
Optimistic Actor-Critic with Parametric Policies: Unifying Sample Efficiency and Practicality

Max Qiushi Lin
Simon Fraser University
maxqslin@gmail.com

Reza Asad
Simon Fraser University
rasad@sfu.ca

Kevin Tan
University of Pennsylvania
kevtan@umich.edu

Haque Ishfaq
Mila and McGill University
haque.ishfaq@mail.mcgill.ca

Csaba Szepesvári
Google DeepMind and University of Alberta
szepi@google.com

Sharan Vaswani
Simon Fraser University
vaswani.sharan@gmail.com

Abstract

Although actor-critic (AC) methods have been successful in practice, their theoretical analyses have several limitations. Specifically, existing theoretical work either sidesteps the exploration problem by making strong assumptions or analyzes impractical methods that require complicated algorithmic modifications. Moreover, the AC methods analyzed for finite-horizon MDPs often construct “implicit” policies without explicitly parameterizing them, further exacerbating the mismatch between theory and practice. To that end, we propose an optimistic AC framework with parametric policies that is both practical and equipped with theoretical guarantees for episodic linear MDPs. In particular, we introduce a tractable regression objective for the actor to train log-linear policies. This enables us to control the error between the parameterized actor and the easier-to-analyze implicit policies induced by natural policy gradient. To train the critic, we use approximate Thompson sampling via Langevin Monte Carlo to obtain optimistic value estimates. This results in a principled, yet flexible exploration scheme without any additional assumptions on the MDP. We prove that our algorithm achieves an $\tilde{O}(\epsilon^{-4})$ sample complexity in the on-policy setting and an $\tilde{O}(\epsilon^{-2})$ complexity in the off-policy setting. Our algorithm matches prior theoretical work in achieving state-of-the-art sample efficiency, while being more aligned with practice.

1 Introduction

Reinforcement learning (RL) is a general framework for sequential decision making under uncertainty and has been successful in various real-world applications, such as robotics [Kober et al., 2013] and aligning language models [Uc-Cetina et al., 2023]. Policy Gradient (PG) methods [Williams, 1992, Sutton et al., 1999, Kakade, 2001, Schulman et al., 2017a] are an important class of algorithms that assume a differentiable parameterization of the policy, and directly optimize the policy parameters using the return from interacting with the environment. PG methods are widely used in practice as they can easily handle function approximation or structured state-action spaces. However, since the environment is typically stochastic in practice, the estimated returns usually have high variance, resulting in poor sample efficiency [Dulac-Arnold et al., 2019].

Actor-critic (AC) methods [Konda and Tsitsiklis, 1999, Peters et al., 2005, Bhatnagar et al., 2009] alleviate this issue by using value-based approaches in conjunction with PG methods. In particular, they utilize a critic that estimates the policy’s value and an actor that performs PG to improve the policy towards obtaining higher returns. These AC methods have been proven to be empirically successful in both on-policy [Schulman et al., 2015, 2017b] and off-policy [Lillicrap et al., 2015, Fujimoto et al., 2018, Haarnoja et al., 2018] settings.

Subsequently, there have been many attempts to provide a theoretical understanding of actor-critic methods, especially in the presence of function approximation [Cai et al., 2020, Zhong and Zhang, 2023, Liu et al., 2023]. However, there are two prevalent issues that result in mismatches between theory and practice: the studied methods either (i) do not consider strategic exploration in a systematic manner or (ii) analyze complicated and impractical variants of the algorithm. In particular, much of the literature makes unrealistic assumptions to avoid dealing with exploration, a central challenge in RL. For instance, existing work on PG methods [Agarwal et al., 2021a, Yuan et al., 2023, Alfano and Rebeschini, 2022, Asad et al., 2025] obtains convergence rates that involve a mismatch ratio between the optimal policy and the initial state distribution. These results are only meaningful if the mismatch ratio is bounded. However, a bounded mismatch ratio indicates that the initial state distribution already provides a good coverage over the state space, thereby sidestepping the exploration problem. Within actor-critic methods, some early analyses make assumptions on the reachability of the state-action space or the coverage of collected data [Abbasi-Yadkori et al., 2019, Neu et al., 2017, Bhandari and Russo, 2024, Agarwal et al., 2021a, Cen et al., 2022, Gaur et al., 2024], which again imply that the state-action space is already relatively easy to explore. Follow-up work [Hong et al., 2023, Fu et al., 2021, Xu et al., 2020, Cayci et al., 2024] assumes a bounded mismatch ratio, while others [Khodadadian et al., 2022, Gaur et al., 2023] require mixing assumptions on the induced Markov chain.

On the other hand, recent work [Cai et al., 2020, Jin et al., 2021, Zanette et al., 2021, Zhong and Zhang, 2023, Agarwal et al., 2023, He et al., 2023, Liu et al., 2023, Sherman et al., 2023, Cassel and Rosenberg, 2024, Tan et al., 2025] tackles the exploration issue directly. However, the algorithms analyzed are significantly different from those implemented in practice. Much of this body of work studies AC methods that use the natural policy gradient (NPG) update for policy optimization. However, the canonical implementation of the NPG update does not consider an explicit policy parameterization. Instead, the update involves constructing “implicit” policies on the fly using all previously stored Q -functions. Consequently, this implementation has a memory complexity that is linear in the number of updates. This drastically deviates from practice, where algorithms typically employ explicitly parameterized complex models as learnable policies and optimize them with gradient descent-based methods. Furthermore, in the off-policy setting, these works require complicated algorithmic modifications, further exacerbating the mismatch between theory and practice. For example, Sherman et al. [2023] adopts a warm-up procedure from Wagenmaker et al. [2022] that does not resemble policy optimization and is difficult to implement in practice. Although Cassel and Rosenberg [2024] can avoid the warm-up phase, it still requires feature contraction techniques that are non-standard in practice, and difficult to extend beyond linear function approximation. The issues above indicate a significant gap between theory and practice for AC methods. Thus, we address the following question:

Can we design a provably sample-efficient, yet practical, actor-critic algorithm with parametric policies for both the on- and off-policy regimes?

Contributions We answer the above question affirmatively, and make the following contributions.

1. General framework with an explicitly parameterized actor. In Section 3, we propose a general optimistic actor-critic framework that employs an explicitly parameterized policy. We analyze this framework in the setting of linear function approximation for both the environment (i.e., linear MDP [Jin et al., 2020]) and the policy (i.e., log-linear policy class). In Section 4, we propose an actor algorithm that learns a log-linear policy by solving a specific regression problem at each iteration. This allows us to directly control the error between the explicitly parametrized policy and the implicit policy induced by NPG. Using this error bound in conjunction with the well-established theoretical results of NPG [Hazan et al., 2016, Szepesvári, 2022] enables us to analyze the performance of the parameterized actor. We show that the proposed algorithm benefits from a substantially improved memory complexity, while retaining similar theoretical guarantees.

2. LMC critic for practical strategic exploration. In Section 5, instead of constructing UCB bonuses [Jin et al., 2020], which are ubiquitous within prior work [Cassel and Rosenberg, 2024, Sherman et al., 2023, Liu et al., 2023, Zhong and Zhang, 2023], we adopt a more practical approach. We employ Langevin Monte Carlo (LMC) [Welling and Teh, 2011] to update the critic parameters at each episode. Unlike UCB-based approaches that require computing confidence sets at every episode, LMC simply perturbs (by Gaussian noise) the gradient descent update on the critic loss. This gradient descent-based approach is both easier to implement [Ishfaq et al., 2025] and to extend to general function approximation [Ishfaq et al., 2024b]. Furthermore, the LMC algorithm directly leads to an optimistic estimate of the Q -function that has similar guarantees as UCB bonuses. Nevertheless, previous works have only successfully designed provably efficient algorithms for solving multi-armed bandits Mazumdar et al. [2020], contextual bandits [Xu et al., 2022], and linear MDPs via value-based methods [Ishfaq et al., 2024a]. Our paper is the first to analyze an LMC based approach in the context of policy optimization.

3. End-to-end theoretical guarantees for actor-critic. In Section 6, we analyze the proposed actor-critic framework in both the on-policy and off-policy settings without making any assumptions on the mismatch ratio or data coverage. In particular, in the on-policy setting, we prove that our method requires $\tilde{O}(1/\epsilon^4)$ samples to learn an ϵ -optimal policy. This matches the result in [Liu et al., 2023] that uses an implicit NPG policy in conjunction with UCB bonuses. On the other hand, we also prove that our framework can attain a sample complexity of $\tilde{O}(1/\epsilon^2)$ in the off-policy setting. This matches the results of Sherman et al. [2023], Cassel and Rosenberg [2024], Tan et al. [2025], but with a far less complicated algorithm design.

We thus demonstrate that our optimistic actor-critic method is both practical and sample-efficient.

2 Preliminaries

In this section, we introduce the episodic linear MDP setting and the log-linear policy class.

Episodic Linear MDP. An episodic MDP is a tuple $\mathcal{M} = (\mathcal{S}, \mathcal{A}, \mathbb{P}, r, H)$ where $\mathcal{S}$ denotes the state space, $\mathcal{A}$ is the action set and $H \in \mathbb{Z}_+$ is the length of the horizon, $\mathbb{P} = \{\mathbb{P}_h\}_{h \in [H]}$ is a set of time-dependent transition kernels, and $r = \{r_h\}_{h \in [H]}$ denotes a sequence of reward functions. We assume that the state space $\mathcal{S}$ is a (possibly infinite) measurable space, whereas $\mathcal{A}$ is a finite set with cardinality $|\mathcal{A}|$. We note that $\mathbb{P}_h(\cdot | s, a) \in \Delta(\mathcal{S})$ is the distribution over states when taking action $a \in \mathcal{A}$ in state $s \in \mathcal{S}$ at step $h \in [H]$, and $r_h(s, a) \in [0, 1]$ is the corresponding reward. Additionally, for any given function $V : \mathcal{S} \rightarrow \mathbb{R}$, we define that $[\mathbb{P}_h V_{h+1}](s, a) := \mathbb{E}_{s' \sim \mathbb{P}_h(\cdot | s, a)} V_{h+1}(s')$.

The agent interacts with the environment by starting at an initial state (w.l.o.g., fixed to be $s_1 \in \mathcal{S}$). At step h , the agent first observes the current state $s_h \in \mathcal{S}$, then takes an action $a_h \in \mathcal{A}$ and receives the reward $r_h(s_h, a_h)$. After that, the agent transitions to $s_{h+1} \sim \mathbb{P}_h(\cdot | s_h, a_h)$. The agent follows a given policy $\pi : [H] \times \mathcal{S} \mapsto \Delta(\mathcal{A})$ in which $\pi_h(\cdot | s) \in \Delta(\mathcal{A})$ is the probability distribution over $\mathcal{A}$ in state s at step h .

To quantify the performance of any given policy π , we define the value function as $V_h^\pi(s) := \mathbb{E}[\sum_{\tau=h}^H r_\tau(s_\tau, a_\tau) | s_h = s, \pi]$, and the corresponding state-action value function is defined as $Q_h^\pi(s, a) := \mathbb{E}_{\pi, \mathbb{P}}[\sum_{\tau=h}^H r_\tau(s_\tau, a_\tau) | s_h = s, a_h = a, \pi]$, where the expectation is with respect to the randomness in the stochastic policy and the transition dynamics. The value function (resp. Q -function) corresponds to the expected cumulative rewards when starting in state s (resp. state-action (s, a)) at step h , and subsequently following the policy π until reaching step H .

We assume that both $\mathbb{P}$ and r are unknown to the agent. In order to efficiently learn these quantities, we consider the linear MDP assumption [Jin et al., 2020] where both the transition kernel and the reward function are assumed to be linear functions of given features.

Definition 2.1 (Linear MDP). *An episodic MDP $\mathcal{M} = (\mathcal{S}, \mathcal{A}, \mathbb{P}, r, H)$ is a linear MDP with a feature map $\phi : \mathcal{S} \times \mathcal{A} \mapsto \mathbb{R}^{d_c}$ if the following holds. There exist H signed measures $\psi_h : \mathcal{S} \rightarrow \mathbb{R}^{d_c}$ and $v_h : \mathbb{R}^{d_c} \rightarrow \mathbb{R}$ such that $\mathbb{P}_h(s' | s, a) = \langle \phi(s, a), \psi_h(s') \rangle$ and $r_h(s, a) = \langle \phi(s, a), v_h \rangle$. It should also satisfy the following constraints: $\|\phi(s, a)\| \leq 1$, $\|\psi_h(s)\| \leq \sqrt{d_c}$, and $\|v_h\| \leq \sqrt{d_c}$ for all h, s , and a . Additionally, for any measurable function $V : \mathcal{S} \rightarrow [0, 1]$, $\|\int_{s \in \mathcal{S}} V(s) \psi_h(s) ds\| \leq \sqrt{d_c}$.*

According to Jin et al. [2020, Proposition 2.3], for a linear MDP and any policy π , Q_h^π is a linear function of the features: for all (h, s, a) , there exists a $w_h \in \mathbb{R}^{d_c}$ such that $Q_h^\pi(s, a) = \langle \phi(s, a), w_h \rangle$.

Learning Objective. For this linear MDP setting, we assume that only ϕ is available to the learner whereas ψ and v are not. The agent sequentially interacts with the environment for T episodes and aims to minimize the *cumulative regret* defined as $\text{Reg}(T) := \sum_{t=1}^T [V_1^*(s_1) - V_1^{\pi^t}(s_1)]$, where $V_1^* := V_1^{\pi^*} := \sup_{\pi} V_1^{\pi}$ is the value function of the optimal policy $\pi^* := \arg \sup_{\pi} V_1^{\pi}(s_1)$. Equivalently, if $\bar{\pi}^T$ denotes the mixture policy that picks a policy among $\{\pi^1, \dots, \pi^T\}$ uniformly randomly, we aim to learn an ϵ -optimal $\bar{\pi}^T$, i.e., its *optimality gap* (OG) is bounded such that

$$\text{OG}(T) := \mathbb{E} [V_1^*(s_1) - V_1^{\bar{\pi}^T}(s_1)] = \frac{\text{Reg}(T)}{T} \leq \tilde{\mathcal{O}}(\epsilon),$$

where the expectation is taken with respect to the randomness of the mixture policy.

Log-Linear Policy. We consider a restricted policy class Π_{lin} consisting of *log-linear policies*. Log-linear policies are represented using the softmax function with linear function approximation. In particular, a log-linear policy is defined as follows: for all $h \in [H]$,

$$\pi_h(a | s, \theta) = \frac{\exp(z_h(s, a | \theta_h))}{\sum_{a' \in \mathcal{A}} \exp(z_h(s, a' | \theta_h))}, \quad (1)$$

where $z_h(s, a | \theta_h) = \langle \varphi(s, a), \theta_h \rangle$ represents the logits parameterized by θ_h , and $\varphi : \mathcal{S} \times \mathcal{A} \mapsto \mathbb{R}^{d_a}$ are policy features given to the learner. W.l.o.g, we assume that $\|\varphi(s, a)\| \leq 1$ for all s and a . For convenience, we use the shorthand $\pi(\theta) : [H] \times \mathcal{S} \rightarrow \Delta(\mathcal{A})$ to refer to the log-linear policy corresponding to the parameters θ .

3 Optimistic Actor-Critic Framework

In this section, we start by introducing our general optimistic actor-critic framework as shown in [Algorithm 1](#). Starting with a uniform policy π^1 , at the beginning of every learning episode $t \in [T]$, the agent interacts with the environment using policy π^t (Line 4). Our framework allows for collecting data from the environment in either an on-policy or off-policy fashion. In the on-policy setting, at episode t , the agent collects N fresh trajectories $\mathcal{D}^t$ by interacting with the environment using the current policy π^t . On the other hand, in the off-policy setting, at episode t , the agent collects only 1 trajectory from the environment using π^t . However, the agent stores all the historical data collected by the previous policies, and hence, $\mathcal{D}^t$ consists of t trajectories, each collected by $\pi^1, \dots, \pi^t$ respectively.

The *critic* uses the collected data and estimates an (optimistic) Q -function via learning the critic parameters $w^{t+1} \in [H] \times \mathbb{R}^{d_c}$ (Line 5). The *actor* then uses the estimated Q -function, and updates the parameters of $\theta^t \in [H] \times \mathbb{R}^{d_a}$ of the log-linear policy (Line 6). The updated log-linear policy is denoted by π^{t+1} (Line 7), and is used to collect data in the next episode.

Algorithm 1 Optimistic Actor-Critic with Parametric Policies

- 1: **Input:** number of update steps T , data collection batch size N (only for on-policy)
 - 2: set $\mathcal{D}^0 \leftarrow \emptyset$, $w_h^1 \leftarrow \mathbf{0}$, $\pi_h^1(\cdot | s) \leftarrow \mathcal{U}(\mathcal{A}) \quad \forall (h, s)$
 - 3: **for** $t = 1, \dots, T - 1$ **do**
 - 4: Collect data: $\mathcal{D}^t \leftarrow \begin{cases} \text{On-Policy:} & \{N \text{ fresh traj.} \stackrel{\text{i.i.d.}}{\sim} \pi^t\} \\ \text{Off-Policy:} & \mathcal{D}^{t-1} \cup \{1 \text{ traj.} \sim \pi^t\} \end{cases}$
 - 5: Update the critic: $w^{t+1} \leftarrow \text{Critic}(\mathcal{D}^t, \pi^t, w^t)$
 - 6: Update the actor: $\theta^{t+1} \leftarrow \text{Actor}(w^{t+1}, \theta^t)$
 - 7: Instantiate the parametric policy: $\pi^{t+1} = \pi(\theta^{t+1})$
 - 8: **Return:** mixture policy $\bar{\pi}^T$
-

Given this general framework, we will next instantiate the actor in [Section 4](#) and the critic in [Section 5](#).

4 Instantiating the Actor: Projected Natural Policy Gradient

In this section, we instantiate the actor using natural policy gradient (NPG) with parametric policies and analyze its behavior. In particular, in [Section 4.1](#), we devise an algorithm that projects the standard NPG update onto the class of realizable policies. In [Section 4.2](#), we analyze and control the errors induced by the projection step. Finally, in [Section 4.3](#), we put everything together and instantiate the complete actor algorithm for the log-linear policy class.

4.1 Projected Natural Policy Gradient

At episode $t \in [T]$, given $\hat{Q}_h^t$, the estimated Q -function, NPG updates the policy as: for each (h, s) ,

$$\pi_h^{t+1}(\cdot | s) \propto \pi_h^t(\cdot | s) \exp(\eta \hat{Q}_h^t(s, \cdot)) \quad (2)$$

with the corresponding normalization across $\mathcal{A}$. Existing work on policy optimization in linear MDPs [Liu et al., 2023, Sherman et al., 2023, Cassel and Rosenberg, 2024] uses NPG to update the actor because of its favorable theoretical properties. Importantly, *these works do not consider any explicit parameterization for the actor*. Directly implementing the update in Eq. (2) requires $\mathcal{O}(|S||\mathcal{A}|)$ memory, and is therefore impractical with large state-action spaces. Consequently, existing works use the following equivalent form of the NPG update: $\pi_h^{t+1}(\cdot | s) \propto \pi_h^1(\cdot | s) \exp(\eta \sum_{i=1}^t \hat{Q}_h^i(s, \cdot))$ and characterize the policy implicitly. In particular, at episode t , for any (h, s, a) , we can compute the policy *on the fly* if we have access to the sum of all the parameterized Q -functions up to episode t . However, in Liu et al. [2023], Sherman et al. [2023], Cassel and Rosenberg [2024], the sum of parameterized Q functions cannot be stored in a succinct manner. Consequently, these existing works require storing *all* the parameterized Q functions, and have a memory complexity linear in $|S|$ or T . Consequently, the resulting algorithm is far from practice that typically uses an explicit (and often sophisticated) actor parameterization.

To alleviate these issues, we aim to compute a policy that is (i) realizable by the explicit actor parameterization and (ii) provably approximates the policy induced by the NPG update in Eq. (2) (referred to as the *implicit policy*). To this end, we use a projected NPG update:

$$\pi_h^{t+1}(\cdot | s) = \text{Proj}_{\Pi} \left[\frac{\pi_h^t(\cdot | s) \exp(\eta \hat{Q}_h^t(s, \cdot))}{\sum_{a'} \pi_h^t(a' | s) \exp(\eta \hat{Q}_h^t(s, a'))} \right],$$

where Proj is the projection operator, which will be instantiated subsequently in Section 4.2.

When theoretically analyzing policy optimization methods, an important intermediate result is the bound on the regret for a specific online linear optimization problem. For the standard NPG update in Eq. (2), this regret can be bounded by $\tilde{\mathcal{O}}(\sqrt{T})$ [Hazan et al., 2016, Szepesvári, 2022]. In the following lemma, we analyze the effect of the projection operator and bound the regret for the projected NPG.

Lemma 4.1. *Given a sequence of linear functions $\{ \langle p^t, g^t \rangle \}_{t \in [T]}$ for a sequence of vectors $\{ g^t \}_{t \in [T]}$ where for any $t \in [T]$, $p^t \in \Delta(\mathcal{A})$, $g^t \in \mathbb{R}^{|\mathcal{A}|}$, and $\|g^t\|_{\infty} \leq H$. Consider $p^{t \in [T]}$ where p^1 is the uniform distribution, and for all $t \in [T]$,*

$$p^{t+1/2} = \arg \min_{p \in \Delta_{\mathcal{A}}} \{ \langle p, -\eta g^t \rangle + \text{KL}(p \| p^t) \}, \quad (3)$$

$$p^{t+1} = \text{Proj}_{\Pi}(p^{t+1/2}). \quad (4)$$

Let $\epsilon^t := \text{KL}(u \| p^{t+1}) - \text{KL}(u \| p^{t+1/2})$ be the projection error induced by Eq. (4). Then, for any comparator $u \in \Delta(\mathcal{A})$, it holds that

$$\sum_{t=1}^T \langle u - p^t, g^t \rangle \leq \frac{\log|\mathcal{A}| + \sum_{t=1}^T \epsilon^t}{\eta} + \frac{\eta H^2 T}{2}.$$

The update in Eq. (3) with $p^t = \pi_h^t(\cdot | s)$ is equivalent to the standard NPG update in Eq. (2) [Xiao, 2022]. Using this lemma for each state s and step h , with $g^t = \hat{Q}_h^t(s, \cdot)$ and an appropriate choice of η gives the following regret bound¹ for the projected NPG:

$$\sum_{t=1}^T \sum_{h=1}^H \max_{s \in \mathcal{S}} \langle \pi_h^*(\cdot | s) - \pi_h^t(\cdot | s), \hat{Q}_h^t(s, \cdot) \rangle \leq \mathcal{O} \left(H^2 \sqrt{\log|\mathcal{A}|} \sqrt{T} + H^2 \sqrt{\bar{\epsilon}} T \right), \quad (5)$$

where $\bar{\epsilon} := \max_{t,s,h} \epsilon_h^t(s)$ is the largest error across all t, s , and h . For the NPG in Eq. (2) without projection, $\bar{\epsilon} = 0$ and the above result recovers the standard regret bound for NPG. The above lemma suggests that by choosing the projection operator carefully and controlling the projection errors, we can bound the regret.

¹This generalized regret bound holds for any other mirror descent-based policy optimization method (e.g., SPMA [Asad et al., 2025] in Appendix B.3), but we discuss NPG within the main text for the ease of exposition.

4.2 Controlling the Projection Error for Log-Linear Policies

To bound the projection error in [Lemma 4.1](#), one could choose that $\text{Proj}_\Pi(p) = \arg \min_p \text{KL}(u \parallel p) - \text{KL}(u \parallel p^{t+1/2})$, and hence directly control ϵ_h^t . However, this results in a non-convex optimization problem. Consequently, we instead choose Proj to minimize the following regression loss in the logit space: $\frac{1}{2} \|z - (z^t + \eta g^t)\|^2$ where z^t is the logit corresponding to p^t such that $p^t \propto \exp(z^t)$. For the projected NPG with log-linear policies, we aim to minimize the sum of such regression losses (across all $(s, a) \in \mathcal{S} \times \mathcal{A}$) at episode t and step h , and obtain the loss function: $\frac{1}{2} \sum_{(s,a) \in \mathcal{S} \times \mathcal{A}} \left[\langle \varphi(s, a), \theta - \hat{\theta}_h^t \rangle - \eta \hat{Q}_h^t(s, a) \right]^2$, where $\hat{Q}_h^t(s, a)$ is the estimated Q -function from the critic. As a regression problem, this actor loss can be easily optimized via gradient descent-based methods.

However, note that the above actor loss requires a minimization over the entire state-action space, which may be impractical. Therefore, we propose to construct a good and preferably small subset $\mathcal{D}_{\text{exp}} \subset \mathcal{S} \times \mathcal{A}$ along with a corresponding distribution $\rho_{\text{exp}} \in \Delta(\mathcal{D}_{\text{exp}})$ that offers good coverage of the feature space. Given $\mathcal{D}_{\text{exp}}$ and ρ_{exp} , we instantiate the actor loss $\tilde{\ell}_h^t(\theta)$:

$$\tilde{\ell}_h^t(\theta) = \frac{1}{2} \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s, a) \left[\langle \varphi(s, a), \theta - \hat{\theta}_h^t \rangle - \eta \hat{Q}_h^t(s, a) \right]^2, \quad (6)$$

In order to construct $\mathcal{D}_{\text{exp}}$ and ρ_{exp} and to show that optimizing the above actor loss can indeed bound the projection error, we require the following assumptions. We assume that the given policy features φ are expressive enough to control the bias when minimizing $\tilde{\ell}_h^t(\theta)$.

