduce the following quantities used in Zhou & Gu (2022) as

$$\widehat{R}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^j \min \left\{ 1, \beta \left\| \widehat{\boldsymbol{\phi}}_{k,h,m} \right\|_{\dot{\boldsymbol{\Sigma}}_{k,m}^{-1}} \right\}, \forall m \in [\overline{M}] \quad (\text{A.10})$$

$$\widetilde{R}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^j \min \left\{ 1, \beta \left\| \widetilde{\boldsymbol{\phi}}_{k,h,m} \right\|_{\dot{\boldsymbol{\Sigma}}_{k,m}^{-1}} \right\}, \forall m \in [\overline{M}] \quad (\text{A.11})$$

$$\widehat{A}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\left[\mathbb{P}\widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) - \widehat{V}_{k,h+1}^{2^m} (s_{h+1}^k) \right], \forall m \in [\overline{M}] \quad (\text{A.12})$$

$$\widetilde{A}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\left[\mathbb{P}\widetilde{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) - \widetilde{V}_{k,h+1}^{2^m} (s_{h+1}^k) \right], \forall m \in [\overline{M}] \quad (\text{A.13})$$

$$\widehat{S}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\mathbb{V}\widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k), \forall m \in [\overline{M}] \quad (\text{A.14})$$

$$\widetilde{S}_m = \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\mathbb{V}\widetilde{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k), \forall m \in [\overline{M}] \quad (\text{A.15})$$

$$I_h^k = \mathbb{1} \left\{ \forall m \in [\overline{M}], \det \left(\dot{\boldsymbol{\Sigma}}_{k,m}^{-1/2} \right) / \det \left(\widehat{\boldsymbol{\Sigma}}_{k,h,m}^{-1/2} \right) \leq 4 \text{ and } \det \left(\dot{\boldsymbol{\Sigma}}_{k,m}^{-1/2} \right) / \det \left(\widetilde{\boldsymbol{\Sigma}}_{k,h,m}^{-1/2} \right) \leq 4 \right\} \quad (\text{A.16})$$

$$G = \sum_{k=1}^K (1 - I_H^k), \quad (\text{A.17})$$

Lemma A.5. Let γ, α , be defined in Algorithm 2, $\{\hat{R}_m\}_{m \in \overline{[M]}}$, $\{\tilde{R}_m\}_{m \in \overline{[M]}}$, $\{\hat{S}_m\}_{m \in \overline{[M]}}$, $\{\tilde{S}_m\}_{m \in \overline{[M]}}$ be defined in (A.10), (A.11), (A.14), (A.15). Then for $m \in \overline{[M-1]}$, we have

$$\hat{R}_m \leq \min \left\{ KH, 4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{\hat{S}_m + 4\hat{R}_m + 2\hat{R}_{m+1} + KH\alpha^2} \right\} \quad (\text{A.18})$$

$$\tilde{R}_m \leq \min \left\{ KH, 4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{\tilde{S}_m + 4\tilde{R}_m + 2\tilde{R}_{m+1} + KH\alpha^2} \right\}, \quad (\text{A.19})$$

where $\iota = \log(1 + KH/(d\lambda\alpha^2))$. For $\hat{R}_{M-1}$ and $\tilde{R}_{M-1}$, we have the trivial bound $\hat{R}_{M-1} \leq KH$ and $\tilde{R}_{M-1} \leq KH$.

Lemma A.6. Let $\{\hat{R}_m\}_{m \in \overline{[M]}}$, $\{\tilde{R}_m\}_{m \in \overline{[M]}}$, $\{\hat{S}_m\}_{m \in \overline{[M]}}$, $\{\tilde{S}_m\}_{m \in \overline{[M]}}$, $\{\hat{A}_m\}_{m \in \overline{[M]}}$, $\{\tilde{A}_m\}_{m \in \overline{[M]}}$, G be defined as (A.10), (A.11), (A.14), (A.15), (A.12), (A.13), (A.17). Then, conditioned on the event $\mathcal{E}$, for $m \in \overline{[M-1]}$, we have

$$\hat{S}_m \leq |\hat{A}_{m+1}| + G + 2^{m+1} (\tilde{R}_0 + 4\hat{R}_0) \quad (\text{A.20})$$

$$\tilde{S}_m \leq |\tilde{A}_{m+1}| + G + 2^{m+1} (K + 2\tilde{R}_0) \quad (\text{A.21})$$

Lemma A.7. Let $\{\hat{S}_m\}_{m \in \overline{[M]}}$, $\{\tilde{S}_m\}_{m \in \overline{[M]}}$, $\{\hat{A}_m\}_{m \in \overline{[M]}}$, $\{\tilde{A}_m\}_{m \in \overline{[M]}}$ be defined as (A.14), (A.15), (A.12), (A.13). Then we have $\mathbb{P}(\mathcal{E}_{A.7}) > 1 - 2M\delta$, with $\mathcal{E}_{A.7}$ be defined as,

$$\mathcal{E}_{A.7} := \left\{ \forall m \in \overline{[M]}, |\hat{A}_m| \leq \min \left\{ \sqrt{2\zeta \hat{S}_m} + \zeta, KH \right\} \text{ and } |\tilde{A}_m| \leq \min \left\{ \sqrt{2\zeta \tilde{S}_m} + \zeta, KH \right\} \right\}, \quad (\text{A.22})$$

where $\zeta = 4\log(4\log(KH)/\delta)$.

Lemma A.8. Let G be defined in (A.17). Then we have

$$G \leq M d \iota, \quad (\text{A.23})$$

where $\iota = \log(1 + KH/(d\lambda\alpha^2))$.

Lemma A.9. (Restatement of Lemma 6.3) For any $0 < \delta < 1$, with probability at least $1 - 4M\delta$, we have

$$\begin{aligned} & \sum_{k=1}^K \hat{V}_{k,1}(s_1^k; \theta_k, \pi_k, r_k) \\ & \leq 896 \max \{ 64\beta^2 d\iota, 2\zeta \} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota \sqrt{M} + 24\sqrt{\zeta M d\iota} + M d\iota \\ & \quad + \left(64 \max \{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \} + 120\beta\sqrt{d\iota} H \alpha^2 \right) \sqrt{K} \end{aligned}$$

where $\iota = \log(1 + KH/(d\lambda\alpha^2))$, $\zeta = 4\log(4\log(KH)/\delta)$.

Proof of Lemma A.9. All the following proofs are conditioned on $\mathcal{E}_{A.1} \cap \mathcal{E}_{A.7}$, which happens with probability at least $1 - 4M\delta$. Firstly, we have

$$\begin{aligned} & \sum_{k=1}^K \hat{V}_{k,1}(s_h^k) \\ & = \sum_{k=1}^K \sum_{h=1}^H \left[I_h^k \left[\hat{V}_{k,h}(s_h^k) - \hat{V}_{k,h+1}(s_{h+1}^k) \right] + (1 - I_h^k) \left[\hat{V}_{k,h}(s_h^k) - \hat{V}_{k,h+1}(s_{h+1}^k) \right] \right] \\ & = \sum_{k=1}^K \left[\sum_{h=1}^H I_h^k u_{k,h}(s_h^k, a_h^k) + \sum_{h=1}^H I_h^k \left[\hat{V}_{k,h}(s_h^k) - u_{k,h}(s_h^k, a_h^k) - \mathbb{P} \hat{V}_{k,h+1}(s_h^k, a_h^k) \right] \right] \\ & \quad + \sum_{h=1}^H I_h^k \left[\mathbb{P} \hat{V}_{k,h+1}(s_h^k, a_h^k) - \hat{V}_{k,h+1}(s_{h+1}^k) \right] + \sum_{k=1}^K \sum_{h=1}^H (1 - I_h^k) \left[\hat{V}_{k,h}(s_h^k) - \hat{V}_{k,h+1}(s_{h+1}^k) \right] \end{aligned}$$

$$\begin{aligned}
 &\leq \underbrace{\sum_{k=1}^K \sum_{h=1}^H I_h^k u_{k,h}(s_h^k, a_h^k)}_{I_1} + \underbrace{\sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{h,k}(s_h^k) - u_{k,h}(s_h^k, a_h^k) - \mathbb{P} \widehat{V}_{k,h+1}(s_h^k, a_h^k) \right]}_{I_2} \\
 &\quad + \underbrace{\sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\mathbb{P} \widehat{V}_{k,h+1}(s_h^k, a_h^k) - \widehat{V}_{h+1,k}(s_{h+1}^k) \right]}_{I_3} + \underbrace{\sum_{k=1}^K (1 - I_{h_k}^k) \widehat{V}_{k,h_k}(s_{h_k}^k)}_{I_4},
 \end{aligned}$$

where h_k is the smallest index such that $I_{h_k}^k = 0$. Following the definition of $u_{k,h}$,

$$I_1 = \sum_{k=1}^K \sum_{h=1}^H I_h^k \min \left\{ 1, \beta \left\| \widetilde{\phi}_{k,h,0} \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right\} = \widetilde{R}_0.$$

By Lemma A.4,

$$I_2 \leq 4 \sum_{k=1}^K \sum_{h=1}^H I_h^k \min \left\{ 1, \beta \left\| \widehat{\phi}_{k,h,0} \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right\} = 4\widehat{R}_0$$

By definitions,

$$\begin{aligned}
 I_3 &= \widehat{A}_0, \\
 I_4 &\leq \sum_{k=1}^K (1 - I_H^k) = G.
 \end{aligned}$$

Thus,

$$\sum_{k=1}^K \widehat{V}_{k,1}(s_h^k) \leq \widetilde{R}_0 + 4\widehat{R}_0 + \widehat{A}_0 + G \quad (\text{A.24})$$

