
Finite-time Convergence Analysis of Actor-Critic with Evolving Reward

Rui Hu¹ Yu Chen¹ Longbo Huang¹

Abstract

Many popular practical reinforcement learning (RL) algorithms employ evolving reward functions, through techniques such as reward shaping, entropy regularization, or curriculum learning, yet their theoretical foundations remain underdeveloped. This paper provides the first finite-time convergence analysis of a single-timescale actor-critic algorithm in the presence of an evolving reward function under Markovian sampling. We consider a setting where the reward parameters may change at each time step, affecting both policy optimization and value estimation. Under standard assumptions, we derive non-asymptotic bounds for both actor and critic errors. Our result shows that an $O(1/\sqrt{T})$ convergence rate is achievable, matching the best-known rate for static rewards, provided the reward parameters evolve slowly enough. This rate is preserved when the reward is updated via a gradient-based rule with bounded gradient and on the same timescale as the actor and critic, offering a theoretical foundation for many popular RL techniques. As a secondary contribution, we introduce a novel analysis of distribution mismatch under Markovian sampling, improving the best-known rate by a factor of $\log^2 T$ in the static-reward case.

1. Introduction

Reinforcement Learning (RL, Sutton et al. 1998) has attracted great research interest in the past decades. On the empirical side, a variety of practical algorithms have been proposed and have demonstrated remarkable success in a wide range of real-world scenarios (Mnih et al., 2013; Lillicrap et al., 2016; Ouyang et al., 2022; DeepSeek-AI et al., 2025). On the theoretical side, great efforts have been made to bridge the gap between empirical practice and theoretical

foundations, providing rigorous convergence guarantees and solid theoretical understandings to the empirically powerful algorithms (Agarwal et al., 2021; Mei et al., 2020; Kumar et al., 2023; Wu et al., 2020; Olshevsky & Ghahserifard, 2023).

Still, a very common setting adopted by many practical RL algorithms has been largely overlooked by theoretical analyses. RL theory is built upon Markov Decision Processes (MDPs), which typically assume the existence of a static underlying reward function, and the goal is to learn a policy that maximizes the expected cumulative reward. However, when applying RL in many real-world scenarios, a significant impediment is the difficulty of designing a reward function that is both learnable and aligns with the desired task. This challenge has led to the development of techniques that utilize evolving rewards. These include:

- **Reward Shaping:** Adding auxiliary rewards to guide the policy to the desired goal (Ng et al., 1999; Pathak et al., 2017; Burda et al., 2019; Zheng et al., 2018; Hu et al., 2020; Mahankali et al., 2024; Ma et al., 2024). The auxiliary rewards can come from prior knowledge, or be learned in a self-supervised manner during the training process.
- **(Adaptive) Entropy or KL Regularization:** Introducing an entropy or KL regularization term to the optimization objective, which is equivalent to modifying the reward according to the current policy (Haarnoja et al., 2017; 2018a;b; Jaques et al., 2019; Stiennon et al., 2020). The regularization factor can be automatically adjusted during training.
- **Curriculum Learning:** Starting with easier tasks (and their associated rewards) and gradually increasing the difficulty (Narvekar et al., 2020).

Intuitively, a slight change of the reward function does not drastically alter the solution of the underlying MDP, allowing an RL algorithm to remain effective. However, when this change is negligible compared to the under-training policy or value function, the algorithm’s effectiveness becomes questionable, as they are closely interconnected. Therefore, to rigorously support the use of these evolving reward techniques, we must answer the following fundamental question precisely:

¹IIS, Tsinghua University, Beijing, China. Correspondence to: Longbo Huang <longbohuang@tsinghua.edu.cn>.

How fast can the reward change while still ensuring the convergence of an RL algorithm?

This paper aims to establish a theoretical foundation for this setting by providing the first finite-time convergence analysis for an actor-critic algorithm with an evolving reward. We focus on a single-sample, single-timescale actor-critic algorithm with linear function approximation for the critic under Markovian sampling. This setting is particularly practical yet challenging, as the non-stationarity from the evolving reward affects both the policy gradient (actor) and the value estimation (critic), creating a complex feedback loop.

Specifically, we bound both the expected actor error and the expected critic error in terms of the number of iterations T and the total change of the reward parameters. From this, we derive conditions on the evolving rate of reward parameters necessary to achieve asymptotic convergence and to maintain the $O(1/\sqrt{T})$ convergence rate as in the static-reward case. Moreover, it turns out that $O(1/\sqrt{T})$ convergence can be achieved if the reward parameter follows a gradient-based update with bounded gradient and the same timescale as the actor and the critic, providing theoretical guarantees to a wide range of practical techniques, including curiosity-driven reward shaping (Pathak et al., 2017), random network distillation methods (Burda et al., 2019), and soft actor-critic with automated entropy adjustment (Haarnoja et al., 2018b).

Our analysis is developed upon the existing finite-time convergence studies of single-timescale actor-critic methods in the static-reward scenario (Chen et al., 2021; Olshevsky & Gharesifard, 2023; Tian et al., 2023; Chen & Zhao, 2023; 2025). To handle the evolving reward, we exploit the Lipschitz continuity of the objective function and the optimal critic parameter with respect to the reward parameter, which relies on the Lipschitz assumption regarding the rewards themselves. In addition, we introduce a novel analysis on the distribution mismatch caused by Markovian sampling, improving the convergence rate by a factor of $\log^2 T$ in the static-reward case.

In summary, our contributions include the following:

- **Important Problem Formulation:** We formalize the problem of Actor-Critic with Evolving Reward, where the reward parameter φ_t (encompassing the true reward and regularization terms) can be updated by an arbitrary oracle at every time step.
- **Novel Non-Asymptotic Results:** Under standard assumptions (Linear function approximated critic, Lipschitz continuity of policy and reward, sufficient exploration), we derive the convergence rate of the single-sample single-timescale actor-critic algorithm with Markovian sampling, and show that it achieves a convergence rate of $O(1/\sqrt{T})$ to a neighborhood of a

stationary point under mild condition on the evolving reward, validating a wide range of practical RL techniques.

- **Interesting Technical Tools:** We establish the necessary assumptions and key properties to analyze the effects of the evolving reward on standard actor-critic algorithms. We also provide a novel analysis of the distribution mismatch caused by Markovian sampling, which independently improves the convergence rate for the static-reward case.

2. Related Work

Our work is developed upon the finite-time analysis of policy gradient methods and actor-critic methods.

Finite-time analysis of Policy Gradient Methods. The finite-time convergence guarantees of policy gradient methods have been studied by (Agarwal et al., 2021; Mei et al., 2020; Xiao, 2022), assuming access to exact gradient oracles. For the stochastic case where the algorithm can only access gradient estimators from sampled trajectories or transitions, convergence results in terms of sample complexity have been established by (Liu et al., 2020; Ding et al., 2022; 2025; Fatkhullin et al., 2023; Mondal & Aggarwal, 2024).

Finite-time analysis of Actor-Critic Methods. The finite-time analysis of actor-critic methods encompasses several variants of the algorithm, including the double-loop setting (Yang et al., 2019; Kumar et al., 2023; Xu et al., 2020a; Cayci et al., 2024; Gaur et al., 2024; Ganesh et al., 2025), the two-timescale setting (Wu et al., 2020; Xu et al., 2020b; Shen et al., 2023; Hong et al., 2023), and the single-timescale setting (Chen et al., 2021; Olshevsky & Gharesifard, 2023; Tian et al., 2023; Chen & Zhao, 2023; 2025; Kumar et al., 2026). The analysis of single-timescale actor-critic methods is particularly relevant to our work. Both (Chen et al., 2021) and (Olshevsky & Gharesifard, 2023) obtain an $O(1/\sqrt{T})$ local convergence rate in discrete spaces, assuming i.i.d. sampling. (Chen & Zhao, 2023) and (Tian et al., 2023) made efforts to tackle the more practical yet challenging Markovian sampling setting. However, the former considers the average-reward scenario instead of the commonly used discounted-reward scenario, while the latter employs Markovian samples for the critic and i.i.d. samples for the actor. (Chen & Zhao, 2025) ultimately resolves the problem of Markovian sampling, obtaining an $O(\log^2 T/\sqrt{T})$ convergence rate, and further extends the analysis to continuous action spaces. When translating these results to global convergence, the above-mentioned local convergence rate of $\tilde{O}(1/\sqrt{T})$ leads to a sample complexity of $\tilde{O}(\epsilon^{-4})$ in order to obtain an ϵ -optimal policy under mild assumptions (Mei et al., 2020). A recent study by

Kumar et al. (2026) improves this bound to $\tilde{O}(\epsilon^{-3})$ under i.i.d. sampling. However, this result is limited to finite state-action spaces and tabular critics, whereas Tian et al. (2023) incorporates an analysis of neural network approximated critics, and the other four studies assume a linear function approximated critic.

While we focus on analyzing the convergence of the single-timescale actor-critic algorithm under evolving rewards, we recognize two related research areas that focus on designing reinforcement learning (RL) algorithms for non-static Markov Decision Processes (MDPs), where both rewards and transitions can change:

Adversarial RL and Non-stationary RL. Tailored for on-line learning, adversarial RL (Even-Dar et al., 2009; Zimin & Neu, 2013; Jin et al., 2020) and non-stationary RL (Cheung et al., 2023; Feng et al., 2023; Mao et al., 2025) aim to minimize cumulative static or dynamic regret through strategic algorithm design in the presence of adversarial rewards and transitions. However, their formulation does not align with the practical scenarios we consider, where the RL policy is trained before executing, and the evolving reward techniques we previously mentioned focus on enhancing convergence during training and improving performance in execution, rather than adapting in an adversarial manner.

Performative RL. Performative RL (Mandal et al., 2023; Rank et al., 2024; Mandal & Radanovic, 2025) examines situations in which the underlying MDP changes in response to the deployed policy. This setting partially overlaps with ours, as performative rewards represent a specific case within our general evolving reward framework. Under certain mild assumptions regarding the performative rewards and transitions, these studies demonstrate the existence of a stable policy and provide finite-time convergence guarantees for their proposed algorithms, which involve iterating between deployment and retraining. However, their approach necessitates finding exact solutions for a saddle point at each retraining step, which is often infeasible in practice.

3. Preliminaries

3.1. Markov Decision Process

We consider an infinite-horizon discounted Markov Decision Process (MDP), defined by the tuple $\mathcal{M} = (\mathcal{S}, \mathcal{A}, \mathcal{P}, r, \gamma)$. Here, $\mathcal{S}$ is the state space, $\mathcal{A}$ is the action space, $\mathcal{P} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ is the transition kernel, $r : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ is the reward function, and $\gamma \in (0, 1)$ is the discount factor.

A policy $\pi : \mathcal{S} \rightarrow \Delta(\mathcal{A})$ maps states to distributions over actions. Starting from an initial state s_0 , at each time step $t = 0, 1, \dots$, an action $a_t \sim \pi(\cdot|s_t)$ is sampled, yielding

a reward $r_t = r(s_t, a_t)$ and transitioning to the next state $s_{t+1} \sim \mathcal{P}(\cdot|s_t, a_t)$. The standard value function $V^\pi(s)$ and action-value function $Q^\pi(s, a)$ represent the expected discounted cumulative reward starting from state s (and action a), respectively, and are defined as

$$V^\pi(s) = \mathbb{E}_{\pi, \mathcal{P}} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s \right], \quad (1)$$

$$Q^\pi(s, a) = \mathbb{E}_{\pi, \mathcal{P}} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s, a_0 = a \right]. \quad (2)$$

Entropy Regularized MDPs. Policy optimization often benefits from entropy regularization to encourage exploration (Haarnoja et al., 2017; 2018a;b). The entropy of a policy π at a state s is defined as $\mathcal{H}(\cdot|s) = -\int_{\mathcal{A}} \pi(a|s) \log \pi(a|s) da$. Let

$$H^\pi(s) = \mathbb{E}_{\pi, \mathcal{P}} \left[\sum_{t=0}^{\infty} \gamma^t \mathcal{H}(\pi(\cdot|s_t)) \mid s_0 = s \right], \quad (3)$$

the regularized (soft) value function is defined as

$$\tilde{V}^\pi(s) = V^\pi(s) + \alpha H^\pi(s),$$

where $\alpha \geq 0$ is a hyper-parameter, and $\alpha = 0$ recovers the standard, unregularized case. This formulation is equivalent to solving the original MDP with a regularized reward function:

$$\tilde{r}(s, a) = r(s, a) - \alpha \log \pi(a|s). \quad (4)$$

Consequently, $\tilde{V}(s)$ and $\tilde{Q}(s, a)$ can be interpreted as the value functions under this regularized reward $\tilde{r}$:

$$\tilde{V}^\pi(s) = \mathbb{E}_{\pi, \mathcal{P}} \left[\sum_{t=0}^{\infty} \gamma^t \tilde{r}(s_t, a_t) \mid s_0 = s \right],$$

$$\tilde{Q}^\pi(s, a) = \mathbb{E}_{\pi, \mathcal{P}} \left[\sum_{t=0}^{\infty} \gamma^t \tilde{r}(s_t, a_t) \mid s_0 = s, a_0 = a \right].$$

Note on KL Regularization. While we focus on entropy regularization for clarity, our analysis also applies to KL regularization against a fixed reference policy π_{ref} . The corresponding regularized reward becomes $\tilde{r}(s, a) = r(s, a) + \alpha \log \pi_{\text{ref}}(a|s) - \alpha \log \pi(a|s)$. Since the term $\alpha \log \pi_{\text{ref}}(a|s)$ can be absorbed into the base reward $r(s, a)$, we will consider only the entropy regularizer in the following analysis without loss of generality.