Assumption 4.1 (Bias). *Suppose φ is the given policy feature for the log-linear policy. Then, it holds that $\inf_{\theta} \sup_{t,h,s,a} \left| \langle \varphi(s, a), \theta \rangle - \eta \hat{Q}_h^t(s, a) \right| \leq \epsilon_{\text{bias}}$.*

In practice, ϵ_{bias} can be controlled by choosing high-dimensional features (e.g., $d_a \gg d_c$) or a sufficiently expressive policy class (e.g., neural network).

Next, we assume the loss $\tilde{\ell}_h^t(\theta)$ is sufficiently minimized.

Assumption 4.2 (Optimization Error). *Suppose θ_h^t is obtained by minimizing $\tilde{\ell}_h^t(\theta)$. Let $\hat{\theta}_h^{t,*} = \arg \min_{\theta} \tilde{\ell}_h^t(\theta)$. Then, it holds that $\sup_{t,h} \left\| \theta_h^t - \hat{\theta}_h^{t,*} \right\| \leq \epsilon_{\text{opt}}$.*

In practice, minimizing $\tilde{\ell}_h^t(\theta)$ by K_t steps of gradient descent ensures that $\epsilon_{\text{opt}} \leq \mathcal{O}(\exp(-K_t))$. Given these two mild assumptions, we can then proceed to bound the projection error. Using [Lemma 4.1](#) for the projected NPG update at state s , step h , and setting $u = \pi^*(\cdot \mid s)$, the projection error $\epsilon_h^t(s)$ can be bounded as follows.

Lemma 4.2. *Under [Assumptions 4.1](#) and [4.2](#), suppose $\bar{\varphi}_G := \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} \|\varphi(s, a)\|_{G^{-1}}$ where $G := \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s, a) \varphi(s, a) \varphi(s, a)^\top$, then, for all (t, h, s) ,*

$$|\epsilon_h^t(s)| \leq \bar{\epsilon} := \sqrt{2} (\bar{\varphi}_G + 1) \epsilon_{\text{bias}} + \sqrt{2} \epsilon_{\text{opt}}.$$

The above lemma is true for any choice of $\mathcal{D}_{\text{exp}}$ and ρ_{exp} , and suggests that if we can control $\bar{\varphi}_G$, the projection error can be bounded. Therefore, we would like to construct a suitable $\mathcal{D}_{\text{exp}}$ and ρ_{exp} to bound $\|\varphi(s, a)\|_{G^{-1}}$ for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and solve the following optimization problem:

$$\begin{aligned} & \inf_{\substack{\mathcal{D}_{\text{exp}} \subset \mathcal{S} \times \mathcal{A} \\ \rho_{\text{exp}} \in \Delta(\mathcal{D}_{\text{exp}})}} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} \|\varphi(s, a)\|_{G^{-1}} \\ & \text{s.t. } G = \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s, a) \varphi(s, a) \varphi(s, a)^\top, \end{aligned}$$

which fits the form of experimental design. Ideally, we would also like $|\mathcal{D}_{\text{exp}}|$ to be relatively small so that the actor parameters can be updated efficiently.

There are standard techniques to solve this problem. The most common approach is the G -optimal design, which involves constructing a *coreset* and bounds $\|\varphi(s, a)\|_{G^{-1}}$. In particular, the Kiefer–Wolfowitz theorem [[Kiefer and Wolfowitz, 1960](#)] guarantees that there exists a coreset such that $\|\varphi(s, a)\|_{G^{-1}} \leq \mathcal{O}(d_a)$ and $|\mathcal{D}_{\text{exp}}| \leq \tilde{\mathcal{O}}(d_a)$. Constructing such a coreset can be achieved using various methods, such as the Frank-Wolfe algorithm [[Frank et al., 1956](#), [Szepesvári, 2022](#)]. We remark that this method only uses the given policy features φ , and does not involve the linear

MDP features. Furthermore, the required coreset can be constructed offline, even before the learning procedure or without any knowledge of the environment (see [Appendix C.1](#) for details). Given access to such a coreset, we can guarantee that $\bar{\epsilon} \leq \mathcal{O}(d_a \epsilon_{\text{bias}} + \epsilon_{\text{opt}})$, and optimizing the actor loss only requires $\mathcal{O}(d_a)$ computation.

Rather than forming a coreset, alternative approaches assume $\varphi = \phi$, and use some limited interaction with the environment to construct $\mathcal{D}_{\text{exp}}$. In particular, under some standard assumptions (e.g., [Wagenmaker and Jamieson, 2022](#), Assumption 1), we can apply methods such as CoverTraj [[Wagenmaker et al., 2022](#)] and OptCov [[Wagenmaker et al., 2022](#)] that bound $\|\varphi(s, a)\|_{G^{-1}}$ and offer similar guarantees. We defer these details to [Appendix C](#).

4.3 Putting Everything Together: Projected NPG with Log-Linear Policies

In [Algorithm 2](#), we instantiate the complete actor algorithm, which uses the projected NPG update for log-linear policies. Unlike the standard NPG update, [Algorithm 2](#) alleviates the necessity of storing past Q -functions, improving the memory complexity to $\mathcal{O}(d_a)$, while enjoying similar theoretical guarantees. Furthermore, the actor parameters are updated by using gradient descent on a properly defined surrogate loss, rendering it closer to the practical implementation of common algorithms (e.g., PPO [[Schulman et al., 2017b](#)]).

We remark that although we focused on the log-linear policies, our theoretical guarantees readily extend to general function approximation when [Assumptions 4.1](#) and [4.2](#) are satisfied and one has access to an exploratory policy [[Hao et al., 2021](#), Definition 1]. In the next section, we instantiate the critic in [Algorithm 1](#).

Algorithm 2 Actor: Projected NPG

- 1: **Input:** critic parameters w^t , policy optimization learning rate η , number of actor updates K_t , actor learning rate α_a^t , subset and distribution of the state-action space $\mathcal{D}_{\text{exp}}$ and ρ_{exp}
 - 2: **for** $h = 1, 2, \dots, H$ **do**
 - 3: $\hat{Q}_h^t(\cdot, \cdot) = \text{clip}_{[0, H-h+1]} \left\{ \max_{m \in [M]} \left\langle \phi(\cdot, \cdot), w_h^{t, m, J_t} \right\rangle \right\}$
 - 4: Define the actor loss $\tilde{\ell}_h^t(\theta)$ using [Eq. \(6\)](#)
 - 5: **for** $k = 1, \dots, K_t$ **do**
 - 6: $\theta_h^{t, k} \leftarrow \theta_h^{t, k-1} - \alpha_a^t \nabla_{\theta} \tilde{\ell}_h^t(\theta_h^{t, k-1})$
 - 7: **Return:** actor parameters for the policy θ^t
-

5 Instantiating the Critic: Langevin Monte Carlo

In this section, we use Langevin Monte Carlo (LMC) to instantiate the critic. We describe the resulting algorithm in [Section 5.1](#), and analyze it in [Section 5.2](#).

The LMC approaches allow for sampling from a posterior distribution and have recently been used in sequential decision-making problems. For example, [Mazumdar et al. \[2020\]](#) achieves optimal instance-dependent regret bounds for multi-armed bandits using Langevin dynamics for approximate Thompson sampling. On the other hand, [Xu et al. \[2022\]](#) uses LMC for contextual bandits, achieving comparable theoretical results to Thompson sampling. More recently, [Ishfaq et al. \[2024a\]](#) leverages LMC for linear MDPs by using it to sample the Q -function from its posterior distribution, achieving the optimal $\tilde{\mathcal{O}}(\sqrt{T})$ regret.

Nevertheless, all existing LMC-based approaches for MDPs, including those for general function approximation [[Ishfaq et al., 2024b](#), [Jorge et al., 2024](#)] use value-based algorithms. To the best of our knowledge, such approaches have never been theoretically analyzed in the context of policy optimization. Next, we incorporate the LMC algorithm into our actor-critic framework and provide the first provable result.

5.1 LMC for Linear MDPs

At episode t , the critic uses the collected dataset $\mathcal{D}^t$ to obtain an optimistic estimate of the Q function. In order to instantiate the critic loss, we consider the dataset $\mathcal{D}^t$ as split into H disjoint subsets $\{\mathcal{D}_h^t\}_{h \in [H]}$, where $\mathcal{D}_h^t$ consists of (s_h, a_h, s_{h+1}) tuples indexed as $\{(s_h^i, a_h^i, s_{h+1}^i)\}_{i=1}^{|\mathcal{D}_h^t|}$ ². The critic

² $|\mathcal{D}^t|$ represents the number of trajectories in $\mathcal{D}^t$ or the number of (s_h, a_h, s_{h+1}) tuples in $\mathcal{D}_h^t$

loss at episode t and step h uses the estimated value function at step $h + 1$, and forms the following ridge regression problem:

$$\mathcal{L}_h^t(w) = \frac{1}{2} \sum_{i=1}^{|\mathcal{D}^t|} \left[r_h(s_h^i, a_h^i) + \widehat{V}_{h+1}^t(s_{h+1}^i) - \langle \phi(s_h^i, a_h^i), w \rangle \right]^2 + \frac{\lambda}{2} \|w\|^2, \quad (7)$$

For each step h , LMC iteratively adds Gaussian noise to the gradient descent updates on $\mathcal{L}_h^t(w)$, and aims to produce approximate samples of the critic parameters from its underlying posterior distribution (Line 6-8). In particular, for an arbitrary loss ℓ , the LMC update can be written as:

$$w^{t+1} = w^t - \alpha^t \nabla_w \ell(w^t) + \sqrt{\alpha^t / \zeta} \nu^t,$$

where α_t is the learning rate, ζ is the inverse temperature parameter, and ν_t is sampled from an isotropic Gaussian distribution. After J_t steps of the LMC update on the critic loss (Lines 6-8 in [Algorithm 3](#)), the resulting critic parameters are used to produce an optimistic sample of the Q -function (Line 9). From a theoretical perspective, we note that it is important to clip Q_h^t appropriately. In order to improve the optimism guarantees of the LMC algorithm, we follow the idea in [Ishfaq et al. \[2021\]](#), and repeat the LMC update M times, taking the maximum over these samples (Line 9). Iterating this procedure backward from $h = H$ to 1, we can obtain the desired critic parameters.

Note that compared to UCB-based approaches, LMC does not require computing confidence sets at every episode. Instead, it simply perturbs gradient descent by injecting Gaussian noise, allowing for a natural extension beyond the linear function approximation setting and rendering it easier to implement in practice. Moreover, our proposed framework, instantiated by the projected NPG and LMC, has less space complexity as stated in the following remark.

Remark 5.1. Existing works (e.g., [Liu et al. \[2023\]](#), [Sherman et al. \[2023\]](#), [Cassel and Rosenberg \[2024\]](#)) that use the standard NPG update ($\pi_h^{t+1}(\cdot|s) \propto \pi_h^1(\cdot|s) \exp(\eta \sum_{i=1}^t \widehat{Q}_h^i(s, \cdot))$) with UCB bonuses require storing $\mathcal{D}^t$ and all the historical Q -function for every previous iteration, resulting in the space complexity of $\mathcal{O}(THd)$. Our proposed framework, instantiated by the projected NPG and LMC, does not have such requirement, and only require $\mathcal{O}(TH + Hd)$, where $\mathcal{O}(TH)$ is for storing $\mathcal{D}^t$ and $\mathcal{O}(Hd)$ is for storing the LMC critic parameters.

Algorithm 3 Critic: LMC

- 1: **Input:** collected data $\mathcal{D}^t$, policy π^{t-1} , number of critic updates J_t , critic learning rate $\alpha_c^{h,t}$, inverse temperature ζ , number of critic samples M
 - 2: $\widehat{V}_{H+1}^t(\cdot) \leftarrow 0$
 - 3: **for** $h = H, H-1, \dots, 1$ **do**
 - 4: Define the critic loss $\mathcal{L}_h^t(w)$ using [Eq. \(7\)](#)
 - 5: $w_h^{t,m,0} \leftarrow w_h^{t-1,m,J_t-1} \quad \forall m \in [M]$
 - 6: **for** $j = 1, \dots, J_t$ **do**
 - 7: $\nu_h^{t,m,j} \leftarrow \mathbf{N}(0, I) \quad \forall m \in [M]$
 - 8: $w_h^{t,m,j} \leftarrow w_h^{t,m,j-1} - \alpha_c^{h,t} \nabla_w \mathcal{L}_h^t(w_h^{t,m,j-1}) + \sqrt{\alpha_c^{h,t} / \zeta} \nu_h^{t,m,j} \quad \forall m \in [M]$
 - 9: $\widehat{Q}_h^t(\cdot, \cdot) = \text{clip}_{[0, H-h+1]} \left\{ \max_{m \in [M]} \left\langle \phi(\cdot, \cdot), w_h^{t,m,J_t} \right\rangle \right\}$
 - 10: $\widehat{V}_h^t(\cdot) = \mathbb{E}_{a \sim \pi^{t-1}(\cdot|s)} \widehat{Q}_h^t(\cdot, a)$
 - 11: **Return:** critic parameters for the estimated Q -function $\{w_h^{t,m,J_t}\}_{(m,h) \in [M] \times [H]}$
-

5.2 Optimism Guarantee and Error Bound

In order to theoretically analyze [Algorithm 3](#), we first define the following model prediction error.

Definition 5.1. Given an estimated Q -function $\widehat{Q}^t$ and the corresponding estimated value function $\widehat{V}^t$, for all (t, h, s, a) , the model prediction error is $\iota_h^t(s, a) := r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) - \widehat{Q}_h^t(s, a)$.

The theoretical analyses in existing works [[Jin et al., 2020](#), [Zhong and Zhang, 2023](#), [Liu et al., 2023](#)] that use UCB bonuses typically proceed by proving an upper bound of 0 on ι_h^t (optimism) and a lower bound of $\tilde{\mathcal{O}}(\sqrt{T})$. The following lemma shows that LMC can offer similar guarantees.

Lemma 5.1. Let $\Lambda_h^t := \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s, a) \phi(s, a)^\top + \lambda I$. With appropriate choices of λ , ζ , J_t , $\alpha_c^{h,t}$, M and for any $\delta \in (0, 1)$, [Algorithm 1](#) with the LMC critic in [Algorithm 3](#) ensures that in both

the on-policy and off-policy settings, for all t, h, s, a and some constant $\Gamma_{\text{LMC}} = \tilde{\mathcal{O}}(H d_c)$, with probability at least $1 - \delta$,

$$-\Gamma_{\text{LMC}} \times \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \leq \iota_h^t(s, a) \leq 0.$$

The exact definition of Γ_{LMC} varies between the on-policy and off-policy settings, although they are both bounded by $\tilde{\mathcal{O}}(H d_c)$ (see [Appendix D](#) for the full version of this lemma). In order to prove this result in the on-policy setting, we use the fact that all the data points in $\mathcal{D}_h^t$ are collected via independent trajectories from the same policy π^t , and are therefore independent and identically distributed. Hence, we can use the self-normalized bounds in [Abbasi-Yadkori et al. \[2011\]](#) to analyze the dependence in h , and prove the corresponding result. In the off-policy setting, since the data points in $\mathcal{D}_h^t$ are collected by different data-dependent policies, these samples are correlated in a complicated manner. Hence, we use the value-aware uniform concentration result from [Jin et al. \[2020\]](#). We remark that this result requires control over the log covering number of the value function class, which is deferred to [Section 6](#).

Therefore, we conclude that, compared to UCB bonuses, LMC offers significant practical advantages while still providing similar theoretical guarantees.

6 Sample Complexity Analysis

In this section, we analyze the sample complexity of [Algorithm 1](#) with the projected NPG actor from [Algorithm 2](#) and the LMC critic from [Algorithm 3](#). [Section 6.1](#) focuses on the on-policy setting, while [Section 6.2](#) addresses the off-policy setting.

6.1 On-Policy Setting

We now present the following theorem that shows that our proposed algorithm achieves a sample complexity of $\tilde{\mathcal{O}}(1/\epsilon^4)$ in the on-policy setting, matching the result in [Liu et al. \[2023\]](#).

Theorem 6.1. *Under [Assumptions 4.1](#) and [4.2](#), consider [Algorithm 1](#) in the on-policy setting with the LMC critic ([Algorithm 3](#)) and the projected NPG actor ([Algorithm 2](#)). For an appropriate choice of the actor and critic parameters, $N = d_c^3 T / (H^2 \log |\mathcal{A}|)$ and $\delta \in (0, 1)$, if $\bar{\epsilon}$ is the projection error in the actor, then, with probability at least $1 - \delta$,*

$$\text{OG}(T) \leq \tilde{\mathcal{O}}\left(\frac{H^2 \sqrt{\log |\mathcal{A}|}}{\sqrt{T}} + H^2 \sqrt{\bar{\epsilon}}\right).$$

Hence, for any $\epsilon > 0$, by setting $T = H^4 \log |\mathcal{A}| / \epsilon^2$, [Algorithm 1](#) returns an $(\epsilon + H^2 \sqrt{\bar{\epsilon}})$ -optimal mixture policy, and therefore requires $T \times N = \tilde{\mathcal{O}}(1/\epsilon^4)$ samples.

Proof sketch. We decompose the difference between $V_1^{\pi^T}(s_1)$ and $V_1^*(s_1)$ into two terms that only depend on either the actor or the critic.

$$\begin{aligned} \mathbb{E}\left[V_1^{\pi^*} - V_1^{\pi^T}(s_1)\right] &= \frac{1}{T} \sum_{t=1}^T \sum_{h=1}^H \underbrace{\mathbb{E}_{\pi^*} \left[\left\langle \pi_h^*(\cdot | s_h) - \pi_h^t(\cdot | s_h), \hat{Q}_h^t(s_h, \cdot) \right\rangle \right]}_{\text{policy optimization (actor) error}} \\ &\quad + \frac{1}{T} \sum_{t=1}^T \sum_{h=1}^H \underbrace{(\mathbb{E}_{\pi^*} [\iota_h^t(s_h, a_h)] - \mathbb{E}_{\pi^t} [\iota_h^t(s_h, a_h)])}_{\text{policy evaluation (critic) error}}. \end{aligned}$$

The policy optimization (actor) error can be bounded using [Eq. \(5\)](#), and the policy evaluation (critic) error is bounded using [Lemma D.2](#). In particular, the lower-bound in [Lemma D.2](#) can be instantiated as $-\iota_h^t(s, a) \leq \mathcal{O}\left(\sqrt{d_c^3 H^4 T \log^2(N/\delta)/N}\right)$ in the on-policy setting. Putting everything together with the chosen value of N leads to the stated sample complexity. $\square$

6.2 Off-Policy Setting

Next, we show that, in the off-policy setting, [Algorithm 1](#) can achieve $\tilde{\mathcal{O}}(1/\epsilon^2)$ sample complexity, matching [Sherman et al. \[2023\]](#), [Cassel and Rosenberg \[2024\]](#).

Theorem 6.2 (Off-Policy Sample Efficiency). *For an appropriate choice of the actor and critic parameters, $\delta \in (0, 1)$, if $\bar{\epsilon}$ is the projection error in the actor, then, with probability at least $1 - \delta$,*

$$\text{OG}(T) \leq \tilde{\mathcal{O}}\left(\frac{d_c^2 \max\{d_a, d_c\} H^2 \sqrt{\log|A|}}{\sqrt{T}} + H^2 \sqrt{\bar{\epsilon}}\right).$$

Hence, for any $\epsilon > 0$, by setting $T = H^4 \log|A|/\epsilon^2$, [Algorithm 1](#) returns a $(\epsilon + H^2 \sqrt{\bar{\epsilon}})$ -optimal mixture policy, and therefore requires $T \times 1 = \tilde{\mathcal{O}}(1/\epsilon^2)$ samples.

Proof sketch. The proof uses a similar regret decomposition to [Theorem 6.1](#). Compared to the on-policy setting, the most significant difference is the bound on the policy evaluation (critic) errors. We use the uniform concentration argument in [Jin et al. \[2020\]](#) to obtain that

$$-\iota_h^t(s, a) \leq \mathcal{O}\left(\sqrt{d_c^3 H^4 T \log(T/\delta)} \left[\sqrt{\log(\mathcal{N}_\Delta(\mathcal{V}))} + T^2 \Delta^2\right]\right),$$

which involves a bound on $\log(\mathcal{N}_\Delta(\mathcal{V}))$, the log covering number of the value function class. The log-covering number is a measure of the complexity of the space of value functions. We show that for an actor with log-linear policies, we can easily bound the log covering number using the following lemma.

Lemma 6.1. *Let Π_{lin} be the policy class induced by [Eq. \(1\)](#) such that $\sup_{\theta, h, s, a} \|z_h(s, a | \theta)\| \leq \bar{Z}$. Let $\mathcal{Q} = \left\{ \min\{\langle \phi(\cdot, \cdot), w \rangle, H \}^+ \mid \|w\| \leq \bar{W} \right\}$ be the Q -function class and $\mathcal{V} = \{ \langle Q(\cdot, \cdot), \pi(\cdot | \cdot, \theta) \rangle_{\mathcal{A}} \mid Q \in \mathcal{Q}, \pi \in \Pi_{lin} \}$ be the corresponding value function class. Then, it holds that $\log \mathcal{N}_\Delta(\mathcal{V}) \leq V$ where $V := d_c \log\left(1 + \frac{4\bar{W} + 4H\sqrt{2\bar{Z}}}{\Delta}\right) + d_a \log\left(1 + \frac{4H\sqrt{2\bar{Z}}}{\Delta}\right)$.*

In particular, we can show $\bar{W} \leq \mathcal{O}(\sqrt{T})$ ([Lemma D.7](#)), and $\bar{Z} \leq \mathcal{O}(\bar{\epsilon}T)$ ([Lemma F.1](#)). Putting everything together and setting $\Delta = \mathcal{O}(1/T^2)$ yields that

$$-\iota_h^t(s, a) \leq \tilde{\mathcal{O}}(\sqrt{d_c^2 \max\{d_a, d_c\} H^4 T}).$$

Following a proof similar to [Theorem 6.1](#) leads to the desired sample complexity. $\square$

In order to control this log covering number, previous work [[Sherman et al., 2023](#), [Cassel and Rosenberg, 2024](#), [Tan et al., 2025](#)] has incorporated various algorithmic tweaks, including reward-free warm-ups, feature contractions, and rare-switching. On the contrary, since our algorithm learns a parametric policy at each iteration, the log covering number of the policy class is bounded by $\mathcal{O}(d_a \log(T))$ without any bespoke tricks.

Furthermore, our proposed framework is also compatible with other policy optimization or policy evaluation methods. For the actor, we can replace the projected NPG with other policy mirror descent-based methods (e.g., SPMA [[Asad et al., 2025](#)]) that can provide a similar bound for the policy optimization error as in [Lemma 4.1](#) (details in [Appendix B.3](#)). We can also instantiate the critic with the UCB bonuses that can provide a similar bound for the policy evaluation error as in [Lemma 5.1](#). In this case, we can recover the same guarantees for sample complexity as [Liu et al. \[2023\]](#) for the on-policy setting and as [Sherman et al. \[2023\]](#), [Cassel and Rosenberg \[2024\]](#) for the off-policy setting.

7 Discussion

We proposed an optimistic actor-critic algorithm with explicitly parameterized policies and a systematic exploration mechanism. In particular, for the actor, we demonstrated that using projected NPG with parametric policies is not only practical, but also equipped with theoretical guarantees. For the critic, we demonstrated that LMC is a principled and easy-to-implement exploration scheme for policy optimization methods. We derived theoretical guarantees in both the on-policy and off-policy settings, showcasing that the proposed actor-critic framework can simultaneously achieve sample efficiency and practicality.

For future work, we aim to investigate the actor-critic methods in more practical setups (e.g., infinite-horizon discounted MDPs) with more general function approximation schemes beyond linear models for both the environment and the policy. It would also be fruitful to further explore the practical implementations of the proposed method and evaluate their performance across standard benchmarks.

Acknowledgments

We would like to thank Xingtu Liu and Yunxiang Li for helpful feedback on the paper. This work was partially supported by the Natural Sciences and Engineering Research Council of Canada (NSERC) Discovery Grant RGPIN-2022-04816, and enabled in part by support provided by the Digital Research Alliance of Canada (alliancecan.ca).

References

- Yasin Abbasi-Yadkori, Dávid Pál, and Csaba Szepesvári. Improved algorithms for linear stochastic bandits. *Advances in Neural Information Processing Systems*, 24, 2011.
- Yasin Abbasi-Yadkori, Peter Bartlett, Kush Bhatia, Nevena Lazic, Csaba Szepesvari, and Gellért Weisz. Politex: Regret bounds for policy iteration using expert prediction. In *International Conference on Machine Learning*, pages 3692–3702. PMLR, 2019.
- Milton Abramowitz and Irene A Stegun. *Handbook of mathematical functions with formulas, graphs, and mathematical tables*, volume 55. US Government Printing Office, 1948.
- Alekh Agarwal, Sham M Kakade, Jason D Lee, and Gaurav Mahajan. On the theory of policy gradient methods: Optimality, approximation, and distribution shift. *Journal of Machine Learning Research*, 22(98):1–76, 2021a.
- Alekh Agarwal, Yujia Jin, and Tong Zhang. VOQL: Towards optimal regret in model-free RL with nonlinear function approximation. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 987–1063. PMLR, 2023.
- Naman Agarwal, Syomantak Chaudhuri, Prateek Jain, Dheeraj Nagaraj, and Praneeth Netrapalli. Online target q-learning with reverse experience replay: Efficiently finding the optimal policy for linear mdps. *arXiv preprint arXiv:2110.08440*, 2021b.
- Carlo Alfano and Patrick Rebeschini. Linear convergence for natural policy gradient with log-linear policy parametrization. *arXiv preprint arXiv:2209.15382*, 2022.
- Carlo Alfano, Rui Yuan, and Patrick Rebeschini. A novel framework for policy mirror descent with general parameterization and linear convergence. *Advances in Neural Information Processing Systems*, 36:30681–30725, 2023.
- Reza Asad, Reza Babanezhad Harikandeh, Issam H Laradji, Nicolas Le Roux, and Sharan Vaswani. Fast convergence of softmax policy mirror ascent. In *International Conference on Artificial Intelligence and Statistics*, pages 3943–3951. PMLR, 2025.
- Jalaj Bhandari and Daniel Russo. Global optimality guarantees for policy gradient methods. *Operations Research*, 72(5):1906–1927, 2024.
- Shalabh Bhatnagar, Richard S Sutton, Mohammad Ghavamzadeh, and Mark Lee. Natural actor-critic algorithms. *Automatica*, 45(11):2471–2482, 2009.
- Qi Cai, Zhuoran Yang, Chi Jin, and Zhaoran Wang. Provably efficient exploration in policy optimization. In *International Conference on Machine Learning*, pages 1283–1294. PMLR, 2020.
- Asaf Cassel and Aviv Rosenberg. Warm-up free policy optimization: Improved regret in linear Markov decision processes. *Advances in Neural Information Processing Systems*, 37:3275–3303, 2024.
- Semih Cayci, Niao He, and R. Srikant. Finite-time analysis of entropy-regularized natural actor-critic algorithm. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856. URL <https://openreview.net/forum?id=BkEqk7pS1I>.
- Shicong Cen, Chen Cheng, Yuxin Chen, Yuting Wei, and Yuejie Chi. Fast global convergence of natural policy gradient methods with entropy regularization. *Operations Research*, 70(4):2563–2578, 2022.
- Gabriel Dulac-Arnold, Daniel Mankowitz, and Todd Hester. Challenges of real-world reinforcement learning. *arXiv preprint arXiv:1904.12901*, 2019.