Substituting (A.20) in Lemma A.6 into (A.18) in Lemma A.5, we have

$$\begin{aligned}
 \widehat{R}_m &\leq 4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{\left| \widehat{A}_{m+1} \right| + G + 2^{m+1} \left(\widetilde{R}_0 + 4\widehat{R}_0 \right) + 4\widehat{R}_m + 2\widehat{R}_{m+1} + KH\alpha^2} \\
 &\leq 2\beta\sqrt{d\iota} \sqrt{\left| \widehat{A}_{m+1} \right| + 2^{m+1} \left(\widetilde{R}_0 + 4\widehat{R}_0 \right) + 4\widehat{R}_m + 2\widehat{R}_{m+1} + \underbrace{4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{G + KH\alpha^2}}_{I_c}}, \quad (\text{A.25})
 \end{aligned}$$

where the second inequality holds due to $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$. Substituting (A.20) in Lemma A.6 into (A.22) in Lemma A.7, we have

$$\begin{aligned}
 \left| \widehat{A}_m \right| &\leq \sqrt{2\zeta} \sqrt{\left| \widehat{A}_{m+1} \right| + G + 2^{m+1} \left(\widetilde{R}_0 + 4\widehat{R}_0 \right) + \zeta} \\
 &\leq \sqrt{2\zeta} \sqrt{\left| \widehat{A}_{m+1} \right| + 2^{m+1} \left(\widetilde{R}_0 + 4\widehat{R}_0 \right) + \sqrt{2\zeta G} + \zeta} \quad (\text{A.26})
 \end{aligned}$$

Substituting (A.21) in Lemma A.6 into (A.19) in Lemma A.5, we have

$$\begin{aligned}
 \widetilde{R}_m &\leq 4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{\left| \widetilde{A}_{m+1} \right| + G + 2^{m+1} \left(K + 2\widetilde{R}_0 \right) + 4\widetilde{R}_m + 2\widetilde{R}_{m+1} + KH\alpha^2} \\
 &\leq 2\beta\sqrt{d\iota} \sqrt{\left| \widetilde{A}_{m+1} \right| + 2^{m+1} \left(K + 2\widetilde{R}_0 \right) + 4\widetilde{R}_m + 2\widetilde{R}_{m+1} + \underbrace{4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota} \sqrt{G + KH\alpha^2}}_{I_c}} \quad (\text{A.27})
 \end{aligned}$$

Substituting (A.21) in Lemma A.6 into (A.22) in Lemma A.7, we have

$$\begin{aligned} |\tilde{A}_m| &\leq \sqrt{2\zeta} \sqrt{|\tilde{A}_{m+1}| + G + 2^{m+1} (K + 2\tilde{R}_0)} + \zeta \\ &\leq \sqrt{2\zeta} \sqrt{|\tilde{A}_{m+1}| + 2^{m+1} (K + 2\tilde{R}_0)} + \sqrt{2\zeta G} + \zeta \end{aligned} \quad (\text{A.28})$$

Thus, calculating (A.27) + (A.28) + $4 \times$ (A.25) + (A.26) and using $\sqrt{a} + \sqrt{b} + \sqrt{c} + \sqrt{d} \leq 2\sqrt{a+b+c+d}$, we have

$$\begin{aligned} &\tilde{R}_m + |\tilde{A}_m| + 4\hat{R}_m + |\hat{A}_m| \\ &\leq 5I_c + 2\sqrt{2\zeta G} + 2\zeta + 2 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{2|\hat{A}_{m+1}| + 2 \cdot 2^{m+1} (\tilde{R}_0 + 4\hat{R}_0)} \\ &\quad + 4\hat{R}_m + 2\hat{R}_{m+1} + 2|\tilde{A}_{m+1}| + 2 \cdot 2^{m+1} (K + 2\tilde{R}_0) + 4\tilde{R}_m + 2\tilde{R}_{m+1} \\ &\leq 5I_c + 2\sqrt{2\zeta G} + 2\zeta + 4 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{(\tilde{R}_m + |\tilde{A}_m| + 4\hat{R}_m + |\hat{A}_m|)} \\ &\quad + (\tilde{R}_{m+1} + |\tilde{A}_{m+1}| + 4\hat{R}_{m+1} + |\hat{A}_{m+1}|) + 2 \cdot 2^{m+1} (K + \tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0|). \end{aligned}$$

Then by Lemma D.3 with $a_m = \tilde{R}_m + |\tilde{A}_m| + 4\hat{R}_m + |\hat{A}_m| \leq 7KH$ and $M = \log(7KH)/\log 2$, $\tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0|$ can be bounded as

$$\begin{aligned} &\tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0| \\ &\leq 22 \cdot 16 \max\{64\beta^2 d\iota, 2\zeta\} + 30I_c + 12\sqrt{\zeta G} + 12\zeta \\ &\quad + 32 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{K + \tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0|} \\ &\leq 352 \max \{64\beta^2 d\iota, 2\zeta\} + 30I_c + 12\sqrt{\zeta G} + 12\zeta + 32 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{K} \\ &\quad + 32 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{\tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0|}. \end{aligned} \quad (\text{A.29})$$

By the fact that $x \leq a\sqrt{x} + b \Rightarrow x \leq 2a^2 + 2b$, (A.29) implies that

$$\begin{aligned} &\tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0| \\ &\leq 896 \max \{64\beta^2 d\iota, 2\zeta\} + 60I_c + 24\sqrt{\zeta G} + 24\zeta + 64 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{K}. \end{aligned} \quad (\text{A.30})$$

Finally, plugging (A.30) into (A.24) and bounding G with Lemma A.8, we have

$$\begin{aligned} &\sum_{k=1}^K \hat{V}_{k,1}(s_h^k) \\ &\leq \tilde{R}_0 + |\tilde{A}_0| + 4\hat{R}_0 + |\hat{A}_0| + G \\ &\leq 896 \max \{64\beta^2 d\iota, 2\zeta\} + 24\zeta + 64 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} \sqrt{K} \\ &\quad + 60 \left(4d\iota + 4\beta\gamma^2 d\iota + 2\beta\sqrt{d\iota}\sqrt{Md\iota + KH\alpha^2} \right) + 24\sqrt{\zeta Md\iota} + Md\iota \\ &\leq 896 \max \{64\beta^2 d\iota, 2\zeta\} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota\sqrt{M} + 24\sqrt{\zeta Md\iota} + Md\iota \\ &\quad + \left(64 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} + 120\beta\sqrt{d\iota H\alpha^2} \right) \sqrt{K} \end{aligned} \quad (\text{A.31})$$

□

A.5. Proof of Main Results

Lemma A.10. (Restatement of Theorem 5.1) For Algorithm 1, set $M = \log(7KH)/\log(2)$, $\{\beta_k\}_{k \geq 1}$ as

$$\begin{aligned} \beta_k = & 12\sqrt{d \log(1 + kH/(\alpha^2 d\lambda)) \log(32(\log(\gamma^2/\alpha) + 1)k^2 H^2/\delta)} \\ & + 30 \log(32(\log(\gamma^2/\alpha) + 1)k^2 H^2/\delta)/\gamma^2 + \sqrt{\lambda}B, \end{aligned}$$

and denote $\beta = \beta_K$, then for any $0 < \delta' < 1$, we have with probability at least $1 - \delta'$, where $\delta' = 4M\delta$, after collecting K trajectories, algorithm 1 returns a policy satisfying the following sub-optimality bound,

$$\begin{aligned} & V_1^*(s_1; r) - V_1(s_1; \theta^*, \hat{\pi}_r, r) \\ & \leq \frac{4}{K} \left(2752 \max \{64\beta^2 d\iota, 2\zeta\} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota \sqrt{M} \right) \\ & \quad + \frac{4}{\sqrt{K}} \left(64 \max \{8\beta\sqrt{d\iota}, \sqrt{2\zeta}\} + 120\beta\sqrt{d\iota} \right), \end{aligned}$$

where $\iota = \log(1 + KH/(d\lambda\alpha^2))$, $\zeta = 4 \log(4 \log(KH)/\delta)$. Moreover, setting $\alpha = H^{-1/2}$, $\gamma = d^{-1/4}$, and $\lambda = d/B^2$, we have the horizon-free suboptimality bound

$$V_1^*(s_1; r) - V_1(s_1; \theta^*, \hat{\pi}_r, r) = \tilde{O} \left(\frac{d^2}{K} + \frac{d}{\sqrt{K}} \right). \quad (\text{A.32})$$

Proof of Theorem A.10. The following proof is conditioned on $\mathcal{E}_{A.1} \cap \mathcal{E}_{A.7}$, which holds with probability at least $1 - 4M\delta = 1 - \delta'$. We have

$$\begin{aligned} & V_1^*(s_1; r) - V_1(s_1; \theta^*, \hat{\pi}_r, r) \\ & \leq 4\hat{V}_{K,1}(s_1; \theta_K, \hat{\pi}_K, r_K) \\ & \leq \frac{4}{K} \sum_{k=1}^K V_{k,1}(s_1; \theta_k, \hat{\pi}_k, r_k) \\ & \leq \frac{4}{K} \left(896 \max \{64\beta^2 d\iota, 2\zeta\} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota \sqrt{M} + 24\sqrt{\zeta M d\iota} + M d\iota \right) \\ & \quad + \frac{4}{\sqrt{K}} \left(64 \max \{8\beta\sqrt{d\iota}, \sqrt{2\zeta}\} + 120\beta\sqrt{d\iota H\alpha^2} \right), \end{aligned}$$

where the first inequality holds due to Lemma 6.1, the second inequality holds due to Lemma 6.2, and the third equality holds due to Lemma A.9. $\square$

Given the regret bound provided in Lemma A.10, we can prove the following sample complexity upper bound.