The goal of reinforcement learning is to find a parameterized policy π_θ that maximizes the objective:

$$J(\theta) = \int_{\mathcal{S}} \rho(s) \tilde{V}^{\pi_\theta}(s) ds, \quad (5)$$

where θ denotes the policy parameter and $\rho \in \Delta(\mathcal{S})$ is the initial distribution.

3.2. Actor-Critic Method

In order to optimize the policy parameter θ , a popular approach is to compute the gradient of the objective $J(\theta)$ and iteratively adjust θ in the direction of $\nabla J(\theta)$. This approach, named the policy gradient method, is based on the policy gradient theorem (Sutton et al., 1999):

$$\begin{aligned}\nabla_{\theta} J(\theta) &= \frac{1}{1-\gamma} \mathbb{E}[\tilde{Q}^{\pi_{\theta}}(s, a) \nabla_{\theta} \log \pi_{\theta}(a|s)] \\ &= \frac{1}{1-\gamma} \mathbb{E}[(\tilde{Q}^{\pi_{\theta}}(s, a) - \tilde{V}^{\pi_{\theta}}(s)) \nabla_{\theta} \log \pi_{\theta}(a|s)],\end{aligned}\quad (6)$$

where s is sampled from the discounted visitation distribution $\nu_{\rho}^{\pi_{\theta}} \in \Delta(\mathcal{S})$ defined as

$$\nu_{\rho}^{\pi_{\theta}}(s) = (1-\gamma) \sum_{t=0}^{\infty} \gamma^t \Pr(s_t = s),$$

with $s_0 \sim \rho(\cdot)$, $a_t \sim \pi_{\theta}(\cdot|s_t)$, $s_{t+1} \sim \mathcal{P}(\cdot|s_t, a_t)$.

Computing this gradient requires an estimator for the $\tilde{V}$ function or $\tilde{Q}$ function associated with the policy π_{θ} . A Monte-Carlo estimator (used by REINFORCE, Williams 1992) suffers from high variance, resulting in slow convergence. Hence, the actor-critic method (Konda & Tsitsiklis, 1999) introduces another trainable model to approximate the true value functions.

Following the setting of prior work (Chen et al., 2021; Olshesky & Ghahesifard, 2023; Chen & Zhao, 2023; 2025), we assume that the critic approximates the regularized value function linearly as

$$\hat{V}_{\omega}(s) = \phi(s)^{\top} \omega,$$

where $\phi: \mathcal{S} \rightarrow \mathbb{R}^d$ is a known feature mapping satisfying $\|\phi(s)\|_2 \leq 1$ for any state s , and $\omega \in \mathbb{R}^d$ is the trainable parameter. Note that $\tilde{Q}^{\pi}(s, a) = \tilde{r}(s, a) + \gamma \mathbb{E}_{s'}[\tilde{V}^{\pi}(s')]$, then with $s' \sim \mathcal{P}(\cdot|s, a)$, the temporal difference (TD) error

$$\begin{aligned}\hat{\delta}(s, a, s') &= r(s, a) - \alpha \log \pi_{\theta}(a|s) \\ &\quad + (\gamma \phi(s') - \phi(s))^{\top} \omega\end{aligned}\quad (7)$$

serves as a biased but low-variance gradient estimator for (6), and the update rule for θ will be

$$\theta_{t+1} \leftarrow \theta_t + \eta_t^{\theta} \hat{\delta}(s_t, a_t, s'_t) \nabla_{\theta} \log \pi_{\theta}(a_t|s_t), \quad (8)$$

where η_t^{θ} denotes the step size for θ at step t .

Also, we need to align $\hat{V}_{\omega}$ with the true value function $\tilde{V}^{\pi_{\theta}}$. As proven by (Haarnoja et al., 2017), $\tilde{V}^{\pi}$ is the unique solution to the soft Bellman equation $V = \mathcal{T}_{\alpha}^{\pi} V$ where the soft Bellman operator $\mathcal{T}_{\alpha}^{\pi}$ is defined as

$$\mathcal{T}_{\alpha}^{\pi} V(s) = \mathbb{E}_{\substack{a \sim \pi(\cdot|s) \\ s' \sim \mathcal{P}(\cdot|a, s)}} [r(s, a) - \alpha \log \pi(a|s) + \gamma V(s')].$$

Hence, a common practice is to adjust the value iteration update $V \leftarrow \mathcal{T}_{\alpha}^{\pi} V$ to a stochastic semi-gradient TD(0) update:

$$\omega_{t+1} \leftarrow \omega_t + \eta_t^{\omega} \hat{\delta}(s_t, a_t, s'_t) \phi(s_t), \quad (9)$$

where η_t^{ω} denotes the step size for ω at step t .

Markovian Sampling. In our single-sample, single-timescale setting, both actor and critic are updated using the same sample tuple (s_t, a_t, s'_t) at each step. Ideally, s_t shall be sampled from the stationary distribution $\nu_{\rho}^{\pi_{\theta}}$, but this is impractical. Instead, we adopt a Markovian sampling scheme (Chen & Zhao, 2025):

$$s_t \sim \hat{\mathcal{P}}(\cdot|s_{t-1}, a_{t-1}), a_t \sim \pi_{\theta_t}(\cdot|s_t), s'_t \sim \mathcal{P}(\cdot|s_t, a_t),$$

where $\hat{\mathcal{P}}(\cdot|s, a) = \gamma \mathcal{P}(\cdot|s, a) + (1-\gamma)\rho(\cdot)$ ensures ergodicity. Let $\hat{\nu}_t(\cdot)$ denote the probability distribution of s_t induced by this process. When the policy is fixed, $\hat{\nu}_t$ will converge to $\nu_{\rho}^{\pi_{\theta}}$ geometrically. Under a slowly changing policy, the distribution mismatch $\|\hat{\nu}_t - \nu_{\rho}^{\pi_{\theta_t}}\|_1$ can be controlled by the magnitude of the policy updates. Note that this constitutes an off-policy learning setting for the critic.

3.3. Actor-Critic with Evolving Reward

Algorithm 1 Actor Critic with Evolving Reward

```
Initialize:  $\theta_0, \omega_0, \varphi_0, \rho, \{\eta_t^{\theta}\}_{t \geq 0}, \{\eta_t^{\omega}\}_{t \geq 0}$ 
Sample  $s_0 \sim \rho$ 
for  $t = 0, 1, \dots, T-1$  do
    Sample  $a_t \sim \pi_{\theta_t}(\cdot|s_t), s'_t \sim \mathcal{P}(\cdot|s_t, a_t), s_{t+1} \sim \hat{\mathcal{P}}(\cdot|s_t, a_t)$ 
     $\hat{\delta}_t \leftarrow \tilde{r}_{\varphi_t, \theta_t}(s_t, a_t) + (\gamma \phi(s'_t) - \phi(s_t))^{\top} \omega_t$ 
     $\theta_{t+1} \leftarrow \theta_t + \eta_t^{\theta} \hat{\delta}_t \nabla_{\theta} \log \pi_{\theta}(a_t|s_t)$ 
     $\omega_{t+1} \leftarrow \text{Proj}_{C_{\omega}}(\omega_t + \eta_t^{\omega} \hat{\delta}_t \phi(s_t))$ 
     $\varphi_{t+1} \leftarrow \text{UpdateReward}(\varphi_t)$ 
end for
```

A central focus of this work is the setting where the regularized reward $\tilde{r}(s, a)$ evolves during training. Depending on the algorithmic design, this evolution can arise from modifications to the base reward $r(s, a)$, the regularization factor α , or the policy π_{θ} itself. To encompass these variables, we introduce a general reward parameter φ to include all factors that determine $\tilde{r}(s, a)$ along with θ , i.e.,

$$\tilde{r}_{\varphi, \theta}(s, a) = r(s, a; \varphi) - \alpha(\varphi) \log \pi_{\theta}(a|s). \quad (10)$$

We denote the corresponding soft value function as $\tilde{V}_{\varphi}^{\pi_{\theta}}(s)$, and the policy objective as $J_{\varphi}(\theta)$. The reward parameter φ is updated concurrently with θ and ω at each time step via an arbitrary update rule, which can either be a pre-defined

schedule or a feedback-driven strategy. For the critic update, we also introduce a projection Proj_{C_ω} to keep the critic norm bounded by C_ω , which is widely adopted in the literature (Wu et al., 2020; Chen et al., 2021; Olshevsky & Ghahesifard, 2023; Chen & Zhao, 2023; 2025). This framework, summarized in Algorithm 1, unifies a wide range of existing techniques, from automated reward shaping (Martin et al., 2017; Pathak et al., 2017; Burda et al., 2019) to adaptive entropy and KL regularization (Haarnoja et al., 2018b). We provide a further literature review of these evolving reward techniques in Appendix A.

4. Main Results

This section presents the finite-time convergence guarantees for the Actor-Critic with Evolving Reward algorithm (Algorithm 1). We begin by stating the standard assumptions required for our analysis, then present the main theorem and a key corollary. Finally, we provide an intuitive proof sketch to elucidate the key technical challenges and innovations.

4.1. Assumptions

Our analysis relies on several standard assumptions in the literature, which we adapt to accommodate the evolving reward setting.

Take the expectation of ω_{t+1} conditioning on ω_t in (9) with respect to the discounted visitation distribution, we have

$$\mathbb{E}[\omega_{t+1}|\omega_t] = \omega_t + \eta_t^\omega (\mathbf{b}_{\varphi,\theta} - \mathbf{A}_\theta \omega_t),$$

where

$$\mathbf{A}_\theta = \mathbb{E}_{\substack{s \sim \nu_\rho^{\pi_\theta}(\cdot) \\ a \sim \pi_\theta(\cdot|s) \\ s' \sim \mathcal{P}(\cdot|s,a)}} [\phi(s)(\phi(s) - \gamma\phi(s'))^\top], \quad (11)$$

$$\mathbf{b}_{\varphi,\theta} = \mathbb{E}_{\substack{s \sim \nu_\rho^{\pi_\theta}(\cdot) \\ a \sim \pi_\theta(\cdot|s)}} [\tilde{r}_{\varphi,\theta}(s,a)\phi(s)]. \quad (12)$$

It has been shown by Sutton et al. (1998) that the TD limiting point $\omega^*(\varphi, \theta)$ satisfies

$$\mathbf{A}_\theta \omega^*(\varphi, \theta) = \mathbf{b}_{\varphi,\theta}. \quad (13)$$

To ensure the existence of ω^* , we need $\mathbf{A}_\theta$ to be non-singular. Fortunately, we can show that $\mathbf{A}_\theta$ is positive definite given sufficient exploration over the state space.

Assumption 4.1 (Sufficient Exploration). Let $\Sigma_\theta = \mathbb{E}_{\nu_\rho^{\pi_\theta}} [\phi(s)\phi(s)^\top]$, then for any $\theta \in \Omega(\theta)$, Σ_θ is positive definite with singular values lower-bounded by $\lambda_\Sigma > 0$.

Proposition 4.2. For any $\theta \in \Omega(\theta)$, $\mathbf{A}_\theta$ is positive definite with singular values lower-bounded by $\lambda = (1 - \sqrt{\gamma})\lambda_\Sigma$.

Proposition 4.2 is widely adopted as an assumption in analyzing TD-learning and actor-critic with linear function approximation (Wu et al., 2020; Chen et al., 2021; Olshevsky

& Ghahesifard, 2023; Chen & Zhao, 2023; 2025). Although the definition of $\mathbf{A}_\theta$ here differs from previous literature where the expectation is taken over the stationary distribution instead of the discounted visitation distribution as in our definition, this nice property can still be induced from the fundamental Assumption 4.1 (Bhandari et al., 2018). Equipped with Proposition 4.2, we can now solve from (13) that $\omega^*(\varphi, \theta) = \mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi,\theta}$, and further imply that $\omega^*(\varphi, \theta)$ is bounded by some constant C_ω since both $\mathbf{A}_\theta^{-1}$ and $\mathbf{b}_{\varphi,\theta}$ can be shown bounded, which justifies the projection operator introduced in Algorithm 1.

As ω^* is bounded, $\hat{V}_{\omega^*}$ is bounded, the linear function approximation error has a uniform upper bound, denoted as ϵ . Formally,

$$\epsilon := \sup_{\theta, \varphi} \sqrt{\mathbb{E}_{s \sim \nu_\rho^{\pi_\theta}(\cdot)} \left[\left(\phi(s)^\top \omega^*(\varphi, \theta) - \tilde{V}^{\pi_\theta}(s) \right)^2 \right]}. \quad (14)$$

The error ϵ will be zero if $\tilde{V}^{\pi_\theta}(\cdot)$ is indeed a linear function for any φ and θ given the feature mapping $\phi(\cdot)$. With a proper $\phi(\cdot)$, ϵ is controllable (Eshwar et al., 2026).

To capture the bias of the TD-gradient estimator for the actor, we need the following bound that controls the error of TD-errors (refer to Appendix D for a detailed proof):

Proposition 4.3. For any $\theta \in \Omega(\theta)$, $\varphi \in \Omega(\varphi)$,

$$\sqrt{\mathbb{E}_{\nu_\rho^{\pi_\theta}, \pi_\theta, \mathcal{P}} \left[\left(\hat{\delta}(s, a, s'; \omega^*) - \tilde{\delta}(s, a, s') \right)^2 \right]} \leq 2\sqrt{2}\epsilon,$$

where $\tilde{\delta}(s, a, s') = r(s, a) - \alpha \log \pi_\theta(a|s) + \gamma \tilde{V}^{\pi_\theta}(s') - \tilde{V}^{\pi_\theta}(s)$.