- Marguerite Frank, Philip Wolfe, et al. An algorithm for quadratic programming. *Naval research logistics quarterly*, 3(1-2):95–110, 1956.
- Zuyue Fu, Zhuoran Yang, and Zhaoran Wang. Single-timescale actor-critic provably finds globally optimal policy. In *International Conference on Learning Representations*, 2021. URL https://openreview.net/forum?id=pqZV_srUvMk.
- Scott Fujimoto, Herke Hoof, and David Meger. Addressing function approximation error in actor-critic methods. In *International conference on machine learning*, pages 1587–1596. PMLR, 2018.
- Mudit Gaur, Amrit Singh Bedi, Di Wang, and Vaneet Aggarwal. On the global convergence of natural actor-critic with two-layer neural network parametrization. *arXiv preprint arXiv:2306.10486*, 2023.
- Mudit Gaur, Amrit Bedi, Di Wang, and Vaneet Aggarwal. Closing the gap: Achieving global convergence (last iterate) of actor-critic under Markovian sampling with neural network parametrization. In *International Conference on Machine Learning*, pages 15153–15179. PMLR, 2024.
- Tuomas Haarnoja, Aurick Zhou, Pieter Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *International conference on machine learning*, pages 1861–1870. Pmlr, 2018.
- Botao Hao, Tor Lattimore, Csaba Szepesvári, and Mengdi Wang. Online sparse reinforcement learning. In *International Conference on Artificial Intelligence and Statistics*, pages 316–324. PMLR, 2021.
- Elad Hazan et al. Introduction to online convex optimization. *Foundations and Trends® in Optimization*, 2(3-4):157–325, 2016.
- Jiafan He, Heyang Zhao, Dongruo Zhou, and Quanquan Gu. Nearly minimax optimal reinforcement learning for linear markov decision processes. In *International Conference on Machine Learning*, pages 12790–12822. PMLR, 2023.
- Mingyi Hong, Hoi-To Wai, Zhaoran Wang, and Zhuoran Yang. A two-timescale stochastic algorithm framework for bilevel optimization: Complexity analysis and application to actor-critic. *SIAM Journal on Optimization*, 33(1):147–180, 2023.
- Haqee Ishfaq, Qiwen Cui, Viet Nguyen, Alex Ayoub, Zhuoran Yang, Zhaoran Wang, Doina Precup, and Lin Yang. Randomized exploration in reinforcement learning with general value function approximation. In *International Conference on Machine Learning*, pages 4607–4616. PMLR, 2021.
- Haqee Ishfaq, Qingfeng Lan, Pan Xu, A Rupam Mahmood, Doina Precup, Anima Anandkumar, and Kamyar Azizzadenesheli. Provable and practical: Efficient exploration in reinforcement learning via langevin monte carlo. In *The Twelfth International Conference on Learning Representations*, 2024a.
- Haqee Ishfaq, Yixin Tan, Yu Yang, Qingfeng Lan, Jianfeng Lu, A Rupam Mahmood, Doina Precup, and Pan Xu. More efficient randomized exploration for reinforcement learning via approximate sampling. In *Reinforcement Learning Conference*, 2024b.
- Haqee Ishfaq, Guangyuan Wang, Sami Nur Islam, and Doina Precup. Langevin soft actor-critic: Efficient exploration through uncertainty-driven critic learning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Chi Jin, Zhuoran Yang, Zhaoran Wang, and Michael I Jordan. Provably efficient reinforcement learning with linear function approximation. In *Conference on learning theory*, pages 2137–2143. PMLR, 2020.
- Chi Jin, Qinghua Liu, and Sobhan Miryoosefi. Bellman eluder dimension: New rich classes of rl problems, and sample-efficient algorithms. *Advances in Neural Information Processing Systems*, 34:13406–13418, 2021.

- Emilio Jorge, Christos Dimitrakakis, and Debabrota Basu. Isoperimetry is all we need: Langevin posterior sampling for rl with sublinear regret. *arXiv preprint arXiv:2412.20824*, 2024.
- Sham M Kakade. A natural policy gradient. *Advances in Neural Information Processing Systems*, 14, 2001.
- Sajad Khodadadian, Thinh T Doan, Justin Romberg, and Siva Theja Maguluri. Finite-sample analysis of two-time-scale natural actor-critic algorithm. *IEEE Transactions on Automatic Control*, 68(6): 3273–3284, 2022.
- Jack Kiefer and Jacob Wolfowitz. The equivalence of two extremum problems. *Canadian Journal of Mathematics*, 12:363–366, 1960.
- Jens Kober, J Andrew Bagnell, and Jan Peters. Reinforcement learning in robotics: A survey. *The International Journal of Robotics Research*, 32(11):1238–1274, 2013.
- Vijay Konda and John Tsitsiklis. Actor-critic algorithms. *Advances in Neural Information Processing Systems*, 12, 1999.
- Tor Lattimore and Csaba Szepesvári. *Bandit algorithms*. Cambridge University Press, 2020.
- Timothy P Lillicrap, Jonathan J Hunt, Alexander Pritzel, Nicolas Heess, Tom Erez, Yuval Tassa, David Silver, and Daan Wierstra. Continuous control with deep reinforcement learning. *arXiv preprint arXiv:1509.02971*, 2015.
- Qinghua Liu, Gellért Weisz, András György, Chi Jin, and Csaba Szepesvári. Optimistic natural policy gradient: a simple efficient policy optimization framework for online rl. *Advances in Neural Information Processing Systems*, 36:3560–3577, 2023.
- Eric Mazumdar, Aldo Pacchiano, Yian Ma, Michael Jordan, and Peter Bartlett. On approximate thompson sampling with langevin algorithms. In *international conference on machine learning*, pages 6797–6807. PMLR, 2020.
- Gergely Neu, Anders Jonsson, and Vicenç Gómez. A unified view of entropy-regularized markov decision processes. *arXiv preprint arXiv:1705.07798*, 2017.
- Ian Osband, Yotam Doron, Matteo Hessel, John Aslanides, Eren Sezener, Andre Saraiva, Katrina McKinney, Tor Lattimore, Csaba Szepesvari, Satinder Singh, et al. Behaviour suite for reinforcement learning. *arXiv preprint arXiv:1908.03568*, 2019.
- Jan Peters, Sethu Vijayakumar, and Stefan Schaal. Natural actor-critic. In *European Conference on Machine Learning*, pages 280–291. Springer, 2005.
- John Schulman, Sergey Levine, Pieter Abbeel, Michael Jordan, and Philipp Moritz. Trust region policy optimization. In *International conference on machine learning*, pages 1889–1897. PMLR, 2015.
- John Schulman, Xi Chen, and Pieter Abbeel. Equivalence between policy gradients and soft q-learning. *arXiv preprint arXiv:1704.06440*, 2017a.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017b.
- Uri Sherman, Alon Cohen, Tomer Koren, and Yishay Mansour. Rate-optimal policy optimization for linear markov decision processes. *arXiv preprint arXiv:2308.14642*, 2023.
- Richard S Sutton, David McAllester, Satinder Singh, and Yishay Mansour. Policy gradient methods for reinforcement learning with function approximation. *Advances in Neural Information Processing Systems*, 12, 1999.
- Csaba Szepesvári. *Algorithms for reinforcement learning*. Springer nature, 2022.
- Kevin Tan, Wei Fan, and Yuting Wei. Actor-critics can achieve optimal sample efficiency. *arXiv preprint arXiv:2505.03710*, 2025.

- Michael J Todd. *Minimum-volume ellipsoids: Theory and algorithms*. SIAM, 2016.
- Victor Uc-Cetina, Nicolás Navarro-Guerrero, Anabel Martin-Gonzalez, Cornelius Weber, and Stefan Wermter. Survey on reinforcement learning for language processing. *Artificial Intelligence Review*, 56(2):1543–1575, 2023.
- Andrew Wagenmaker and Kevin G Jamieson. Instance-dependent near-optimal policy identification in linear mdps via online experiment design. *Advances in Neural Information Processing Systems*, 35:5968–5981, 2022.
- Andrew J Wagenmaker, Yifang Chen, Max Simchowitz, Simon Du, and Kevin Jamieson. Reward-free rl is no harder than reward-aware rl in linear markov decision processes. In *International Conference on Machine Learning*, pages 22430–22456. PMLR, 2022.
- Max Welling and Yee W Teh. Bayesian learning via stochastic gradient langevin dynamics. In *Proceedings of the 28th international conference on machine learning (ICML-11)*, pages 681–688, 2011.
- Ronald J Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8(3):229–256, 1992.
- Lin Xiao. On the convergence rates of policy gradient methods. *Journal of Machine Learning Research*, 23(282):1–36, 2022.
- Pan Xu, Hongkai Zheng, Eric V Mazumdar, Kamyar Azizzadenesheli, and Animashree Anandkumar. Langevin monte carlo for contextual bandits. In *International Conference on Machine Learning*, pages 24830–24850. PMLR, 2022.
- Tengyu Xu, Zhe Wang, and Yingbin Liang. Improving sample complexity bounds for (natural) actor-critic algorithms. *Advances in Neural Information Processing Systems*, 33:4358–4369, 2020.
- Rui Yuan, Simon S Du, Robert M Gower, Alessandro Lazaric, and Lin Xiao. Linear convergence of natural policy gradient methods with log-linear policies. In *International Conference on Learning Representations*, 2023.
- Andrea Zanette, Ching-An Cheng, and Alekh Agarwal. Cautiously optimistic policy optimization and exploration with linear function approximation. In *Conference on Learning Theory*, pages 4473–4525. PMLR, 2021.
- Han Zhong and Tong Zhang. A theoretical analysis of optimistic proximal policy optimization in linear markov decision processes. *Advances in Neural Information Processing Systems*, 36: 73666–73690, 2023.

Contents of Appendix

A	Notation	16
B	Analyses for the Actor	16
B.1	Generalized OMD Regret (Proof of Lemma 4.1)	16
B.2	Projection Error (Proof of Lemma 4.2)	17
B.3	Instantiating the Actor with SPMA	19
B.4	Technical Tools	21
C	Constructing $\mathcal{D}_{\text{exp}}$ via Experimental Design	22
C.1	Kiefer–Wolfowitz Theorem and G-Experimental Design	22
C.2	Exploratory Policy and Minimum Eigenvalue	22
D	Analyses for the Critic	23
D.1	Proof of Lemma 5.1	23
D.1.1	Preliminary Properties	24
D.1.2	Main Analysis	25
D.2	Proofs of Preliminary Properties	27
D.2.1	Proof of Lemma D.3	27
D.2.2	Proof of Lemma D.4	28
D.2.3	Proof of Lemma D.5	29
D.2.4	Proof of Lemma D.6	32
D.2.5	Proof of Lemma D.7	32
D.2.6	Proof of Lemma D.8	34
D.2.7	Proof of Lemma D.9	36
D.3	Technical Tools	38
E	Sample Complexity in the On-Policy Setting	39
E.1	Proof of Good Event	39
E.2	Proof of Theorem 6.1	39
E.3	Technical Tools	42
F	Sample Complexity in the Off-Policy Setting	43
F.1	Covering Number (Proof of Lemma 6.1)	43
F.2	Proof of Good Event	44
F.3	Proof of Theorem 6.2	45
F.4	Technical Tools	47
G	Experiments	47
G.1	Environment Setup	47
G.2	Coreset Construction	47
G.3	Hyperparameters	48
G.4	Experimental Results	48
G.5	Additional Results	49

A Notation

Notation for Problem Setting and Algorithm Design

Table 1: Notation for Problem Setting and Algorithm Design

Notation	Meaning
Problem Definition	
$\mathcal{S}, \mathcal{A}$	state space and action space
H, h	horizon length (total number of steps), current index of step
$r \in \mathbb{R}^{ \mathcal{S} \times \mathcal{A} }$	reward function
$\mathbb{P} \in \mathbb{R}^{ \mathcal{S} \times \mathcal{A} \times \mathcal{S} }$	transition probability
$\phi : \mathcal{S} \times \mathcal{A} \mapsto \mathbb{R}^{d_c}$	features for the linear MDP environment
$\varphi : \mathcal{S} \times \mathcal{A} \mapsto \mathbb{R}^{d_a}$	features for the learnable policy
Algorithm Design	
T, t	total number of learning episodes, index of current episode
$\mathcal{D}^t, \mathcal{D}_h^t$	collected data of at episode t , split data at h -th step (subset of $\mathcal{D}^t$)
N	number of samples collected for on-policy learning
$w \in \mathbb{R}^{d_c}$	learnable critic parameters
J	number of critic updates
α_c	critic learning rate
ν	noise vector for LMC sampled from the standard normal distribution
ζ	inverse temperature for the LMC critic loss
M	number of samples for the critic parameters
$\mathcal{D}_{\text{exp}}, \rho_{\text{exp}}$	subset of $\mathcal{S} \times \mathcal{A}$, distribution over $\mathcal{D}_{\text{exp}}$
$\theta \in \mathbb{R}^{d_a}$	learnable actor parameters
K	number of actor updates
α_a	actor learning rate
η	policy optimization learning rate

Additional Notation Throughout this paper, we use subscripts to represent the index of the step within the horizon of the episodic MDP and superscripts to denote the index of the episode for learning. For example, V_h^t means the value function for the h -th step derived at the learning episode t . In some cases, where the subscripts are omitted, it represents a set of H functions for all steps $h \in [H]$ (e.g., $V^t := \{V_h^t\}_{h \in [H]}$). $|\mathcal{D}^t|$ represents the number of trajectories in $\mathcal{D}^t$ or the number of (s_h, a_h, s_{h+1}) tuples in $\mathcal{D}_h^t$. Additionally, for any vector $v \in \mathbb{R}^d$ and any matrix $M \in \mathbb{R}^{d \times d}$, we denote $\|v\|_M = \sqrt{v^\top M v}$.

B Analyses for the Actor

B.1 Generalized OMD Regret (Proof of Lemma 4.1)

Proof of Lemma 4.1. Given the update of $p^{t+1/2}$ and the fact that $\Delta(\mathcal{A})$ is a convex set, we have the following optimality condition:

$$\left\langle u - p^{t+1/2}, -\eta g^t + \log(p^{t+1/2}) - \log(p^t) \right\rangle \geq 0. \quad (8)$$

Then, for each $t \in [T]$, we have that

$$\begin{aligned}
\langle u - p^t, \eta g^t \rangle &= \langle u - p^{t+1/2}, \eta g^t \rangle + \langle p^{t+1/2} - p^t, \eta g^t \rangle \\
&= \langle u - p^{t+1/2}, \eta g^t - \log(p^{t+1/2}) + \log(p^t) \rangle + \langle u - p^{t+1/2}, \log(p^{t+1/2}) - \log(p^t) \rangle \\
&\quad + \langle p^{t+1/2} - p^t, \eta g^t \rangle \\
&\stackrel{(i)}{\leq} \langle u - p^{t+1/2}, \log(p^{t+1/2}) - \log(p^t) \rangle + \langle p^{t+1/2} - p^t, \eta g^t \rangle \\
&\stackrel{(ii)}{=} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \langle p^{t+1/2} - p^t, \eta g^t \rangle
\end{aligned}$$

$$\begin{aligned}
&\stackrel{\text{(iii)}}{\leq} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \frac{1}{2} \|p^{t+1/2} - p^t\|_1^2 + \frac{1}{2} \|\eta g^t\|_\infty^2 \\
&\stackrel{\text{(iv)}}{\leq} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \text{KL}(p^{t+1/2} \parallel p^t) + \frac{\eta^2 H^2}{2} \\
&\leq \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) + \frac{\eta^2 H^2}{2} \\
&\leq \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \text{KL}(u \parallel p^{t+1}) - \text{KL}(u \parallel p^{t+1/2}) + \frac{\eta^2 H^2}{2} \\
&= \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \epsilon^t + \frac{\eta^2 H^2}{2}.
\end{aligned}$$

(i) drops the first term due to the optimality condition from Eq. (8). (ii) applies the three-point property of Bregman divergence (Lemma B.2) by setting $x = u$, $y = p^{t+1/2}$, and $z = p^t$. (iii) follows from the Hölder's inequality and then the Young's inequality (i.e., $\langle u, v \rangle \leq \|u\|_1 \|v\|_\infty \leq \|u\|_1^2/2 + \|v\|_\infty^2/2$), and (iv) applies the Pinsker's inequality and $|g^t| \leq H$. Summing up the above inequality from $t = 1$ to T yields that

$$\begin{aligned}
\sum_{t=1}^T \langle u - p^t, \eta g^t \rangle &= \sum_{t=1}^T \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \sum_{t=1}^T \epsilon^t + \frac{\eta^2 H^2 T}{2} \\
&= \text{KL}(u \parallel p^1) - \text{KL}(u \parallel p^{T+1}) + \sum_{t=1}^T \epsilon^t + \frac{\eta^2 H^2 T}{2} \\
&\stackrel{\text{(v)}}{\leq} \text{KL}(u \parallel p^1) + \sum_{t=1}^T \epsilon^t + \frac{\eta^2 H^2 T}{2} \\
&\leq \sum_{a \in \mathcal{A}} u(a) \log(u(a)) - \sum_{a \in \mathcal{A}} u(a) \log(p^1(a)) + \sum_{t=1}^T \epsilon^t + \frac{\eta^2 H^2 T}{2} \\
&\stackrel{\text{(vi)}}{\leq} \log |\mathcal{A}| + \sum_{t=1}^T \epsilon^t + \frac{\eta^2 H^2 T}{2}.
\end{aligned}$$

(v) follows from the fact that KL-divergence is non-negative, and (vi) stands because the first term is negative, and for the second term, p^1 is a uniform distribution. Dividing both side by η , we have that

$$\sum_{t=1}^T \langle u - p^t, g^t \rangle \leq \frac{\log |\mathcal{A}| + \sum_{t=1}^T \epsilon^t}{\eta} + \frac{\eta H^2 T}{2}.$$

This concludes the proof. $\square$

B.2 Projection Error (Proof of Lemma 4.2)

Proof of Lemma 4.2. First, we define Φ as the log-sum-exp mirror map and Φ^* as negative entropy, its Fenchel conjugate. Based on this, for any softmax policy π , we can also define its logit as $z := \nabla \Phi^*(\pi) = (\nabla \Phi)^{-1}(\pi)$. Consequently, $\pi = \nabla \Phi(z)$. Additionally, for any two softmax policies π, π' and their corresponding logits z, z' , it holds that, $D_\Phi(z, z') = \text{KL}(\pi', \pi)$.

Since we are using the log-linear policy class, we have $z_h^t(s, a) = \langle \varphi(s, a), \hat{\theta}_h^t \rangle$ for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ where $\hat{\theta}_h^t$ represents the parameters we attain at episode t . Therefore, for any $s \in \mathcal{S}$,

$$\begin{aligned}
\epsilon_h^t(s) &= \text{KL}(\pi_h^*(\cdot \mid s) \parallel \pi_h^{t+1}(\cdot \mid s)) - \text{KL}(\pi_h^*(\cdot \mid s) \parallel \pi_h^{t+1/2}(\cdot \mid s)) \\
&= D_\Phi(z_h^{t+1}(\cdot \mid s), z_h^*(s, \cdot)) - \text{KL}(\pi_h^*(\cdot \mid s) \parallel \pi_h^{t+1/2}(\cdot \mid s)) \\
&\stackrel{\text{(i)}}{=} \left\langle \nabla \Phi(z_h^{t+1}(s, \cdot)) - \nabla \Phi(z_h^*(s, \cdot)), z_h^{t+1}(s, \cdot) - z_h^{t+1/2}(\cdot \mid s) \right\rangle - \text{KL}(\pi_h^*(\cdot \mid s) \parallel \pi_h^{t+1/2}(\cdot \mid s)) \\
&= \left\langle \pi_h^{t+1}(s, \cdot) - \pi_h^*(s, \cdot), z_h^{t+1}(s, \cdot) - z_h^{t+1/2}(\cdot \mid s) \right\rangle - \text{KL}(\pi_h^*(\cdot \mid s) \parallel \pi_h^{t+1/2}(\cdot \mid s))
\end{aligned}$$

$$\begin{aligned}
&\stackrel{\text{(ii)}}{\leq} \left\langle \pi_h^{t+1}(\cdot | s) - \pi_h^*(\cdot | s), z_h^{t+1}(s, \cdot) - z_h^{t+1/2}(s, \cdot) \right\rangle \\
&\stackrel{\text{(iii)}}{\leq} \left\| \pi_h^{t+1}(\cdot | s) - \pi_h^*(\cdot | s) \right\|_2 \left\| z_h^{t+1}(s, \cdot) - z_h^{t+1/2}(s, \cdot) \right\|_2 \\
&\leq \sqrt{2} \left\| z_h^{t+1}(s, \cdot) - z_h^{t+1/2}(s, \cdot) \right\|_2 \\
&\stackrel{\text{(iv)}}{=} \sqrt{2} \left\| \left\langle \varphi(s, \cdot), \hat{\theta}_h^{t+1} - \hat{\theta}_h^t \right\rangle - \eta \hat{Q}_h^t(s, \cdot) \right\|_2 \\
&\stackrel{\text{(v)}}{\leq} \sqrt{2} \left| \left\langle \varphi(s, \cdot), \hat{\theta}_h^{t+1} - \hat{\theta}_h^t \right\rangle - \eta \hat{Q}_h^t(s, \cdot) \right|.
\end{aligned}$$

(i) follows from the three-point property of Bregman divergence ([Lemma B.2](#)) by setting $x = z_h^{t+1/2}(\cdot | s)$, $y = z_h^{t+1}(\cdot | s)$, and $z = z_h^*(\cdot | s)$ where z^* is the logit of π^* . (ii) is based on the fact that KL-divergence is non-negative. (iii) uses the Cauchy-Schwarz inequality. (iv) uses the NPG update. (v) holds because $\|\cdot\|_2 \leq \|\cdot\|_1$.

Since the actor is designed to minimize the ridge regression in [Algorithm 2](#), the minimizer can be written as

$$\hat{\theta}_h^{t,*} = \arg \min_{\theta_h} \frac{1}{2} \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho(s,a) \left[\langle \varphi(s,a), \theta_h \rangle - \hat{Z}_h^t(s,a) \right]^2,$$

where $\hat{Z}_h^t(s,a) := \left\langle \varphi(s, \cdot), \hat{\theta}_h^t(s, \cdot) \right\rangle + \eta \hat{Q}_h^t(s,a)$ for all $t \in [T]$. We define $\hat{\theta}_h^{t,*}$ as the minimizer, and it has the following explicit solution:

$$\hat{\theta}_h^{t,*} = G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \hat{Z}_h^t(s',a') \varphi(s',a') \right],$$

where $G := \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho(s,a) \varphi(s,a) \varphi(s,a)^\top \in \mathbb{R}^{d_a \times d_a}$.