Lemma A.11. (Restatement of Corollary 5.2) Under the same conditions as in Theorem A.10, Algorithm 1 has sample complexity of

$$\begin{aligned} m(\varepsilon, \delta') = & \frac{16}{\varepsilon^2} \left(64 \max \{8\beta\sqrt{d\iota}, \sqrt{2\zeta}\} + 120\beta\sqrt{d\iota H\alpha^2} \right)^2 \\ & + \frac{8}{\varepsilon} \left(2752 \max \{64\beta^2 d\iota, 2\zeta\} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota \sqrt{M} \right) \end{aligned} \quad (\text{A.33})$$

Moreover, setting $\alpha = H^{-1/2}$, $\gamma = d^{-1/4}$, and $\lambda = d/B^2$, we have the horizon-free sample complexity bound

$$m(\varepsilon, \delta') = \tilde{O} \left(\frac{d^2}{\varepsilon^2} \right).$$

Proof of Lemma A.11. (A.33) is derived directly from Lemma A.10 by setting the suboptimality to ε and solving the K . $\square$

Lemma A.12 (Restatement of Corollary 5.4). When rescaling the assumption $\sum_{h=1}^H r_h(s_h, a_h) \leq 1$ to $\sum_{h=1}^H r_h(s_h, a_h) \leq H$, under the same conditions as Lemma A.10, Algorithm 1 has sample complexity of

$$m(\varepsilon, \delta') = \frac{16H^2}{\varepsilon^2} \left(64 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} + 120\beta\sqrt{d\iota H\alpha^2} \right)^2 + \frac{8H}{\varepsilon} \left(2752 \max \{ 64\beta^2 d\iota, 2\zeta \} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota\sqrt{M} \right) \quad (\text{A.34})$$

Moreover, setting $\alpha = H^{-1/2}$, $\gamma = d^{-1/4}$, and $\lambda = d/B^2$, we have the horizon-free sample complexity bound

$$m(\varepsilon, \delta') = \tilde{O} \left(\frac{H^2 d^2}{\varepsilon^2} \right).$$

Proof of Lemma A.12. (A.34) is a direct result of Lemma A.11. Let $r'_h(s_h, a_h) = r_h(s_h, a_h)/H$, then $\sum_{h=1}^H r'_h(s_h, a_h) \leq 1$. Thus the sample complexity of achieving an ε/H -optimal policy for reward r'_h with probability $1 - \delta'$ is

$$m(\varepsilon, \delta') = \frac{16H^2}{\varepsilon^2} \left(64 \max \left\{ 8\beta\sqrt{d\iota}, \sqrt{2\zeta} \right\} + 120\beta\sqrt{d\iota H\alpha^2} \right)^2 + \frac{8H}{\varepsilon} \left(2752 \max \{ 64\beta^2 d\iota, 2\zeta \} + 24\zeta + 240d\iota + 240\beta\gamma^2 d\iota + 120\beta d\iota\sqrt{M} \right). \quad (\text{A.35})$$

Since $r_h(s_h, a_h) = Hr'_h(s_h, a_h)$, for the same policy, the suboptimality for rewards r_h is H times the suboptimality for rewards r'_h . Thus, the ε/H -optimal policy for r'_h is a ε -optimal policy for r_h . Therefore, the sample complexity of achieving an ε -optimal policy for reward r_h with probability $1 - \delta'$ is $m(\varepsilon, \delta')$. $\square$

B. Proof of Lemmas in Appendix A

B.1. Proof of Lemma A.1

Lemma B.1. (Theorem 4.3 in Zhou & Gu (2022)) Let $\{\mathcal{G}_k\}_{k=1}^\infty$ be a filtration, and $\{\mathbf{x}_k, \eta_k\}_{k \geq 1}$ be a stochastic process such that $\mathbf{x}_k \in \mathbb{R}^d$ is $\mathcal{G}_k$ -measurable and $\eta_k \in \mathbb{R}$ is $\mathcal{G}_{k+1}$ -measurable. Let $L, \sigma, \lambda, \varepsilon > 0, \boldsymbol{\mu}^* \in \mathbb{R}^d$. For $k \geq 1$, let $y_k = \langle \boldsymbol{\mu}^*, \mathbf{x}_k \rangle + \eta_k$ and suppose that $\eta_k, \mathbf{x}_k$ also satisfy

$$\mathbb{E}[\eta_k | \mathcal{G}_k] = 0, \mathbb{E}[\eta_k^2 | \mathcal{G}_k] \leq \sigma^2, |\eta_k| \leq R, \|\mathbf{x}_k\|_2 \leq L \quad (\text{B.1})$$

For $k \geq 1$, let $\mathbf{Z}_k = \lambda \mathbf{I} + \sum_{i=1}^k \mathbf{x}_i \mathbf{x}_i^\top$, $\mathbf{b}_k = \sum_{i=1}^k y_i \mathbf{x}_i$, $\boldsymbol{\mu}_k = \mathbf{Z}_k^{-1} \mathbf{b}_k$, and

$$\beta_k = 12\sqrt{\sigma^2 d \log(1 + kL^2/(d\lambda)) \log(32(\log(R/\varepsilon) + 1)k^2/\delta)} + 24 \log(32(\log(R/\varepsilon) + 1)k^2/\delta) \max_{1 \leq i \leq k} \left\{ |\eta_i| \min \left\{ 1, \|\mathbf{x}_i\|_{\mathbf{Z}_{i-1}^{-1}} \right\} \right\} + 6 \log(32(\log(R/\varepsilon) + 1)k^2/\delta) \varepsilon. \quad (\text{B.2})$$

Then, for any $0 < \delta < 1$, we have with probability at least $1 - \delta$ that,

$$\forall k \geq 1, \left\| \sum_{i=1}^k \mathbf{x}_i \eta_i \right\|_{\mathbf{Z}_k^{-1}} \leq \beta_k, \quad \|\boldsymbol{\mu}_k - \boldsymbol{\mu}^*\|_{\mathbf{Z}_k} \leq \beta_k + \sqrt{\lambda} \|\boldsymbol{\mu}^*\|_2$$

Lemma B.2. Let $\tilde{V}_{k,h}, \hat{V}_{k,h}, \hat{\Sigma}_{k,m}, \dot{\Sigma}_{k,m}, \tilde{\theta}_{k,m}, \hat{\theta}_{k,m}, \tilde{\phi}_{k,h,m}, \hat{\phi}_{k,h,m}$ be defined in Algorithm 1, for any $k \in [K]$, $h \in [H]$, $m \in [M]$. We have

$$\left| \nabla \hat{V}_{k,h+1}^{2^m}(s_h^k, a_h^k) - \hat{\nabla} \hat{V}_{k,h+1}^{2^m}(s_h^k, a_h^k) \right| \leq \min \left\{ 1, \left\| \hat{\phi}_{k,h,m+1} \right\|_{\hat{\Sigma}_{k,m+1}^{-1}} \left\| \hat{\theta}_{k,m+1} - \theta^* \right\|_{\dot{\Sigma}_{k,m+1}} \right\}$$

$$+ \min \left\{ 1, 2 \left\| \hat{\phi}_{k,h,m} \right\|_{\hat{\Sigma}_{k,m}^{-1}} \left\| \hat{\theta}_{k,m} - \theta^* \right\|_{\hat{\Sigma}_{k,m}} \right\}, \quad (\text{B.3})$$

and

$$\begin{aligned} & \left| \mathbb{V}_{k,h+1}^{2^m}(s_h^k, a_h^k) - \tilde{\mathbb{V}}_{k,h+1}^{2^m}(s_h^k, a_h^k) \right| \\ & \leq \min \left\{ 1, \left\| \tilde{\phi}_{k,h,m+1} \right\|_{\tilde{\Sigma}_{k,m+1}^{-1}} \left\| \tilde{\theta}_{k,m+1} - \theta^* \right\|_{\tilde{\Sigma}_{k,m+1}} \right\} \\ & \quad + \min \left\{ 1, 2 \left\| \tilde{\phi}_{k,h,m} \right\|_{\tilde{\Sigma}_{k,m}^{-1}} \left\| \tilde{\theta}_{k,m} - \theta^* \right\|_{\tilde{\Sigma}_{k,m}} \right\}. \end{aligned} \quad (\text{B.4})$$