Assumption 4.4. There exist constants L and S such that for any $\theta \in \Omega(\theta)$, $s \in \mathcal{S}$, $a \in \mathcal{A}$,

$$\|\nabla_\theta \log \pi_\theta(a|s)\|_2 \leq L, \quad \|\nabla_\theta^2 \log \pi_\theta(a|s)\|_2 \leq S.$$

Assumption 4.4 is standard in the literature of policy gradient and actor-critic (Wu et al., 2020; Chen et al., 2021; Olshevsky & Ghahesifard, 2023; Chen & Zhao, 2023; Tian et al., 2023), which further implies the following proposition that the policy π_θ is Lipschitz continuous with respect to θ (refer to Appendix D for a detailed proof):

Proposition 4.5 (Lipschitz Continuity of Policy). For any $\theta_1, \theta_2 \in \Omega(\theta)$ and $s \in \mathcal{S}$,

$$\|\pi_{\theta_1}(\cdot|s) - \pi_{\theta_2}(\cdot|s)\|_1 \leq L\|\theta_1 - \theta_2\|_2.$$

Apart from the policy, we also need the Lipschitz continuity and smoothness of the regularized reward to control the bias caused by the evolving reward. These properties can be derived from the following basic assumptions.

Assumption 4.6 (Bounded Reward). For any $\varphi \in \Omega(\varphi)$ and $(s, a) \in \mathcal{S} \times \mathcal{A}$, $r(s, a; \varphi)$ is bounded by $C_R > 0$.

Assumption 4.7 (Bounded Regularization Factor). For any $\varphi \in \Omega(\varphi)$, $\alpha(\varphi)$ is upper-bounded by $C_\alpha > 0$.

Assumption 4.8. For any $\theta \in \Omega(\theta)$, $\varphi \in \Omega(\varphi)$, and $s \in \mathcal{S}$,

$$\mathbb{E}_{\pi_\theta}[-\log \pi_\theta(a|s)] \leq C_{H,1}, \quad \mathbb{E}_{\pi_\theta}[(\log \pi_\theta(a|s))^2] \leq C_{H,2}.$$

Assumption 4.9. There exists constants D_R and D_α such that for any $\varphi \in \Omega(\varphi)$ and $(s, a) \in \mathcal{S} \times \mathcal{A}$,

$$\|\nabla_\varphi r(s, a; \varphi)\|_2 \leq D_R, \quad \|\nabla_\varphi \alpha(\varphi)\|_2 \leq D_\alpha.$$

Remark. It is important to note that Assumption 4.8 holds inherently for finite action spaces. For infinite action spaces, this assumption also applies to a broad class of policies, including Gaussian policies and exponential policies, as long as their variance remains bounded. This is because entropy regularization implicitly constrains the policy to maintain bounded entropy, which in turn restricts the variance. If entropy were unbounded, it can result in infinite negative rewards, which is practically avoided.

Proposition 4.10 (Properties of Regularized Reward). *There exist constants $C, D > 0$ such that for any $\theta \in \Omega(\theta)$, $\varphi \in \Omega(\varphi)$, and $s \in \mathcal{S}$, the expected regularized reward satisfies:*

1. **Boundedness:** $|\mathbb{E}_{a \sim \pi_\theta(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)]| \leq C$;
2. **Bounded Variance:** $\mathbb{E}_{a \sim \pi_\theta(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)^2] \leq C^2$;
3. **Lipschitz in θ :** $\|\nabla_\theta \mathbb{E}_{a \sim \pi_\theta(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \leq CL$;
4. **Smoothness in θ :** $\|\nabla_\theta^2 \mathbb{E}_{a \sim \pi_\theta(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \leq C(L^2 + S)$;
5. **Lipschitz in φ :** $\|\nabla_\varphi \mathbb{E}_{a \sim \pi_\theta(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \leq D$,

where $C = C_R + C_\alpha (1 + C_{H,1} \vee \sqrt{C_{H,2}})$, $D = D_R + C_{H,1} D_\alpha$.

Parts 1 and 2 of Proposition 4.10 are natural extensions of the standard bounded reward assumption, which can be directly derived from Assumption 4.6, 4.7 and 4.8.

Parts 3 and 4 of Proposition 4.10 can be implied from Part 1 and Assumption 4.4. Since

$$\begin{aligned} \nabla_\theta \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)] &= \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a) \nabla_\theta \log \pi_\theta(a|s)], \\ \nabla_\theta^2 \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)] &= \mathbb{E}[(\tilde{r}_{\varphi, \theta}(s, a) - \alpha) \nabla_\theta \log \pi_\theta(a|s) \nabla_\theta \log \pi_\theta(a|s)^\top] \\ &\quad + \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a) \nabla_\theta^2 \log \pi_\theta(a|s)], \end{aligned}$$

the Lipschitz continuity and smoothness w.r.t. θ can be derived from the boundedness of $\mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)]$, $\nabla_\theta \log \pi_\theta(a|s)$, $\nabla_\theta^2 \log \pi_\theta(a|s)$ and α .

Part 5 of Proposition 4.10 can be derived from Assumption 4.8 and 4.9. It is the most critical for handling evolving rewards. It guarantees that small changes in the reward parameters φ (e.g., from reward shaping or entropy adjustment) lead to proportionally small changes in the expected reward. This allows the algorithm to track the evolving learning objective rather than being destabilized by it.

In essence, Proposition 4.10 ensures that the regularized reward function changes in a controlled and predictable manner as the policy and reward parameters evolve. This is crucial for analyzing non-stationary learning dynamics.

A detailed derivation of Proposition 4.10 can be found in Appendix D.

4.2. Main Results

With the assumptions above, we are ready to present our finite-time analysis of Algorithm 1. We measure the performance of Algorithm 1 using the following time-averaged errors over the second half of the T iterations:

$$\textbf{Actor Error: } G_T = \frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E} \|\nabla_\theta J_{\varphi_t}(\theta_t)\|_2^2$$

$$\textbf{Critic Error: } W_T = \frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E} \|\omega_t - \omega_t^*\|_2^2, \text{ where } \omega_t^* = \omega^*(\varphi_t, \theta_t)$$

$$\textbf{Reward Variation: } F_T = \frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2$$

The average squared gradient norm of the actor's objective over the second half of the training steps characterizes the convergence of the policy under the algorithm. Furthermore, as shown by (Agarwal et al., 2021; Mei et al., 2020; Ding et al., 2022), the stationarity of the policy often implies its optimality when the policy parameterization is specified. Thus, this metric can naturally serve as an indicator of the optimality gap in executing.

Theorem 4.11. *Consider Algorithm 1 with $\eta_t^\theta = \frac{c_\theta}{\sqrt{t}}$ and $\eta_t^\omega = \frac{c_\omega}{\sqrt{t}}$, where the ratio $\frac{c_\theta}{c_\omega}$ is chosen to be sufficiently small such that $\frac{c_\theta}{c_\omega} \leq \frac{\lambda}{4LS_\omega} \wedge \frac{1}{16LS_\omega}$. Under the above mentioned assumptions, the following bounds hold:*

$$\begin{aligned} G_T &= O\left(\frac{1}{\sqrt{T}}\right) + O\left(F_T \sqrt{T}\right) + O\left(\sqrt{\frac{F_T}{T}}\right) + O(\epsilon) \\ W_T &= O\left(\frac{1}{\sqrt{T}}\right) + O\left(F_T \sqrt{T}\right) + O\left(\sqrt{\frac{F_T}{T}}\right) + O(\epsilon) \end{aligned}$$

Interpretation of Theorem 4.11:

- **Static-Reward Case** ($F_T \equiv 0$): The algorithm achieves the canonical $O(1/\sqrt{T})$ convergence rate for both actor and critic, matching the best-known rate for single-timescale actor-critic methods under i.i.d. sampling (Chen et al., 2021; Olshevsky & Ghahserifard, 2023). Our analysis, by carefully handling the

Markovian sampling, improves upon previous works by eliminating a $\log^2 T$ factor compared to (Tian et al., 2023; Chen & Zhao, 2023; 2025).

- **Evolving-Reward Case** ($F_T > 0$): The convergence rate depends critically on the total variation of the reward parameters, F_T . For the errors to converge to zero asymptotically, we require $F_T = o(1/\sqrt{T})$. To preserve the $O(1/\sqrt{T})$ rate, we need the stronger condition $F_T = O(1/T)$. This means the reward function must change slowly enough for the actor-critic algorithm to track it effectively.

To further connect our theoretical findings with practical applications, we present the following key corollary, demonstrating that a common class of reward update rules satisfies this stringent condition. Roughly speaking, as long as the reward is updated at a rate comparable to the actor’s learning rate, stability is theoretically guaranteed.

Corollary 4.12. *If the reward parameter adopts a gradient-based update rule, i.e.*

$$\varphi_{t+1} \leftarrow \varphi_t + \eta_t^\varphi h_\varphi(t),$$

then given $\mathbb{E}\|h_\varphi(t)\|_2^2 \leq C_\varphi^2$ and $\eta_t^\varphi = \frac{c_\varphi}{\sqrt{t}}$ where C_φ and c_φ are constants, we have $F_K = O(1/T)$, and hence

$$G_T = O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon), \quad W_T = O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon).$$

Here, the step size for updating the reward parameter is of the same order as the actor’s and the critic’s, and the requirements on the gradient norm can be achieved by applying gradient clipping, a technique that is very common in practice. Therefore, Corollary 4.12 provides a solid theoretical foundation for a wide range of empirical practice of RL, including learning-based reward shaping methods (Pathak et al., 2017; Burda et al., 2019; Zheng et al., 2018; Hu et al., 2020; Mahankali et al., 2024; Ma et al., 2024) as well as soft actor-critic with automated entropy adjustment (Haarnoja et al., 2018b).

The proof of Corollary 4.12 can be found in Appendix D.

4.3. Proof Sketch of the Main Theorem

The proof of Theorem 4.11 proceeds in three interconnected steps, which we outline below. A rigorous proof is provided in Appendix C. The key innovations lie in (1) rigorously analyzing the impact of evolving rewards by establishing Lipschitz continuity properties of policy objective $J_\varphi(\theta)$ and optimal critic parameter $\omega^*(\varphi, \theta)$ w.r.t φ , and (2) providing a novel analysis on the distribution mismatch induced by Markovian sampling through deriving the following key proposition (refer to Appendix D for a detailed proof):

Proposition 4.13. *Following the Markovian sampling strategy described in Algorithm 1, we have*

$$\mathbb{E}\|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1 \leq LC_\delta L_\nu \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^\theta + \gamma^t \|\rho - \nu_\rho^{\pi_{\theta_0}}\|_1$$

for any $t \geq 0$, where C_δ and L_ν are constants (refer to Appendix B for a formal definition).

This analysis does not rely on the mixing time of the ergodic Markov chain. Instead, it directly utilizes the contraction properties of the induced operator acting on state distributions, which is a stronger property than the ergodicity, hence resulting in a tighter bound on the distribution mismatch.

Step 1: Bounding the Actor Error. The primary challenge introduced by an evolving reward is that the policy optimization objective $J_{\varphi_t}(\theta_t)$ changes at every time step t . To address this, we first show that $J_\varphi(\theta)$ is D_J -Lipschitz with respect to the reward parameter φ (Lemma B.1), which allows us to bound the change in the objective function by the change in φ :

$$\mathbb{E}[J_{\varphi_{t+1}}(\theta_{t+1})] - J_{\varphi_t}(\theta_t) \geq \mathbb{E}[J_{\varphi_t}(\theta_{t+1})] - J_{\varphi_t}(\theta_t) - D_J \mathbb{E}\|\varphi_{t+1} - \varphi_t\|_2$$

We then analyze the improvement in the objective for a fixed reward parameter. A Taylor expansion of $J_{\varphi_t}(\theta)$ around θ_t yields a bound on the squared policy gradient norm, $\|\nabla_\theta J_{\varphi_t}(\theta_t)\|_2^2$. This bound involves several error terms:

- **I_1 (Approximation Error):** This term arises from the bias introduced by linear function approximation and is bounded by $O(\epsilon)$.
- **I_2 (Critic Error):** This term captures the error from using an estimated critic ω_t instead of the optimal critic ω_t^* and is bounded by $O(\|\omega_t - \omega_t^*\|_2)$.
- **I_3 (Markovian Noise):** This term quantifies the error due to sampling states from the Markovian distribution $\hat{\nu}_t$ rather than the true stationary distribution $\nu_\rho^{\pi_{\theta_t}}$. Proposition 4.13 provides a tighter bound on this distribution mismatch, which is crucial for improving the overall convergence rate.

After summing over iterations and applying a telescoping series, we obtain the following inequality for the actor error (Theorem C.1):

$$(1 - \gamma)G_T \leq 2L\sqrt{G_T W_T} + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon) \quad (15)$$

Step 2: Bounding the Critic Error. The critic update must track a moving target: the optimal parameter $\omega_t^* = \omega^*(\varphi_t, \theta_t)$ changes with both the policy parameter θ_t and the reward parameter φ_t . We analyze the evolution of the critic error $\|\omega_{t+1} - \omega_{t+1}^*\|_2^2$. A central decomposition yields:

$$\begin{aligned} \mathbb{E} \|\omega_{t+1} - \omega_{t+1}^*\|_2^2 &\leq \|\omega_t - \omega_t^*\|_2^2 \\ &\quad + 2\mathbb{E} \left\| \hat{\delta}(s_t, a_t, s'_t) \phi(s_t) \right\|_2^2 \\ &\quad + 2\mathbb{E} \|\omega_t^* - \omega_{t+1}^*\|_2^2 \\ &\quad + 2\eta_t \mathbb{E} \left\langle \omega_t - \omega_t^*, \hat{\delta}(s_t, a_t, s'_t) \phi(s_t) \right\rangle \\ &\quad + 2\mathbb{E} \left\langle \omega_t - \omega_t^*, \omega_t^* - \omega_{t+1}^* \right\rangle, \end{aligned}$$

where $\mathbb{E} \left\langle \omega_t - \omega_t^*, \hat{\delta}(s_t, a_t, s'_t) \phi(s_t) \right\rangle$ is further decomposed into three components:

- J_1 : This term is zero by the definition of ω_t^* .
- $J_2(\text{Contraction})$: This term provides a negative contribution $-\lambda \|\omega_t - \omega_t^*\|_2^2$, ensuring the critic error contracts towards zero.
- $J_3(\text{Markovian Noise})$: Similar to I_3 in the actor analysis, this term is bounded using Proposition 4.13.