Suppose $\theta_h^{t,*}$ is the minimizer of the regression loss over the entire state-action space, $\hat{\theta}_h^{t,*}$ is the minimizer over the coreset, and $\hat{\theta}_h^t$ is the parameters produced by the actor after K_t rounds of gradient descent as shown in [Algorithm 2](#). Then, for any arbitrary $(s,a) \in \mathcal{S} \times \mathcal{A}$, using the triangular inequality, we have that

$$\begin{aligned}
&\left| \left\langle \varphi(s,a), \hat{\theta}_h^t \right\rangle - \hat{Z}_h^t(s,a) \right| \\
&\leq \left| \left\langle \varphi(s,a), \theta_h^{t,*} \right\rangle - \hat{Z}_h^t(s,a) \right| + \left| \left\langle \varphi(s,a), \hat{\theta}_h^t \right\rangle - \left\langle \varphi(s,a), \theta_h^{t,*} \right\rangle \right| \\
&= \epsilon_{\text{bias}} + \left| \left\langle \varphi(s,a), \hat{\theta}_h^t \right\rangle - \left\langle \varphi(s,a), \theta_h^{t,*} \right\rangle \right| \\
&\leq \epsilon_{\text{bias}} + \left| \left\langle \varphi(s,a), \hat{\theta}_h^t \right\rangle - \left\langle \varphi(s,a), \hat{\theta}_h^{t,*} \right\rangle \right| + \left| \left\langle \varphi(s,a), \hat{\theta}_h^{t,*} \right\rangle - \left\langle \varphi(s,a), \theta_h^{t,*} \right\rangle \right| \\
&= \epsilon_{\text{bias}} + \epsilon_{\text{opt}} + \left| \left\langle \varphi(s,a), \hat{\theta}_h^{t,*} - \theta_h^{t,*} \right\rangle \right|.
\end{aligned}$$

Therefore, it suffices to bound $\left| \left\langle \varphi(s,a), \hat{\theta}_h^{t,*}(s,a) - \theta_h^{t,*}(s,a) \right\rangle \right|$. To do that, we first define $\Upsilon(s',a') := \hat{Z}_h^t(s',a') - \left\langle \varphi(s',a'), \theta_h^{t,*} \right\rangle$ for any $(s',a') \in \mathcal{D}_{\text{exp}}$. Then, we have that

$$\begin{aligned}
\hat{\theta}_h^{t,*} &= G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \left[\Upsilon(s',a') + \left\langle \varphi(s',a'), \theta_h^{t,*} \right\rangle \right] \varphi(s',a') \right] \\
&= G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \varphi(s',a') \varphi(s',a')^\top \right] \theta_h^{t,*} \\
&\quad + G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \Upsilon(s',a') \varphi(s',a') \right]
\end{aligned}$$

$$= \theta_h^{t,*} + G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \Upsilon(s',a') \varphi(s',a') \right].$$

This implies that

$$\widehat{\theta}_h^{t,*} - \theta_h^{t,*} = G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \Upsilon(s',a') \varphi(s',a') \right].$$

Hence, for any arbitrary $(s, a) \in \mathcal{S} \times \mathcal{A}$,

$$\begin{aligned} \left| \left\langle \varphi(s, a), \widehat{\theta}_h^{t,*} - \theta_h^{t,*} \right\rangle \right| &= \left| \sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \Upsilon(s',a') \varphi(s, a)^\top G^{-1} \varphi(s',a') \right| \\ &\stackrel{\text{(vi)}}{\leq} \sum_{(s',a') \in \mathcal{D}_{\text{exp}}} |\Upsilon(s',a')| \rho(s',a') |\varphi(s, a)^\top G^{-1} \varphi(s',a')| \\ &\leq \left(\max_{(s',a') \in \mathcal{D}_{\text{exp}}} |\Upsilon(s',a')| \right) \sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') |\varphi(s, a)^\top G^{-1} \varphi(s',a')| \\ &\leq \epsilon_{\text{bias}} \sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') |\varphi(s, a)^\top G^{-1} \varphi(s',a')| \\ &= \epsilon_{\text{bias}} \sqrt{\left(\mathbb{E}_{(s',a') \sim \rho} |\varphi(s, a)^\top G^{-1} \varphi(s',a')| \right)^2} \\ &\stackrel{\text{(vii)}}{\leq} \epsilon_{\text{bias}} \sqrt{\mathbb{E}_{(s',a') \sim \rho} |\varphi(s, a)^\top G^{-1} \varphi(s',a')|^2} \\ &= \epsilon_{\text{bias}} \sqrt{\varphi(s, a)^\top G^{-1} \left[\sum_{(s',a') \in \mathcal{D}_{\text{exp}}} \rho(s',a') \varphi(s',a') \varphi(s',a')^\top \right] G^{-1} \varphi(s, a)} \\ &= \epsilon_{\text{bias}} \|\varphi(s, a)\|_{G^{-1}}. \end{aligned}$$

(vi) applies the Cauchy-Schwarz inequality, and (vii) follows from Jensen's inequality.

Putting everything together, we have that

$$\begin{aligned} \left| \left\langle \varphi(s, a), \widehat{\theta}_h^t \right\rangle - \widehat{Z}_h^t(s, a) \right| &\leq (\|\varphi(s, a)\|_{G^{-1}} + 1) \epsilon_{\text{bias}} + \epsilon_{\text{opt}} \\ &\leq (\overline{\varphi}_G + 1) \epsilon_{\text{bias}} + \epsilon_{\text{opt}}. \end{aligned}$$

Recall that $\epsilon_h^t(s) \leq \sqrt{2} \left\| \left\langle \varphi(s, \cdot), \widehat{\theta}_h^{t+1}(s, \cdot) \right\rangle - \widehat{Z}_h^t(s, \cdot) \right\|_2$. Therefore, for any $s \in \mathcal{S}$,

$$\epsilon_h^t(s) \leq \sqrt{2} \left| \left\langle \varphi(s, \cdot), \widehat{\theta}_h^t \right\rangle - \widehat{Z}_h^t(s, a) \right| \leq \sqrt{2} (\overline{\varphi}_G + 1) \epsilon_{\text{bias}} + \sqrt{2} \epsilon_{\text{opt}}.$$

This concludes the proof. $\square$

B.3 Instantiating the Actor with SPMA

Lemma 4.2 can not only be applied to NPG but also other mirror descent-based policy optimization methods such as TRPO [Schulman et al. \(2015\)](#), AMPO [\(Alfano et al., 2023\)](#), and SPMA [\(Asad et al., 2025\)](#). In this section, as an example, we show that the projected variant of SPMA (projected SPMA) is also compatible with our framework and can enjoy similar sample complexity guarantees as projected NPG. We can instantiate the actor in [Algorithm 1](#) with the projected SPMA by setting the actor loss in [Algorithm 2](#) as the following:

$$\begin{aligned} \tilde{\ell}_h^t(\theta) &= \frac{1}{2} \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s, a) \left[\langle \varphi(s, a), \theta \rangle - \widehat{Z}_h^t(s, a) \right]^2, \\ \text{where } \widehat{Z}_h^t(s, a) &:= \left\langle \varphi(s, a), \widehat{\theta}_h^t \right\rangle + \log \left(1 + \eta A^{\pi^t}(s, \cdot) \right). \end{aligned}$$

Equivalently, the projected SPMA update can be expressed as follows. For any $s \in \mathcal{S}$, $\pi^1(\cdot | s)$ is a uniform distribution, and

$$\begin{aligned}\pi^{t+1/2}(\cdot | s) &= \arg \min_{p \in \Delta_{\mathcal{A}}} \left\{ \left\langle \pi^t(\cdot | s), -\log(1 + \eta A^{\pi^t}(s, \cdot)) \right\rangle + \text{KL}(p \parallel \pi^t(\cdot | s)) \right\}, \\ \pi^{t+1}(\cdot | s) &= \text{Proj}_{\Pi}(\pi^{t+1/2}(\cdot | s)).\end{aligned}$$

Hence, we introduce the following alternative lemma to show that [Lemma 4.1](#) also holds for the projected SPMA.

Lemma B.1. *Given a sequence of linear functions $\{\langle p^t, g^t \rangle\}_{t \in [T]}$ for a sequence of vectors $\{g^t\}_{t \in [T]}$ where for any $t \in [T]$, $p^t \in \Delta(\mathcal{A})$, $g^t \in \mathbb{R}^{|\mathcal{A}|}$, and $g^t(a) \in [0, H]$ for all $a \in \mathcal{A}$. Consider p^1 is the uniform distribution, and for all $t \in [T]$,*

$$\begin{aligned}p^{t+1/2} &= \arg \min_{p \in \Delta_{\mathcal{A}}} \left\{ \langle p, -\log(1 + \eta (g^t - \langle p^t, g^t \rangle \mathbf{1})) \rangle + \text{KL}(p \parallel p^t) \right\}, \\ p^{t+1} &= \text{Proj}_{\Pi}(p^{t+1/2}),\end{aligned}$$

where $\mathbf{1} \in \mathbb{R}^{|\mathcal{A}|}$ is an all-one vector. Let $\epsilon^t := \text{KL}(u \parallel p^{t+1}) - \text{KL}(u \parallel p^{t+1/2})$ be the projection error induced by [Eq. \(4\)](#). If $\eta \leq \frac{1}{2H}$, then for any comparator $u \in \Delta(\mathcal{A})$, it holds that

$$\sum_{t=1}^T \langle u - p^t, g^t \rangle \leq \frac{\log|\mathcal{A}| + \sum_{t=1}^T \epsilon^t}{\eta} + \frac{3\eta H^2 T}{2}.$$

Proof of Lemma B.1. We first denote that $d^t = \log(1 + \eta (g^t - \langle p^t, g^t \rangle \mathbf{1}))$ for all $t \in [T]$. Then, for all $a \in \mathcal{A}$, since $\eta \leq \frac{1}{2H}$ and $g^t(a) - \langle p^t, g^t \rangle \in [-H, H]$, we have $\eta (g^t(a) - \langle p^t, g^t \rangle) > -\frac{1}{2}$ and therefore

$$\begin{aligned}d^t(a) &\stackrel{(i)}{\leq} \eta (g^t(a) - \langle p^t, g^t \rangle) \leq \eta H, \\ d^t(a) &\stackrel{(ii)}{\geq} \eta (g^t(a) - \langle p^t, g^t \rangle) - \eta^2 (g^t(a) - \langle p^t, g^t \rangle)^2 \geq \eta (g^t(a) - \langle p^t, g^t \rangle) - \eta H^2,\end{aligned}$$

where (i) follows from $\log(1+x) \leq x$ for all $x > -1$, and (ii) holds because $\log(1+x) \geq x - x^2$ for all $x > -\frac{1}{2}$.

Given the update of $p^{t+1/2}$ and the fact that $\Delta(\mathcal{A})$ is a convex set, we have the following optimality condition:

$$\left\langle u - p^{t+1/2}, -d^t + \log(p^{t+1/2}) - \log(p^t) \right\rangle \geq 0. \quad (9)$$

Then, for all $t \in [T]$, we have that

$$\begin{aligned}\langle u - p^t, d^t \rangle &= \langle u - p^{t+1/2}, d^t \rangle + \langle p^{t+1/2} - p^t, d^t \rangle \\ &= \langle u - p^{t+1/2}, d^t - \log(p^{t+1/2}) + \log(p^t) \rangle + \langle u - p^{t+1/2}, \log(p^{t+1/2}) - \log(p^t) \rangle \\ &\quad + \langle p^{t+1/2} - p^t, d^t \rangle \\ &\stackrel{(iii)}{\leq} \langle u - p^{t+1/2}, \log(p^{t+1/2}) - \log(p^t) \rangle + \langle p^{t+1/2} - p^t, d^t \rangle \\ &\stackrel{(iv)}{=} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \langle p^{t+1/2} - p^t, d^t \rangle \\ &\stackrel{(v)}{\leq} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \frac{1}{2} \|p^{t+1/2} - p^t\|_1^2 + \frac{1}{2} \|d^t\|_2^2 \\ &\stackrel{(vi)}{\leq} \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) - \text{KL}(p^{t+1/2} \parallel p^t) + \text{KL}(p^{t+1/2} \parallel p^t) + \frac{\eta^2 H^2}{2} \\ &\leq \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1/2}) + \frac{\eta^2 H^2}{2} \\ &\leq \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \text{KL}(u \parallel \pi^{t+1}) - \text{KL}(u \parallel p^{t+1/2}) + \frac{\eta^2 H^2}{2}\end{aligned}$$

$$= \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \epsilon^t + \frac{\eta^2 H^2}{2}.$$

(iii) drops the first term due to the optimality condition from [Eq. \(9\)](#). (iv) applies the three-point property of Bregman divergence ([Lemma B.2](#)) by setting $x = u$, $y = p^{t+1/2}$, and $z = p^t$. (v) follows from the Hölder's inequality and then the Young's inequality (i.e., $\langle u, v \rangle \leq \|u\|_1 \|v\|_\infty \leq \|u\|_1^2/2 + \|v\|_\infty^2/2$), and (vi) applies the Pinkster's inequality and $\|d^t\|_\infty \leq H$.

Moreover, we have that

$$\langle u - p^t, d^t \rangle \stackrel{\text{(vii)}}{\geq} \langle u - p^t, \eta (g^t - \langle p^t, g^t \rangle \mathbf{1}) \rangle - \eta H^2 \stackrel{\text{(viii)}}{\geq} \langle u - p^t, \eta g^t \rangle - \eta H^2,$$

where (vii) comes from the fact that $d^t(a) \geq \eta (g^t(a) - \langle p^t, g^t \rangle)$ for all $a \in \mathcal{A}$, and (viii) follows from the fact that $\langle p^t, g^t \rangle \geq 0$ since $p^t \in \Delta(\mathcal{A})$ and $g^t(a) \in [0, H]$ for all $a \in \mathcal{A}$. This implies that

$$\begin{aligned} \langle u - p^t, \eta g^t \rangle &\leq \langle u - p^t, d^t \rangle + \eta H^2 \\ &\leq \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \epsilon^t + \frac{3\eta^2 H^2}{2}. \end{aligned}$$

Summing up the above inequality from $t = 1$ to T yields that

$$\begin{aligned} \sum_{t=1}^T \langle u - p^t, \eta g^t \rangle &= \sum_{t=1}^T \text{KL}(u \parallel p^t) - \text{KL}(u \parallel p^{t+1}) + \sum_{t=1}^T \epsilon^t + \frac{3\eta^2 H^2 T}{2} \\ &= \text{KL}(u \parallel p^1) - \text{KL}(u \parallel p^{T+1}) + \sum_{t=1}^T \epsilon^t + \frac{3\eta^2 H^2 T}{2} \\ &\stackrel{\text{(ix)}}{\leq} \text{KL}(u \parallel p^1) + \sum_{t=1}^T \epsilon^t + \frac{3\eta^2 H^2 T}{2} \\ &\leq \sum_{a \in \mathcal{A}} u(a) \log(u(a)) - \sum_{a \in \mathcal{A}} u(a) \log(p^1(a)) + \sum_{t=1}^T \epsilon^t + \frac{3\eta^2 H^2 T}{2} \\ &\stackrel{\text{(x)}}{\leq} \log |\mathcal{A}| + \sum_{t=1}^T \epsilon^t + \frac{3\eta^2 H^2 T}{2}. \end{aligned}$$

(ix) follows from the fact that KL-divergence is non-negative, and (x) stands because the first term is negative, and for the second term, p^1 is a uniform distribution. Dividing both side by η , we have that

$$\sum_{t=1}^T \langle u - p^t, g^t \rangle \leq \frac{\log |\mathcal{A}| + \sum_{t=1}^T \epsilon^t}{\eta} + \frac{3\eta H^2 T}{2}.$$

This concludes the proof. $\square$

In order to obtain a meaningful regret bound, we should set $\eta = \min \left\{ \frac{1}{2H}, \sqrt{\frac{2(\log |\mathcal{A}| + \bar{\epsilon} T)}{3H^2 T}} \right\}$.

Furthermore, we introduce the alternative version of [Assumption 4.1](#).

Assumption B.1 (Bias). *Let φ be the policy feature. Then,*

$$\inf_{\theta} \sup_{t, h, s, a} \left| \langle \varphi(s, a), \theta \rangle - \log(1 + \eta A_h^{\pi^t}(s, a)) \right| \leq \epsilon_{\text{bias}}.$$

Therefore, under [Assumptions 4.2](#) and [B.1](#), we can easily prove that [Lemma 4.2](#) also holds for the projected SPMA, and consequently, all the sample complexity guarantees for the projected NPG should also hold.

B.4 Technical Tools

Lemma B.2 (Three-Point Property of Bregman Divergence). *Suppose $X \subseteq \mathbb{R}^d$ is closed and convex. Consider a strictly convex function $\Phi : X \rightarrow \mathbb{R}$. For all $x \in X$ and $y, z \in \text{int}X$,*

$$D_\Phi(x, y) + D_\Phi(y, z) - D_\Phi(x, z) = \langle \nabla \Phi(z) - \nabla \Phi(y), x - y \rangle.$$

C Constructing $\mathcal{D}_{\text{exp}}$ via Experimental Design

In this section, we introduce various methods of experimental design to bound $\bar{\varphi}_G$ defined in Lemma 4.2. The experimental design problem can be written as

$$\begin{aligned} & \inf_{\substack{\mathcal{D}_{\text{exp}} \in \mathcal{S} \times \mathcal{A} \\ \rho_{\text{exp}} \in \Delta(\mathcal{D}_{\text{exp}})}} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} \|\varphi(s,a)\|_{G^{-1}} \\ \text{s.t. } & G = \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s,a) \varphi(s,a) \varphi(s,a)^\top. \end{aligned}$$

In Appendix C.1, we consider constructing a coreset for the policy features. The Kiefer–Wolfowitz theorem guarantees that there exists a coreset that can ensure that $\bar{\varphi}_G$ is bounded, and that such a coreset has a small $O(d)$ size. Such a coreset can be formed using G-experimental design. In Appendix C.2, we consider using the linear MDP features as the policy features and constructing $\mathcal{D}_{\text{exp}}$ through limited interaction with the environment.

C.1 Kiefer–Wolfowitz Theorem and G-Experimental Design

We first introduce the Kiefer–Wolfowitz theorem (Kiefer and Wolfowitz, 1960) which guarantees that there exists a coreset $\mathcal{D}_{\text{exp}}$ and its corresponding distribution ρ_{exp} that can be used to bound $\bar{\varphi}_G$.

Proposition C.1 (Kiefer–Wolfowitz). *Let $G := \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \rho_{\text{exp}}(s,a) \varphi(s,a) \varphi(s,a)^\top$ be the covariance matrix for any $\mathcal{D}_{\text{exp}} \subset \mathcal{S} \times \mathcal{A}$ and $\rho \in \Delta(\mathcal{D}_{\text{exp}})$. There exists a coreset $\mathcal{D}_{\text{exp}}$ and a distribution ρ_{exp} such that*

$$\sup_{(s,a) \in \mathcal{D}_{\text{exp}}} \|\varphi(s,a)\|_{G^{-1}} \leq 2d_a \quad \text{and} \quad |\mathcal{D}_{\text{exp}}| \leq 4d_a \log \log(d_a + 4) + 28.$$

Note that the size of $\mathcal{D}_{\text{exp}}$ is also bounded by $\tilde{\mathcal{O}}(d_a)$, suggesting that the computation cost of calculating the actor loss over $\mathcal{D}_{\text{exp}}$ is inexpensive. The problem of constructing such a coreset is often framed as G-experimental design, and it can typically be solved using numerous efficient approximation algorithms such as the Franke-Wolfe algorithm (Frank et al., 1956) as mentioned in Todd (2016); Lattimore and Szepesvári (2020). Using $\mathcal{D}_{\text{exp}}$ and ρ_{exp} produced by such methods to construct the actor loss in Algorithm 2 offers the guarantees that $\bar{\varphi}_G \leq \mathcal{O}(d_a)$, which is consequently used to bound the projection error in Lemma 4.2 as $\bar{\epsilon} \leq \mathcal{O}(d_a \epsilon_{\text{bias}} + \epsilon_{\text{opt}})$.

We remark that the coreset construction can be done before the learning process in the actor-critic algorithm since it is independent of the linear MDP environment. However, these algorithms typically require traversing through all the policy features in $\mathcal{S} \times \mathcal{A}$, which is not ideal for large state-action spaces.

C.2 Exploratory Policy and Minimum Eigenvalue

Alternatively, we can choose to use the linear MDP features as the policy features (i.e., $\varphi = \phi$) and construct $\mathcal{D}_{\text{exp}}$ via interacting with the environment. Note that bounding $\bar{\varphi}_G$ is equivalent to controlling $\|\phi(s,a)\|_{G^{-1}}$ for all $(s,a) \in \mathcal{S} \times \mathcal{A}$. Consequently, given that $\|\phi(s,a)\|_2 \leq 1$ by the linear MDP assumption and since

$$\|\phi(s,a)\|_{G^{-1}} \leq \frac{\|\phi(s,a)\|_2}{\lambda_{\min}(G)} = \frac{1}{\lambda_{\min}(G)},$$

we only need a well-conditioned covariance matrix G that has a positive minimum eigenvalue.

Several existing works (Hao et al., 2021; Agarwal et al., 2021b) assume access to an exploratory (not necessarily optimal) policy π_{exp} that is able to collect such covariance matrices with minimum eigenvalue bounded away from 0. Given that, we can directly apply π_{exp} to roll-out trajectories and collect observations, which can be used to construct $\mathcal{D}_{\text{exp}}$ and the corresponding covariance G .

However, in practice, we rarely have access to such an oracle policy. Consequently, Wagenmaker et al. (2022) proposed a reward-free approach, CoverTraj, that can effectively collect such observations without assuming access to an exploratory policy. In particular, the CoverTraj algorithm offers the following theoretical guarantee.

Proposition C.2 (Wagenmaker et al. 2022, Theorem 4). Fix $h \in [H]$ and $\gamma \in [0, 1]$. Suppose there exists a problem-dependent constant $\epsilon_{\mathcal{M}} > 0$ such that $\sup_{\pi \in \Pi} \lambda_{\min}(\mathbb{E}_{\pi}[\phi(s, a) \phi(s, a)^{\top}]) \geq \epsilon_{\mathcal{M}}$. Running K rounds of CoverTraj to collect $\mathcal{D}_{\text{exp}} = \{(s_h^{\tau}, a_h^{\tau})\}_{\tau=1}^K$ where

$$K = \tilde{\mathcal{O}}\left(\frac{1}{\epsilon_{\mathcal{M}}} \cdot \max\left\{\frac{d_c}{\gamma^2}, d_c^4 H^3, \log^3\left(\frac{1}{\delta}\right)\right\}\right)$$

ensures that for any $\delta \in (0, 1)$, with probability of at least $1 - \delta$,

$$\lambda_{\min}(G) \geq \frac{\epsilon_{\mathcal{M}}}{\gamma^2},$$

where $G = \sum_{(s,a) \in \mathcal{D}_{\text{exp}}} \phi(s, a) \phi(s, a)^{\top}$.

Note that CoverTraj does not utilize the reward function of the MDP and merely use the transition kernel when interacting with the environment. Alternatively, Wagenmaker and Jamieson (2022) provides another approach, OptCov, that utilizes regret minimization algorithms to construct the desired covariance matrix. According to Wagenmaker and Jamieson (2022, Theorem 9), OptCov can also offer a similar guarantee of the minimum eigenvalue ensuring that

$$\lambda_{\min}(G) \geq \max\left\{d_c \log\left(\frac{1}{\delta}\right), \epsilon_{\mathcal{M}}\right\}.$$

To conclude, the Frank-Wolfe algorithm can be used to form a coresset and subsequently bound $\bar{\varphi}_G$ for any given policy features. If we use the linear MDP features as the policy features, we can construct $\mathcal{D}_{\text{exp}}$ by interacting with the environment. Either having access to an exploratory policy or running CoverTraj or OptCov can offer guarantees on the minimum eigenvalues of the covariance matrix, which will consequently control $\bar{\varphi}_G$.

D Analyses for the Critic

D.1 Proof of Lemma 5.1

In order to prove Lemma 5.1, we introduce the following “good” event for the estimated value function.

Lemma D.1 (Good Event). *There exists some $C_{\delta} > 0$ such that for any fixed $\delta \in (0, 1)$, the following event,*

$$\mathcal{E}_{\delta} := \left\{ \forall (t, h) \in [T] \times [H] : \left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s, a) [\hat{V}_{h+1}^t(s') - \mathbb{P}_h \hat{V}_{h+1}^t(s, a)] \right\|_{(\Lambda_h^t)^{-1}} \leq C_{\delta} H \sqrt{d_c} \right\},$$

holds with probability at least $1 - \delta$ (i.e., $\Pr(\mathcal{E}_{\delta}) \geq 1 - \delta$).

The exact definition of C_{δ} varies between the on-policy and the off-policy settings. We will prove that $\Pr(\mathcal{E}_{\delta}) \geq 1 - \delta$ for the on-policy and off-policy setting in Appendix E and Appendix F respectively where C_{δ} will be set to C_{δ}^{on} in Lemma E.1 and C_{δ}^{off} in Lemma F.2.

Next, conditioned on the above event, we present a formal version of Lemma 5.1, which provides an upper and a lower bound for the model prediction error induced by the LMC critic.

Lemma D.2 (Formal version of Lemma 5.1). *Consider Algorithm 1 with the LMC critic from Algorithm 3. Conditioned on $\mathcal{E}_{\delta}$ defined in Lemma D.1, if we choose that $\lambda = 1$, $\zeta = (2H\sqrt{d_c}C_{\delta} + 8/3)^{-2}$, $\alpha_c^{h,t} = 1/(2\lambda_{\max}(\Lambda_h^t))$, $J_t \geq 2\kappa_t \log(1/\sigma)$, and $M = \log(HT/\delta)/\log(1/(1-c))$ where $\kappa_t = \max_{h \in [H]} \lambda_{\max}(\Lambda_h^t)/\lambda_{\min}(\Lambda_h^t)$, $\sigma = 1/(4H(|\mathcal{D}^t| + 1)\sqrt{d_c})$, and $c = 1/(2\sqrt{2e\pi})$, then, for all $(t, h, s, a) \in [T] \times [H] \times \mathcal{S} \times \mathcal{A}$ and for any $\delta \in (0, 1)$, with probability at least $1 - \delta$,*

$$-\Gamma_{\text{LMC}} \times \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \leq \ell_h^t(s, a) \leq 0, \quad (10)$$

where $\Gamma_{\text{LMC}} = C_{\delta} H \sqrt{d_c} + \frac{4}{3} \sqrt{\frac{2d_c \log(1/\delta)}{3\zeta}} + \frac{4}{3}$.

D.1.1 Preliminary Properties

In this section, we introduce some useful properties of LMC and state the supporting lemmas that will be helpful in proving the above result.