Proof of Lemma B.2. The proof follows the proof of Lemma C.1 in (Zhou et al., 2021b). We first prove (B.3), and the proof of (B.4) is similar. We have

$$\begin{aligned} & \left| [\hat{\mathbb{V}}_{k,h} \hat{V}_{k,h+1}^{2^m}](s_h^k, a_h^k) - [\mathbb{V}_{k,h} \hat{V}_{k,h+1}](s_h^k, a_h^k) \right| \\ & = \left| [\langle \hat{\phi}_{k,h,m+1}, \hat{\theta}_{k,m+1} \rangle]_{[0,1]} - \langle \hat{\phi}_{k,h,m+1}, \theta^* \rangle \right. \\ & \quad \left. + (\langle \hat{\phi}_{k,h,m}, \theta^* \rangle)^2 - [\langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle]_{[0,1]}^2 \right| \\ & \leq \underbrace{\left| [\langle \hat{\phi}_{k,h,m+1}, \hat{\theta}_{k,m+1} \rangle]_{[0,1]} - \langle \hat{\phi}_{k,h,m+1}, \theta^* \rangle \right|}_{I_1} \\ & \quad + \underbrace{\left| (\langle \hat{\phi}_{k,h,m}, \theta^* \rangle)^2 - [\langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle]_{[0,1]}^2 \right|}_{I_2} \end{aligned} \quad (\text{B.5})$$

where the inequality holds due to triangle inequality. We have $I_1 \leq 1$ since both terms in I_1 lie in the interval $[0, 1]$. Furthermore,

$$\begin{aligned} I_1 & \leq \left| [\langle \hat{\phi}_{k,h,m+1}, \hat{\theta}_{k,m+1} \rangle]_{[0,1]} - \langle \hat{\phi}_{k,h,m+1}, \theta^* \rangle \right| \\ & = \left| [\langle \hat{\phi}_{k,h,m+1}, \hat{\theta}_{k,m+1} - \theta^* \rangle]_{[0,1]} \right| \\ & \leq \left\| \phi_{k,h,m+1} \right\|_{\hat{\Sigma}_{k,m+1}^{-1}} \left\| \hat{\theta}_{k,m+1} - \theta^* \right\|_{\hat{\Sigma}_{k,m+1}}, \end{aligned}$$

where the first inequality holds due to $\langle \hat{\phi}_{k,h,m+1}(s_h^k, a_h^k), \theta^* \rangle \in [0, 1]$, the second inequality holds due to Cauchy-Schwarz inequality. Thus, we obtain

$$I_1 \leq \min \{ 1, \left\| \phi_{k,h,m+1} \right\|_{\hat{\Sigma}_{k,m+1}^{-1}} \left\| \hat{\theta}_{k,m+1} - \theta^* \right\|_{\hat{\Sigma}_{k,m+1}} \} \quad (\text{B.6})$$

For I_2 , we have

$$\begin{aligned} I_2 & = \left| (\langle \hat{\phi}_{k,h,m}(s_h^k, a_h^k), \theta^* \rangle) - [\langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle]_{[0,1]} \right| \cdot \left| (\langle \hat{\phi}_{k,h,m}(s_h^k, a_h^k), \theta^* \rangle) + [\langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle]_{[0,1]} \right| \\ & \leq 2 \left| (\langle \hat{\phi}_{k,h,m}(s_h^k, a_h^k), \theta^* \rangle) - \langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle \right| \\ & \leq 2 \left\| \phi_{k,h,m}(s_h^k, a_h^k) \right\|_{\hat{\Sigma}_{k,m}^{-1}} \left\| \hat{\theta}_{k,m} - \theta^* \right\|_{\hat{\Sigma}_{k,m}} \end{aligned}$$

where the first inequality holds due to that both $\langle \hat{\phi}_{k,h,m}(s_h^k, a_h^k), \theta^* \rangle$ and $[\langle \hat{\phi}_{k,h,m}, \hat{\theta}_{k,m} \rangle]_{[0,1]}$ lie in the interval $[0, 1]$, and the second inequality holds due to Cauchy-Schwarz inequality. Since I_2 belongs to the interval $[0, 1]$, we have

$$I_2 \leq \min \{ 1, 2 \left\| \phi_{k,h,m}(s_h^k, a_h^k) \right\|_{\hat{\Sigma}_{k,m}^{-1}} \left\| \hat{\theta}_{k,m} - \theta^* \right\|_{\hat{\Sigma}_{k,m}} \} \quad (\text{B.7})$$

Substituting (B.6) and (B.7) into (B.5), we obtain (B.3). The proof of B.4 is nearly identical to the proof of (B.5). The only difference is to replace $\hat{\phi}$ with $\tilde{\phi}$, $\hat{\theta}$ with $\tilde{\theta}$, $\hat{\Sigma}$ with $\tilde{\Sigma}$. $\square$

Proof of Lemma A.1. The proof follows Lemma C.2 in Zhou & Gu (2022). Symbols we used here may have small intuitively understandable modification compared to Algorithm 2 since we have to distinguish between Algorithm 2 applied to $\tilde{V}_{k,h}$ and $\hat{V}_{k,h}$. We first prove that Equation (A.1) holds with high probability. By definitions,

$$\begin{aligned}\hat{\sigma}_{k,h,m}^2 &= \max \left\{ \gamma^2 \left\| \hat{\phi}_{k,h,m} \right\|_{\hat{\Sigma}_{k,h,m}^{-1}}, \left[\hat{V}_{k,m} \hat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) + \hat{E}_{k,h,m}, \alpha^2 \right\} \\ \hat{\sigma}_{k,h,m}^2 &= \max \left\{ \gamma^2 \left\| \hat{\phi}_{k,h,M-1} \right\|_{\hat{\Sigma}_{k,h,M-1}^{-1}}, 1, \alpha^2 \right\}.\end{aligned}$$

We define $\mathcal{C}_{k,m}$ as

$$\hat{\mathcal{C}}_{k,m} := \{ \theta : \|\theta - \hat{\theta}_{k,m}\|_{\hat{\Sigma}_{k,m}} \leq \beta_k \}.$$

For each m , let

$$\begin{aligned}\mathbf{x}_{k,h,m} &= \hat{\sigma}_{k,h,m}^{-1} \hat{\phi}_{k,h,m}, \\ \eta_{k,h,m} &= \hat{\sigma}_{k,h,m}^{-1} \mathbb{1}\{\theta^* \in \hat{\mathcal{C}}_{k,m} \cap \hat{\mathcal{C}}_{k,m+1}\} [\hat{V}_{k,h+1}^{2^m} (s_{h+1}^k) - \langle \hat{\phi}_{k,h,m}, \theta^* \rangle], \\ \eta_{k,h,M-1} &= \hat{\sigma}_{k,h,M-1}^{-1} [\hat{V}_{k,h+1}^{2^{M-1}} - \langle \hat{\phi}_{k,h,M-1}, \theta^* \rangle], \\ \mathcal{G}_{k,h} &= \mathcal{F}_{k,h}, \\ \mu^* &= \theta^*.\end{aligned}$$

We have

$$\mathbb{E}[\eta_{k,h,m} | \mathcal{G}_{k,h}] = 0, \quad \|\mathbf{x}_{k,h,m}\|_2 \leq \hat{\sigma}_{k,h,m}^{-1} \leq 1/\alpha, \quad |\eta_{k,h,m}| \leq 1/\alpha$$

Since $\mathbb{1}\{\theta^* \in \hat{\mathcal{C}}_{k,m} \cap \hat{\mathcal{C}}_{k,m+1}\}$ is $\mathcal{G}_{k,h}$ -measurable, then we can bound the variance for $m \in \overline{[M]}$ as follows:

$$\begin{aligned}\mathbb{E}[\eta_{k,h,m}^2 | \mathcal{G}_{k,h}] &= \hat{\sigma}_{k,h,m}^{-2} \mathbb{1}\{\theta^* \in \hat{\mathcal{C}}_{k,m} \cap \hat{\mathcal{C}}_{k,m+1}\} [\hat{V}_{k,h+1}^{2^m} (s_h^k, a_h^k) \\ &\leq \hat{\sigma}_{k,h,m}^{-2} \mathbb{1}\{\theta^* \in \hat{\mathcal{C}}_{k,m} \cap \hat{\mathcal{C}}_{k,m+1}\} \left[\hat{V}_{k,h+1}^{2^m} (s_h^k, a_h^k) \right. \\ &\quad \left. + \min \left\{ 1, \left\| \hat{\phi}_{k,h,m+1} \right\|_{\hat{\Sigma}_{k,m+1}^{-1}} \left\| \hat{\theta}_{k,m+1} - \theta^* \right\|_{\hat{\Sigma}_{k,m+1}} \right\} \right. \\ &\quad \left. + \min \left\{ 1, 2 \left\| \hat{\phi}_{k,h,m} \right\|_{\hat{\Sigma}_{k,m}^{-1}} \left\| \hat{\theta}_{k,m} - \theta^* \right\|_{\hat{\Sigma}_{k,m}} \right\} \right] \\ &\leq \hat{\sigma}_{k,h,m}^{-2} \left[\hat{V}_{k,h+1}^{2^m} (s_h^k, a_h^k) + \min \left\{ 1, \beta_k \left\| \hat{\phi}_{k,h,m+1} \right\|_{\hat{\Sigma}_{k,m+1}^{-1}} \right\} \right. \\ &\quad \left. + \min \left\{ 1, 2\beta_k \left\| \hat{\phi}_{k,h,m} \right\|_{\hat{\Sigma}_{k,m}^{-1}} \right\} \right] \\ &\leq 1,\end{aligned}$$

where the first inequality holds due to Lemma B.2, the second inequality holds due to the definition of the indicator function, and the third inequality holds due to the definition of $\hat{\sigma}_{k,h,m}^{-2}$. For $m = M - 1$, we have $\mathbb{E}[\eta_{k,h,m}^2 | \mathcal{G}_{k,h}] \leq 1$ directly by the definition of $\hat{\sigma}_{k,h,m}^2$. For any $m \in \overline{[M]}$, we have

$$|\eta_{k,h,m}| \max\{1, \|\mathbf{x}_{k,h,m}\|_{\hat{\Sigma}_{k,h-1,m}^{-1}}\} \leq \hat{\sigma}_{k,h,m}^{-2} \|\hat{\phi}_{k,h,m}\|_{\hat{\Sigma}_{k,h-1,m}^{-1}} \leq 1/\gamma^2,$$

where the first inequality follows from the definition of $\eta_{k,h,m}$ and $\mathbf{x}_{k,h,m}$, and the second inequality follows from the definition of $\hat{\sigma}_{k,h,m}$. Let

$$y_{k,h,m} = \langle \mu^*, \mathbf{x}_{k,h,m} \rangle + \eta_{k,h,m},$$

$$\begin{aligned}
 \mathbf{Z}_{k,m} &= \lambda \mathbf{I} + \sum_{i=1}^k \sum_{h=1}^H \mathbf{x}_{i,h,m} \mathbf{x}_{i,h,m}^\top = \hat{\Sigma}_{k+1,m}, \\
 \mathbf{b}_{k,m} &= \sum_{i=1}^k \sum_{h=1}^H \mathbf{x}_{i,h,m} y_{i,h,m}, \\
 \boldsymbol{\mu}_{k,m} &= \mathbf{Z}_{k,m}^{-1} \mathbf{b}_{k,m}, \\
 \varepsilon &= 1/\gamma^2.
 \end{aligned}$$