The critical difference from the static-reward case is the presence of terms involving $\omega_t^* - \omega_{t+1}^*$. We bound these by establishing the Lipschitz continuity of ω^* with respect to both θ and φ (Lemma B.5). This introduces terms proportional to $\mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2$ and $\mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2$ into the bound. After summation, we derive the following inequality for the critic error (Theorem C.2):

$$\begin{aligned} \frac{1}{1-\gamma} W_T &\leq 2S\omega \frac{c\theta}{c\omega} \sqrt{G_T W_T} + O(F_T \sqrt{T}) \\ &\quad + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon) \quad (16) \end{aligned}$$

Step 3: Solving the System of Inequalities Steps 1 and 2 result in a system of two inequalities (15 and 16) that couple the actor error G_T and the critic error W_T . To solve this system, we use the algebraic inequality

$$2\sqrt{G_T W_T} \leq \frac{1-\gamma}{2L} G_T + \frac{2L}{1-\gamma} W_T.$$

By substituting this into the inequalities and choosing the step-size ratio $\frac{c\theta}{c\omega}$ to be sufficiently small, we can decouple the two errors and obtain a bound of $O(1/\sqrt{T}) + O(F_T \sqrt{T}) + O(\sqrt{F_T/T}) + O(\epsilon)$ for both G_T and W_T , thus completing the proof.

5. Conclusion

In this work, we have undertaken a systematic theoretical investigation of actor-critic methods in the presence of evolving rewards—a setting that mirrors the reality of many practical RL algorithms but has been largely overlooked by theoretical analyses. We formulated the problem, established necessary assumptions, and provided the first finite-time convergence guarantees for a single-timescale actor-critic algorithm under Markovian sampling.

Our analysis demonstrates that the single-timescale actor-critic algorithm is remarkably robust to reward non-stationarity. The canonical $O(1/\sqrt{T})$ convergence rate can be maintained for both the actor and critic, provided the reward parameters evolve at a controlled pace. A key corollary confirms that gradient-based reward updates—a common pattern in algorithms that learn intrinsic rewards or adapt regularization strengths—satisfy this condition, thereby providing a solid theoretical foundation for their empirical success. Furthermore, our novel technique for bounding distribution mismatch under Markovian sampling yields a tighter analysis, improving upon prior rates by a factor of $\log^2 T$ even when the reward is static.

This work opens several avenues for future research. Extending the analysis to nonlinear function approximation, particularly with neural networks, is a critical next step. Furthermore, exploring the implications of our theoretical findings for the design of more effective and provably stable reward-shaping algorithms presents an exciting direction for both theoretical and applied work. Finally, this analysis lays a foundational stone for a deeper theoretical understanding of reinforcement learning with dynamic objectives due to evolving reward, shifting initial distribution or transition probabilities.

Acknowledgements

This work was supported by the National Natural Science Foundation of China Grant 52494974.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. There are many potential societal consequences of our work, none of which we feel must be specifically highlighted here.

References

- Agarwal, A., Kakade, S. M., Lee, J. D., and Mahajan, G. On the theory of policy gradient methods: Optimality, approximation, and distribution shift. *J. Mach. Learn. Res.*, 22:98:1–98:76, 2021. URL <https://jmlr.org>

- [rg/papers/v22/19-736.html](http://papers.v22/19-736.html).
- Ahmed, Z., Roux, N. L., Norouzi, M., and Schuurmans, D. Understanding the impact of entropy on policy optimization. In Chaudhuri, K. and Salakhutdinov, R. (eds.), *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA*, volume 97 of *Proceedings of Machine Learning Research*, pp. 151–160. PMLR, 2019. URL <http://proceedings.mlr.press/v97/ahmed19a.html>.
- Asmuth, J., Littman, M. L., and Zinkov, R. Potential-based shaping in model-based reinforcement learning. In Fox, D. and Gomes, C. P. (eds.), *Proceedings of the Twenty-Third AAAI Conference on Artificial Intelligence, AAAI 2008, Chicago, Illinois, USA, July 13-17, 2008*, pp. 604–609. AAAI Press, 2008. URL <http://www.aaai.org/Library/AAAI/2008/aaai08-096.php>.
- Bhandari, J., Russo, D., and Singal, R. A finite time analysis of temporal difference learning with linear function approximation. In Bubeck, S., Perchet, V., and Rigollet, P. (eds.), *Conference On Learning Theory, COLT 2018, Stockholm, Sweden, 6-9 July 2018*, volume 75 of *Proceedings of Machine Learning Research*, pp. 1691–1692. PMLR, 2018. URL <http://proceedings.mlr.press/v75/bhandari18a.html>.
- Burda, Y., Edwards, H., Storkey, A. J., and Klimov, O. Exploration by random network distillation. In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*. OpenReview.net, 2019. URL <https://openreview.net/forum?id=H1lJjN5Ym>.
- Cayci, S., He, N., and Srikant, R. Finite-time analysis of entropy-regularized neural natural actor-critic algorithm. *Trans. Mach. Learn. Res.*, 2024, 2024. URL <https://openreview.net/forum?id=BkEqk7pS1I>.
- Chen, T., Sun, Y., and Yin, W. Closing the gap: Tighter analysis of alternating stochastic gradient methods for bilevel problems. In Ranzato, M., Beygelzimer, A., Dauphin, Y. N., Liang, P., and Vaughan, J. W. (eds.), *Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*, pp. 25294–25307, 2021. URL <https://papers.nips.cc/paper/2021/hash/d4dd111a4fd973394238aca5c05bebe3-Abstract.html>.
- Chen, X. and Zhao, L. Finite-time analysis of single-timescale actor-critic. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*, 2023. URL http://papers.nips.cc/paper_files/paper/2023/hash/160adf2dc118a920e7858484b92a37d8-Abstract-Conference.html.
- Chen, X. and Zhao, L. On the convergence of continuous single-timescale actor-critic. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=pV7hSmGJXP>.
- Cheung, W. C., Simchi-Levi, D., and Zhu, R. Nonstationary reinforcement learning: The blessing of (more) optimism. *Manage. Sci.*, 69(10):5722–5739, October 2023. ISSN 0025-1909. doi: 10.1287/mnsc.2023.4704. URL <https://doi.org/10.1287/mnsc.2023.4704>.
- DeepSeek-AI, Guo, D., Yang, D., Zhang, H., Song, J., Zhang, R., Xu, R., Zhu, Q., Ma, S., Wang, P., Bi, X., Zhang, X., Yu, X., Wu, Y., Wu, Z. F., Gou, Z., Shao, Z., Li, Z., Gao, Z., Liu, A., Xue, B., Wang, B., Wu, B., Feng, B., Lu, C., Zhao, C., Deng, C., Zhang, C., Ruan, C., Dai, D., Chen, D., Ji, D., Li, E., Lin, F., Dai, F., Luo, F., Hao, G., Chen, G., Li, G., Zhang, H., Bao, H., Xu, H., Wang, H., Ding, H., Xin, H., Gao, H., Qu, H., Li, H., Guo, J., Li, J., Wang, J., Chen, J., Yuan, J., Qiu, J., Li, J., Cai, J. L., Ni, J., Liang, J., Chen, J., Dong, K., Hu, K., Gao, K., Guan, K., Huang, K., Yu, K., Wang, L., Zhang, L., Zhao, L., Wang, L., Zhang, L., Xu, L., Xia, L., Zhang, M., Zhang, M., Tang, M., Li, M., Wang, M., Li, M., Tian, N., Huang, P., Zhang, P., Wang, Q., Chen, Q., Du, Q., Ge, R., Zhang, R., Pan, R., Wang, R., Chen, R. J., Jin, R. L., Chen, R., Lu, S., Zhou, S., Chen, S., Ye, S., Wang, S., Yu, S., Zhou, S., Pan, S., and Li, S. S. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *CoRR*, abs/2501.12948, 2025. doi: 10.48550/ARXIV.2501.12948. URL <https://doi.org/10.48550/arXiv.2501.12948>.
- Devlin, S. and Kudenko, D. Dynamic potential-based reward shaping. In *Proceedings of the 11th International Conference on Autonomous Agents and Multiagent Systems - Volume 1, AAMAS '12*, pp. 433–440, Richland, SC, 2012. International Foundation for Autonomous Agents and Multiagent Systems. ISBN 0981738117.
- Ding, Y., Zhang, J., and Lavaei, J. On the global optimum convergence of momentum-based policy gradient. In Camps-Valls, G., Ruiz, F. J. R., and Valera, I. (eds.), *International Conference on Artificial Intelligence and Statistics, AISTATS 2022, 28-30 March 2022, Virtual Event*, volume 151 of *Proceedings of Machine Learning Research*, pp. 1910–1934. PMLR, 2022. URL <https://proceedings.mlr.press/v151/ding22a.html>.

- Ding, Y., Zhang, J., Lee, H., and Lavaei, J. Beyond exact gradients: Convergence of stochastic soft-max policy gradient methods with entropy regularization. *IEEE Trans. Autom. Control.*, 70(8):5129–5144, 2025. doi: 10.1109/TAC.2025.3540965. URL <https://doi.org/10.1109/TAC.2025.3540965>.
- Eshwar, S., Thoppe, G., Barua, A., Gopalan, A., and Dalal, G. Monotone and conservative policy iteration beyond the tabular case. In *The 29th International Conference on Artificial Intelligence and Statistics*, 2026. URL <https://openreview.net/forum?id=TPSX51x8oR>.
- Even-Dar, E., Kakade, S. M., and Mansour, Y. Online markov decision processes. *Mathematics of Operations Research*, 34(3):726–736, 2009. doi: 10.1287/moor.1090.0396. URL <https://doi.org/10.1287/moor.1090.0396>.
- Fan, Y., Watkins, O., Du, Y., Liu, H., Ryu, M., Boutilier, C., Abbeel, P., Ghavamzadeh, M., Lee, K., and Lee, K. Reinforcement learning for fine-tuning text-to-image diffusion models. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*, 2023. URL http://papers.nips.cc/paper_files/paper/2023/hash/fc65fab891d83433bd3c8d966edde311-Abstract-Conference.html.
- Fatkhullin, I., Barakat, A., Kireeva, A., and He, N. Stochastic policy gradient methods: Improved sample complexity for fisher-non-degenerate policies. In Krause, A., Brunskill, E., Cho, K., Engelhardt, B., Sabato, S., and Scarlett, J. (eds.), *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pp. 9827–9869. PMLR, 2023. URL <https://proceedings.mlr.press/v202/fatkhullin23a.html>.
- Feng, S., Yin, M., Huang, R., Wang, Y., Yang, J., and Liang, Y. Non-stationary reinforcement learning under general function approximation. In Krause, A., Brunskill, E., Cho, K., Engelhardt, B., Sabato, S., and Scarlett, J. (eds.), *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pp. 9976–10007. PMLR, 2023. URL <https://proceedings.mlr.press/v202/feng23e.html>.
- Ganesh, S., Chen, J., Mondal, W. U., and Aggarwal, V. Order-optimal global convergence for actor-critic with general policy and neural critic parametrization. In *The 41st Conference on Uncertainty in Artificial Intelligence*, 2025. URL <https://openreview.net/forum?id=HPxE1IejA5>.
- Gaur, M., Bedi, A., Wang, D., and Aggarwal, V. Closing the gap: Achieving global convergence (Last iterate) of actor-critic under Markovian sampling with neural network parametrization. In Salakhutdinov, R., Kolter, Z., Heller, K., Weller, A., Oliver, N., Scarlett, J., and Berkenkamp, F. (eds.), *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pp. 15153–15179. PMLR, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/gaur24a.html>.
- Gupta, D., Chandak, Y., Jordan, S. M., Thomas, P. S., and da Silva, B. C. Behavior alignment via reward function optimization. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*, 2023. URL http://papers.nips.cc/paper_files/paper/2023/hash/a5357781c204d4412e44ed9cbcd08d5-Abstract-Conference.html.
- Haarnoja, T., Tang, H., Abbeel, P., and Levine, S. Reinforcement learning with deep energy-based policies. In Precup, D. and Teh, Y. W. (eds.), *Proceedings of the 34th International Conference on Machine Learning, ICML 2017, Sydney, NSW, Australia, 6-11 August 2017*, volume 70 of *Proceedings of Machine Learning Research*, pp. 1352–1361. PMLR, 2017. URL <http://proceedings.mlr.press/v70/haarnoja17a.html>.
- Haarnoja, T., Zhou, A., Abbeel, P., and Levine, S. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In Dy, J. G. and Krause, A. (eds.), *Proceedings of the 35th International Conference on Machine Learning, ICML 2018, Stockholm, Sweden, July 10-15, 2018*, volume 80 of *Proceedings of Machine Learning Research*, pp. 1856–1865. PMLR, 2018a. URL <http://proceedings.mlr.press/v80/haarnoja18b.html>.
- Haarnoja, T., Zhou, A., Hartikainen, K., Tucker, G., Ha, S., Tan, J., Kumar, V., Zhu, H., Gupta, A., Abbeel, P., and Levine, S. Soft actor-critic algorithms and applications. *CoRR*, abs/1812.05905, 2018b. URL <http://arxiv.org/abs/1812.05905>.
- Hong, M., Wai, H.-T., Wang, Z., and Yang, Z. 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. doi: 10.1137/20M1387341. URL <https://doi.org/10.1137/20M1387341>.
- Hu, Y., Wang, W., Jia, H., Wang, Y., Chen, Y., Hao, J., Wu, F., and Fan, C. Learning to utilize shaping rewards: A new approach of reward shaping. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/b710915795b9e9c02cf10d6d2bdb688c-Abstract.html>.
- Jaques, N., Ghandeharioun, A., Shen, J. H., Ferguson, C., Lapedriza, À., Jones, N., Gu, S., and Picard, R. W. Way off-policy batch deep reinforcement learning of implicit human preferences in dialog. *CoRR*, abs/1907.00456, 2019. URL <http://arxiv.org/abs/1907.00456>.
- Jin, C., Jin, T., Luo, H., Sra, S., and Yu, T. Learning adversarial Markov decision processes with bandit feedback and unknown transition. In III, H. D. and Singh, A. (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 4860–4869. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/jin20c.html>.
- Konda, V. R. and Tsitsiklis, J. N. Actor-critic algorithms. In Solla, S. A., Leen, T. K., and Müller, K. (eds.), *Advances in Neural Information Processing Systems 12, [NIPS Conference, Denver, Colorado, USA, November 29 - December 4, 1999]*, pp. 1008–1014. The MIT Press, 1999. URL <http://papers.nips.cc/paper/1786-actor-critic-algorithms>.
- Kumar, H., Koppel, A., and Ribeiro, A. On the sample complexity of actor-critic method for reinforcement learning with function approximation. *Mach. Learn.*, 112(7):2433–2467, 2023. doi: 10.1007/S10994-023-06303-2. URL <https://doi.org/10.1007/s10994-023-06303-2>.
- Kumar, N., Agrawal, P., Ramponi, G., Levy, K. Y., and Mannor, S. On the convergence of single-timescale actor-critic. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2026. URL <https://openreview.net/forum?id=OixkIljSZD>.
- Lillicrap, T. P., Hunt, J. J., Pritzel, A., Heess, N., Erez, T., Tassa, Y., Silver, D., and Wierstra, D. Continuous control with deep reinforcement learning. In Bengio, Y. and LeCun, Y. (eds.), *4th International Conference on Learning Representations, ICLR 2016, San Juan, Puerto Rico, May 2-4, 2016, Conference Track Proceedings*, 2016. URL <http://arxiv.org/abs/1509.02971>.
- Liu, Y., Zhang, K., Basar, T., and Yin, W. An improved analysis of (variance-reduced) policy gradient and natural policy gradient methods. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/56577889b3c1cd083b6d7b32d32f99d5-Abstract.html>.
- Lobel, S., Bagaria, A., and Konidaris, G. Flipping coins to estimate pseudocounts for exploration in reinforcement learning. In Krause, A., Brunskill, E., Cho, K., Engelhardt, B., Sabato, S., and Scarlett, J. (eds.), *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pp. 22594–22613. PMLR, 2023. URL <https://proceedings.mlr.press/v202/lobel23a.html>.
- Ma, H., Sima, K., Vo, T. V., Fu, D., and Leong, T. Reward shaping for reinforcement learning with an assistant reward agent. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net, 2024. URL <https://openreview.net/forum?id=a3XFF0PGLU>.
- Machado, M. C., Bellemare, M. G., and Bowling, M. Count-based exploration with the successor representation. In *The Thirty-Fourth AAAI Conference on Artificial Intelligence, AAAI 2020, New York, NY, USA, February 7-12, 2020*, pp. 5125–5133. AAAI Press, 2020. doi: 10.1609/AAAI.V34I04.5955. URL <https://doi.org/10.1609/aaai.v34i04.5955>.
- Mahankali, S., Hong, Z., Sekhari, A., Rakhlin, A., and Agrawal, P. Random latent exploration for deep reinforcement learning. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net, 2024. URL <https://openreview.net/forum?id=Y9qzwNlKVU>.
- Mandal, D. and Radanovic, G. Performative reinforcement learning with linear markov decision process. In Li, Y., Mandt, S., Agrawal, S., and Khan, E. (eds.), *Proceedings of The 28th International Conference on Artificial Intelligence and Statistics*, volume 258 of *Proceedings of Machine Learning Research*, pp. 3232–3240. PMLR, 03–05 May 2025. URL <https://proceedings.mlr.press/v258/mandal25a.html>.
- Mandal, D., Triantafyllou, S., and Radanovic, G. Performative reinforcement learning. In Krause, A., Brunskill, E.,