First, we obtain the derivative of the critic loss defined in [Algorithm 3](#).

$$\begin{aligned}\nabla L_h^t(w_h) &= \Lambda_h^t w_h - b_h^t, \text{ where} \\ \Lambda_h^t &:= \sum_{(s,a) \in \mathcal{D}_h^t} \phi(s,a) \phi(s,a)^\top + \lambda I, \\ b_h^t &:= \sum_{(s,a,s') \in \mathcal{D}_h^t} \left[r_h(s,a) + \widehat{V}_{h+1}^t(s') \right] \phi(s,a).\end{aligned}\tag{11}$$

Consequently, by setting $\nabla L_h^t(w_h) = 0$, we get the minimizer of $L_h^t(w_h)$ as

$$\widehat{w}_h^t := (\Lambda_h^t)^{-1} b_h^t.\tag{12}$$

We now introduce the following lemma, showing that the noisy gradient descent performed by the LMC critic ensures that the sampled critic parameter w follows a Gaussian distribution.

Lemma D.3 ([Ishfaq et al. 2024a](#), Proposition B.1). *Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). For any $(t, h, m) \in [T] \times [H] \times [M]$, the sampled parameters w_h^{t,m,J_t} follows a Gaussian distribution $\mathbf{N}(\mu_h^{t,m,J_t}, \Sigma_h^{t,m,J_t})$. The mean and the covariance are defined as*

$$\mu_h^{t,J_t} = A_t^{J_t} \dots A_1^{J_1} w^{1,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \widehat{w}_h^i,\tag{13}$$

$$\Sigma_h^{t,J_t} = \frac{1}{\zeta} \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t},\tag{14}$$

where $A_t := I - \alpha_c^t \Lambda_h^t$ for all $t \in [T]$.

Since w_h^{t,m,J_t} follows the Gaussian distribution of $\mathbf{N}(\mu_h^{t,m,J_t}, \Sigma_h^{t,m,J_t})$, $\langle \phi_h(s,a), w_h^{t,m,J_t} \rangle$ also follows the Gaussian distribution of $\mathbf{N}(\phi_h(s,a)^\top \mu_h^{t,m,J_t}, \phi_h(s,a)^\top \Sigma_h^{t,m,J_t} \phi_h(s,a))$. Therefore, we introduce the following lemmas to bound the terms related to the mean and variance.

Lemma D.4. *Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). If we follow the hyperparameter choices of [Lemma D.2](#), then for any $(s,a) \in \mathcal{S} \times \mathcal{A}$,*

$$\left| \langle \phi(s,a), (\mu_h^{t,J_t} - \widehat{w}_h^t) \rangle \right| \leq \frac{4}{3} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}.$$

Lemma D.5. *Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). If we follow the hyperparameter choices of [Lemma D.2](#), then for any $(s,a) \in \mathcal{S} \times \mathcal{A}$,*

$$\frac{1}{2\sqrt{6}\zeta} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \leq \|\phi(s,a)\|_{\Sigma_h^{t,m,J_t}} \leq \frac{4}{3} \sqrt{\frac{2}{3\zeta}} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}.$$

Additionally, we outline the necessary supporting lemmas that are useful for bounding the model prediction error. Recall that $|\mathcal{D}^t| := \sup_{h \in [H]} |\mathcal{D}_h^t|$ represents the number of trajectories in $\mathcal{D}^t$ or the number of (s_h, a_h, s_{h+1}) tuples in $\mathcal{D}_h^t$, where $|\mathcal{D}^t| = N$ in the on-policy setting, and $|\mathcal{D}^t| = t$ in the off-policy setting.

Lemma D.6. *Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). For any $(t, h) \in [T] \times [H]$, it holds that*

$$\|\widehat{w}_h^t\|_2 \leq 2H \sqrt{d_c |\mathcal{D}^t| / \lambda}.$$

Lemma D.7. *Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). If we follow the hyperparameter choices of [Lemma D.2](#), then for any $(t, m, h) \in [T] \times [M] \times [H]$ and for any $\delta \in (0, 1)$, with probability at least $1 - \delta$,*

$$\left\| w_h^{t,m,J_t} \right\|_2 \leq \overline{W}_\delta^t := \frac{16}{3} H \sqrt{d_c |\mathcal{D}^t|} + \sqrt{\frac{2 d_c^3 t}{3 \zeta \delta}}.$$

Lemma D.8. Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). If we follow the hyperparameter choices of [Lemma D.2](#), then for any $(t, m, h, s, a) \in [T] \times [M] \times [H] \times \mathcal{S} \times \mathcal{A}$ and for any $\delta \in (0, 1)$, with probability at least $1 - \delta$,

$$\left| \left\langle \phi(s, a), \hat{w}_h^t - w_h^{t, m, J_t} \right\rangle \right| \leq \left(\frac{8}{3} \sqrt{\frac{2 d_c \log(1/\delta)}{3 \zeta}} + \frac{4}{3} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.$$

Lemma D.9. Consider [Algorithm 1](#) with the LMC critic from [Algorithm 3](#). Conditioned on $\mathcal{E}_\delta$ defined in [Lemma D.1](#), if we follow the hyperparameter choices of [Lemma D.2](#), then for any $(t, h, s, a) \in [T] \times [H] \times \mathcal{S} \times \mathcal{A}$ and for any $\delta \in (0, 1)$, it holds that

$$\left| \left\langle \phi(s, a), \hat{w}_h^t \right\rangle - r_h(s, a) - \mathbb{P}_h \hat{V}_{h+1}^t(s, a) \right| \leq 3 C_\delta H \sqrt{d_c} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.$$

D.1.2 Main Analysis

We will use the above lemmas to complete the main proof in this section.

Proof of [Lemma D.2](#).

Optimism (RHS of [Eq. \(10\)](#)) Using the definition of the model prediction error, we need to show that with high probability, $\hat{Q}_h^t(s, a) \geq r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a)$. Recall that $\hat{Q}_h^t(s, a) = \min \left\{ \left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle, H - h + 1 \right\}$. Since $r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) \leq H - h + 1$, when $\left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle > H - h + 1$, the statement is trivially true. Thus, we only need to consider the case when $\left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle \leq H - h + 1$ and thus $\hat{Q}_h^t(s, a) = \left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle$.

Based on the mean and covariance matrix defined in [Lemma D.3](#), we have that $\left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle$ follows the distribution $\mathbf{N} \left(\phi(s, a)^\top \mu_h^{t, J_t}, \phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a) \right)$.

In order to prove that $\hat{Q}_h^t(s, a) \geq r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a)$, we consider the following variable $X_t := \frac{r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) - \left\langle \phi(s, a), \mu_h^{t, J_t} \right\rangle}{\sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)}}$ and will next show that $|X_t| \leq 1$. First, we have that

$$\begin{aligned} & \left| r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) - \left\langle \phi(s, a), \mu_h^{t, J_t} \right\rangle \right| \\ & \stackrel{(i)}{\leq} \left| r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) - \left\langle \phi(s, a), \hat{w}_h^t \right\rangle \right| + \left| \left\langle \phi(s, a), \hat{w}_h^t - \mu_h^{t, J_t} \right\rangle \right| \\ & \stackrel{(ii)}{\leq} C_\delta H \sqrt{d_c} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \frac{4}{3} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\ & = \left(C_\delta H \sqrt{d_c} + \frac{4}{3} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}, \end{aligned}$$

where (i) uses the triangular inequality, and (ii) is implied by [Lemmas D.4](#) and [D.9](#). Therefore,

$$\begin{aligned} |X_t| &= \left| \frac{r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) - \left\langle \phi(s, a), \mu_h^{t, J_t} \right\rangle}{\sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)}} \right| \\ &\leq \sqrt{\zeta} \left(2 H \sqrt{d_c} C_\delta + 8/3 \right). \end{aligned}$$

Since we choose $\zeta = (2 H \sqrt{d_c} C_\delta + 8/3)^{-2}$, we have that $|X_t| \leq 1$. Then, using [Lemma D.12](#), we can get that

$$\begin{aligned} & \Pr \left(\left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle \geq r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) \right) \\ &= \Pr \left(\frac{\left\langle \phi(s, a), w_h^{t, m, J_t} \right\rangle - \left\langle \phi(s, a), \mu_h^{t, J_t} \right\rangle}{\sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)}} \geq \frac{r_h(s, a) + \mathbb{P}_h \hat{V}_{h+1}^t(s, a) - \left\langle \phi(s, a), \mu_h^{t, J_t} \right\rangle}{\sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)}} \right) \end{aligned}$$

$$\begin{aligned}
&= \Pr \left(\frac{\langle \phi(s, a), w_h^{t, m, J_t} \rangle - \langle \phi(s, a), \mu_h^{t, J_t} \rangle}{\sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)}} \geq X_t \right) \\
&\geq \frac{1}{2\sqrt{2\pi}} \exp(-X_t^2/2) \\
&\geq \frac{1}{2\sqrt{2e\pi}}.
\end{aligned}$$

The above result holds for any $m \in [M]$. Since we have M parallel critic parameters, it holds that

$$\begin{aligned}
&\Pr(\exists(s, a) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^t(s, a) \leq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)) \\
&= \Pr\left(\exists(s, a) \in \mathcal{S} \times \mathcal{A} : \max_{m \in [M]} \widehat{Q}_h^{t, m}(s, a) \leq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right) \\
&= \Pr\left(\exists(s, a) \in \mathcal{S} \times \mathcal{A} : \forall m \in [M], \widehat{Q}_h^{t, m}(s, a) \leq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right) \\
&\leq \Pr\left(\forall m \in [M], \exists(s^m, a^m) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^{t, m}(s^m, a^m) \leq r_h(s^m, a^m) + \mathbb{P}_h \widehat{V}_{h+1}^t(s^m, a^m)\right) \\
&= \prod_{m=1}^M \Pr\left(\exists(s, a) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^{t, m}(s, a) \leq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right) \\
&= \prod_{m=1}^M \left(1 - \Pr\left(\forall(s, a) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^{t, m}(s, a) \geq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right)\right) \\
&= \prod_{m=1}^M \left(1 - \Pr\left(\forall(s, a) \in \mathcal{S} \times \mathcal{A} : \langle \phi(s, a), w_h^{t, m, J_t} \rangle \geq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right)\right) \\
&\leq \left(1 - \frac{1}{2\sqrt{2e\pi}}\right)^M.
\end{aligned}$$

This further implies that

$$\begin{aligned}
&\Pr(\forall(s, a) \in \mathcal{S} \times \mathcal{A} : \iota_h^t(s, a) \leq 0) \\
&= \Pr\left(\forall(s, a) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^t(s, a) \geq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right) \\
&= 1 - \Pr\left(\exists(s, a) \in \mathcal{S} \times \mathcal{A} : \widehat{Q}_h^t(s, a) \leq r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)\right) \\
&= 1 - \left(1 - \frac{1}{2\sqrt{2e\pi}}\right)^M.
\end{aligned}$$

Let $1 - \left(1 - \frac{1}{2\sqrt{2e\pi}}\right)^M \geq 1 - \delta/(HT)$, which yields that $M = \log(HT/\delta)/\log(1/(1-c))$ where $c = 1/(2\sqrt{2e\pi})$. Therefore, we have that

$$\Pr(\iota_h^t(s, a) \leq 0, \forall(s, a) \in \mathcal{S} \times \mathcal{A}) \geq 1 - \frac{\delta}{HT}.$$

Applying union bound over $[H]$ and $[T]$, we have that $\iota_h^t(s, a) \leq 0$ with probability $1 - \delta$.

Error Bound (LHS of Eq. (10)) We can lower bound ι_h^t as follows.

$$\begin{aligned}
-\iota_h^t(s, a) &= \widehat{Q}_h^t(s, a) - r_h(s, a) - \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) \\
&= \min \left\{ \max_{m \in [M]} \langle \phi(s, a), w_h^{t, m, J_t} \rangle, H - h + 1 \right\}^+ - r_h(s, a) - \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) \\
&\leq \max_{m \in [M]} \langle \phi(s, a), w_h^{t, m, J_t} \rangle - r_h(s, a) - \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) \\
&= \max_{m \in [M]} \langle \phi(s, a), w_h^{t, m, J_t} \rangle - \langle \phi(s, a), \widehat{w}_h^t \rangle + \langle \phi(s, a), \widehat{w}_h^t \rangle - r_h(s, a) - \mathbb{P}_h \widehat{V}_{h+1}^t(s, a)
\end{aligned}$$

$$\begin{aligned}
&\leq \left| \max_{m \in [M]} \langle \phi(s, a), w_h^{t,m,J_t} \rangle - \langle \phi(s, a), \hat{w}_h^t \rangle \right| + \left| \langle \phi(s, a), \hat{w}_h^t \rangle - r_h(s, a) - \mathbb{P}_h \hat{V}_{h+1}^t(s, a) \right| \\
&\stackrel{\text{(iii)}}{\leq} \left(C_\delta H \sqrt{d_c} + \frac{4}{3} \sqrt{\frac{2 d_c \log(1/\delta)}{3 \zeta}} + \frac{4}{3} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}},
\end{aligned}$$

where (iii) is derived from [Lemmas D.8](#) and [D.9](#).

This concludes the proof. $\square$

D.2 Proofs of Preliminary Properties

D.2.1 Proof of [Lemma D.3](#)

Proof. For any $(t, m) \in [T] \times [M]$, the critic update rule at j -th round can be written as

$$w_h^{t,m,j} = w_h^{t,m,j-1} - \alpha_c^{h,t} \nabla L_h^t(w_h^{t,m,j-1}) + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \nu_h^{t,m,j}.$$

Considering $j = J_t$ and plugging in [Eq. \(11\)](#), we have that

$$\begin{aligned}
w_h^{t,m,J_t} &= w_h^{t,m,J_t-1} - \alpha_c^{h,t} \left(\Lambda_h^t w_h^{t,m,J_t-1} - b_h^t \right) + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \nu_h^{t,m,J_t} \\
&= (I - \alpha_c^{h,t} \Lambda_h^t) w_h^{t,m,J_t-1} + \alpha_c^{h,t} b_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \nu_h^{t,m,J_t} \\
&\stackrel{\text{(i)}}{=} (I - \alpha_c^{h,t} \Lambda_h^t)^{J_t} w_h^{t,m,0} + \sum_{l=0}^{J_t-1} (I - \alpha_c^{h,t} \Lambda_h^t)^l \left(\alpha_c^{h,t} b_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \nu_h^{t,m,J_t-l} \right) \\
&= (I - \alpha_c^{h,t} \Lambda_h^t)^{J_t} w_h^{t,m,0} + \alpha_c^{h,t} \sum_{l=0}^{J_t-1} (I - \alpha_c^{h,t} \Lambda_h^t)^l b_h^t \\
&\quad + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \sum_{l=0}^{J_t-1} (I - \alpha_c^{h,t} \Lambda_h^t)^l \nu_h^{t,m,J_t-l} \\
&\stackrel{\text{(ii)}}{=} A_t^{J_t} w_h^{t,m,0} + \alpha_c^{h,t} \sum_{l=0}^{J_t-1} A_t^l \Lambda_h^t \hat{w}_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \sum_{l=0}^{J_t-1} A_t^l \nu_h^{t,m,J_t-l} \\
&\stackrel{\text{(iii)}}{=} A_t^{J_t} w_h^{t,m,0} + (I - A_t) \left(A_t^0 + A_t^1 + \dots + A_t^{J_t-1} \right) \hat{w}_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \sum_{l=0}^{J_t-1} A_t^l \nu_h^{t,m,J_t-l} \\
&\stackrel{\text{(iv)}}{=} A_t^{J_t} w_h^{t,m,0} + (I - A_t^{J_t}) \hat{w}_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \sum_{l=0}^{J_t-1} A_t^l \nu_h^{t,m,J_t-l}.
\end{aligned}$$

(i) comes from telescoping the previous equation from $l = 0$ to $J_t - 1$. (ii) uses the definition that $A_t = I - \alpha_c^{h,t} \Lambda_h^t$ and $b_h^t = \Lambda_h^t \hat{w}_h^t$. (iii) uses the definition of A_t . (iv) follows from $I + A + \dots + A^{n-1} = (I - A^n)(I - A)^{-1}$. Since we set $\alpha_c^{h,t} = 1/(2 \lambda_{\max}(\Lambda_h^t))$, A_t satisfies $I \succ A_t \succ 0$ for all $t \in [T]$. Note that we warm-start the parameters from the previous episode and set $w_h^{t,m,0} = w_h^{t-1,m,J_{t-1}}$. Therefore, by telescoping the above equation from $i = 0$ to t , we further have that

$$\begin{aligned}
w_h^{t,m,J_t} &= A_t^{J_t} w_h^{t-1,m,J_{t-1}} + (I - A_t^{J_t}) \hat{w}_h^t + \sqrt{\alpha_c^{h,t} \zeta^{-1}} \sum_{l=0}^{J_t-1} A_t^l \nu_h^{t,m,J_t-l} \\
&= A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \hat{w}_h^i \\
&\quad + \sum_{i=1}^t \sqrt{\alpha_c^{h,i} \zeta^{-1}} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} \sum_{l=0}^{J_i-1} A_i^l \nu_h^{i,J_i-l}.
\end{aligned}$$

Note that if $\xi \sim \mathbf{N}(0, I_{d \times d})$, then we have that $A\xi + \mu \sim \mathbf{N}(\mu, A A^\top)$ for any $A \in \mathbb{R}^{d \times d}$ and $\mu \in \mathbb{R}^d$. This implies that w_h^{t,m,J_t} follows the Gaussian distribution $N(\mu_h^{t,m,J_t}, \Sigma_h^{t,m,J_t})$, where

$$\mu_h^{t,m,J_t} = A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \hat{w}_h^i.$$

We then derive the covariance matrix Σ_h^{t,m,J_t} . For any $i \in [t]$, we denote that $\mathcal{A}_{i+1} = A_t^{J_t} \dots A_{i+1}^{J_{i+1}}$. Therefore,

$$\begin{aligned} \sqrt{\alpha_c^i \zeta^{-1}} \mathcal{A}_{i+1} \sum_{l=0}^{J_i-1} A_i^l \nu_h^{i,J_i-l} &= \sum_{l=0}^{J_i-1} \sqrt{\alpha_c^i \zeta^{-1}} \mathcal{A}_{i+1} A_i^l \nu_h^{i,J_i-l} \\ &\sim \mathbf{N}\left(0, \sum_{l=0}^{J_i-1} \alpha_c^i \zeta^{-1} \mathcal{A}_{i+1} A_i^l (\mathcal{A}_{i+1} A_i^l)^\top\right) \sim \mathbf{N}\left(0, \alpha_c^i \zeta^{-1} \mathcal{A}_{i+1} \left(\sum_{l=0}^{J_i-1} A_i^{2l}\right) \mathcal{A}_{i+1}^\top\right). \end{aligned}$$

This further implies that

$$\begin{aligned} \Sigma_h^{t,m,J_t} &= \sum_{i=1}^t \alpha_c^i \zeta^{-1} \mathcal{A}_{i+1} \left(\sum_{l=0}^{J_i-1} A_i^{2l}\right) \mathcal{A}_{i+1}^\top \\ &= \sum_{i=1}^t \alpha_c^i \zeta^{-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} \left(\sum_{l=0}^{J_i-1} A_i^{2l}\right) A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \\ &\stackrel{(v)}{=} \sum_{i=1}^t \alpha_c^i \zeta^{-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{2J_i}) (I - A_i^2)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \\ &= \sum_{i=1}^t \alpha_c^i \zeta^{-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{2J_i}) (\Lambda_h^i) (\Lambda_h^i)^{-1} (I - A_i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \\ &\stackrel{(vi)}{=} \sum_{i=1}^t \zeta^{-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{2J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t}. \end{aligned}$$

(v) uses the fact that $I + A + \dots + A^{n-1} = (I - A^n)(I - A)^{-1}$, and (vi) uses the fact that $\alpha_c^{h,t} \Lambda_h^t = I - A_t$.

This concludes the proof. $\square$

D.2.2 Proof of Lemma D.4

Proof. Using Lemma D.3, we first have that

$$\begin{aligned} \mu_h^{t,J_t} &= A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \hat{w}_h^i \\ &= A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} \hat{w}_h^i - \sum_{i=1}^t A_t^{J_t} \dots A_i^{J_i} \hat{w}_h^i \\ &= A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^{t-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (\hat{w}_h^i - \hat{w}_h^{i+1}) - A_t^{J_t} \dots A_1^{J_1} \hat{w}_h^1 + \hat{w}_h^t \\ &= A_t^{J_t} \dots A_1^{J_1} (w_h^{1,m,0} - \hat{w}_h^1) + \sum_{i=1}^{t-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (\hat{w}_h^i - \hat{w}_h^{i+1}) + \hat{w}_h^t. \end{aligned}$$

This implies that

$$\begin{aligned} &\left| \left\langle \phi(s, a), \left(\mu_h^{t,J_t} - \hat{w}_h^t \right) \right\rangle \right| \\ &= \phi(s, a)^\top A_t^{J_t} \dots A_1^{J_1} (w_h^{1,m,0} - \hat{w}_h^1) + \phi(s, a)^\top \sum_{i=1}^{t-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (\hat{w}_h^i - \hat{w}_h^{i+1}) \end{aligned}$$

$$\begin{aligned}
& \stackrel{(i)}{=} \left| \phi(s, a)^\top \sum_{i=0}^{t-1} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (\hat{w}_h^i - \hat{w}_h^{i+1}) \right| \\
& = \left| \sum_{i=0}^{t-1} \phi(s, a)^\top A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (\hat{w}_h^i - \hat{w}_h^{i+1}) \right| \\
& \stackrel{(ii)}{\leq} \sum_{i=0}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_2 \|\hat{w}_h^i - \hat{w}_h^{i+1}\|_2 \\
& \stackrel{(iii)}{\leq} \sum_{i=0}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_2 (\|\hat{w}_h^i\|_2 + \|\hat{w}_h^{i+1}\|_2) \\
& \stackrel{(iv)}{\leq} 4H \sqrt{d_{\mathbf{c}} |\mathcal{D}^t| / \lambda} \sum_{i=0}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_2 \\
& \stackrel{(v)}{\leq} 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}} / \lambda} \sum_{i=0}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
& \stackrel{(vi)}{\leq} 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}}} \sum_{i=0}^{t-1} \sigma^{t-i} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
& \stackrel{(vii)}{\leq} 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}}} \left(\sum_{i=0}^{t-1} \sigma^{t-i} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
& \leq 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}}} \left(\sum_{i=0}^{t-1} \sigma^{t-i} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
& = 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}}} \left(\sum_{i=1}^{t-1} \sigma^i \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
& \stackrel{(viii)}{\leq} 4H (|\mathcal{D}^t| + 1) \sqrt{d_{\mathbf{c}}} \left(\frac{\sigma}{1 - \sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
& = \left(\frac{1}{1 - \sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
& \leq \frac{4}{3} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

For (i), we choose $w_h^{1,m,0} = \mathbf{0}$ and denote that $\hat{w}_h^0 = \mathbf{0}$. (ii) comes from $A_i \prec (1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j)) I$ and the Hölder's inequality. (iii) uses the triangular inequality. (iv) uses [Lemma D.6](#). (v) uses the fact that $\|\phi(s, a)\| \leq \sqrt{|\mathcal{D}^t| + 1} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}}$. (vi) hold because we set $\lambda = 1$ and uses [Lemma D.16](#) by setting $J_j \geq \kappa_j \log(1/\sigma)$ where $\sigma = 1/(4H(|\mathcal{D}^t| + 1)\sqrt{d_{\mathbf{c}}})$. (vii) follows from $\|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \leq \|\phi(s, a)\|_2 \leq \sqrt{|\mathcal{D}^t| + 1} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}$. (viii) follows from $\sum_{i=1}^t \sigma^i \leq \sum_{i=1}^\infty \sigma^i \leq \sigma/(1 - \sigma)$.

This concludes the proof. $\square$

D.2.3 Proof of [Lemma D.5](#)

Proof. We first bound the RHS. Using [Lemma D.3](#), we have that

$$\begin{aligned}
& \phi(s, a)^\top \sum_h^{t, J_t} \phi(s, a) \\
& = \frac{1}{\zeta} \sum_{i=1}^t \phi(s, a)^\top A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A^{2J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \phi(s, a)
\end{aligned}$$

$$\begin{aligned}
&\stackrel{(i)}{=} \frac{1}{\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (I - A^{2J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \\
&\stackrel{(ii)}{\leq} \frac{2}{3\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} \left((\Lambda_h^i)^{-1} - A_i^{J_i} (\Lambda_h^i)^{-1} A_i^{J_i} \right) \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{2}{3\zeta} \left(\sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) - \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_i (\Lambda_h^i)^{-1} \mathcal{A}_i^\top \phi(s, a) \right) \\
&\stackrel{(iii)}{\leq} \frac{2}{3\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{2}{3\zeta} \left(\|\phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2 + \sum_{i=1}^{t-1} \|\mathcal{A}_{i+1}^\top \phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2 \right) \\
&\leq \frac{2}{3\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 + \frac{2}{3\zeta} \sum_{i=1}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_c \lambda_{\min}(\Lambda_h^j) \right)^{2J_j} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2.
\end{aligned}$$

For (i), we denote $\mathcal{A}_{i+1} = A_i^{J_t} \dots A_{i+1}^{J_{i+1}}$. (ii) follows from $I + A_i \succeq \frac{3}{2}I$ since we set $\alpha_c^{h,j} = 1/(2\lambda_{\max}(\Lambda_h^j))$. In particular, it is trivial that A and $(\Lambda_h^t)^{-1}$ are commuting matrices. Hence,

$$\begin{aligned}
A^{2J_i} (\Lambda_h^i)^{-1} &= A^{2J_i-1} (I - \alpha_c^{h,t} \Lambda_h^t) (\Lambda_h^t)^{-1} \\
&= A^{2J_i-1} (\Lambda_h^t)^{-1} (I - \alpha_c^{h,t} \Lambda_h^t) \\
&= A^{2J_i-1} (\Lambda_h^t)^{-1} A \\
&\vdots \\
&= A^{J_i} (\Lambda_h^t)^{-1} A^{J_i}.
\end{aligned}$$

(iii) follows from the fact that $\sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_i (\Lambda_h^i)^{-1} \mathcal{A}_i^\top \phi(s, a) > 0$. Therefore,

$$\begin{aligned}
&\left\| \phi(s, a)^\top \left(\Sigma_h^{t, J_t} \right)^{1/2} \right\|_2 = \sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)} \\
&\stackrel{(iv)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2}{3\zeta}} \sum_{i=1}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_c \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
&\stackrel{(v)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2}{3\zeta}} \sum_{i=1}^{t-1} \sigma^{t-i} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
&\stackrel{(vi)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2(|\mathcal{D}^t|+1)}{3\zeta}} \left(\sum_{i=1}^{t-1} \sigma^{t-i} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{(vii)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2(|\mathcal{D}^t|+1)}{3\zeta}} \left(\frac{\sigma}{1-\sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \frac{1}{4} \sqrt{\frac{2}{3\zeta}} \left(\frac{1}{1-\sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \left(\sqrt{\frac{2}{3\zeta}} + \frac{1}{3} \sqrt{\frac{2}{3\zeta}} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \frac{4}{3} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

(iv) follows from the fact that $\sqrt{a+b} \leq a+b$ for all $a, b > 0$. (v) uses [Lemma D.16](#) by setting $J_j \geq 2\kappa_j \log(1/\sigma)$ where $\sigma = 1/(4H(|\mathcal{D}^t|+1)\sqrt{d_c})$. (vi) follows from $\|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \leq \|\phi(s, a)\|_2 \leq \sqrt{|\mathcal{D}^t|+1} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}$. (vii) follows from $\sum_{i=1}^t \sigma^{t-i} \leq \sum_{i=1}^\infty \sigma^i \leq \sigma/(1-\sigma)$.