Then, by Lemma B.1, for each $m \in [\overline{M}]$, with probability at least $1 - \delta$, $\forall k \in [K + 1]$,

$$\begin{aligned}
 \|\boldsymbol{\mu}_{k-1,m} - \boldsymbol{\theta}^*\|_{\hat{\Sigma}_{k,m}} &\leq 12\sqrt{d \log(1 + kH/(\alpha^2 d \lambda)) \log(32(\log(\gamma^2/\alpha) + 1)k^2 H^2/\delta)} \\
 &\quad + 30 \log(32(\log(\gamma^2/\alpha) + 1)k^2 H^2/\delta)/\gamma^2 + \sqrt{\lambda} B \\
 &= \beta_k
 \end{aligned} \tag{B.8}$$

Define the event that (B.8) happens for all k and m as $\hat{\mathcal{E}}$. Conditioned on $\hat{\mathcal{E}}$, the following properties hold:

- For $k = 1$, $m \in [\overline{M}]$, by the definition of $\hat{\boldsymbol{\theta}}_{1,m}$ and $\hat{\Sigma}_{1,m}$, we have $\|\boldsymbol{\theta}^* - \hat{\boldsymbol{\theta}}_{1,m}\|_{\hat{\Sigma}_{1,m}} = \|\boldsymbol{\theta}^*\|_{\lambda \mathbf{I}} \leq \sqrt{\lambda} B = \beta_1$, which implies

$$\boldsymbol{\theta}^* \in \hat{\mathcal{C}}_{1,m} \tag{B.9}$$

- For $k \in [K]$ and $m = M - 1$, we directly have $\boldsymbol{\mu}_{k,M-1} = \hat{\boldsymbol{\theta}}_{k+1,M-1}$, which implies

$$\boldsymbol{\theta}^* \in \hat{\mathcal{C}}_{k+1,M-1}. \tag{B.10}$$

- For $k \in [K]$ and $m \in [\overline{M} - 1]$, we have

$$\boldsymbol{\theta}^* \in \hat{\mathcal{C}}_{k,m} \cap \hat{\mathcal{C}}_{k,m+1} \Rightarrow y_{k,h,m} = \hat{\sigma}^{-1} \hat{V}_{k,h+1}^{2^m}(s_{h+1}^k) \Rightarrow \boldsymbol{\mu}_{k,m} = \hat{\boldsymbol{\theta}}_{k+1,m} \Rightarrow \boldsymbol{\theta}^* \in \hat{\mathcal{C}}_{k+1,m}. \tag{B.11}$$

Therefore, by induction based on initial conditions (B.9) and (B.10), induction rule (B.11), we have for $k \in [K]$ and $m \in [M]$, $\boldsymbol{\theta}^* \in \hat{\mathcal{C}}_{k,m}$. Taking the union bound gives that (A.1) happens with probability at least $1 - M\delta$. We can use the nearly identical argument to prove that (A.2) holds with probability at least $1 - M\delta$. The only difference is to replace $\hat{\sigma}$ with $\tilde{\sigma}$, $\hat{\phi}$ with $\tilde{\phi}$, $\hat{V}$ with $\tilde{V}$, $\hat{\Sigma}$ with $\tilde{\Sigma}$, $\hat{\boldsymbol{\theta}}$ with $\tilde{\boldsymbol{\theta}}$. By taking the union bound, we obtain that with probability at least $1 - 2M\delta$, Equations (A.1) (A.2) both hold.

For (A.3) and (A.4), we have

$$\begin{aligned}
 \|\boldsymbol{\theta}_k - \boldsymbol{\theta}^*\|_{\hat{\Sigma}_{k,0}} &\leq \|\boldsymbol{\theta}_k - \hat{\boldsymbol{\theta}}_{k,m}\|_{\hat{\Sigma}_{k,0}} + \|\hat{\boldsymbol{\theta}}_{k,m} - \boldsymbol{\theta}^*\|_{\hat{\Sigma}_{k,0}} \leq 2\beta_k, \\
 \|\boldsymbol{\theta}_k - \boldsymbol{\theta}^*\|_{\tilde{\Sigma}_{k,0}} &\leq \|\boldsymbol{\theta}_k - \tilde{\boldsymbol{\theta}}_{k,m}\|_{\tilde{\Sigma}_{k,0}} + \|\tilde{\boldsymbol{\theta}}_{k,m} - \boldsymbol{\theta}^*\|_{\tilde{\Sigma}_{k,0}} \leq 2\beta_k
 \end{aligned}$$

□

B.2. Proof of Lemma A.2

Proof of Lemma A.2. We prove this inequality by induction. Suppose

$$\begin{aligned}
 &V_{h+1}(s_{h+1}; \boldsymbol{\theta}_K, \pi, r) - V_{h+1}(s_{h+1}; \boldsymbol{\theta}^*, \pi, r) \\
 &= \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}} W_{h+1}(\{(\mathbb{P}_K - \mathbb{P})V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r)\}),
 \end{aligned} \tag{B.12}$$

which is true for $h = H$. Then, we have

$$V_h(s_h; \boldsymbol{\theta}_K, \pi, r) - V_h(s_h; \boldsymbol{\theta}^*, \pi, r)$$

$$\begin{aligned}
 &= \min \{1, r_h(s_h, a_h) + \mathbb{P}_K V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) - (r_h(s_h, a_h) + \mathbb{P} V_{h+1}(s_h, a_h; \boldsymbol{\theta}^*, \pi, r))\} \\
 &= \min \{1, \mathbb{P}_K V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) - \mathbb{P} V_{h+1}(s_h, a_h; \boldsymbol{\theta}^*, \pi, r)\} \\
 &= \min \{1, (\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) + \mathbb{P}(V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) - V_{h+1}(s_h, a_h; \boldsymbol{\theta}^*, \pi, r))\} \\
 &= \min \{1, (\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) \\
 &\quad + \mathbb{E}_{s_{h+1} \sim \mathbb{P}(\cdot | s_h, a_h)} \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}} W_{h+1}(\{(\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h; \boldsymbol{\theta}_K, \pi, r)\})\} \\
 &= \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} \min \{1, (\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r) + W_{h+1}(\{(\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r)\})\} \\
 &= \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} W_h(\{(\mathbb{P}_K - \mathbb{P}) V_{h+1}(s_h, a_h; \boldsymbol{\theta}_K, \pi, r)\}).
 \end{aligned}$$

The first equality holds due to that $V_h(s_h; \boldsymbol{\theta}_K, \pi, r)$ and $V_h(s_h; \boldsymbol{\theta}^*, \pi, r)$ both belong to $[0, 1]$, the third equality holds due to (B.12), and the forth equality holds due to that $\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} = \mathbb{E}_{s_{h+1} \sim \mathbb{P}(\cdot | s_h, a_h)} \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}}$. Thus, by induction, we obtain the desired result (A.5). $\square$

B.3. Proof of Lemma A.3

Proof of Lemma A.3. We first prove (B.13) by induction.

$$\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_1} W_1(\{u_{K,h}(s_h, \pi(s_h)); \boldsymbol{\theta}_K, \pi, r\}) \leq \widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi, r). \quad (\text{B.13})$$

Suppose

$$\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}} W_{h+1}(\{u_{K,h}(s_h, \pi(s_h)); \boldsymbol{\theta}_K, \pi, r\}) \leq \widehat{V}_{K,h+1}(s_1; \boldsymbol{\theta}_K, \pi, r), \quad (\text{B.14})$$

which is true for $h = H$. Then,

$$\begin{aligned}
 &\widehat{V}_{K,h}(s_h; \boldsymbol{\theta}_K, \pi, r) - \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} W_h(\{u_{K,h}(s_h, \pi(s_h)); \boldsymbol{\theta}_K, \pi, r\}) \\
 &\geq \min \left\{ 0, u_{K,h}(s_h, \pi(s_h); \boldsymbol{\theta}_K, \pi, r) + 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} \right. \\
 &\quad \left. + \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}^\top(s_h, \pi(s_h)) \boldsymbol{\theta}_K - \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} W_h(\{u_{K,h}(s_h, \pi(s_h)); \boldsymbol{\theta}_K, \pi, r\}) \right\} \\
 &\geq \min \left\{ 0, u_{K,h}(s_h, \pi(s_h); \boldsymbol{\theta}_K, \pi, r) + 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} \right. \\
 &\quad \left. + \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}^\top(s_h, \pi(s_h)) \boldsymbol{\theta}_K - u_{K,h}(s_h, \pi(s_h); \boldsymbol{\theta}_K, \pi, r) \right. \\
 &\quad \left. - \mathbb{E}_{s_{h+1} \sim \mathbb{P}(\cdot | s_h, \pi(s_h))} \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}} W_{h+1}(\{u_{K,h}(s_h, \pi(s_h)); \boldsymbol{\theta}_K, \pi, r\}) \right\} \\
 &\geq \min \left\{ 0, 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} + \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}^\top(s_h, \pi(s_h)) \boldsymbol{\theta}_K \right. \\
 &\quad \left. - \mathbb{E}_{s_{h+1} \sim \mathbb{P}(\cdot | s_h, \pi(s_h))} \widehat{V}_{K,h+1}(s_{h+1}; \boldsymbol{\theta}_K, \pi, r) \right\} \\
 &\geq \min \left\{ 0, 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} + \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}^\top(s_h, \pi(s_h)) (\boldsymbol{\theta}_K - \boldsymbol{\theta}^*) \right\} \\
 &\geq \min \left\{ 0, 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} - 2\beta \left\| \phi_{\widehat{V}_{K,h+1}(\cdot; \boldsymbol{\theta}_K, \pi, r)}(s_h, \pi(s_h)) \right\|_{\dot{\Sigma}_{K,0}} \right\} \\
 &\geq 0,
 \end{aligned}$$