- Cho, K., Engelhardt, B., Sabato, S., and Scarlett, J. (eds.), *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pp. 23642–23680. PMLR, 2023. URL <https://proceedings.mlr.press/v202/mandal23a.html>.
- Mao, W., Zhang, K., Zhu, R., Simchi-Levi, D., and Basar, T. Model-free nonstationary reinforcement learning: Near-optimal regret and applications in multiagent reinforcement learning and inventory control. *Manag. Sci.*, 71(2): 1564–1580, 2025. doi: 10.1287/MNSC.2022.02533. URL <https://doi.org/10.1287/mnsc.2022.02533>.
- Martin, J., Sasikumar, S. N., Everitt, T., and Hutter, M. Count-based exploration in feature space for reinforcement learning. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence, IJCAI’17*, pp. 2471–2478. AAAI Press, 2017. ISBN 9780999241103.
- Mei, J., Xiao, C., Szepesvári, C., and Schuurmans, D. On the global convergence rates of softmax policy gradient methods. In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13-18 July 2020, Virtual Event*, volume 119 of *Proceedings of Machine Learning Research*, pp. 6820–6829. PMLR, 2020. URL <http://proceedings.mlr.press/v119/mei20b.html>.
- Memarian, F., Goo, W., Lioutikov, R., Niekum, S., and Topcu, U. Self-supervised online reward shaping in sparse-reward environments. In *IEEE/RSJ International Conference on Intelligent Robots and Systems, IROS 2021, Prague, Czech Republic, September 27 - Oct. 1, 2021*, pp. 2369–2375. IEEE, 2021. doi: 10.1109/IROS51168.2021.9636020. URL <https://doi.org/10.1109/IROS51168.2021.9636020>.
- Mguni, D., Jafferjee, T., Wang, J., Nieves, N. P., Song, W., Tong, F., Taylor, M. E., Yang, T., Dai, Z., Chen, H., Zhu, J., Shao, K., Wang, J., and Yang, Y. Learning to shape rewards using a game of two partners. In Williams, B., Chen, Y., and Neville, J. (eds.), *Thirty-Seventh AAAI Conference on Artificial Intelligence, AAAI 2023, Washington, DC, USA, February 7-14, 2023*, pp. 11604–11612. AAAI Press, 2023. doi: 10.1609/AAAI.V37I10.26371. URL <https://doi.org/10.1609/aaai.v37i10.26371>.
- Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., and Riedmiller, M. A. Playing atari with deep reinforcement learning. *CoRR*, abs/1312.5602, 2013. URL <http://arxiv.org/abs/1312.5602>.
- Mondal, W. U. and Aggarwal, V. Improved sample complexity analysis of natural policy gradient algorithm with general parameterization for infinite horizon discounted reward markov decision processes. In Dasgupta, S., Mandt, S., and Li, Y. (eds.), *International Conference on Artificial Intelligence and Statistics, 2-4 May 2024, Palau de Congressos, Valencia, Spain*, volume 238 of *Proceedings of Machine Learning Research*, pp. 3097–3105. PMLR, 2024. URL <https://proceedings.mlr.press/v238/u-mondal24a.html>.
- Narvekar, S., Peng, B., Leonetti, M., Sinapov, J., Taylor, M. E., and Stone, P. Curriculum learning for reinforcement learning domains: A framework and survey. *J. Mach. Learn. Res.*, 21:181:1–181:50, 2020. URL <http://jmlr.org/papers/v21/20-212.html>.
- Ng, A. Y., Harada, D., and Russell, S. J. Policy invariance under reward transformations: Theory and application to reward shaping. In *Proceedings of the Sixteenth International Conference on Machine Learning, ICML ’99*, pp. 278–287, San Francisco, CA, USA, 1999. Morgan Kaufmann Publishers Inc. ISBN 1558606122.
- Olshevsky, A. and Gharesifard, B. A small gain analysis of single timescale actor critic. *SIAM J. Control. Optim.*, 61(2):980–1007, 2023. doi: 10.1137/22M1483335. URL <https://doi.org/10.1137/22m1483335>.
- Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C. L., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., Schulman, J., Hilton, J., Kelton, F., Miller, L., Simens, M., Askell, A., Welinder, P., Christiano, P. F., Leike, J., and Lowe, R. Training language models to follow instructions with human feedback. In Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds.), *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, 2022. URL http://papers.nips.cc/paper_files/paper/2022/hash/blefde53be364a73914f58805a001731-Abstract-Conference.html.
- Pathak, D., Agrawal, P., Efros, A. A., and Darrell, T. Curiosity-driven exploration by self-supervised prediction. In Precup, D. and Teh, Y. W. (eds.), *Proceedings of the 34th International Conference on Machine Learning, ICML 2017, Sydney, NSW, Australia, 6-11 August 2017*, volume 70 of *Proceedings of Machine Learning Research*, pp. 2778–2787. PMLR, 2017. URL <http://proceedings.mlr.press/v70/pathak17a.html>.
- Pathak, D., Gandhi, D., and Gupta, A. Self-supervised exploration via disagreement. In Chaudhuri, K. and Salakhutdinov, R. (eds.), *Proceedings of the 36th International*

- Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA, volume 97 of *Proceedings of Machine Learning Research*, pp. 5062–5071. PMLR, 2019. URL <http://proceedings.mlr.press/v97/pathak19a.html>.
- Ramesh, A. A., Kirsch, L., van Steenkiste, S., and Schmidhuber, J. Exploring through random curiosity with general value functions. In Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds.), *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, 2022. URL http://papers.nips.cc/paper_files/paper/2022/hash/76e57c3c6b3e06f332a4832ddd6a9a12-Abstract-Conference.html.
- Rank, B., Triantafyllou, S., Mandal, D., and Radanovic, G. Performative reinforcement learning in gradually shifting environments. In Kiyavash, N. and Mooij, J. M. (eds.), *Proceedings of the Fortieth Conference on Uncertainty in Artificial Intelligence*, volume 244 of *Proceedings of Machine Learning Research*, pp. 3041–3075. PMLR, 15–19 Jul 2024. URL <https://proceedings.mlr.press/v244/rank24a.html>.
- Shao, Z., Wang, P., Zhu, Q., Xu, R., Song, J., Zhang, M., Li, Y. K., Wu, Y., and Guo, D. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *CoRR*, abs/2402.03300, 2024. doi: 10.48550/ARXIV.2402.03300. URL <https://doi.org/10.48550/arXiv.2402.03300>.
- Shen, H., Zhang, K., Hong, M., and Chen, T. Towards understanding asynchronous advantage actor-critic: Convergence and linear speedup. *Trans. Sig. Proc.*, 71: 2579–2594, January 2023. ISSN 1053-587X. doi: 10.1109/TSP.2023.3268475. URL <https://doi.org/10.1109/TSP.2023.3268475>.
- Stadie, B. C., Zhang, L., and Ba, J. Learning intrinsic rewards as a bi-level optimization problem. In Adams, R. P. and Gogate, V. (eds.), *Proceedings of the Thirty-Sixth Conference on Uncertainty in Artificial Intelligence, UAI 2020, virtual online, August 3-6, 2020*, volume 124 of *Proceedings of Machine Learning Research*, pp. 111–120. AUAI Press, 2020. URL <http://proceedings.mlr.press/v124/stadie20a.html>.
- Stiennon, N., Ouyang, L., Wu, J., Ziegler, D. M., Lowe, R., Voss, C., Radford, A., Amodei, D., and Christiano, P. F. Learning to summarize from human feedback. *CoRR*, abs/2009.01325, 2020. URL <https://arxiv.org/abs/2009.01325>.
- Sun, H., Han, L., Yang, R., Ma, X., Guo, J., and Zhou, B. Exploit reward shifting in value-based deep-rl: Optimistic curiosity-based exploration and conservative exploitation via linear reward shaping. In Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds.), *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, 2022. URL http://papers.nips.cc/paper_files/paper/2022/hash/f600d1a3f6a63f782680031f3ce241a7-Abstract-Conference.html.
- Sutton, R. S., Barto, A. G., et al. *Reinforcement learning: An introduction*, volume 1. MIT press Cambridge, 1998.
- Sutton, R. S., McAllester, D., Singh, S., and Mansour, Y. Policy gradient methods for reinforcement learning with function approximation. In Solla, S., Leen, T., and Müller, K. (eds.), *Advances in Neural Information Processing Systems*, volume 12. MIT Press, 1999. URL https://proceedings.neurips.cc/paper_files/paper/1999/file/464d828b85b0bed98e80ade0a5c43b0f-Paper.pdf.
- Tian, H., Olshevsky, A., and Paschalidis, Y. Convergence of actor-critic with multi-layer neural networks. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems*, volume 36, pp. 9279–9321. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/1dc9fbdb6b4d9955ad377cb983232c9f-Paper-Conference.pdf.
- Trott, A., Zheng, S., Xiong, C., and Socher, R. Keeping your distance: Solving sparse reward tasks using self-balancing shaped rewards. In Wallach, H. M., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E. B., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems 2019, NeurIPS 2019, December 8-14, 2019, Vancouver, BC, Canada*, pp. 10376–10386, 2019. URL <https://proceedings.neurips.cc/paper/2019/hash/64c26b2a2dcf068c49894bd07e0e6389-Abstract.html>.
- Wiewiora, E. Potential-based shaping and q-value initialization are equivalent. *J. Artif. Int. Res.*, 19(1):205–208, September 2003. ISSN 1076-9757.
- Williams, R. J. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Mach. Learn.*, 8:229–256, 1992. doi: 10.1007/BF00992696. URL <https://doi.org/10.1007/BF00992696>.