We then proceed to bound the LHS. Using the definition of Σ_h^{t, J_t} from [Eq. \(14\)](#), we have

$$\begin{aligned}
& \phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a) \\
&= \sum_{i=1}^t \frac{1}{\zeta} \phi(s, a)^\top A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A^{2J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \phi(s, a) \\
&\stackrel{\text{(iii)}}{\geq} \frac{1}{2\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (I - A^{2J_i}) (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{1}{2\zeta} \sum_{i=1}^t \frac{1}{2\zeta} \phi(s, a)^\top \mathcal{A}_{i+1} \left((\Lambda_h^i)^{-1} - A_t^{J_t} (\Lambda_h^i)^{-1} A_t^{J_t} \right) \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{1}{2\zeta} \sum_{i=1}^{t-1} \phi(s, a)^\top \mathcal{A}_{i+1} \left((\Lambda_h^i)^{-1} - (\Lambda_h^{i+1})^{-1} \right) \mathcal{A}_{i+1}^\top \phi(s, a) \\
&\quad - \frac{1}{2\zeta} \phi(s, a)^\top A_t^{J_t} \dots A_1^{J_1} (\Lambda_h^1)^{-1} A_1^{J_1} \dots A_t^{J_t} \phi(s, a) + \frac{1}{2\zeta} \phi(s, a)^\top (\Lambda_h^t)^{-1} \phi(s, a),
\end{aligned}$$

where (iii) follows from $(I + A_t)^{-1} \succeq \frac{1}{2} I$ for all $t \in [T]$.

$$\begin{aligned}
& \left| \phi(s, a)^\top \mathcal{A}_{i+1} \left((\Lambda_h^i)^{-1} - (\Lambda_h^{i+1})^{-1} \right) \mathcal{A}_{i+1}^\top \phi(s, a) \right| \\
&\leq \left| \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \right| \\
&\quad + \left| \langle \phi(s, a), \mathcal{A}_{i+1} (\Lambda_h^{i+1})^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \rangle \right| \\
&\leq \left\| \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1/2} \right\|^2 + \left\| \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^{i+1})^{-1/2} \right\|^2 \\
&= \prod_{j=i+1}^t \left(1 - \alpha_c^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{2J_j} \left(\|\phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2 + \|\phi(s, a)\|_{(\Lambda_h^{i+1})^{-1}}^2 \right) \\
&\leq 2 \prod_{j=i+1}^t \left(1 - \alpha_c^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{2J_j} \|\phi(s, a)\|_2^2,
\end{aligned}$$

where we used $0 < \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \leq \|\phi(s, a)\|_2$. Therefore, we have that

$$\begin{aligned}
& \phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a) \\
&\geq \frac{1}{2\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 - \frac{1}{2\zeta} \prod_{i=1}^t \left(1 - \alpha_c^{h,j} \lambda_{\min}(\Lambda_h^i) \right)^{2J_i} \|\phi(s, a)\|_2^2 \\
&\quad - \frac{1}{\zeta} \sum_{i=1}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_c^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{2J_j} \|\phi(s, a)\|_2^2 \\
&\stackrel{\text{(iv)}}{\geq} \frac{1}{2\zeta} \left(\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 - \sigma^t \|\phi(s, a)\|_2^2 - \sum_{i=1}^{t-1} 2\sigma^i \|\phi(s, a)\|_2^2 \right) \\
&\stackrel{\text{(v)}}{\geq} \frac{1}{2\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 \left(1 - (|\mathcal{D}^t|+1)\sigma^t - 2(|\mathcal{D}^t|+1) \sum_{i=1}^{t-1} \sigma^i \right) \\
&\geq \frac{1}{2\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 \left(1 - \sigma^{t-1} - \frac{1}{2(1-\sigma)} \right) \\
&\geq \frac{1}{2\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 \left(1 - \frac{1}{4} - \frac{2}{3} \right)
\end{aligned}$$

$$= \frac{1}{24\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2,$$

where (iv) uses [Lemma D.16](#) by setting $J_j \geq 2\kappa_j \log(1/\sigma)$ where $\sigma = 1/(4H(|\mathcal{D}^t| + 1)\sqrt{d_c})$, and (v) use $\|\phi(s, a)\|_2 \leq \sqrt{|\mathcal{D}^t| + 1} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}$.

This concludes the proof. $\square$

D.2.4 Proof of [Lemma D.6](#)

Proof. Given the definition of $\hat{w}_h^t$ in [Eq. \(12\)](#), we have that

$$\begin{aligned} \|\hat{w}_h^t\| &= \left\| (\Lambda_h^t)^{-1} \sum_{(s,a) \in \mathcal{D}_h^t} [r_h(s, a) + \hat{V}_{h+1}^t(s)] \cdot \phi(s, a) \right\| \\ &\leq \sqrt{\frac{|\mathcal{D}^t|}{\lambda}} \left(\sum_{(s,a) \in \mathcal{D}_h^t} \left\| [r_h(s, a) + \hat{V}_{h+1}^t(s)] \cdot \phi(s, a) \right\|_{(\Lambda_h^t)^{-1}}^2 \right)^{1/2} \\ &\leq 2H \sqrt{\frac{|\mathcal{D}^t|}{\lambda}} \left(\sum_{(s,a) \in \mathcal{D}_h^t} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 \right)^{1/2} \\ &\leq 2H \sqrt{d_c |\mathcal{D}^t| / \lambda}, \end{aligned}$$

where the first inequality follows from [Lemma D.15](#), the second inequality is due to the fact that $V_h^t \in [0, H]$ and the reward function is bounded by 1, and the last inequality follows from [Lemma D.10](#). $\square$

D.2.5 Proof of [Lemma D.7](#)

Proof. From [Lemma D.3](#), we know w_h^{t,m,J_t} follows Gaussian distribution $\mathbf{N}(\mu_h^{t,J_t}, \Sigma_h^{t,J_t})$. Therefore, we have that

$$\|w_h^{t,m,J_t}\|_2 = \|\mu_h^{t,J_t} + \xi_h^{t,J_t}\|_2 \leq \underbrace{\|\mu_h^{t,J_t}\|_2}_{\text{(I)}} + \underbrace{\|\xi_h^{t,J_t}\|_2}_{\text{(II)}},$$

where $\xi_h^{t,J_t} \sim \mathbf{N}(0, \Sigma_h^{t,J_t})$. We first start by bounding Term (I). Given [Lemma D.3](#), by setting $w_h^{1,m,0} = \mathbf{0}$, we can obtain that

$$\begin{aligned} \|\mu_h^{t,J_t}\|_2 &= \left\| A_t^{J_t} \dots A_1^{J_1} w_h^{1,m,0} + \sum_{i=1}^t A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \hat{w}_h^i \right\|_2 \\ &\leq \sum_{i=1}^t \left\| A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \hat{w}_h^i \right\|_2 \\ &\stackrel{\text{(i)}}{\leq} \sum_{i=1}^t \left\| A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \right\|_2 \|\hat{w}_h^i\|_2 \\ &\stackrel{\text{(ii)}}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \left\| A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A_i^{J_i}) \right\|_2 \\ &\stackrel{\text{(iii)}}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \|A_t\|_2^{J_t} \dots \|A_{i+1}\|_2^{J_{i+1}} \left\| (I - A_i^{J_i}) \right\|_2 \\ &\stackrel{\text{(iv)}}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \prod_{j=i+1}^t \left(1 - \alpha_c^{h,j} \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \left(\|I\|_2 + \|A_i^{J_i}\|_2 \right) \end{aligned}$$

$$\begin{aligned}
&\stackrel{(v)}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j)\right)^{J_j} \left(\|I\|_2 + \|A_i\|_2^{J_i}\right) \\
&\stackrel{(vi)}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j)\right)^{J_j} \left(1 + (1 - \alpha_{\mathbf{c}}^i \lambda_{\min}(\Lambda_h^i))^{J_i}\right) \\
&\leq 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \left(\prod_{j=i+1}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j)\right)^{J_j} + \prod_{j=i}^t \left(1 - \alpha_{\mathbf{c}}^{h,j} \lambda_{\min}(\Lambda_h^j)\right)^{J_j} \right) \\
&\stackrel{(vii)}{\leq} 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t \left(\prod_{j=i+1}^t (1 - 1/(2\kappa_j))^{J_j} + \prod_{j=i}^t (1 - 1/(2\kappa_j))^{J_j} \right),
\end{aligned}$$

(i) uses the definition of the matrix norm (i.e., $\|A\|_2 := \max_x \frac{\|Ax\|_2}{\|x\|_2} \implies \|Ax\|_2 \leq \|A\|_2 \|x\|_2$). (ii) uses [Lemma D.6](#) and sets $\lambda = 1$. (iii) and (v) come from the submultiplicativity of matrix norm. (iv) and (vi) use the fact that $\|A\|_2 \leq \lambda_{\max}(A)$, and (iv) also uses the triangular inequality. (vii) uses the fact that we set $\alpha_{\mathbf{c}}^{h,j} = 1/(2\lambda_{\max}(\Lambda_h^j))$ and denotes that $\kappa_j = \max_{h \in [H]} \lambda_{\max}(\Lambda_h^j)/\lambda_{\min}(\Lambda_h^j)$.

Using [Lemma D.16](#), we can set $J_j \geq 2\kappa_j \log(1/\sigma)$ where $\sigma = 1/(4H(|\mathcal{D}^t| + 1)\sqrt{d_c})$ and further get that

$$\begin{aligned}
\|\mu_h^{t,J_t}\|_2 &\leq 2H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=1}^t (\sigma^{t-i} + \sigma^{t-i+1}) \\
&\leq 4H \sqrt{d_c |\mathcal{D}^t|} \sum_{i=0}^{\infty} \sigma^i \\
&= 4H \sqrt{d_c |\mathcal{D}^t|} \left(\frac{1}{1-\sigma} \right) \\
&= \frac{16}{3} H \sqrt{d_c |\mathcal{D}^t|}.
\end{aligned}$$

Next, we continue to bound Term (II). Since $\xi_h^{t,J_t} \sim \mathcal{N}(0, \Sigma_h^{t,J_t})$, using [Lemma D.11](#), we have that

$$\Pr\left(\left\|\xi_h^{t,J_t}\right\|_2 \leq \sqrt{\frac{1}{\delta} \text{Tr}(\Sigma_h^{t,J_t})}\right) \geq 1 - \delta.$$

Recall from [Lemma D.3](#) that

$$\Sigma_h^{t,J_t} = \sum_{i=1}^t \frac{1}{\zeta} A_t^{J_t} \dots A_{i+1}^{J_{i+1}} \left(I - A_i^{2J_i}\right) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t}.$$

Therefore, we can use [Lemma D.13](#) and derive that

$$\begin{aligned}
\text{Tr}(\Sigma_h^{t,J_t}) &= \sum_{i=1}^t \frac{1}{\zeta} \text{Tr}\left(A_t^{J_t} \dots A_{i+1}^{J_{i+1}} \left(I - A_i^{2J_i}\right) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t}\right) \\
&\leq \sum_{i=1}^t \frac{1}{\zeta} \text{Tr}(A_t^{J_t}) \dots \text{Tr}(A_{i+1}^{J_{i+1}}) \text{Tr}(I - A_i^{2J_i}) \times \\
&\quad \text{Tr}((\Lambda_h^i)^{-1}) \text{Tr}((I + A_i)^{-1}) \text{Tr}(A_{i+1}^{J_{i+1}}) \dots \text{Tr}(A_t^{J_t}).
\end{aligned}$$

To bound each term, we first have,

$$\begin{aligned}
\text{Tr}(A_i^{J_i}) &\leq \text{Tr}\left((1 - \alpha_{\mathbf{c}}^i \lambda_{\min}(\Lambda_h^i))^{J_i} I\right) \\
&\leq d_c (1 - \alpha_{\mathbf{c}}^i \lambda_{\min}(\Lambda_h^i))^{J_i}
\end{aligned}$$

$$\leq d_c \sigma \leq 1,$$

where the first inequality follows from the fact that $A_i^{J_i} \prec (1 - \alpha_c^t \lambda_{\min}(\Lambda_h^t))^{J_i} I$. Similarly, since we set $0 < \alpha_c^{h,j} < 1/(2 \lambda_{\max}(\Lambda_j))$, we have $A_i^{J_i} \succ \frac{1}{2^{J_i}} I$ and therefore,

$$\text{Tr}(I - A_i^{2^{J_i}}) \leq \left(1 - \frac{1}{2^{2^{J_i}}}\right) d_c < d_c.$$

Similarly, since we set $0 < \alpha_c^{h,j} < 1/(2 \lambda_{\max}(\Lambda_j))$ and thus $I + A_i \succ \frac{3}{2} I$, we have that

$$\text{Tr}((I + A_i)^{-1}) \leq \frac{2}{3} d_c.$$

Additionally, since all eigenvalues of Λ_h^i are greater than or equal to 1,

$$\text{Tr}((\Lambda_h^i)^{-1}) \leq d_c \cdot 1 = d_c.$$

Finally, we have that

$$\text{Tr}(\Sigma_h^{t,J_t}) \leq \sum_{i=1}^t \frac{1}{\zeta} \cdot \frac{2}{3} \cdot d_c^3 = \frac{2 d_c^3}{3 \zeta} t.$$

Therefore, using [Lemma D.11](#), we have that

$$\Pr\left(\left\|\xi_h^{t,J_t}\right\|_2 \leq \sqrt{\frac{1}{\delta} \cdot \frac{2 d_c^3}{3 \zeta} T}\right) \geq \Pr\left(\left\|\xi_h^{t,J_t}\right\|_2 \leq \sqrt{\frac{1}{\delta} \text{Tr}(\Sigma_h^{t,J_t})}\right) \geq 1 - \delta.$$

Putting everything together, with probability at least $1 - \delta$, we can obtain that

$$\left\|w_h^{t,m,J_t}\right\|_2 \leq \overline{W}_\delta := \frac{16}{3} H \sqrt{d_c |\mathcal{D}^t|} + \sqrt{\frac{2 d_c^3 t}{3 \zeta \delta}}.$$

This concludes the proof. $\square$

D.2.6 Proof of [Lemma D.8](#)

Proof. To start, we decompose the LHS using the triangle inequality,

$$\left|\left\langle \phi(s, a), w_h^{t,m,J_t} - \widehat{w}_h^t \right\rangle\right| \leq \underbrace{\left|\left\langle \phi(s, a), w_h^{t,m,J_t} - \mu_h^{t,J_t} \right\rangle\right|}_{\text{(I)}} + \underbrace{\left|\left\langle \phi(s, a), \mu_h^{t,J_t} - \widehat{w}_h^t \right\rangle\right|}_{\text{(II)}},$$

where μ_h^{t,J_t} is defined in [Eq. \(13\)](#). To bound Term (I), we first apply Hölder's inequality and obtain that

$$\left|\left\langle \phi(s, a), w_h^{t,m,J_t} - \mu_h^{t,J_t} \right\rangle\right| \leq \left\|\phi(s, a)^\top (\Sigma_h^{t,J_t})^{1/2}\right\|_2 \left\|(\Sigma_h^{t,J_t})^{-1/2} (w_h^{t,m,J_t} - \mu_h^{t,J_t})\right\|_2.$$

Since $w_h^{t,m,J_t} \sim \mathbf{N}(\mu_h^{t,J_t}, \Sigma_h^{t,J_t})$, we know that $(\Sigma_h^{t,J_t})^{-1/2} (w_h^{t,m,J_t} - \mu_h^{t,J_t}) \sim \mathbf{N}(0, I_{d_c \times d_c})$. Therefore,

$$\Pr\left(\left\|(\Sigma_h^{t,J_t})^{-1/2} (w_h^{t,m,J_t} - \mu_h^{t,J_t})\right\|_2 \geq 2 \sqrt{d_c \log(1/\delta)}\right) \leq \delta^2.$$

Then, we continue to bound $\left\|\phi(s, a)^\top (\Sigma_h^{t,J_t})^{1/2}\right\|_2$.

$$\begin{aligned} & \phi(s, a)^\top \Sigma_h^{t,J_t} \phi(s, a) \\ &= \frac{1}{\zeta} \sum_{i=1}^t \phi(s, a)^\top A_t^{J_t} \dots A_{i+1}^{J_{i+1}} (I - A^{2^{J_i}}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} A_{i+1}^{J_{i+1}} \dots A_t^{J_t} \phi(s, a) \end{aligned}$$

$$\begin{aligned}
&\stackrel{(i)}{=} \frac{1}{\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (I - A^{2J_i}) (\Lambda_h^i)^{-1} (I + A_i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \\
&\stackrel{(ii)}{\leq} \frac{2}{3\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} \left((\Lambda_h^i)^{-1} - A_i^{J_i} (\Lambda_h^i)^{-1} A_i^{J_i} \right) \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{2}{3\zeta} \left(\sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) - \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_i (\Lambda_h^i)^{-1} \mathcal{A}_i^\top \phi(s, a) \right) \\
&\stackrel{(iii)}{\leq} \frac{2}{3\zeta} \sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_{i+1} (\Lambda_h^i)^{-1} \mathcal{A}_{i+1}^\top \phi(s, a) \\
&= \frac{2}{3\zeta} \left(\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 + \sum_{i=1}^{t-1} \|\mathcal{A}_{i+1}^\top \phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2 \right) \\
&\leq \frac{2}{3\zeta} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}^2 + \frac{2}{3\zeta} \sum_{i=1}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_c \lambda_{\min}(\Lambda_h^j) \right)^{2J_j} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}}^2.
\end{aligned}$$

For (i), we use the denotation that $\mathcal{A}_{i+1} = A_t^{J_t} \dots A_{i+1}^{J_{i+1}}$. (ii) follows from $I + A_i \succ \frac{3}{2}I$ since we set $\alpha_c^{h,j} = 1/(2\lambda_{\max}(\Lambda_h^j))$. (iii) follows from the fact that $\sum_{i=1}^t \phi(s, a)^\top \mathcal{A}_i (\Lambda_h^i)^{-1} \mathcal{A}_i^\top \phi(s, a) > 0$. Therefore,

$$\begin{aligned}
&\left\| \phi(s, a)^\top \left(\Sigma_h^{t, J_t} \right)^{1/2} \right\|_2 = \sqrt{\phi(s, a)^\top \Sigma_h^{t, J_t} \phi(s, a)} \\
&\stackrel{(iv)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2}{3\zeta}} \sum_{i=1}^{t-1} \prod_{j=i+1}^t \left(1 - \alpha_c \lambda_{\min}(\Lambda_h^j) \right)^{J_j} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
&\stackrel{(v)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2}{3\zeta}} \sum_{i=1}^{t-1} \sigma^{t-i} \|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \\
&\stackrel{(vi)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2(|\mathcal{D}^t| + 1)}{3\zeta}} \left(\sum_{i=1}^{t-1} \sigma^{t-i} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{(vii)}{\leq} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \sqrt{\frac{2(|\mathcal{D}^t| + 1)}{3\zeta}} \left(\frac{\sigma}{1 - \sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} + \frac{1}{4} \sqrt{\frac{2}{3\zeta}} \left(\frac{1}{1 - \sigma} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \left(\sqrt{\frac{2}{3\zeta}} + \frac{1}{3} \sqrt{\frac{2}{3\zeta}} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&\leq \frac{4}{3} \sqrt{\frac{2}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

(iv) follows from the fact that $\sqrt{a+b} \leq a+b$ for all $a, b > 0$. (v) uses [Lemma D.16](#) by setting $J_j \geq \kappa_j \log(1/\sigma)$ where $\sigma = 1/(4\bar{H}(|\mathcal{D}^t| + 1)\sqrt{d_c})$. (vi) follows from $\|\phi(s, a)\|_{(\Lambda_h^i)^{-1}} \leq \|\phi(s, a)\|_2 \leq \sqrt{|\mathcal{D}^t| + 1} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}$. (vii) follows from $\sum_{i=1}^t \sigma^{t-i} \leq \sum_{i=1}^\infty \sigma^i \leq \sigma/(1 - \sigma)$. Therefore, we have

$$\begin{aligned}
&\Pr \left(\left| \left\langle \phi(s, a), w_h^{t, m, J_t} - \mu_h^{t, J_t} \right\rangle \right| \geq \frac{8}{3} \sqrt{\frac{2d_c \log(1/\delta)}{3\zeta}} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right) \\
&\leq \Pr \left(\left\| \phi(s, a)^\top \left(\Sigma_h^{t, J_t} \right)^{1/2} \right\|_2 \left\| \left(\Sigma_h^{t, J_t} \right)^{-1/2} \left(w_h^{t, m, J_t} - \mu_h^{t, J_t} \right) \right\|_2 \right) \\
&\geq 2\sqrt{d_c \log(1/\delta)} \left\| \phi(s, a)^\top \left(\Sigma_h^{t, J_t} \right)^{1/2} \right\|_2
\end{aligned}$$

$$= \Pr\left(\left\|\left(\Sigma_h^{t,J_t}\right)^{-1/2}\left(w_h^{t,m,J_t} - \mu_h^{t,J_t}\right)\right\|_2 \geq 2\sqrt{d_c \log(1/\delta)}\right) = \delta^2 \leq \delta.$$

This implies that

$$\Pr\left(\left|\left\langle\phi(s,a), w_h^{t,m,J_t} - \mu_h^{t,J_t}\right\rangle\right| \leq \frac{8}{3}\sqrt{\frac{2d_c \log(1/\delta)}{3\zeta}}\|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}\right) \geq 1 - \delta.$$

Putting everything together, with probability at least $1 - \delta$,

$$\begin{aligned} \left|\left\langle\phi(s,a), w_h^{t,m,J_t} - \hat{w}_h^t\right\rangle\right| &\leq \left|\left\langle\phi(s,a), w_h^{t,m,J_t} - \mu_h^{t,J_t}\right\rangle\right| + \left|\left\langle\phi(s,a), \mu_h^{t,J_t} - \hat{w}_h^t\right\rangle\right| \\ &\leq \left(\frac{8}{3}\sqrt{\frac{2d_c \log(1/\delta)}{3\zeta}} + \frac{4}{3}\right)\|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}. \end{aligned}$$

□

D.2.7 Proof of Lemma D.9

Proof. Recall that $\mathbb{P}_h V(s,a) := \mathbb{E}_{s' \sim \mathbb{P}_h(\cdot|s,a)} V(s')$ and $\mathbb{P}_h(\cdot | s,a) = \langle \phi(s,a), \psi_h(\cdot) \rangle$ due to the linear MDP assumption (Definition 2.1). We also denote that $\hat{\Psi}_h^t := \langle \psi_h, \hat{V}_{h+1}^t \rangle_s$ and thus $\mathbb{P}_h \hat{V}_{h+1}^t(s,a) = \langle \phi(s,a), \hat{\Psi}_h^t \rangle$. Then, we have that

$$\begin{aligned} \mathbb{P}_h \hat{V}_{h+1}^t(s,a) &= \langle \phi(s,a), \hat{\Psi}_h^t \rangle \\ &= \phi(s,a)^\top (\Lambda_h^t)^{-1} \Lambda_h^t \hat{\Psi}_h^t \\ &= \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \phi(s,a)^\top + \lambda I \right) \hat{\Psi}_h^t \\ &= \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) (\mathbb{P}_h \hat{V}_{h+1}^t)(s,a) + \lambda \hat{\Psi}_h^t \right). \end{aligned}$$

This further implies that

$$\begin{aligned} &\langle \phi(s,a), \hat{w}_h^t \rangle - r_h(s,a) - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \\ &= \phi(s,a)^\top (\Lambda_h^t)^{-1} \sum_{(s,a,s') \in \mathcal{D}_h^t} \left[r_h(s,a) + \hat{V}_{h+1}^t(s') \right] \cdot \phi(s,a) - r_h(s,a) \\ &\quad - \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) (\mathbb{P}_h \hat{V}_{h+1}^t)(s,a) + \lambda \hat{\Psi}_h^t \right) \\ &= \underbrace{\phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \left[\hat{V}_{h+1}^t(s') - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \right] \right)}_{\text{(I)}} \\ &\quad + \underbrace{\phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) \right) - r_h(s,a)}_{\text{(II)}} - \underbrace{\lambda \phi(s,a)^\top (\Lambda_h^t)^{-1} \hat{\Psi}_h^t}_{\text{(III)}}. \end{aligned}$$

We first start by bounding Term (I). With probability at least $1 - \delta$, it holds that

$$\phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \left[\hat{V}_{h+1}^t(s') - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \right] \right)$$

$$\begin{aligned}
&\stackrel{(i)}{\leq} \left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \left[\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a) \right] \right\|_{(\Lambda_h^t)^{-1}} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{(ii)}{\leq} C_\delta H \sqrt{d_c} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}},
\end{aligned}$$

where (i) follows from the Cauchy-Schwarz inequality, and (ii) follows from the good event defined in [Lemma D.1](#).