where the first inequality holds due to the definition of $\widehat{V}_{K,h}$, the second inequality holds due to the definition of $W_h(\cdot)$ and $\mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_h} = \mathbb{E}_{s_{h+1} \sim \mathbb{P}(\cdot | s_h, \pi(s_h))} \mathbb{E}_{\text{traj} \sim (\pi, \mathbb{P}) | \text{traj}_{h+1}}$, the third inequality holds due to B.14, the fifth inequality holds due to Lemma A.1. Thus, by induction, B.13 holds. Thanks to the optimism of $\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K)$, we have

$$\widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi, r) \leq \widehat{V}_{K,1}(s_1; \boldsymbol{\theta}_K, \pi_K, r_K),$$

which concludes the proof. $\square$

B.4. Proof of Lemma A.4

Proof of Lemma A.4. For the equation (A.8), we have

$$\begin{aligned}
 & \widehat{V}_{k,h}(s_h^k) - u_{k,h}(s_h^k, a_h^k) - \mathbb{P}\widehat{V}_{k,h+1}(s_h^k, a_h^k) \\
 & \leq \min \left\{ 1, 2\beta \left\| \widehat{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} + \widehat{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \boldsymbol{\theta}_k - \widehat{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \boldsymbol{\theta} \right\} \\
 & = \min \left\{ 1, 2\beta \left\| \widehat{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} + \widehat{\phi}_{k,h,0}^\top(s_h^k, a_h^k) (\boldsymbol{\theta}_k - \boldsymbol{\theta}) \right\} \\
 & \leq \min \left\{ 1, 2\beta \left\| \widehat{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} + \left\| \widehat{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} \|\boldsymbol{\theta}_k - \boldsymbol{\theta}\|_{\widehat{\Sigma}_{k,0}} \right\} \\
 & \leq \min \left\{ 1, 4\beta \left\| \widehat{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} \right\} \\
 & \leq 4 \min \left\{ 1, \beta \left\| \widehat{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widehat{\Sigma}_{k,0}^{-1}} \right\}
 \end{aligned}$$

where the first inequality holds due to that each term lies in the interval $[0, 1]$, the second inequality holds due to Cauchy-Schwartz inequality, and the third inequality holds due to lemma A.1. For the equation (A.9), we have

$$\begin{aligned}
 & \widetilde{V}_{k,h}(s_h^k) - r_{k,h}(s_h^k, a_h^k) - \mathbb{P}\widetilde{V}_{k,h+1}(s_h^k, a_h^k) \\
 & \leq \min \left\{ 1, \widetilde{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \boldsymbol{\theta}_k - \widetilde{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \boldsymbol{\theta} \right\} \\
 & = \min \left\{ 1, \widetilde{\phi}_{k,h,0}^\top(s_h^k, a_h^k) (\boldsymbol{\theta}_k - \boldsymbol{\theta}) \right\} \\
 & \leq \min \left\{ 1, \left\| \widetilde{\phi}_{k,h,0}^\top(s_h^k, a_h^k) \right\|_{\widetilde{\Sigma}_{k,0}^{-1}} \|\boldsymbol{\theta}_k - \boldsymbol{\theta}\|_{\widetilde{\Sigma}_{k,0}} \right\} \\
 & \leq \min \left\{ 1, 2\beta \left\| \widetilde{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widetilde{\Sigma}_{k,0}^{-1}} \right\} \\
 & \leq 2 \min \left\{ 1, \beta \left\| \widetilde{\phi}_{k,h,0}(s_h^k, a_h^k) \right\|_{\widetilde{\Sigma}_{k,0}^{-1}} \right\},
 \end{aligned}$$

where the first inequality holds due to that each term lies in the interval $[0, 1]$, the second inequality holds due to the Cauchy-Schwartz inequality, and the third inequality holds due to Lemma A.1. $\square$

B.5. Proof of Lemma A.5

Lemma B.3 (Lemma B.1, Zhou & Gu (2022)). Let $\{\sigma_k, \beta_k\}_{k \geq 1}$ be a sequence of non-negative numbers, $\alpha, \gamma > 0$, $\{\mathbf{x}_k\}_{k \geq 1} \subset \mathbb{R}^d$ and $\|\mathbf{x}_k\|_2 \leq L$. Let $\{\mathbf{Z}_k\}_{k \geq 1}$ and $\{\bar{\sigma}_k\}_{k \geq 1}$ be recursively defined as follows: $\mathbf{Z}_1 = \lambda \mathbf{I}$

$$\forall k \geq 1, \bar{\sigma}_k = \max \left\{ \sigma_k, \alpha, \gamma \|\mathbf{x}_k\|_{\mathbf{Z}_k^{-1}}^{1/2} \right\}, \mathbf{Z}_{k+1} = \mathbf{Z}_k + \mathbf{x}_k \mathbf{x}_k^\top / \bar{\sigma}_k^2.$$

Let $\iota = \log(1 + KL^2 / (d\lambda\alpha^2))$. Then we have

$$\sum_{k=1}^K \min \left\{ 1, \beta_k \|\mathbf{x}_k\|_{\mathbf{Z}_k^{-1}} \right\} \leq 2d\iota + 2 \max_{k \in [K]} \beta_k \gamma^2 d\iota + 2\sqrt{d\iota} \sqrt{\sum_{k=1}^K \beta_k^2 (\sigma_k^2 + \alpha^2)}.$$

Proof of Lemma A.5. The proof is nearly identical to the proof of Lemma C.5 in Zhou & Gu (2022). The only difference is to replace $\widehat{\Sigma}_{k,m}$ with $\widehat{\Sigma}_{k,m}$ (or $\widetilde{\Sigma}_{k,m}$), $\widetilde{\Sigma}_{k,h,m}$ with $\widehat{\Sigma}_{k,h,m}$ (or still $\widetilde{\Sigma}_{k,h,m}$), $\phi_{k,h,m}$ with $\widehat{\phi}_{k,h,m}$ (or $\widetilde{\phi}_{k,h,m}$). $\square$

B.6. Proof of Lemma A.6

Proof of Lemma A.6. The proof follows the proof of Lemma 25 in Zhang et al. (2021c) and Lemma C.6 in Zhou & Gu (2022). We have,

$$\begin{aligned}
 \widehat{S}_m &= \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^{m+1}} \right] (s_h^k, a_h^k) - \left(\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) \right)^2 \right] \\
 &= \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^{m+1}} \right] (s_h^k, a_h^k) - \widehat{V}_{k,h+1}^{2^{m+1}} (s_{h+1}^k) \right] + \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{k,h}^{2^{m+1}} (s_h^k) \right. \\
 &\quad \left. - \left(\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) \right)^2 \right] + \sum_{k=1}^K \sum_{h=1}^H I_h^k \left(\widehat{V}_{k,h+1}^{2^{m+1}} (s_{h+1}^k) - \widehat{V}_{k,h}^{2^{m+1}} (s_h^k) \right) \\
 &\leq \widehat{A}_{m+1} + \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{k,h}^{2^{m+1}} (s_h^k) - \left(\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) \right)^2 \right] + \sum_{k=1}^K I_{h_k}^k \widehat{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k), \tag{B.15}
 \end{aligned}$$

where h_k is the largest index satisfying $I_h^k = 1$. For the second term, we have

$$\begin{aligned}
 &\sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{k,h}^{2^{m+1}} (s_h^k) - \left(\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) \right)^2 \right] \\
 &\leq \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{k,h}^{2^{m+1}} (s_h^k) - \left(\left[\mathbb{P} \widehat{V}_{k,h+1} \right] (s_h^k, a_h^k) \right)^{2^{m+1}} \right] \\
 &= \sum_{k=1}^K \sum_{h=1}^H I_h^k \left(\widehat{V}_{k,h} (s_h^k) - \left[\mathbb{P} \widehat{V}_{k,h+1} \right] (s_h^k, a_h^k) \right) \prod_{i=0}^m \left(\widehat{V}_{k,h}^{2^i} (s_h^k) + \left[\mathbb{P} \widehat{V}_{k,h+1} \right] (s_h^k, a_h^k)^{2^i} \right) \\
 &\leq 2^{m+1} \sum_{k=1}^K \sum_{h=1}^H I_h^k \max \left\{ \widehat{V}_{k,h} (s_h^k) - \left[\mathbb{P} \widehat{V}_{k,h+1} \right] (s_h^k, a_h^k), 0 \right\} \\
 &\leq 2^{m+1} \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[u_{k,h} (s_h^k, a_h^k) + 4 \min \left\{ 1, \beta \left\| \widehat{\phi}_{k,h,0} \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right\} \right] \\
 &\leq 2^{m+1} \left(\widetilde{R}_0 + 4\widehat{R}_0 \right), \tag{B.16}
 \end{aligned}$$