- Wu, Y., Zhang, W., Xu, P., and Gu, Q. A finite-time analysis of two time-scale actor-critic methods. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/cc9b3c69b56df284846bf2432f1cba90-Abstract.html>.
- Xiao, L. On the convergence rates of policy gradient methods. *J. Mach. Learn. Res.*, 23:282:1–282:36, 2022. URL <https://jmlr.org/papers/v23/22-0056.html>.
- Xu, T., Wang, Z., and Liang, Y. Improving sample complexity bounds for (natural) actor-critic algorithms. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020a. URL <https://proceedings.neurips.cc/paper/2020/hash/2e1b24a664f5e9c18f407b2f9c73e821-Abstract.html>.
- Xu, T., Wang, Z., and Liang, Y. Non-asymptotic convergence analysis of two time-scale (natural) actor-critic algorithms. *CoRR*, abs/2005.03557, 2020b. URL <https://arxiv.org/abs/2005.03557>.
- Yang, K., Tao, J., Lyu, J., and Li, X. Exploration and anti-exploration with distributional random network distillation. In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net, 2024. URL <https://openreview.net/forum?id=rIrpzmqRBk>.
- Yang, Z., Chen, Y., Hong, M., and Wang, Z. Provably global convergence of actor-critic: A case for linear quadratic regulator with ergodic cost. In Wallach, H. M., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E. B., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems 2019, NeurIPS 2019, December 8-14, 2019, Vancouver, BC, Canada*, pp. 8351–8363, 2019. URL <https://proceedings.neurips.cc/paper/2019/hash/9713faa264b94e2bf346a1bb52587fd8-Abstract.html>.
- Zheng, Z., Oh, J., and Singh, S. On learning intrinsic rewards for policy gradient methods. In Bengio, S., Wallach, H. M., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 31: Annual Conference on Neural Information Processing Systems 2018, NeurIPS 2018, December 3-8, 2018, Montréal, Canada*, pp. 4649–4659, 2018. URL <https://proceedings.neurips.cc/paper/2018/hash/51de85ddd068f0bc787691d356176df9-Abstract.html>.
- Zimin, A. and Neu, G. Online learning in episodic markovian decision processes by relative entropy policy search. In Burges, C., Bottou, L., Welling, M., Ghahramani, Z., and Weinberger, K. (eds.), *Advances in Neural Information Processing Systems*, volume 26. Curran Associates, Inc., 2013. URL https://proceedings.neurips.cc/paper_files/paper/2013/file/68053af2923e00204c3ca7c6a3150cf7-Paper.pdf.

A. Literature Review of Evolving Reward Techniques

The idea of modifying rewards to improve learning is long-standing. Potential-based reward shaping, introduced and developed by Ng et al. (1999), Wiewiora (2003), Asmuth et al. (2008), and Devlin & Kudenko (2012), defines the shaped reward as $\gamma\Phi(s') - \Phi(s)$ where $\Phi(\cdot)$ is a potential function to guarantee policy invariance. More recent work, however, modifies the reward to balance the exploration and exploitation behavior of the policy. This sorts of work include randomly perturbing the reward function (Mahankali et al., 2024), various design of explicit exploration bonus like count-based methods (Martin et al., 2017; Machado et al., 2020; Lobel et al., 2023), curiosity-driven methods (Pathak et al., 2017; 2019; Ramesh et al., 2022; Sun et al., 2022) and random network distillation (Burda et al., 2019; Yang et al., 2024), as well as fully self-supervised intrinsic rewards (Zheng et al., 2018; Stadie et al., 2020; Memarian et al., 2021; Mguni et al., 2023; Ma et al., 2024) or incorporation of prior knowledge (Trott et al., 2019; Hu et al., 2020; Gupta et al., 2023) that enhance the performance of the resulting policy in terms of the original reward.

Another series of work that result in evolving rewards is the entropy or KL regularization. Entropy regularization (Haarnoja et al., 2018a;b; Ahmed et al., 2019) is commonly used technique to encourage exploration and avoid near-deterministic suboptimal policy. KL regularization is common in fine-tuning RL policies, especially in training Large Language Models (Ouyang et al., 2022; Shao et al., 2024) and Diffusion Models (Fan et al., 2023). These methods result in evolving reward because the regularization term $-\alpha\mathcal{H}(\pi(\cdot|s))$ (or $-\alpha d_{\text{KL}}(\pi(\cdot|s)||\pi_{\text{ref}}(\cdot|s))$) is equivalent to a penalty term $-\alpha \log \pi(a|s)$ (or $-\alpha \log \frac{\pi(a|s)}{\pi_{\text{ref}}(a|s)}$) added to the reward function $r(s, a)$. Hence, the reward function will change because of the under-training policy $\pi(\cdot|s)$, the adaptive regularization factor α , or the change of the reference policy π_{ref} .

Besides, curriculum learning (Narvekar et al., 2020) is also closely related, as it inherently involve a sequence of evolving learning objectives (and thus rewards).

B. Preliminary Lemmas

Lemma B.1. *There exist constants C_J , L_J , S_J and D_J such that for any $\theta \in \Omega(\theta)$, $s \in \mathcal{S}$, $\tilde{V}_{\varphi}^{\pi_{\theta}}(s)$ is C_J -bounded, L_J -Lipschitz and S_J -smooth w.r.t. θ , and D_J -Lipschitz w.r.t. φ , where $C_J = O((1 - \gamma)^{-1})$, $L_J = O((1 - \gamma)^{-2})$, $S_J = O((1 - \gamma)^{-3})$, $D_J = O((1 - \gamma)^{-1})$.*

Corollary B.2. *There exist constants L_{ν} and S_{ν} such that for any $\theta \in \Omega(\theta)$, $\nu_{\rho}^{\pi_{\theta}}(\cdot)$ is L_{ν} -Lipschitz and S_{ν} -smooth w.r.t. θ in terms of $\|\cdot\|_1$, where $L_{\nu} = O((1 - \gamma)^{-1})$, $S_{\nu} = O((1 - \gamma)^{-2})$.*

Lemma B.3. *There exist constants L_A and S_A such that A_{θ} is L_A -Lipschitz and S_A -smooth w.r.t. θ , where $L_A = O((1 - \gamma)^{-1})$, $S_A = O((1 - \gamma)^{-2})$.*

Lemma B.4. *There exist constants C_b , L_b , S_b and D_b such that $b_{\theta, \varphi}$ is C_b -bounded, L_b -Lipschitz and S_b -smooth w.r.t. θ and D_b -Lipschitz w.r.t. φ , where $C_b = O(1)$, $L_b = O((1 - \gamma)^{-1})$, $S_b = O((1 - \gamma)^{-2})$, $D_b = O(1)$.*

Lemma B.5. *There exist constants C_{ω} , L_{ω} and S_{ω} and D_{ω} such that $\omega^*(\varphi, \theta)$ is C_{ω} -bounded, L_{ω} -Lipschitz and S_{ω} -smooth w.r.t. θ and D_{ω} -Lipschitz w.r.t. φ , where $C_{\omega} = O(\lambda^{-1})$, $L_{\omega} = O((1 - \gamma)^{-1}\lambda^{-2})$, $S_{\omega} = O((1 - \gamma)^{-2}\lambda^{-3})$, $D_{\omega} = O(\lambda^{-1})$.*

Lemma B.6. *There exists a constant C_{δ} such that for any $\theta \in \Omega(\theta)$, $\varphi \in \Omega(\varphi)$ and $\|\omega\|_2 \leq C_{\omega}$,*

$$\mathbb{E}_{\nu, \pi_{\theta}, \mathcal{P}}[\hat{\delta}(s, a, s')^2] \leq C_{\delta}^2,$$

where $C_{\delta} = O(\lambda^{-1})$.

C. Proof of Main Theorem

C.1. Step 1: Bounding the Actor Error

Theorem C.1 (Actor Update). *Tate $\eta_t^{\theta} = \frac{c_{\theta}}{\sqrt{t}}$ where c_{θ} is a constant, then*

$$G_T \leq \frac{2L}{1 - \gamma} \sqrt{G_T W_T} + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon).$$

Proof. We first bound the change of the objective function by the change of the reward parameter:

$$\begin{aligned}\mathbb{E}[J_{\varphi_{t+1}}(\boldsymbol{\theta}_{t+1})] - J_{\varphi_t}(\boldsymbol{\theta}_t) &= \mathbb{E}[J_{\varphi_{t+1}}(\boldsymbol{\theta}_{t+1}) - J_{\varphi_t}(\boldsymbol{\theta}_{t+1})] + \mathbb{E}[J_{\varphi_t}(\boldsymbol{\theta}_{t+1})] - J_{\varphi_t}(\boldsymbol{\theta}_t) \\ &\geq \mathbb{E}[-|J_{\varphi_{t+1}}(\boldsymbol{\theta}_{t+1}) - J_{\varphi_t}(\boldsymbol{\theta}_{t+1})|] + \mathbb{E}[J_{\varphi_t}(\boldsymbol{\theta}_{t+1})] - J_{\varphi_t}(\boldsymbol{\theta}_t) \\ &\geq -D_J \mathbb{E}\|\varphi_{t+1} - \varphi_t\|_2 + \mathbb{E}[J_{\varphi_t}(\boldsymbol{\theta}_{t+1})] - J_{\varphi_t}(\boldsymbol{\theta}_t).\end{aligned}$$

For simplicity, we denote $\xi = (s, a, s')$ and

$$\begin{aligned}h_\theta(\boldsymbol{\theta}, \boldsymbol{\omega}, \boldsymbol{\varphi}, \xi) &= (\tilde{r}_{\boldsymbol{\varphi}, \boldsymbol{\theta}}(s, a) + (\gamma\phi(s') - \phi(s))^\top \boldsymbol{\omega}) \nabla_\theta \log \pi_\theta(a|s), \\ \bar{h}_\theta(\boldsymbol{\theta}, \boldsymbol{\omega}, \boldsymbol{\varphi}, \nu) &= \mathbb{E}_{\nu, \pi_{\theta, \mathcal{P}}} [h_\theta(\boldsymbol{\theta}, \boldsymbol{\omega}, \boldsymbol{\varphi}, \xi)],\end{aligned}$$

thereby

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t + \eta_t^\theta h_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \xi_t), \quad \mathbb{E}[\boldsymbol{\theta}_{t+1}|\boldsymbol{\theta}_t] = \boldsymbol{\theta}_t + \eta_t^\theta \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \hat{\nu}_t).$$

Then we apply the Taylor expansion on $J_{\varphi_t}(\boldsymbol{\theta})$ around $\boldsymbol{\theta}_t$. The second-order term is propotional to $\eta_t^{\theta^2}$ as the gradient has bounded variance, while the first-order term can be decomposed into three parts, associated with approximation error, critic error and Markovian noise, respectively.

$$\begin{aligned}\mathbb{E}[J_{\varphi_t}(\boldsymbol{\theta}_{t+1})] - J_{\varphi_t}(\boldsymbol{\theta}_t) &\geq \mathbb{E}\langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), \boldsymbol{\theta}_{t+1} - \boldsymbol{\theta}_t \rangle - \frac{S_J}{2} \mathbb{E}\|\boldsymbol{\theta}_{t+1} - \boldsymbol{\theta}_t\|_2^2 \\ &\geq \eta_t^\theta \langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \hat{\nu}_t) \rangle - \frac{S_J \eta_t^{\theta^2}}{2} \mathbb{E}\|h_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \xi_t)\|_2^2 \\ &\geq \eta_t^\theta \langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), (1 - \gamma) \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t) \rangle \\ &\quad + \eta_t^\theta \langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t^*, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) - (1 - \gamma) \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t) \rangle \\ &\quad + \eta_t^\theta \langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) - \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t^*, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) \rangle \\ &\quad + \eta_t^\theta \langle \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t), \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \hat{\nu}_t) - \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) \rangle \\ &\quad - \frac{L^2 C_\delta^2 S_J \eta_t^{\theta^2}}{2} \\ &\geq \eta_t^\theta (1 - \gamma) \|\nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t)\|_2^2 \\ &\quad - \eta_t^\theta \|\nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t)\|_2 \underbrace{\|\bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t^*, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) - (1 - \gamma) \nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t)\|_2}_{I_1} \\ &\quad - \eta_t^\theta \|\nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t)\|_2 \underbrace{\|\bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}}) - \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t^*, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}})\|_2}_{I_2} \\ &\quad - \eta_t^\theta \|\nabla_\theta J_{\varphi_t}(\boldsymbol{\theta}_t)\|_2 \underbrace{\|\bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \hat{\nu}_t) - \bar{h}_\theta(\boldsymbol{\theta}_t, \boldsymbol{\omega}_t, \boldsymbol{\varphi}_t, \nu_\rho^{\pi_{\theta_t}})\|_2}_{I_3} \\ &\quad - \frac{L^2 C_\delta^2 S_J \eta_t^{\theta^2}}{2}\end{aligned}\tag{17}$$

I_1 is associated with the approximation error. Using Proposition 4.3, we have

$$\begin{aligned}I_1 &= \left\| \mathbb{E}_{\nu_\rho^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[\left(\tilde{r}_{\boldsymbol{\varphi}_t, \boldsymbol{\theta}_t}(s, a) + \gamma \widehat{V}_{\boldsymbol{\omega}_t^*}(s') - \widehat{V}_{\boldsymbol{\omega}_t^*}(s) \right) \right. \right. \\ &\quad \left. \left. - \left(\tilde{r}_{\boldsymbol{\varphi}_t, \boldsymbol{\theta}_t}(s, a) + \gamma \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s') - \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s) \right) \right] \nabla_\theta \log \pi_\theta(a|s) \right\|_2 \\ &= \left\| \mathbb{E}_{\nu_\rho^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[\left(\gamma (\widehat{V}_{\boldsymbol{\omega}_t^*}(s') - \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s')) - (\widehat{V}_{\boldsymbol{\omega}_t^*}(s) - \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s)) \right) \nabla_\theta \log \pi_\theta(a|s) \right] \right\|_2 \\ &\leq L \sqrt{\mathbb{E}_{\nu_\rho^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[\left(\gamma (\widehat{V}_{\boldsymbol{\omega}_t^*}(s') - \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s')) - (\widehat{V}_{\boldsymbol{\omega}_t^*}(s) - \tilde{V}_{\boldsymbol{\varphi}_t}^{\pi_{\theta_t}}(s)) \right)^2 \right]} \\ &\leq 2\sqrt{2}L\epsilon.\end{aligned}$$