Next, we continue to bound Term (II). We observe that

$$\begin{aligned}
&\phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) \right) - r_h(s,a) \\
&\stackrel{(iii)}{=} \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) \right) - \phi(s,a)^\top \theta_h \\
&= \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) - \Lambda_h^t \theta_h \right) \\
&= \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) - \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \phi(s,a)^\top \theta_h - \lambda \theta_h \right) \\
&\stackrel{(iv)}{=} \phi(s,a)^\top (\Lambda_h^t)^{-1} \left(\sum_{(s,a,s') \in \mathcal{D}_h^t} r_h(s,a) \phi(s,a) - \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) r_h(s,a) - \lambda \theta_h \right) \\
&= -\lambda \phi(s,a)^\top (\Lambda_h^t)^{-1} \theta_h \\
&\stackrel{(v)}{\leq} \lambda \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \|\theta_h\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{(vi)}{\leq} \sqrt{\lambda d_c} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

(iii) and (iv) follow from the definition $r_h(s,a) = \langle \phi(s,a), \theta_h \rangle$. (v) applies the Cauchy-Schwarz inequality. (vi) follows from $\|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \leq \sqrt{1/\lambda} \|\phi(s,a)\|_2$ and $\|\theta_h\|_2 \leq \sqrt{d_c}$ ([Definition 2.1](#)).

Lastly, we derive the bound for Term (III).

$$\begin{aligned}
\lambda \phi(s,a)^\top (\Lambda_h^t)^{-1} \widehat{\Psi}_h^t &\stackrel{(vii)}{\leq} \lambda \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \left\| \widehat{\Psi}_h^t \right\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{(viii)}{\leq} \sqrt{\lambda} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \left\| \widehat{\Psi}_h^t \right\|_2 \\
&\leq \sqrt{\lambda} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \left\| \left\langle \psi_h, \widehat{V}_{h+1}^t \right\rangle_S \right\|_2 \\
&= H \sqrt{\lambda} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}} \left\| \int_{s \in \mathcal{S}} \psi_h(s) \left(\widehat{V}_{h+1}^t(s)/H \right) ds \right\|_2 \\
&\stackrel{(xiv)}{\leq} H \sqrt{\lambda d_c} \|\phi(s,a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

(vii) applies the Cauchy-Schwarz inequality. (viii) follows from $\left\| \widehat{\Psi}_h^t \right\|_{(\Lambda_h^t)^{-1}} \leq \sqrt{\lambda} \left\| \widehat{\Psi}_h^t \right\|_2$. (xiv) comes from the assumption that $\left\| \int_{s \in \mathcal{S}} \psi_h(s) \left(\widehat{V}_{h+1}^t(s)/H \right) ds \right\|_2 \leq \sqrt{d_c}$ ([Definition 2.1](#)).

Putting everything together and setting $\lambda = 1$, we have with probability at least $1 - \delta$,

$$\left| \langle \phi(s,a), \widehat{w}_h^t \rangle - r_h(s,a) - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a) \right|$$

$$\begin{aligned}
&\leq \left(C_\delta H \sqrt{d_c} + \sqrt{\lambda d_c} + H \sqrt{\lambda d_c} \right) \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \\
&= 3 C_\delta H \sqrt{d_c} \|\phi(s, a)\|_{(\Lambda_h^t)^{-1}}.
\end{aligned}$$

This concludes the proof. $\square$

D.3 Technical Tools

Lemma D.10 (Jin et al. 2020, Lemma D.1). *Let $\Lambda = \lambda I + \sum_{i=1}^t \phi_i \phi_i^\top$, where $\phi_i \in \mathbb{R}^d$ and $\lambda > 0$. Then,*

$$\sum_{i=1}^t \phi_i^\top (\Lambda)^{-1} \phi_i \leq d.$$

Lemma D.11 (Ishfaq et al. 2024a, Lemma E.1). *Given a multivariate normal distribution $X \sim \mathcal{N}(0, \Sigma_{d \times d})$, for any $\delta \in (0, 1]$, it holds that*

$$\Pr\left(\|X\|_2 \leq \sqrt{\frac{1}{\delta} \text{Tr}(\Sigma)}\right) \geq 1 - \delta.$$

Lemma D.12 (Abramowitz and Stegun 1948). *Suppose X is a Gaussian random variable $X \sim \mathcal{N}(\mu, \sigma^2)$, where $\sigma > 0$. For $z \in [0, 1]$, it holds that*

$$\Pr(X > \mu + z\sigma) \geq \frac{e^{-z^2/2}}{\sqrt{8\pi}} \quad \text{and} \quad \Pr(X < \mu - z\sigma) \geq \frac{e^{-z^2/2}}{\sqrt{8\pi}}.$$

Additionally, for any $z \geq 1$,

$$\frac{e^{-z^2/2}}{2z\sqrt{\pi}} \leq \Pr(|X - \mu| > z\sigma) \leq \frac{e^{-z^2/2}}{z\sqrt{\pi}}.$$

Lemma D.13. *If A and B are positive semi-definite square matrices of the same size, then*

$$[\text{Tr}(AB)]^2 \leq \text{Tr}(A^2) \text{Tr}(B^2) \leq [\text{Tr}(A)]^2 [\text{Tr}(B)]^2.$$

Lemma D.14. *Given two symmetric positive semi-definite square matrices A and B such that $A \succeq B$, it holds that $\|A\|_2 \geq \|B\|_2$.*

Proof. Note that $A - B$ is also positive semi-definite. Then, we have that

$$\|B\|_2 = \sup_{\|x\|=1} x^\top Bx \leq \sup_{\|x\|=1} (x^\top Bx + x^\top (A - B)x) = \sup_{\|x\|=1} x^\top Ax = \|A\|_2.$$

$\square$

Lemma D.15. *Let $A \in \mathbb{R}^{d \times d}$ be a positive definite matrix where its largest eigenvalue $\lambda_{\max}(A) \leq \lambda$. Given that $v_1, \dots, v_n$ are n vectors in $\mathbb{R}^d$, it holds that*

$$\left\| A \sum_{i=1}^n v_i \right\| \leq \sqrt{\lambda n \sum_{i=1}^n \|v_i\|_A^2}.$$

Lemma D.16. *Let Λ be a positive definite matrix and $\kappa = \frac{\lambda_{\max}(\Lambda)}{\lambda_{\min}(\Lambda)}$ be the condition number of Λ . If $\Lambda \succ I$ and $J \geq 2\kappa \log(1/\sigma)$, then, for any $\sigma > 0$,*

$$(1 - 1/(2\kappa))^J < \sigma.$$

Proof. The statement is equivalent to proving that

$$J \geq \frac{\log(1/\sigma)}{\log\left(\frac{1}{1-1/(2\kappa)}\right)}.$$

Since $\kappa \geq 1$ and for any $x \in (0, 1)$, $e^{-x} > 1 - x$, we have that

$$e^{-1/(2\kappa)} > 1 - 1/(2\kappa) \implies \log\left(\frac{1}{1 - 1/(2\kappa)}\right) \geq \frac{1}{2\kappa}.$$

Therefore, we have that

$$J \geq 2\kappa \log(1/\sigma) \geq \frac{\log(1/\sigma)}{\log\left(\frac{1}{1 - 1/(2\kappa)}\right)},$$

which concludes the proof. $\square$

E Sample Complexity in the On-Policy Setting

E.1 Proof of Good Event

Lemma E.1. *Consider [Algorithm 1](#) in the on-policy setting with $\lambda = 1$. Then, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$, it holds that*

$$\left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) [\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a)] \right\|_{(\Lambda_h^t)^{-1}} \leq C_\delta^{\text{on}} H \sqrt{d_c},$$

where $C_\delta^{\text{on}} = \log(N/\delta)$.

Proof of Lemma E.1. Recall that $\mathbb{P}_h \widehat{V}_{h+1}^t(s,a) = \mathbb{E}_{s' \sim \mathbb{P}_h} [\widehat{V}_{h+1}^t(s')]$. Thus, $\mathbb{E}[\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a)] = 0$. Also, $|\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a)| \leq H$. Therefore, $\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a)$ is zero-mean and H -sub Gaussian. Given that, we can invoke [Lemma E.3](#).

$$\begin{aligned} & \left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) [\widehat{V}_{h+1}^t(s') - \mathbb{P}_h \widehat{V}_{h+1}^t(s,a)] \right\|_{(\Lambda_h^t)^{-1}} \\ & \leq \sqrt{2} H \sqrt{\log \left[\frac{\det(\Lambda_h^t)^{1/2} \det(\Lambda_h^0)^{-1/2}}{\delta} \right]} \\ & = \sqrt{2} H \sqrt{\log \left[\left(\frac{N+\lambda}{\lambda} \right)^{d/2} \right] - \log(\delta)} \\ & = \sqrt{2} H \sqrt{\frac{d_c}{2} \log(N/\delta)} \\ & = H \sqrt{d_c \log(N/\delta)}, \end{aligned}$$

where the first equality follows from [Lemma E.4](#), and the second equality holds by setting $\lambda = 1$.

This concludes the proof. $\square$

E.2 Proof of Theorem 6.1

Using [Lemma E.1](#), we can instantiate [Lemma D.2](#) in the on-policy setting with

$$\begin{aligned} \Gamma_{\text{LMC}}^{\text{on}} &= C_\delta^{\text{on}} H \sqrt{d_c} + \frac{4}{3} \sqrt{\frac{2 d_c \log(1/\delta)}{3 \zeta}} + \frac{4}{3} \\ &= H \sqrt{d_c \log(N/\delta)} + \frac{4}{3} \sqrt{\frac{2 d_c \log(1/\delta)}{3 \zeta}} + \frac{4}{3}. \end{aligned}$$

Proof of Theorem 6.1. The optimal gap for the mixture policy can be written as

$$\mathbb{E}\left[V_1^\star(s_1) - V_1^{\bar{\pi}^T}(s_1)\right] = \frac{1}{T} \sum_{t=1}^T \left(V_1^\star(s_1) - V_1^{\pi^t}(s_1)\right).$$

Then, to decompose the above summation, we have that

$$\sum_{t=1}^T \left(V_1^\star(s_1) - V_1^{\pi^t}(s_1)\right) = \sum_{t=1}^T \left(V_1^\star(s_1) - \widehat{V}_1^t(s_1)\right) + \sum_{t=1}^T \left(\widehat{V}_1^t(s_1) - V_1^{\pi^t}(s_1)\right).$$

We can further decompose the first term by invoking Lemma E.2 with $\pi = \pi^\star$ and obtain that

$$\begin{aligned} V_1^\star(s_1) - \widehat{V}_1^t(s_1) &= \sum_{h=1}^H \mathbb{E}_{\pi^\star} \left[\left\langle \pi_h^\star(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right] + \sum_{h=1}^H \mathbb{E}_{\pi^\star} \left[r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) - \widehat{Q}_h^t(s, a) \right]. \end{aligned}$$

Similarly, we can decompose the second term by invoking Lemma E.2 with $\pi = \pi^t$ and get that

$$\begin{aligned} \widehat{V}_1^t(s_1) - V_1^{\pi^t}(s_1) &= \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\left\langle \pi_h^t(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right] - \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) - \widehat{Q}_h^t(s, a) \right] \\ &= - \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[r_h(s, a) + \mathbb{P}_h \widehat{V}_{h+1}^t(s, a) - \widehat{Q}_h^t(s, a) \right]. \end{aligned}$$

Therefore, using the definition of the model prediction error ι in Definition 5.1, we have that

$$\begin{aligned} \sum_{t=1}^T \left(V_1^\star(s_1) - V_1^{\pi^t}(s_1)\right) &= \underbrace{\sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^\star} \left[\left\langle \pi_h^\star(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right]}_{\text{(I) policy optimization (actor) error}} + \underbrace{\sum_{t=1}^T \sum_{h=1}^H \left(\mathbb{E}_{\pi^\star} [l_h^t(s, a)] - \mathbb{E}_{\pi^t} [l_h^t(s, a)] \right)}_{\text{(II) policy evaluation (critic) error}}. \end{aligned}$$

Policy optimization error. We first start by bounding Term (I), the policy optimization error.

$$\begin{aligned} \text{Term (I)} &= \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{s \sim \pi^\star} \left[\left\langle \pi_h^\star(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right] \\ &= \sum_{h=1}^H \mathbb{E}_{s \sim \pi^\star} \left(\sum_{t=1}^T \left\langle \pi_h^\star(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right) \\ &\leq H \max_{(h,s) \in [H] \times \mathcal{S}} \left(\sum_{t=1}^T \left\langle \pi_h^\star(\cdot | s) - \pi_h^t(\cdot | s), \widehat{Q}_h^t(s, \cdot) \right\rangle \right) \\ &\stackrel{(i)}{\leq} H \left(\frac{\log |\mathcal{A}| + \sum_{t=1}^T \|\epsilon_h^t(\cdot)\|_\infty}{\eta} + \frac{\eta H^2 T}{2} \right) \\ &\stackrel{(ii)}{\leq} H^2 \sqrt{(\log |\mathcal{A}| + \bar{\epsilon} T)/2} \sqrt{T} \\ &\stackrel{(iii)}{\leq} \mathcal{O} \left(H^2 \sqrt{\log |\mathcal{A}|} \sqrt{T} + H^2 \sqrt{\bar{\epsilon} T} \right). \end{aligned}$$

(i) follows from Lemma 4.1 with $u = \pi_h^\star(\cdot | s)$. (ii) is obtained by setting $\eta = \frac{\sqrt{2(\log |\mathcal{A}| + \bar{\epsilon} T)}}{H \sqrt{T}}$. (iii) is based on that for all $a, b \geq 0$, $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$.

Policy evaluation error. Then, we continue to bound Term (II), the policy evaluation error.

$$\begin{aligned}
\text{Term (II)} &= \sum_{t=1}^T \sum_{h=1}^H (\mathbb{E}_{\pi^*}[\iota_h^t(s, a)] - \mathbb{E}_{\pi^t}[\iota_h^t(s, a)]) \\
&\stackrel{\text{(iv)}}{\leq} - \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t}[\iota_h^t(s, a)] \\
&\stackrel{\text{(v)}}{\leq} \Gamma_{\text{LMC}}^{\text{on}} \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] \\
&\leq \Gamma_{\text{LMC}}^{\text{on}} T \max_{t \in [T]} \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right].
\end{aligned}$$

(iv) and (v) both follow from [Lemma D.2](#), where (iv) is based on the optimism guarantee (RHS of [Eq. \(10\)](#)), while (v) is based on the error bound (LHS of [Eq. \(10\)](#)).

Bounding the sum of bonuses. Since $\Gamma_{\text{LMC}}^{\text{on}}$ is bounded, it suffices to bound $\mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right]$. Note that $\Lambda_h^t = \sum_{(s, a, s') \in \mathcal{D}_h^t} \phi(s, a) \phi(s, a)^\top + \lambda I$, and $\mathcal{D}_h^t$ only depends on π^t in the on-policy setting. (This is not true for the off-policy setting since Λ_h^t would depend on $\{\pi^1, \dots, \pi^t\}$.) We then index each data point in $\mathcal{D}_h^t$ as $\{(s_h^i, a_h^i, s_{h+1}^i)\}_{i \in [N]}$. Let $\Lambda_h^{t,i} = \left(\sum_{j=1}^i \phi(s_h^j, a_h^j) \phi(s_h^j, a_h^j)^\top + \lambda I \right)$. Then, we have that

$$\begin{aligned}
\sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] &\stackrel{\text{(vi)}}{\leq} \frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^{t,i})^{-1}} \right] \\
&= \frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}} \\
&\quad + \underbrace{\frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \mathbb{E}_{\substack{s_h \sim \mathbb{P}(\cdot | s_{h-1}^i, a_{h-1}^i) \\ a_h \sim \pi_h^t(\cdot | s_h)}} \left[\|\phi(s, a)\|_{(\Lambda_h^{t,i})^{-1}} \right] - \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}}}_{:= \mathcal{M}_{i,h}^{\text{on}}}, \quad (15)
\end{aligned}$$

where (vi) follows from the fact that $\Lambda_h^{t,i} \preceq \Lambda_h^t$.

Applying the elliptical potential lemma. For the first term of [Eq. \(15\)](#), we have that

$$\begin{aligned}
&\frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}} \\
&= \frac{1}{N} \sum_{h=1}^H \sum_{i=1}^N \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}} \\
&\stackrel{\text{(vii)}}{\leq} \frac{1}{N} \sum_{h=1}^H \sqrt{N} \left(\sum_{i=1}^N \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}}^2 \right)^{1/2} \\
&\stackrel{\text{(viii)}}{\leq} \mathcal{O} \left(\sqrt{\frac{d_c H^2 \log(N/\delta)}{N}} \right).
\end{aligned}$$

(vii) applies the Cauchy-Schwarz inequality, and (viii) follows the elliptical argument from [Lemma E.5](#).

A martingale difference sequence. For the second term of [Eq. \(15\)](#), since for a fixed $i \in [N]$, $\{\mathcal{M}_{i,h}^{\text{on}}\}_{h \in [H]}$ forms a martingale sequence adapted to the filtration,

$$\mathcal{F}_{i,h}^{\text{on}} = \{(s_\tau^i, a_\tau^i)\}_{\tau \in [h-1]},$$

such that $\mathbb{E}[\mathcal{M}_{i,h}^{\text{on}} \mid \mathcal{F}_{i,h}^{\text{on}}] = 0$, where the expectation is with respect to the randomness in the policy and the environment at step h . Since $|\mathcal{M}_{i,h}^{\text{on}}| \leq 1$, we can apply the Azuma–Hoeffding inequality and obtain that

$$\Pr\left(\sum_{i=1}^N \sum_{h=1}^H \mathcal{M}_{i,h}^{\text{on}} \geq m\right) \geq \exp\left(\frac{-m^2}{2HN}\right).$$

Setting $m = \sqrt{2HN \log(1/\delta)}$ and using a union bound over $i \in [N]$, with probability at least $1 - \delta$, it holds that

$$\frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \mathcal{M}_{i,h}^{\text{on}} \leq \sqrt{\frac{2H \log(1/\delta)}{N}} \leq \mathcal{O}\left(\sqrt{\frac{H \log(1/\delta)}{N}}\right).$$

Putting everything together. Therefore, we have that

$$\begin{aligned} & \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] \\ &= \frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \|\phi(s_h^i, a_h^i)\|_{(\Lambda_h^{t,i})^{-1}} + \frac{1}{N} \sum_{i=1}^N \sum_{h=1}^H \mathcal{M}_{i,h}^{\text{on}} \leq \mathcal{O}\left(\sqrt{\frac{d_c H^2 \log(N/\delta)}{N}}\right). \end{aligned}$$

It further implies that, with probability at least $1 - \delta$,

$$\begin{aligned} \text{Term (II)} &\leq \Gamma_{\text{LMC}}^{\text{on}} T \max_{t \in [T]} \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] \\ &\stackrel{\text{(ix)}}{\leq} \mathcal{O}\left(\sqrt{\frac{d_c^3 H^4 \log^2(N/\delta)}{N}} T\right) \\ &\leq \tilde{\mathcal{O}}\left(H^2 \sqrt{\log|\mathcal{A}|} \sqrt{T}\right), \end{aligned}$$

where (ix) comes from setting $N = d_c^3 T / \log|\mathcal{A}|$.

Finally, putting everything together, with probability at least $1 - \delta$,

$$\mathbb{E}\left[V_1^*(s_1) - V_1^{\pi^T}(s_1)\right] = \frac{1}{T}(\text{Term (I)} + \text{Term (II)}) = \tilde{\mathcal{O}}\left(\frac{H^2 \sqrt{\log|\mathcal{A}|}}{\sqrt{T}} + H^2 \sqrt{\epsilon}\right).$$

This concludes the proof. $\square$

E.3 Technical Tools

Lemma E.2 (Extended Value Difference). *Given any $\pi, \pi' \in \Delta(\mathcal{A} \mid \mathcal{S}, H)$ and any Q -function $\hat{Q} \in \mathbb{R}^{H \times |\mathcal{S}| \times |\mathcal{A}|}$, we define $\hat{V}_h(\cdot) = \mathbb{E}_{a \sim \pi'_h(\cdot, \cdot)} \hat{Q}_h(\cdot, a)$ for any $h \in [H]$. Then,*

$$\begin{aligned} & \hat{V}_1(s_1) - V_1^\pi(s_1) \\ &= \sum_{h=1}^H \mathbb{E}_{s \sim \pi} \left[\left\langle \pi'_h(s, \cdot) - \pi_h(s, \cdot), \hat{Q}_h(s, \cdot) \right\rangle \right] \\ &\quad + \sum_{h=1}^H \mathbb{E}_{(s,a) \sim \pi} \left[\hat{Q}_h(s, a) - r_h(s, a) - \sum_{s' \in \mathcal{S}} \mathbb{P}_h(s' \mid s, a) \hat{V}_{h+1}(s') \right]. \end{aligned}$$

Lemma E.3 (Concentration of Self-Normalized Processes ([Abbasi-Yadkori et al., 2011](#), Theorem 1)). *Let $\{x_t\}_{t=1}^\infty$ be a real-valued stochastic process with the correspond filtration $\{\mathcal{F}_t\}_{t=0}^\infty$ such that x_t is $\mathcal{F}_{t-1}$ -measurable, and x_t is conditionally σ -sub-Gaussian for some $\sigma > 0$, i.e.,*

$$\forall \lambda \in \mathbb{R}, \quad \mathbb{E}[\exp(\lambda x_t) \mid \mathcal{F}_{t-1}] = \exp(\lambda^2 \sigma^2 / 2).$$

Let $\{\phi_t\}_{t=1}^\infty$ be an $\mathbb{R}^d$ -valued stochastic process such that ϕ_t is $\mathcal{F}_{t-1}$ -measurable. Assume Λ_0 is a $d \times d$ positive definite matrix, and let $\Lambda_t = \Lambda_0 + \sum_{i=1}^t \phi_i \phi_i^\top$. Then, for any $\delta > 0$, with probability at least $1 - \delta$, for all $t \geq 0$, it holds that

$$\left\| \sum_{i=1}^t \phi_i x_i \right\|_{\Lambda_t^{-1}}^2 \leq 2 \sigma^2 \log \left[\frac{\det(\Lambda_t)^{1/2} \det(\Lambda_0)^{-1/2}}{\delta} \right].$$

Lemma E.4 (Determinant-Trace Inequality (Abbasi-Yadkori et al., 2011, Lemma 10)). Suppose $X_1, X_2, \dots, X_t \in \mathbb{R}^d$ and for any $s \in [t]$, $\|X_s\|_2 \leq L$. Let $\Lambda_t = \lambda I + \sum_{s=1}^t X_s X_s^\top$ for some $\lambda > 0$. Then, for all t , it holds that

$$\det(\Lambda_t) \leq (\lambda + t L^2/d)^d.$$

Lemma E.5 (Abbasi-Yadkori et al. 2011, Lemma 11). Suppose $X_1, X_2, \dots, X_t \in \mathbb{R}^d$ and for any $s \in [t]$, $\|X_s\|_2 \leq L$. Let $\Lambda_t = \Lambda_0 + \sum_{s=1}^t X_s X_s^\top$ and $\lambda_{\min}(\Lambda_0) \geq \max\{1, L^2\}$. Then, for all t , it holds that

$$\log \left(\frac{\det(\Lambda_t)}{\det(\Lambda_0)} \right) \leq \sum_{s=1}^t \|X_s\|_{(\Lambda_t)^{-1}}^2 \leq 2 \log \left(\frac{\det(\Lambda_t)}{\det(\Lambda_0)} \right).$$

F Sample Complexity in the Off-Policy Setting

F.1 Covering Number (Proof of Lemma 6.1)

We first present a bound for the norm of the logit.

Lemma F.1. Consider Algorithm 1 with the NPG actor in Algorithm 2. Then, under Assumptions 4.1 and 4.2, for all $(t, h, s, a) \in [T] \times [H] \times \mathcal{S} \times \mathcal{A}$, it holds that

$$\left| \left\langle \varphi(s, a), \theta_h^{t, K_t}(s, a) \right\rangle \right| \leq (\bar{\epsilon} + \eta H) t,$$

where $\bar{\epsilon}$ is defined in Lemma 4.2.