where the first inequality holds due to using $\mathbb{E}X^2 \geq (\mathbb{E}X)^2$ recursively, the first equality holds due to the fact $x^{2^{m+1}} - y^{2^{m+1}} = (x - y) \prod_{i=0}^m (x^{2^i} + y^{2^i})$, the second inequality holds due to $\widehat{V}_{k,h}$ belongs to the interval $[0, 1]$, the third inequality holds due to Lemma A.4, and the last inequality holds due to $u_{k,h}(s_h^k, a_h^k) = \beta \left\| \phi_{V_{h+1}(\cdot; \theta_k, \pi_k, r_k)}(s_h^k, a_h^k) \right\|_{\dot{\Sigma}_{k,0}} = \beta \left\| \widetilde{\phi}_{k,h,0} \right\|_{\dot{\Sigma}_{k,0}}$. If $h_K \leq H$, we have $I_{h_k}^k \widehat{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k) \leq 1 = 1 - I_H^k$, and if $h_K = H$, $I_{h_k}^k \widehat{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k) = 0 = 1 - I_H^k$, which both give

$$\sum_{k=1}^K I_{h_k}^k \widehat{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k) \leq \sum_{k=1}^K (1 - I_H^k) = G \tag{B.17}$$

Substituting Equations (B.15), (B.16), (B.17) into (B.15), we can get (A.20). For Equation (A.15), similarly, we have

$$\widetilde{S}_m \leq \widetilde{A}_{m+1} + \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widetilde{V}_{k,h}^{2^{m+1}} (s_h^k) - \left(\left[\mathbb{P} \widetilde{V}_{k,h+1}^{2^m} \right] (s_h^k, a_h^k) \right)^2 \right] + \sum_{k=1}^K I_{h_k}^k \widetilde{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k), \tag{B.18}$$

$$\sum_{k=1}^K I_{h_k}^k \widetilde{V}_{k,h_k+1}^{2^{m+1}} (s_{h_k+1}^k) \leq \sum_{k=1}^K (1 - I_H^k) = G. \tag{B.19}$$

And we have

$$\begin{aligned}
 & \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\widehat{V}_{k,h}^{2^{m+1}}(s_h^k) - \left(\mathbb{P} \widehat{V}_{k,h+1}^{2^m}(s_h^k, a_h^k) \right)^2 \right] \\
 & \leq 2^{m+1} \sum_{k=1}^K \sum_{h=1}^H I_h^k \max \left\{ \widetilde{V}_{k,h}(s_h^k) - \left[\mathbb{P} \widetilde{V}_{k,h+1} \right](s_h^k, a_h^k), 0 \right\} \\
 & \leq 2^{m+1} \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[r_{k,h}(s_h^k, s_h^k) + \min \left\{ 1, 2\beta \left\| \widetilde{\phi}_{k,h,0} \right\|_{\dot{\Sigma}_{k,0}^{-1}} \right\} \right] \\
 & \leq 2^{m+1} (K + 2\widetilde{R}_0)
 \end{aligned} \tag{B.20}$$

where the first inequality holds similar to the derivation of (B.16), second inequality follows Lemma A.4, and the third inequality holds due to $\sum_{h=1}^H r_{k,h}(s_h^k, a_h^k) \leq 1$. Plugging Equations (B.19) (B.20) into B.18, we obtain A.21 $\square$

B.7. Proof of Lemma A.7

Proof of Lemma A.7. The proof follows the proof of Lemma 25 in Zhang et al. (2021c) and Lemma C.7 in Zhou & Gu (2022). We use Lemma D.2 for $\widehat{A}_m$ and $\widetilde{A}_m$ for each m . To avoid confusion, we write ϵ, δ in Lemma D.2 as ϵ', δ' .

Let $x_{k,h} = I_h^k \left[\left[\mathbb{P} \widehat{V}_{k,h+1}^{2^m} \right](s_h^k, a_h^k) - \widehat{V}_{k,h+1}^{2^m}(s_{h+1}^k) \right]$, $n = KH$, $\epsilon' = \sqrt{\log(1/\delta')}$, and $\delta' = \delta/(4 \log(KH))$. Thus, $\mathbb{E}[\widehat{x}_{k,h} | \mathcal{F}_{k,h}] = 0$ and $\mathbb{E}[\widehat{x}_{k,h}^2 | \mathcal{F}_{k,h}] = I_h^k \left[\mathbb{V} \widehat{V}_{k,h+1}^{2^m} \right](s_h^k, a_h^k)$. Therefore, for each $m \in \overline{[M]}$, with probability at least $1 - \delta$, we have

$$|\widehat{A}_m| = \left| \sum_{k=1}^K \sum_{h=1}^H x_{k,h} \right| \leq \sqrt{2\zeta \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\mathbb{V} \widehat{V}_{k,h+1}^{2^m} \right](s_h^k, a_h^k)} + \zeta.$$

Similarly, let $x_{k,h} = I_h^k \left[\left[\mathbb{P} \widetilde{V}_{k,h+1}^{2^m} \right](s_h^k, a_h^k) - \widetilde{V}_{k,h+1}^{2^m}(s_{h+1}^k) \right]$, $n = KH$, $\epsilon' = \sqrt{\log(1/\delta')}$, and $\delta' = \delta/(4 \log(KH))$. With probability at least $1 - \delta$, we have

$$|\widetilde{A}_m| = \left| \sum_{k=1}^K \sum_{h=1}^H x_{k,h} \right| \leq \sqrt{2\zeta \sum_{k=1}^K \sum_{h=1}^H I_h^k \left[\mathbb{V} \widetilde{V}_{k,h+1}^{2^m} \right](s_h^k, a_h^k)} + \zeta.$$

Taking union bound over $m \in \overline{[M]}$ completes the proof. $\square$

B.8. Proof of Lemma A.8

Proof of Lemma A.8. By the fact that $\det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) < \det \left(\widehat{\Sigma}_{k,H,m}^{-1/2} \right)$ and $\det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) < \det \left(\widetilde{\Sigma}_{k,H,m}^{-1/2} \right)$, we have

$$\begin{aligned}
 (1 - I_H^k) = 1 & \Leftrightarrow \exists m \in \overline{[M]}, \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) / \det \left(\widehat{\Sigma}_{k,H,m}^{-1/2} \right) > 4 \text{ or } \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) / \det \left(\widetilde{\Sigma}_{k,H,m}^{-1/2} \right) > 4 \\
 & \Rightarrow \exists m \in \overline{[M]}, \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) / \det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) > 4 \text{ or } \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) / \det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) > 4
 \end{aligned} \tag{B.21}$$

Let $\widehat{\mathcal{D}}_m$ and $\widetilde{\mathcal{D}}_m$ denote the indices k such that

$$\begin{aligned}
 \widehat{\mathcal{D}}_m &:= \left\{ k \in [K] : \det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) / \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) > 16 \right\} \\
 \widetilde{\mathcal{D}}_m &:= \left\{ k \in [K] : \det \left(\dot{\Sigma}_{k+1,m}^{-1/2} \right) / \det \left(\dot{\Sigma}_{k,m}^{-1/2} \right) > 16 \right\}
 \end{aligned}$$

Then we have

$$G \leq \left| \bigcup_{m=0}^{M-1} \widehat{\mathcal{D}}_m \cup \bigcup_{m=0}^{M-1} \widetilde{\mathcal{D}}_m \right| \leq \sum_{m=0}^{M-1} |\widehat{\mathcal{D}}_m| + \sum_{m=0}^{M-1} |\widetilde{\mathcal{D}}_m|$$

For each m , we have

$$2|\widehat{\mathcal{D}}_m| < \sum_{k \in \widehat{\mathcal{D}}_m} \log 16 < \sum_{k \in \widehat{\mathcal{D}}_m} \log \left(\det \left(\dot{\widehat{\Sigma}}_{k+1,m} \right) / \det \left(\dot{\widehat{\Sigma}}_{k,m} \right) \right) \leq \sum_{k=1}^K \log \left(\det \left(\dot{\widehat{\Sigma}}_{k+1,m} \right) / \det \left(\dot{\widehat{\Sigma}}_{k,m} \right) \right)$$

Furthermore, since $\det \left(\dot{\widehat{\Sigma}}_{K+1,m} \right) \leq \left(\text{tr} \left(\dot{\widehat{\Sigma}}_{K+1,m} \right) / d \right)^d$ and $\text{tr} \left(\dot{\widehat{\Sigma}}_{K+1,m} \right) \leq \text{tr} (\lambda I) +$

$$\sum_{k,h} \left\| \widehat{\phi}_{k,h,m} \right\|_2^2 / \widehat{\sigma}_{k,h,m}^2 \leq d\lambda + KH/\alpha^2$$

$$\begin{aligned} \sum_{k=1}^K \log \left(\det \left(\dot{\widehat{\Sigma}}_{k+1,m} \right) / \det \left(\dot{\widehat{\Sigma}}_{k,m} \right) \right) &= \log \left(\det \left(\dot{\widehat{\Sigma}}_{K+1,m} \right) / \det \left(\dot{\widehat{\Sigma}}_{1,m} \right) \right) \\ &\leq d \left(\log \left(\lambda + KH/(d\alpha^2) \right) - \log(\lambda) \right) \end{aligned}$$

Therefore $|\widehat{\mathcal{D}}_m|$ is upper bounded by

$$|\widehat{\mathcal{D}}_m| < d/2 \log(1 + KH/(d\lambda\alpha^2)).$$

And same for $|\widetilde{\mathcal{D}}_m|$. Taking summation gives the upper bound of G . $\square$

C. Proof of Lower Bound

Reward-free exploration is more difficult than non-reward-free MDP by definitions since we can easily solve non-reward-free MDP by ignoring its reward and executing reward-free exploration. Thus, we will start with acquiring lower bounds under non-reward-free MDP settings and then obtain sample complexity lower bounds of reward-free exploration. The proof follows ideas of Zhou & Gu (2022) and Chen et al. (2021).