I_2 is associated with the critic error. We have

$$\begin{aligned}
 I_2 &= \left\| \mathbb{E}_{\nu_{\rho}^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[\left(\left(\tilde{r}_{\varphi_t, \theta_t}(s, a) + \gamma \hat{V}_{\omega_t}(s') - \hat{V}_{\omega_t}(s) \right) \right. \right. \right. \\
 &\quad \left. \left. \left. - \left(\tilde{r}_{\varphi_t, \theta_t}(s, a) + \gamma \hat{V}_{\omega_t^*}(s') - \hat{V}_{\omega_t^*}(s) \right) \right) \nabla_{\theta} \log \pi_{\theta}(a|s) \right] \right\|_2 \\
 &= \left\| \mathbb{E}_{\nu_{\rho}^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[(\gamma \phi(s') - \phi(s))^{\top} (\omega_t - \omega_t^*) \nabla_{\theta} \log \pi_{\theta}(a|s) \right] \right\|_2 \\
 &\leq L \sqrt{\mathbb{E}_{\nu_{\rho}^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} \left[((\gamma \phi(s') - \phi(s))^{\top} (\omega_t - \omega_t^*))^2 \right]} \\
 &\leq 2L \|\omega_t - \omega_t^*\|_2.
 \end{aligned}$$

I_3 is associated with the Markovian noise. We have

$$\begin{aligned}
 I_3 &= \left\| \int_S ds (\dot{\nu}_t(s) - \nu_{\rho}^{\pi_{\theta_t}}(s)) \mathbb{E}_{a, s'} \left[\hat{\delta}(s, a, s') \nabla_{\theta} \log \pi_{\theta}(a|s) \right] \right\|_2 \\
 &\leq LC_{\delta} \|\dot{\nu}_t - \nu_{\rho}^{\pi_{\theta_t}}\|_1
 \end{aligned}$$

Combining the above, we can derive from (17) that

$$\begin{aligned}
 (1 - \gamma) \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2^2 &\leq \frac{1}{\eta_t^{\theta}} (\mathbb{E}[J_{\varphi_{t+1}}(\theta_{t+1})] - J_{\varphi_t}(\theta_t)) + \frac{D_J \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2}{\eta_t^{\theta}} \\
 &\quad + \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2 (I_1 + I_2 + I_3) + \frac{L^2 C_{\delta}^2 S_J \eta_t^{\theta}}{2} \\
 &\leq \frac{1}{\eta_t^{\theta}} (\mathbb{E}[J_{\varphi_{t+1}}(\theta_{t+1})] - J_{\varphi_t}(\theta_t)) + \frac{D_J \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2}{\eta_t^{\theta}} \\
 &\quad + 2\sqrt{2}LL_J\epsilon + 2L\|\omega_t - \omega_t^*\|_2 \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2 \\
 &\quad + LC_{\delta}L_J \|\dot{\nu}_t - \nu_{\rho}^{\pi_{\theta_t}}\|_1 + \frac{L^2 C_{\delta}^2 S_J}{2} \eta_t^{\theta}.
 \end{aligned}$$

Summing over iterations, we have

$$\begin{aligned}
 (1 - \gamma) \sum_{t=T/2}^{T-1} \mathbb{E} \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2^2 &\leq \underbrace{\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\theta}} (\mathbb{E}[J_{\varphi_{t+1}}(\theta_{t+1})] - \mathbb{E}[J_{\varphi_t}(\theta_t)])}_{S_1} + \underbrace{D_J \sum_{t=T/2}^{T-1} \frac{\mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2}{\eta_t^{\theta}}}_{S_2} \\
 &\quad + \underbrace{\sqrt{2}TLL_J\epsilon + 2L \sum_{t=T/2}^{T-1} \|\omega_t - \omega_t^*\|_2 \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2}_{S_3} \\
 &\quad + \underbrace{LC_{\delta}L_J \sum_{t=T/2}^{T-1} \mathbb{E} \|\dot{\nu}_t - \nu_{\rho}^{\pi_{\theta_t}}\|_1}_{S_4} + \underbrace{\frac{L^2 C_{\delta}^2 S_J}{2} \sum_{t=T/2}^{T-1} \eta_t^{\theta}}_{S_5}. \tag{18}
 \end{aligned}$$

For S_1 , by applying the telescoping skill, we have

$$\begin{aligned}
 S_1 &= \sum_{t=T/2+1}^{T-1} \left(\frac{1}{\eta_{t-1}^{\theta}} - \frac{1}{\eta_t^{\theta}} \right) \mathbb{E}[J_{\varphi_t}(\theta_t)] + \frac{\mathbb{E}[J_{\varphi_T}(\theta_T)]}{\eta_{T-1}^{\theta}} - \frac{\mathbb{E}[J_{\varphi_{T/2}}(\theta_{T/2})]}{\eta_{T/2}^{\theta}} \\
 &\leq \sum_{t=T/2+1}^{T-1} \left(\frac{1}{\eta_t^{\theta}} - \frac{1}{\eta_{t-1}^{\theta}} \right) C_J + \frac{C_J}{\eta_{T-1}^{\theta}} + \frac{C_J}{\eta_{T/2}^{\theta}} \\
 &= \frac{2C_J}{\eta_{T-1}^{\theta}} \\
 &= O(\sqrt{T}).
 \end{aligned}$$

For S_2 , by applying the Cauchy-Schwartz inequality, we have

$$\begin{aligned} S_2 &\leq \sqrt{\sum_{t=T/2}^{T-1} \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2} \sqrt{\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\theta^2}}} \\ &= \sqrt{\frac{TF_T}{2} \sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\theta^2}}} \\ &= O\left(\sqrt{TF_T}\right), \end{aligned}$$

where the last equality is due to the fact that

$$\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\theta^2}} = \frac{1}{c_{\theta}^2} \sum_{t=T/2+1}^T \frac{1}{t} = \frac{1}{c_{\theta}^2} (H_T - H_{T/2}) \sim \frac{\ln 2}{c_{\theta}^2}.$$

For S_3 , by applying the Cauchy-Schwartz inequality, we have

$$\begin{aligned} S_3 &\leq \sqrt{\sum_{t=T/2}^{T-1} \mathbb{E} \|\omega_t - \omega_t^*\|_2^2} \sqrt{\sum_{t=T/2}^{T-1} \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2^2} \\ &= \frac{T}{2} \sqrt{\frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E} \|\omega_t - \omega_t^*\|_2^2} \sqrt{\frac{1}{T/2} \sum_{t=T/2}^{T-1} \|\nabla_{\theta} J_{\varphi_t}(\theta_t)\|_2^2} \\ &= \frac{T}{2} \sqrt{G_T W_T}. \end{aligned}$$

For S_4 , by Proposition 4.13, we have

$$\begin{aligned} S_4 &\leq \sum_{t=T/2}^{T-1} \left(LC_{\delta} L_{\nu} \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^{\theta} + \gamma^t \|\rho - \nu_{\rho}^{\pi_{\theta_0}}\|_1 \right) \\ &= LC_{\delta} L_{\nu} \sum_{t=0}^{T-1} \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^{\theta} + 2 \sum_{t=T/2}^{T-1} \gamma^t \\ &\leq LC_{\delta} L_{\nu} \sum_{t=0}^{T-1} \frac{\eta_t^{\theta}}{1-\gamma} + \frac{2}{\gamma^{T/2}(1-\gamma)} \\ &= \frac{LC_{\delta} L_{\nu}}{c_{\theta}(1-\gamma)} \sum_{t=1}^T \frac{1}{\sqrt{t}} + \frac{2}{\gamma^{T/2}(1-\gamma)} \\ &= O\left(\sqrt{T}\right) \end{aligned}$$

For S_5 , we have

$$S_5 = \frac{1}{c_{\theta}} \sum_{t=T/2+1}^T \frac{1}{\sqrt{t}} = O\left(\sqrt{T}\right).$$

Plug S_1, S_2, S_3, S_4 and S_5 into (18) and divide both sides by $(1-\gamma)\frac{T}{2}$, we obtain

$$G_T \leq \frac{2L}{1-\gamma} \sqrt{G_T W_T} + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon),$$

thus completes the proof. $\square$

C.2. Step 2: Bounding the Critic Error

Theorem C.2 (Critic Update). Take $\eta_t^\theta = \frac{c_\theta}{\sqrt{t}}$, $\eta_t^\omega = \frac{c_\omega}{\sqrt{t}}$ where c_θ and c_ω are constants such that $\frac{c_\theta}{c_\omega} \leq \frac{\lambda}{4LS_\omega}$, then

$$W_T \leq 2(1-\gamma)S_\omega \frac{c_\theta}{c_\omega} \sqrt{W_T G_T} + O(F_T \sqrt{T}) + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon).$$

Proof. For simplicity, we denote $\xi = (s, a, s')$ and

$$\begin{aligned} h_\omega(\theta, \omega, \varphi, \xi) &= (\tilde{r}_{\varphi, \theta}(s, a) + (\gamma\phi(s') - \phi(s))^\top \omega) \phi(s), \\ \bar{h}_\omega(\theta, \omega, \varphi, \nu) &= \mathbb{E}_{\nu, \pi_{\theta, \mathcal{P}}} [h_\omega(\theta, \omega, \varphi, \xi)]. \end{aligned}$$

According to the critic update rule (9), we have

$$\begin{aligned} \|\omega_{t+1} - \omega_{t+1}^*\|_2^2 &= \|\Pi_{C_\omega}(\omega_t + \eta_t^\omega h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t)) - \Pi_{C_\omega}(\omega_{t+1}^*)\|_2^2 \\ &\leq \|\omega_t + \eta_t^\omega h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) - \omega_{t+1}^*\|_2^2 \\ &= \|(\omega_t - \omega_t^*) + \eta_t^\omega h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) + (\omega_t^* - \omega_{t+1}^*)\|_2^2 \\ &= \|\omega_t - \omega_t^*\|_2^2 + \|\eta_t^\omega h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) + (\omega_t^* - \omega_{t+1}^*)\|_2^2 \\ &\quad + 2\langle \omega_t - \omega_t^*, \eta_t^\omega h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) + (\omega_t^* - \omega_{t+1}^*) \rangle \\ &\leq \|\omega_t - \omega_t^*\|_2^2 + 2\eta_t^{\omega^2} \|h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t)\|_2^2 + 2\|\omega_t^* - \omega_{t+1}^*\|_2^2 \\ &\quad + 2\eta_t^\omega \langle \omega_t - \omega_t^*, h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) \rangle + 2\langle \omega_t - \omega_t^*, \omega_t^* - \omega_{t+1}^* \rangle. \end{aligned} \quad (19)$$

To capture the evolution of the critic error $\|\omega_t - \omega_t^*\|_2^2$, we need to bound the other four terms on the right-hand side of (19). By taking the expectation, the two quadratic terms can be bounded by

$$\mathbb{E} \|h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t)\|_2^2 = \mathbb{E}_{\hat{\nu}_t, \pi_{\theta_t}, \mathcal{P}} \left\| \hat{\delta}(s_t, a_t, s'_t) \phi(s_t) \right\|_2^2 \leq C_\delta^2 \quad (20)$$

and

$$\begin{aligned} \mathbb{E} \|\omega_t^* - \omega_{t+1}^*\|_2^2 &\leq \mathbb{E} \left[(L_\omega \|\theta_{t+1} - \theta_t\|_2 + D_\omega \|\varphi_{t+1} - \varphi_t\|_2)^2 \right] \\ &\leq 2L_\omega^2 \eta_t^{\theta^2} \mathbb{E} \|h_\theta(\theta_t, \omega_t, \varphi_t, \xi_t)\|_2^2 + 2D_\omega^2 \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2 \\ &\leq 2L^2 C_\delta^2 L_\omega^2 \eta_t^{\theta^2} + 2D_\omega^2 \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2. \end{aligned} \quad (21)$$

The expectation of first inner-product term can be decomposed as the follows:

$$\begin{aligned} \mathbb{E} \langle \omega_t - \omega_t^*, h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) \rangle &= \langle \omega_t - \omega_t^*, \bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \hat{\nu}_t) \rangle \\ &= \underbrace{\langle \omega_t - \omega_t^*, \bar{h}_\omega(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) \rangle}_{J_1} \\ &\quad + \underbrace{\langle \omega_t - \omega_t^*, \bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) - \bar{h}_\omega(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) \rangle}_{J_2} \\ &\quad + \underbrace{\langle \omega_t - \omega_t^*, \bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \hat{\nu}_t) - \bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) \rangle}_{J_3}. \end{aligned}$$

According to the definition of ω_t^* , it should be a stationary point of the update rule, hence we have

$$J_1 = 0.$$

J_2 is associated with the critic error. We have

$$\begin{aligned} J_2 &= \left\langle \omega_t - \omega_t^*, \mathbb{E}_{\nu_\rho^{\pi_{\theta_t}}, \pi_{\theta_t}, \mathcal{P}} [(\gamma\phi(s') - \phi(s))^\top (\omega_t - \omega_t^*) \phi(s)] \right\rangle \\ &= \langle \omega_t - \omega_t^*, -A_{\theta_t}(\omega_t - \omega_t^*) \rangle \\ &\leq -\lambda \|\omega_t - \omega_t^*\|_2^2. \end{aligned}$$