Proof of Lemma F.1. We will prove this by induction. When $t = 0$, since we set $\theta_h^0 = \mathbf{0}$, the statement is trivially true. For $t \geq 1$, assume that the statement stands true for $t-1$. Since Algorithm 1 optimizes the actor loss up to some errors that are assumed to be bounded, using the triangular inequality, we have that

$$\begin{aligned} & \left| \left\langle \varphi(s, a), \theta_h^{t, K_t}(s, a) \right\rangle \right| \\ &= \left| \left\langle \varphi(s, a), \theta_h^{t, K_t} - \hat{\theta}_h^{t, \star} \right\rangle \right| + \left| \left\langle \varphi(s, a), \hat{\theta}_h^{t, \star}(s, a) \right\rangle \right| \\ &\leq \epsilon_{\text{opt}} + \left| \left\langle \varphi(s, a), \hat{\theta}_h^{t, \star}(s, a) \right\rangle \right| \\ &\leq \epsilon_{\text{opt}} + \left| \left\langle \varphi(s, a), \hat{\theta}_h^{t, \star}(s, a) - \theta_h^{t-1, K_{t-1}}(s, a) \right\rangle - \eta \hat{Q}_h^t(s, a) \right| \\ &\quad + \left| \left\langle \varphi(s, a), \theta_h^{t-1, K_{t-1}}(s, a) \right\rangle + \eta \hat{Q}_h^t(s, a) \right| \\ &\leq \epsilon_{\text{opt}} + \epsilon_{\text{bias}} + \left| \left\langle \varphi(s, a), \theta_h^{t-1, K_{t-1}}(s, a) \right\rangle + \eta \hat{Q}_h^t(s, a) \right| \\ &\leq \bar{\epsilon} + \left| \left\langle \varphi(s, a), \theta_h^{t-1, K_{t-1}}(s, a) \right\rangle \right| + \left| \eta \hat{Q}_h^t(s, a) \right| \\ &\leq (\bar{\epsilon} + \eta H) t, \end{aligned}$$

where $\hat{\theta}_h^{t, \star}$ denotes the optimal actor parameters when optimizing over $\mathcal{D}_{\text{exp}}$ and ρ_{exp} , $\left\langle \varphi(s, a), \theta_h^{t-1, K_{t-1}}(s, a) \right\rangle + \eta \hat{Q}_h^t(s, a)$ is the optimization target in the actor loss of the projected NPG, and the last inequality uses the inductive hypothesis.

This concludes the proof. $\square$

Proof of Lemma 6.1. Consider any $Q, Q' \in \mathcal{Q}$ such that $Q(\cdot, \cdot) = \min\{\langle \phi(\cdot, \cdot), w \rangle, H\}^+$ and $Q'(\cdot, \cdot) = \min\{\langle \phi(\cdot, \cdot), w' \rangle, H\}^+$. Therefore, we have that

$$\begin{aligned} \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} |Q(s,a) - Q'(s,a)| &\leq \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} |\langle \phi(s,a), w - w' \rangle| \\ &\leq \sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} \|\phi(s,a)\| \|w - w'\| \\ &\leq 2\bar{W}, \end{aligned}$$

where the first inequality uses the Cauchy-Schwarz inequality, and the second inequality uses Definition 2.1, the triangular inequality, and the definition of $\bar{W}$.

Consider any $\pi, \pi' \in \Pi_{\text{lin}}$ such that $\pi(\cdot | s) \propto \exp(\langle \phi(s, \cdot), \theta \rangle)$ and $\pi'(\cdot | s) \propto \exp(\langle \phi(s, \cdot), \theta' \rangle)$. By invoking Lemma F.6 and using Lemma F.1, we can observe that for a fixed $s \in \mathcal{S}$,

$$\sup_{a \in \mathcal{A}} |\pi(s,a) - \pi'(s,a)| \leq \|\pi(s, \cdot) - \pi'(s, \cdot)\|_1 \leq 2 \sqrt{\sup_a |\langle \phi(s,a), \theta - \theta' \rangle|} \leq 2\sqrt{2\bar{Z}}.$$

Taking the sup over $\mathcal{S}$, we get that

$$\sup_{(s,a) \in \mathcal{S} \times \mathcal{A}} |\pi(s,a) - \pi'(s,a)| \leq 2\sqrt{2\bar{Z}}.$$

Therefore, we can bound the log covering number of the value function class as follows.

$$\begin{aligned} \log \mathcal{N}_\Delta(\mathcal{V}) &\leq \log \mathcal{N}_{\Delta/2}(\mathcal{Q}) + \log \mathcal{N}_{\Delta/(2H)}(\Pi_{\text{lin}}) \\ &\leq d_c \log\left(1 + \frac{4\bar{W}}{\Delta}\right) + d_a \log\left(1 + \frac{8H\sqrt{2\bar{Z}}}{\Delta}\right), \end{aligned}$$

where the first inequality follows from Lemma F.3, and the second inequality uses Lemma F.5.

This concludes the proof. $\square$

F.2 Proof of Good Event

Lemma F.2. Consider Algorithm 1 in the off-policy setting with $\lambda = 1$. Then, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$, it holds that

$$\left\| \sum_{(s,a) \in \mathcal{D}_h^t} \phi(s,a) \left[\hat{V}_{h+1}^t(s) - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \right] \right\|_{(\Lambda_h^t)^{-1}} \leq C_\delta^{\text{off}} H \sqrt{d_c},$$

where

$$\begin{aligned} C_\delta^{\text{off}} &= 3 \sqrt{\frac{1}{2} \log(T+1) + \log\left(\frac{2\sqrt{2}T}{H}\right) + \log \frac{2}{\delta} + \mathbb{V}}, \\ \mathbb{V} &= d_c \log\left(1 + \frac{4\bar{W} + 4H\sqrt{2\bar{Z}}}{\Delta}\right) + d_a \log\left(1 + \frac{4H\sqrt{2\bar{Z}}}{\Delta}\right), \\ \bar{W} &= \frac{16}{3} H \sqrt{d_c T} + \sqrt{\frac{2d_c^3 T}{3\zeta\delta}}, \quad \bar{Z} = (\bar{\epsilon} + \eta H) T. \end{aligned}$$

Proof of Lemma F.2. Since $\hat{V}(\cdot) \in [0, H]$, we can invoke Lemma F.4. Then, we have that for any $\Delta > 0$, with probability at least $1 - \delta$,

$$\left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \left[\hat{V}_{h+1}^t(s') - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \right] \right\|_{(\Lambda_h^t)^{-1}}$$

$$\begin{aligned}
&\leq \left(4 H^2 \left[\frac{d_c}{2} \log \left(\frac{T + \lambda}{\lambda} \right) + d_c \log \left(\frac{\mathcal{N}_\Delta(\mathcal{V})}{\Delta} \right) + \log \frac{2}{\delta} \right] + \frac{8 T^2 \Delta^2}{\lambda} \right)^{1/2} \\
&\leq 2 H \left[\frac{d_c}{2} \log \left(\frac{T + \lambda}{\lambda} \right) + d_c \log \left(\frac{\mathcal{N}_\Delta(\mathcal{V})}{\Delta} \right) + \log \frac{2}{\delta} \right]^{1/2} + \frac{2 \sqrt{2} T \Delta}{\sqrt{\lambda}}.
\end{aligned}$$

Setting $\lambda = 1$, $\Delta = \frac{H}{2\sqrt{2}T}$, we have that with probability at least $1 - \delta$,

$$\begin{aligned}
&\left\| \sum_{(s,a,s') \in \mathcal{D}_h^t} \phi(s,a) \left[\hat{V}_{h+1}^t(s') - \mathbb{P}_h \hat{V}_{h+1}^t(s,a) \right] \right\|_{(\Lambda_h^t)^{-1}} \\
&\leq 2 H \sqrt{d_c} \left[\frac{1}{2} \log(T+1) + \log \left(\frac{\mathcal{N}_\Delta(\mathcal{V})}{\frac{H}{2\sqrt{2}T}} \right) + \log \frac{2}{\delta} \right]^{1/2} + H \\
&\leq 3 H \sqrt{d_c} \left[\frac{1}{2} \log(T+1) + \log \left(\frac{2\sqrt{2}T}{H} \right) + \log \frac{2}{\delta} + \mathcal{V} \right]^{1/2},
\end{aligned}$$

where the last inequality uses [Lemma 6.1](#) to bound the log covering number.

This concludes the proof. $\square$

E.3 Proof of [Theorem 6.2](#)

We first instantiate [Lemma D.2](#) in the off-policy setting. Given the above good event, we have that

$$\Gamma_{\text{LMC}}^{\text{off}} = C_\delta^{\text{off}} H \sqrt{d_c} + \frac{4}{3} \sqrt{\frac{2 d_c \log(1/\delta)}{3 \zeta}} + \frac{4}{3} \leq \tilde{\mathcal{O}}(H d_c).$$

Proof of [Theorem 6.2](#). Following the proof of [Theorem 6.1](#) ([Appendix E.2](#)), we can use the same regret decomposition as follows.

$$\begin{aligned}
&\sum_{t=1}^T \left(V_1^*(s_1) - V_1^{\pi^t}(s_1) \right) \\
&= \underbrace{\sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^*} \left[\left\langle \pi_h^*(\cdot | s) - \pi_h^t(\cdot | s), \hat{Q}_h^t(s, \cdot) \right\rangle \right]}_{\text{(I) policy optimization (actor) error}} + \underbrace{\sum_{t=1}^T \sum_{h=1}^H \left(\mathbb{E}_{\pi^*} [\ell_h^t(s, a)] - \mathbb{E}_{\pi^t} [\ell_h^t(s, a)] \right)}_{\text{(II) policy evaluation (critic) error}}.
\end{aligned}$$

Term (I) can be bounded the same way as in [Appendix E.2](#). Hence, it suffices to bound Term (II).

$$\begin{aligned}
\text{Term (II)} &= \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^*} [\ell_h^t(s, a)] - \mathbb{E}_{\pi^t} [\ell_h^t(s, a)] \\
&\stackrel{(i)}{\leq} - \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} [\ell_h^t(s, a)] \\
&\stackrel{(ii)}{\leq} \Gamma_{\text{LMC}}^{\text{off}} \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\left\| \phi(s_h^t, a_h^t) \right\|_{(\Lambda_h^t)^{-1}} \right].
\end{aligned}$$

(i) and (ii) both follow from [Lemma D.2](#), where (i) is based on the optimism guarantee (RHS of [Eq. \(10\)](#)), while (ii) is based on the error bound (LHS of [Eq. \(10\)](#)).

Bounding the sum of bonuses. Since $\Gamma_{\text{LMC}}^{\text{off}}$ is bounded, it suffices to bound $\mathbb{E}_{\pi^t} \left[\left\| \phi(s, a) \right\|_{(\Lambda_h^t)^{-1}} \right]$.

We then index each data point in $\mathcal{D}_h^t$ as $\{(s_h^i, a_h^i, s_{h+1}^i)\}_{i \in [T]}$ and get that $\Lambda_h^t = \sum_{i=1}^T \phi(s_h^i, a_h^i) \phi(s_h^i, a_h^i)^\top + \lambda I$. Then, we have that

$$\sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\left\| \phi(s, a) \right\|_{(\Lambda_h^t)^{-1}} \right]$$

$$\begin{aligned}
&= \sum_{i=1}^T \sum_{h=1}^H \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}} \\
&\quad + \underbrace{\sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\substack{s_h \sim \mathbb{P}(\cdot | s_{h-1}^t, a_{h-1}^t) \\ a_h \sim \pi_h^t(\cdot | s_h)}} \left[\|\phi(s_h, a_h)\|_{(\Lambda_h^t)^{-1}} \right] - \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}}}_{:= \mathcal{M}_{t,h}^{\text{off}}} . \quad (16)
\end{aligned}$$

Applying the elliptical potential lemma. For the first term of Eq. (16), we have that

$$\begin{aligned}
\sum_{t=1}^T \sum_{h=1}^H \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}} &= \sum_{h=1}^H \sum_{t=1}^T \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}} \\
&\stackrel{\text{(iii)}}{\leq} \sum_{h=1}^H \sqrt{T} \left(\sum_{t=1}^T \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}}^2 \right)^{1/2} \\
&\stackrel{\text{(iv)}}{\leq} \mathcal{O} \left(\sqrt{d_c H^2 T \log(T/\delta)} \right).
\end{aligned}$$

(iii) applies the Cauchy-Schwarz inequality, and (iv) follows the elliptical potential argument from Lemma E.5.

A martingale difference sequence. For the second term of Eq. (16), since $\{\mathcal{M}_{t,h}^{\text{off}}\}_{(t,h) \in [T] \times [H]}$ forms a martingale sequence adapted to the filtration,

$$\mathcal{F}_{t,h}^{\text{off}} = \{(s_\tau^i, a_\tau^i)\}_{(i,\tau) \in [t-1] \times [H]} \cup \{(s_\tau^t, a_\tau^t)\}_{\tau \in [h-1]},$$

such that $\mathbb{E}[\mathcal{M}_{t,h}^{\text{off}} | \mathcal{F}_{t,h}^{\text{off}}] = 0$. Since $|\mathcal{M}_{t,h}^{\text{off}}| \leq 1$, we can apply the Azuma–Hoeffding inequality and obtain that

$$\Pr \left(\sum_{t=1}^T \sum_{h=1}^H \mathcal{M}_{t,h}^{\text{off}} \geq m \right) \geq \exp \left(\frac{-m^2}{2HT} \right).$$

Setting $m = \sqrt{2HT \log(1/\delta)}$, with probability at least $1 - \delta$, it holds that

$$\sum_{t=1}^T \sum_{h=1}^H \mathcal{M}_{t,h}^{\text{off}} \leq \sqrt{2HT \log(1/\delta)} \leq \mathcal{O} \left(\sqrt{HT \log(1/\delta)} \right).$$

Therefore, we have that

$$\begin{aligned}
&\sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] \\
&= \sum_{t=1}^T \sum_{h=1}^H \|\phi(s_h^t, a_h^t)\|_{(\Lambda_h^t)^{-1}} + \sum_{t=1}^T \sum_{h=1}^H \mathcal{M}_{t,h}^{\text{off}} \leq \mathcal{O} \left(\sqrt{d_c H^2 T \log(T/\delta)} \right).
\end{aligned}$$

It further implies that, with probability at least $1 - \delta$,

$$\text{Term (II)} \leq \Gamma_{\text{LMC}}^{\text{off}} \sum_{t=1}^T \sum_{h=1}^H \mathbb{E}_{\pi^t} \left[\|\phi(s, a)\|_{(\Lambda_h^t)^{-1}} \right] \leq \tilde{\mathcal{O}} \left(\sqrt{d_c^2 \max\{d_c, d_a\} H^4 T} \right).$$

Putting everything together. Therefore, we have that with probability at least $1 - \delta$,

$$\mathbb{E} \left[V_1^*(s_1) - V_1^{\pi^T}(s_1) \right] = \frac{1}{T} (\text{Term (I)} + \text{Term (II)}) = \tilde{\mathcal{O}} \left(\frac{H^2 \sqrt{d_c^2 \max\{d_c, d_a\} \log |\mathcal{A}|}}{\sqrt{T}} + H^2 \sqrt{\epsilon} \right).$$

This concludes the proof. $\square$

F.4 Technical Tools

Lemma F.3 (Zhong and Zhang, 2023, Lemma B.1). Consider the value function class $\mathcal{V} = \{Q(\cdot, \cdot), \hat{\pi}(\cdot | \cdot)\}_{\mathcal{A}} \mid Q \in \mathcal{Q}, \hat{\pi} \in \Pi\}$. Then, it holds that

$$\mathcal{N}_\Delta(\mathcal{V}) \leq \mathcal{N}_{\Delta/2}(\mathcal{Q}) \cdot \mathcal{N}_{\Delta/(2H)}(\Pi).$$

Lemma F.4 (Value-Aware Uniform Concentration (Jin et al., 2020, Lemma D.4)). Let $\{s_t\}_{t=1}^\infty$ be a stochastic process on the state space $\mathcal{S}$ with the correspond filtration $\{\mathcal{F}_t\}_{t=0}^\infty$ such that s_t is $\mathcal{F}_{t-1}$ -measurable. Let $\{\phi_t\}_{t=1}^\infty$ be an $\mathbb{R}^d$ -valued stochastic process such that ϕ_t is $\mathcal{F}_{t-1}$ -measurable, and $\|\phi_t\| \leq 1$. Let $\Lambda_t = I + \sum_{s=1}^t \phi_s \phi_s^\top$. Assume $\mathcal{V}$ is a value function class such that $\sup_{s \in \mathcal{S}} |V(s)| \leq H$. Then, for any $\delta > 0$, with probability at least $1 - \delta$, for all $t \geq 0$ and any $V \in \mathcal{V}$, it holds that

$$\left\| \sum_{i=1}^t \phi_i \{V(s_i) - \mathbb{E}[V(s_i) \mid \mathcal{F}_{i-1}]\} \right\|_{\Lambda_t^{-1}}^2 \leq 4H^2 \left[\frac{d}{2} \log\left(\frac{t+\lambda}{\lambda}\right) + \log\left(\frac{\mathcal{N}_\Delta}{\delta}\right) \right] + \frac{8t^2 \Delta^2}{\lambda},$$

where $\mathcal{N}_\Delta$ is the Δ -covering number of $\mathcal{V}$ with the distance measured by $\text{dist}(V, V') = \sup_{s \in \mathcal{S}} |V(s) - V'(s)|$.

Lemma F.5 (Covering Number of Euclidean Ball). For any $\Delta > 0$, the Δ -covering number, $\mathcal{N}_\Delta$, of the Euclidean ball of radius $B > 0$ in $\mathbb{R}^d$ satisfies that

$$\mathcal{N}_\Delta \leq \left(1 + \frac{2B}{\Delta}\right)^d.$$

Lemma F.6 (Zhong and Zhang 2023, Lemma B.3). For $\pi, \pi' \in \Delta(\mathcal{A})$ and $Z, Z' : \mathcal{A} \rightarrow \mathbb{R}^+$, if $\pi(\cdot) \propto \exp(Z(\cdot))$ and $\pi'(\cdot) \propto \exp(Z'(\cdot))$, then it holds that

$$\|\pi - \pi'\|_1 \leq 2\sqrt{\|Z - Z'\|_\infty}.$$

G Experiments

In this section, we evaluate our actor-critic algorithm on a Random MDP and a linear MDP version of the Deep Sea (Osband et al., 2019) in the off-policy setting.

G.1 Environment Setup

Our experimental setup is an extension of Ishfaq et al. (2024a). In particular, we extend the prior off-policy setting to test our actor-critic framework on a Random MDP and a linear MDP version of the Deep Sea (Osband et al., 2019). In particular, we use the linear MDP features as the policy features, and $d := d_c = d_a$ represents the feature dimension for both the actor and the critic parameters.

For the Deep Sea environment, we use a $N \times N$ grid with $N = 10$ where the agent always starts at $(0, 0)$ and can move either bottom-right or bottom-left, receiving rewards of 0 and $-0.01/N$ respectively. Reaching the bottom-right corner yields a reward of 1. Furthermore, we generate the actor and critic features by projecting each state-action pair uniformly between $[0, d-1]$, which recovers one-hot encoded features for $d = |\mathcal{S}| \times |\mathcal{A}|$. Given the true transition probabilities and rewards, and following the linear MDP assumption in Definition 2.1, we solve for ψ_h and v_h^* via least squares and obtain the corresponding $\mathbb{P}_h$ and r_h .

For the Random MDP environment, we consider 15 states and 5 actions and set $d = 30$. For each state $s \in \mathcal{S}$, we generate $\psi_h(s) \in \mathbb{R}^d$ uniformly at random in $[0, 1]$ and construct tile coded features. Following Definition 2.1 and using least squares (similar to Deep Sea), we obtain the probability transitions. The agent always starts from state 0, receiving a small reward of 0.1 upon taking action 0, and obtains the maximum reward when reaching the final state and taking action 1. All other state-action pairs yield zero reward. Using the same procedure for Deep Sea, we compute r_h and ensure linearity of the MDP.

G.2 Coreset Construction

To construct the coreset, we follow the offline G-experimental design outlined in Algorithm 4. In particular, in each iteration, this greedy iterative algorithm traverses the entire state-action space and adds a data point to the coreset that has the highest marginal gain $g(s, a) = \|\varphi(s, a)\|_{G^{-1}}$. For a specific threshold ϵ_G , the algorithm only terminates when $g_{\max} = \max_{s, a \in (\mathcal{S} \times \mathcal{A})} g(s, a) \leq \epsilon_G$, hence giving us direct control over $\sup_{(s, a) \in \mathcal{S} \times \mathcal{A}} \|\varphi(s, a)\|_{G^{-1}}$. In practice, we find that it often selects too many data points, so we cap the coreset at 80% of the total data.

Algorithm 4 Coreset Construction via G-Experimental Design

```

1: Input: features  $\varphi : \mathcal{S} \times \mathcal{A} \mapsto \mathbb{R}^{d_a}$ , threshold  $\epsilon_G \in \mathbb{R}$ 
2: Initialize:  $G = I_{d_a \times d_a}$ ,  $\mathcal{D}_{\text{exp}} = \emptyset$ ,  $g_{\text{max}} = \infty$ 
3: while  $g_{\text{max}} > \epsilon_G$  do
4:    $g_{\text{max}} = 0$ 
5:   for  $(s, a) \in \mathcal{S} \times \mathcal{A}$  do
6:      $g(s, a) = \|\varphi(s, a)\|_{G^{-1}}$ 
7:     if  $g_{\text{max}} < g(s, a)$  then
8:        $(s^*, a^*) = (s, a)$ 
9:        $g_{\text{max}} = g(s, a)$ 
10:     $\mathcal{D}_{\text{exp}} = \mathcal{D}_{\text{exp}} \cup \{(s^*, a^*)\}$ 
11:     $G = G + \varphi(s^*, a^*) \varphi(s^*, a^*)^\top$ 

```

G.3 Hyperparameters

In Table 2, we list the hyperparameters we tested across all experiments. For log-linear policies, the actor loss in Eq. (6) admits a closed-form solution, allowing us to avoid tuning α_a and K_t by minimizing the objective exactly. We swept the hyperparameters and picked the best combination for each of the considered methods to report the results.

Table 2: Hyperparameter sweep in our experiments.

Hyperparameter	LMC	LMC-NPG-IMP	LMC-NPG-EXP
Policy Optimization Learning Rate (η)	$\times$	[0.1, 1, 10, 100]	[0.1, 1, 10, 100]
Inverse Temperature (ζ^{-1})		[10^{-2} , 10^{-3} , 10^{-4} , 10^{-5}]	
Number of Critic Updates (J_t)		100	
Critic Learning Rate (α_c)		[10^{-2} , 10^{-3} , 10^{-4} , 10^{-5}]	
Number of Episodes (T)		600	
Horizon Length (H)		100	

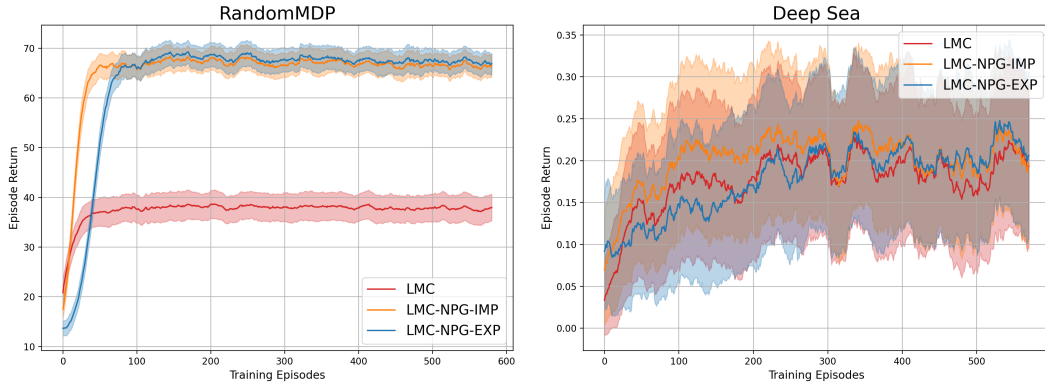
G.4 Experimental Results

Figure 1: Comparison of LMC-NPG-EXP (our proposed framework), LMC-NPG-IMP (memory-intensive variant), and LMC (value-based baseline) in the Random MDP and the Deep Sea.

We denote by LMC-NPG-EXP our proposed framework with an explicit log-linear policy parameterization that uses LMC for policy evaluation and projected NPG for policy optimization on a coreset. We denote by LMC-NPG-IMP an idealized implicit variant of NPG that does not have an explicit actor parameterization and maintains an implicit policy by storing all parameterized Q functions (and hence requires significantly more memory). As a baseline, we include the value-based algorithm LMC (Ishfaq et al., 2024a). Following the protocol of Ishfaq et al. (2024a), each algorithm is run with 20 random seeds. We sweep hyperparameters as discussed in Appendix G.3 and report the best performance with 95% confidence intervals.

For the random linear MDP, Figure 1(left) indicates that LMC-NPG-EXP closely matches LMC-NPG-IMP while outperforming the value-based baseline, LMC. For the Deep Sea, Figure 1(right) showcases that LMC-NPG-EXP can achieve comparable performance with LMC-NPG-IMP and LMC.

G.5 Additional Results

Ablation on Feature Dimensions. When using LMC-NPG-EXP, which employs LMC for policy evaluation to optimize an explicitly parameterized log-linear policy over a coresct, we study the impact of the feature dimension ($d := d_c = d_a$). The results in Figure 2 show that larger feature dimensions d for both the actor and the critic can substantially improve the performance of the proposed framework.

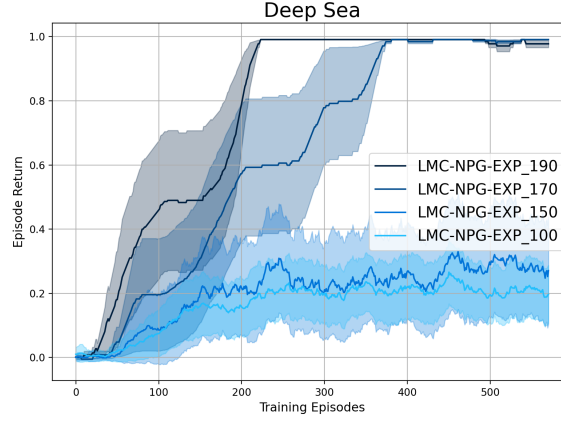


Figure 2: Effect of feature dimension d in the Deep Sea..