As noted in Section 5.2, we will consider the *hard-to-learn linear mixture MDPs* constructed in Zhou & Gu (2022). The state space $\mathcal{S} = \{x_1, x_2, x_3\}$ and the action space $\mathcal{A} = \{\mathbf{a}\} = \{-1, 1\}^{d-1}$. The reward function satisfies $r(x_1, \cdot) = r(x_2, \cdot) = 0$, and $r(x_3, \cdot) = \frac{1}{H}$. The transition probability satisfies $\mathbb{P}(x_2 | x_1, \mathbf{a}) = 1 - (\delta + \langle \boldsymbol{\mu}, \mathbf{a} \rangle)$ and $\mathbb{P}(x_3 | x_1, \mathbf{a}) = \delta + \langle \boldsymbol{\mu}, \mathbf{a} \rangle$, where $\delta = 1/6$ and $\boldsymbol{\mu} \in \{-\Delta, \Delta\}^{d-1}$ with $\Delta = \sqrt{\delta/K'}/(4\sqrt{2})$. The transition kernel is formulated as

$$\phi(s' | s, \mathbf{a}) = \begin{cases} (\alpha(1-\delta), -\beta\mathbf{a}^\top)^\top, & s = x_1, s' = x_2; \\ (\alpha\delta, \beta\mathbf{a}^\top)^\top, & s = x_1, s' = x_3; \\ (\alpha, \mathbf{0}^\top)^\top, & s \in \{x_2, x_3\}, s' = s; \\ \mathbf{0}, & \text{otherwise.} \end{cases}$$

$$\boldsymbol{\theta} = (1/\alpha, \boldsymbol{\mu}^\top/\beta)^\top.$$

The following lemma from Zhou & Gu (2022) lower bounds the regret for linear mixture MDP.

Lemma C.1. (Theorem 5.4 in Zhou & Gu (2022)) Let $B > 1$. Then for any algorithm, when $K' \geq \max \{3d^2, (d-1)/(192(B-1))\}$, there exists a B -bounded linear mixture MDP satisfying Assumptions 3.2 such that its expected regret $\mathbb{E}[\text{Regret}(K')]$ is lower bounded by $\Omega \left(d\sqrt{K'}/(16\sqrt{3}) \right)$.

Given Lemma C.1, we will use the regret lower bound of non-reward-free linear mixture MDPs to derive the sample coomplexity lower bound.

Lemma C.2. Suppose $B > 1$. Then for any algorithm $\text{ALG}_{\text{NonFree}}$ solving non-reward-free linear mixture MDP problems satisfying assumption 3.2, there exist a linear mixture $\mathcal{M}$ such that $\text{ALG}_{\text{NonFree}}$ needs to collect at least $\frac{Cd^2}{\varepsilon^2}$ episodes to output an ε -policy with probability at least $1 - \delta$. Here C is an absolute constant.

Proof of Lemma C.2. For any algorithm $\text{ALG}_{\text{NonFree}}$, we construct an algorithm $\text{ALG}'_{\text{NonFree}}$ executing totally $K_1 = cK$ episodes, where c is a constant integer larger than 1. The first K episodes of $\text{ALG}'_{\text{NonFree}}$ are the same as $\text{ALG}_{\text{NonFree}}$, and the rest episodes keep executing the policy at the end of episode K . By Lemma C.1, we have

$$\sum_{k=1}^{K_1} \mathbb{E} [V(s_1; \theta^*, \pi^*, r) - V(s_1; \theta^*, \pi_k, r)] \geq \frac{c'd\sqrt{K_1}}{16\sqrt{3}}, \quad (\text{C.1})$$

for some constant c' . In addition, based on the construction of *the hard-to-learn MDPs*, where $K' = K_1$, the per-episode regret is upper bounded by

$$\mathbb{E} [V(s_1; \theta^*, \pi^*, r) - V(s_1; \theta^*, \pi_k, r)] \leq \frac{d}{4\sqrt{3K_1}}. \quad (\text{C.2})$$

Thus, calculating (C.1) - $(K_1 - K) \times (\text{C.2})$, and choosing $c = \max \{5/c', 2\}$, we have

$$\sum_{k=K+1}^{K_1} \mathbb{E} [V(s_1; \theta^*, \pi^*, r) - V(s_1; \theta^*, \pi_k, r)] \geq \frac{d\sqrt{K}}{16\sqrt{3}c}.$$

Since the policies in episode $K + 1$ to episode K_1 are same to π_K , we have

$$\mathbb{E} [V(s_1; \theta^*, \pi^*, r) - V(s_1; \theta^*, \pi_K, r)] \geq \frac{d}{16\sqrt{3cKc}}.$$

Suppose the $\text{ALG}_{\text{NonFree}}$ return a ε -optimal policy with probability $1 - \delta$. Then,

$$(1 - \delta)\varepsilon + \delta \frac{d}{4\sqrt{3cK}} \geq \frac{d}{16\sqrt{3cKc}}.$$

Setting $\delta < \min\{1, 1/(4c)\}$, by solving the inequality, we have $K \geq \frac{Cd^2}{\varepsilon^2}$ for some constant C . $\square$

Since reward-free MDP is more difficult than non-reward-free MDP, Lemma C.2 directly indicates Theorem 5.6.

Proof of Theorem 5.6. We will prove the theorem by contradiction. Assume all reward-free linear mixture MDPs can be solved with sample complexity of $o(\frac{d^2}{\varepsilon^2})$. Then, for any non-reward-free MDP $\mathcal{M}$, there exists an algorithm $\text{ALG}'(\varepsilon, \delta)$ learning its reward-free counterpart $\mathcal{M}'$ with sample complexity of $o(\frac{d^2}{\varepsilon^2})$. We define ALG solving $\mathcal{M}$ as follows: it collects K episodes of data and outputs the policy in the same way as ALG' by ignoring the rewards. Then ALG can also (ε, δ) learning $\mathcal{M}$ with sample complexity of $o(\frac{d^2}{\varepsilon^2})$, which contradicts Theorem C.2. $\square$

Corollary 5.8 can be viewed as an direct result of Theorem 5.6.

Proof of Corollary 5.8. The hard-to-learn case we consider here is basically same as we consider in Theorem 5.6, except replacing reward function with $r(x_1, \cdot) = r(x_2, \cdot) = 0$, and $r(x_3, \cdot) = 1$. Since the reward here is H times the reward in Theorem 5.6, the suboptimality is also H times. Therefore, ε/H -optimal policy in Theorem 5.6 is a ε -optimal policy here. According to Theorem 5.6, the sample complexity required to achieve such a policy with probability at least $1 - \delta$ is $\Omega(\frac{H^2 d^2}{\varepsilon^2})$. $\square$

D. Auxiliary Lemmas

Lemma D.1. Suppose $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{d \times d}$ are two positive definite matrices satisfying $\mathbf{A} \succeq \mathbf{B}$, then for any $\mathbf{x} \in \mathbb{R}^d$, $\|\mathbf{x}\|_{\mathbf{A}} \leq \|\mathbf{x}\|_{\mathbf{B}} \sqrt{\det(\mathbf{A}) / \det(\mathbf{B})}$

Lemma D.2 (Lemma 11, Zhang et al. 2021d). Let $M > 0$ be a constant. Let $\{x_i\}_{i=1}^n$ be a stochastic process, $\mathcal{G}_i = \sigma(x_1, \dots, x_i)$ be the σ -algebra of $x_1, \dots, x_i$. Suppose $\mathbb{E}[x_i | \mathcal{G}_{i-1}] = 0$, $|x_i| \leq M$ and $\mathbb{E}[x_i^2 | \mathcal{G}_{i-1}] < \infty$ almost surely. Then, for any $\delta, \varepsilon > 0$, we have

$$\begin{aligned} \mathbb{P} \left(\left| \sum_{i=1}^n x_i \right| \leq 2 \sqrt{2 \log(1/\delta) \sum_{i=1}^n \mathbb{E}[x_i^2 | \mathcal{G}_{i-1}] + 2 \sqrt{\log(1/\delta) \varepsilon} + 2M \log(1/\delta)} \right) \\ > 1 - 2 \left(\log(M^2 n / \varepsilon^2) + 1 \right) \delta \end{aligned} \quad (\text{D.1})$$

Lemma D.3. (Lemma 12 in Zhang et al. (2021c)) Let $\lambda_1, \lambda_2, \lambda_4 > 0, \lambda_3 \geq 1$ and $\kappa = \max\{\log_2 \lambda_1, 1\}$. Let $a_1, \dots, a_\kappa$ be non-negative real numbers such that $a_i \leq \min\left\{\lambda_1, \lambda_2 \sqrt{a_i + a_{i+1} + 2^{i+1} \lambda_3} + \lambda_4\right\}$ for any $1 \leq i \leq \kappa$. Let $a_{\kappa+1} = \lambda_1$. Then we have $a_1 \leq 22\lambda_2^2 + 6\lambda_4 + 4\lambda_2 \sqrt{2\lambda_3}$.