J_3 is associated with the Markovian noise. We have

$$\begin{aligned} J_3 &\leq \|\omega_t - \omega_t^*\|_2 \|\bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \hat{\nu}_t) - \bar{h}_\omega(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}})\|_2 \\ &\leq 2C_\omega \left\| \int_{\mathcal{S}} ds (\hat{\nu}_t(s) - \nu_\rho^{\pi_{\theta_t}}(s)) \mathbb{E}_{a \sim \pi_{\theta_t}(\cdot|s), s' \sim \mathcal{P}(\cdot|s,a)} [\hat{\delta}(s, a, s') \phi(s)] \right\|_2 \\ &\leq 2C_\omega C_\delta \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1 \end{aligned}$$

Hence, we have

$$\mathbb{E} \langle \omega_t - \omega_t^*, h_\omega(\theta_t, \omega_t, \varphi_t, \xi_t) \rangle \leq -\lambda \|\omega_t - \omega_t^*\|_2^2 + 2C_\omega C_\delta \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1. \quad (22)$$

The analysis of the last term in (19) is similar to the analysis of the actor error. We first leverage the Lipschitz continuity of $\omega^*(\varphi, \theta)$ with respect to φ , then apply the Taylor expansion of $\omega^*(\varphi, \theta)$ with respect to θ around θ_t :

$$\begin{aligned} \mathbb{E} \langle \omega_t - \omega_t^*, \omega_t^* - \omega_{t+1}^* \rangle &= \mathbb{E} \langle \omega_t - \omega_t^*, \omega^*(\varphi_t, \theta_t) - \omega^*(\varphi_{t+1}, \theta_{t+1}) \rangle \\ &\quad + \mathbb{E} \langle \omega_t - \omega_t^*, \omega^*(\varphi_{t+1}, \theta_t) - \omega^*(\varphi_{t+1}, \theta_{t+1}) \rangle \\ &= \mathbb{E} \langle \omega_t - \omega_t^*, \omega^*(\varphi_t, \theta_t) - \omega^*(\varphi_{t+1}, \theta_t) \rangle \\ &\quad + \mathbb{E} \langle \omega_t - \omega_t^*, \omega^*(\varphi_t, \theta_t) - \omega^*(\varphi_{t+1}, \theta_t) - \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\theta_t - \theta_{t+1}) \rangle \\ &\quad + \mathbb{E} \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\theta_t - \theta_{t+1}) \rangle \\ &\leq D_\omega \|\omega_t - \omega_t^*\|_2 \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2 + \frac{S_\omega}{2} \|\omega_t - \omega_t^*\|_2 \mathbb{E} \|\theta_{t+1} - \theta_t\|_2^2 \\ &\quad + \mathbb{E} \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\theta_t - \theta_{t+1}) \rangle \\ &\leq L^2 C_\delta^2 C_\omega S_\omega \eta_t^{\theta^2} + 2C_\omega D_\omega \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2 \\ &\quad + \underbrace{\mathbb{E} \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\theta_t - \theta_{t+1}) \rangle}_I, \end{aligned}$$

where the last inequality uses the facts that

$$\mathbb{E} \|\theta_t - \theta_{t+1}\|_2^2 \leq L^2 C_\delta^2 \eta_t^{\theta^2} \quad \text{and} \quad \|\omega_t - \omega_t^*\|_2^2 \leq 2C_\omega.$$

The remaining inner product term I can then be decomposed into terms related to I_1 , I_2 and I_3 .

$$\begin{aligned} I &= -\eta_t^\theta \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top \bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \xi_t) \rangle \\ &= -(1-\gamma) \eta_t^\theta \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top \nabla_\theta J_{\varphi_t}(\theta_t) \rangle \\ &\quad - \eta_t^\theta \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\bar{h}_\theta(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) - (1-\gamma) \nabla_\theta J_{\varphi_t}(\theta_t)) \rangle \\ &\quad - \eta_t^\theta \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) - \bar{h}_\theta(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}})) \rangle \\ &\quad - \eta_t^\theta \langle \omega_t - \omega_t^*, \nabla_\theta \omega^*(\varphi_{t+1}, \theta_t)^\top (\bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \hat{\nu}_t) - \bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}})) \rangle \\ &\leq L^2 C_\delta^2 C_\omega S_\omega \eta_t^{\theta^2} + 2C_\omega D_\omega \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2 \\ &\quad + (1-\gamma) S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \|\nabla_\theta J_{\varphi_t}(\theta_t)\|_2 \\ &\quad + S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \underbrace{\|\bar{h}_\theta(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) - (1-\gamma) \nabla_\theta J_{\varphi_t}(\theta_t)\|}_{I_1} \\ &\quad + S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \underbrace{\|\bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}}) - \bar{h}_\theta(\theta_t, \omega_t^*, \varphi_t, \nu_\rho^{\pi_{\theta_t}})\|}_{I_2} \\ &\quad + S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \underbrace{\|\bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \hat{\nu}_t) - \bar{h}_\theta(\theta_t, \omega_t, \varphi_t, \nu_\rho^{\pi_{\theta_t}})\|}_{I_3}, \end{aligned}$$

where

$$I_1 \leq 2\sqrt{2}L\epsilon, \quad I_2 \leq 2L\|\omega_t - \omega_t^*\|_2, \quad I_3 \leq LC_\delta \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1.$$

Hence, we have

$$\begin{aligned}
 \langle \omega_t - \omega_t^*, \omega_t^* - \omega_{t+1}^* \rangle &\leq L^2 C_\delta^2 C_\omega S_\omega \eta_t^{\theta^2} + 2C_\omega D_\omega \|\varphi_{t+1} - \varphi_t\|_2 \\
 &\quad + (1 - \gamma) S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \|\nabla_\theta J_{\varphi_t}(\theta_t)\|_2 \\
 &\quad + 4\sqrt{2} L C_\omega S_\omega \epsilon \eta_t^\theta + 2L S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2^2 \\
 &\quad + 2L C_\omega C_\delta S_\omega \eta_t^\theta \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1.
 \end{aligned} \tag{23}$$

Plugging (20), (21), (22) and (23) into (19) gives

$$\begin{aligned}
 \mathbb{E} \|\omega_{t+1} - \omega_{t+1}^*\|_2^2 &\leq (1 - 2\lambda \eta_t^\omega + 4L S_\omega \eta_t^\theta) \|\omega_t - \omega_t^*\|_2^2 \\
 &\quad + 2(1 - \gamma) S_\omega \eta_t^\theta \|\omega_t - \omega_t^*\|_2 \|\nabla_\theta J_t(\theta_t)\|_2 \\
 &\quad + 2D_\omega^2 \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2 + 4C_\omega D_\omega \mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2 \\
 &\quad + 4C_\omega C_\delta (L S_\omega \eta_t^\theta + \eta_t^\omega) \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1 \\
 &\quad + 4L^2 C_\delta^2 L_\omega^2 \eta_t^{\theta^2} + 8\sqrt{2} L C_\omega S_\omega \epsilon \eta_t^\theta.
 \end{aligned}$$

Note that $\frac{\eta_t^\theta}{\eta_t^\omega} = \frac{c_\theta}{c_\omega} \leq \frac{\lambda}{4L S_\omega}$, thus

$$\begin{aligned}
 \lambda \sum_{t=T/2}^{T-1} \mathbb{E} \|\omega_t - \omega_t^*\|_2^2 &\leq \underbrace{\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^\omega} \left(\mathbb{E} \|\omega_t - \omega_t^*\|_2^2 - \mathbb{E} \|\omega_{t+1} - \omega_{t+1}^*\|_2^2 \right)}_{S_1} \\
 &\quad + \underbrace{2(1 - \gamma) S_\omega \frac{c_\theta}{c_\omega} \sum_{t=T/2}^{T-1} \|\omega_t - \omega_t^*\|_2 \|\nabla_\theta J_t(\theta_t)\|_2}_{S_2} \\
 &\quad + \underbrace{2D_\omega^2 \sum_{t=T/2}^{T-1} \frac{\mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2^2}{\eta_t^\omega}}_{S_3} + \underbrace{2C_\omega D_\omega \sum_{t=T/2}^{T-1} \frac{\mathbb{E} \|\varphi_{t+1} - \varphi_t\|_2}{\eta_t^\omega}}_{S_4} \\
 &\quad + \underbrace{(1 + \lambda/4) C_\omega C_\delta L_J \sum_{t=T/2}^{T-1} \mathbb{E} \|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1}_{S_5} \\
 &\quad + \underbrace{4L^2 C_\delta^2 L_\omega^2 \frac{c_\theta}{c_\omega} \sum_{t=T/2}^{T-1} \eta_t^\theta + 8\sqrt{2} L C_\omega S_\omega \frac{c_\theta}{c_\omega} \epsilon}_{S_6}
 \end{aligned} \tag{24}$$

For S_1 , by applying the telescoping skill, we have

$$\begin{aligned}
 S_1 &= \sum_{t=T/2+1}^{T-1} \left(\frac{1}{\eta_t^\omega} - \frac{1}{\eta_{t-1}^\omega} \right) \mathbb{E} \|\omega_t - \omega_t^*\|_2^2 - \frac{\mathbb{E} \|\omega_T - \omega_T^*\|_2^2}{\eta_{T-1}^\omega} + \frac{\mathbb{E} \|\omega_{T/2} - \omega_{T/2}^*\|_2^2}{\eta_{T/2}^\omega} \\
 &\leq \sum_{t=T/2+1}^{T-1} \left(\frac{1}{\eta_t^\omega} - \frac{1}{\eta_{t-1}^\omega} \right) 2C_\omega + \frac{2C_\omega}{\eta_{T-1}^\omega} + \frac{2C_\omega}{\eta_{T/2}^\omega} \\
 &= \frac{4C_\omega}{\eta_{T-1}^\omega} \\
 &= O(\sqrt{T}).
 \end{aligned}$$

For S_2 , by applying the Cauchy-Schwartz inequality, we have

$$\begin{aligned}
 S_2 &\leq \sqrt{\sum_{t=T/2}^{T-1} \mathbb{E} \|\boldsymbol{\omega}_t - \boldsymbol{\omega}_t^*\|_2^2} \sqrt{\sum_{t=T/2}^{T-1} \|\nabla_{\boldsymbol{\theta}} J_{\boldsymbol{\varphi}_t}(\boldsymbol{\theta}_t)\|_2^2} \\
 &= \frac{T}{2} \sqrt{\frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E} \|\boldsymbol{\omega}_t - \boldsymbol{\omega}_t^*\|_2^2} \sqrt{\frac{1}{T/2} \sum_{t=T/2}^{T-1} \|\nabla_{\boldsymbol{\theta}} J_{\boldsymbol{\varphi}_t}(\boldsymbol{\theta}_t)\|_2^2} \\
 &= \frac{T}{2} \sqrt{W_T G_T}.
 \end{aligned}$$

For S_3 , note that η_t^ω is decreasing as t grows, we have

$$S_3 = \sum_{t=T/2}^{T-1} \frac{\mathbb{E} \|\boldsymbol{\varphi}_{t+1} - \boldsymbol{\varphi}_t\|_2^2}{\eta_t^\omega} \leq \frac{1}{\eta_{T-1}^\omega} \sum_{t=T/2}^{T-1} \mathbb{E} \|\boldsymbol{\varphi}_{t+1} - \boldsymbol{\varphi}_t\|_2^2 = O\left(F_T T \sqrt{T}\right)$$

For S_4 , by applying the Cauchy-Schwartz inequality, we have

$$\begin{aligned}
 S_4 &\leq \sqrt{\sum_{t=T/2}^{T-1} \mathbb{E} \|\boldsymbol{\varphi}_{t+1} - \boldsymbol{\varphi}_t\|_2^2} \sqrt{\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\omega^2}}} \\
 &= \sqrt{\frac{T F_T}{2} \sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\omega^2}}} \\
 &= O\left(\sqrt{T F_T}\right),
 \end{aligned}$$

where the last equality is due to the fact that

$$\sum_{t=T/2}^{T-1} \frac{1}{\eta_t^{\omega^2}} = \frac{1}{c_\omega^2} \sum_{t=T/2+1}^T \frac{1}{t} = \frac{1}{c_\omega^2} (H_T - H_{T/2}) \sim \frac{\ln 2}{c_\omega^2}.$$

For S_5 , by Proposition 4.13, we have

$$\begin{aligned}
 S_5 &\leq \sum_{t=T/2}^{T-1} \left(LC_\delta L_\nu \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^\theta + \gamma^t \|\rho - \nu_\rho^{\pi_{\boldsymbol{\theta}_0}}\|_1 \right) \\
 &= LC_\delta L_\nu \sum_{t=0}^{T-1} \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^\theta + 2 \sum_{t=T/2}^{T-1} \gamma^t \\
 &\leq LC_\delta L_\nu \sum_{t=0}^{T-1} \frac{\eta_t^\theta}{1-\gamma} + \frac{2}{\gamma^{T/2}(1-\gamma)} \\
 &= \frac{LC_\delta L_\nu}{c_\theta(1-\gamma)} \sum_{t=1}^T \frac{1}{\sqrt{t}} + \frac{2}{\gamma^{T/2}(1-\gamma)} \\
 &= O\left(\sqrt{T}\right)
 \end{aligned}$$

For S_6 , we have

$$S_6 = \frac{1}{c_\theta} \sum_{t=T/2+1}^T \frac{1}{\sqrt{t}} = O\left(\sqrt{T}\right).$$

Plug S_1, S_2, S_3, S_4, S_5 and S_6 into (24) and divide both sides by $\frac{\lambda T}{2}$, we obtain

$$W_T \leq 2(1-\gamma)S_\omega \frac{c_\theta}{c_\omega} \sqrt{W_T G_T} + O\left(F_T \sqrt{T}\right) + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon),$$

thus completes the proof. $\square$

C.3. Step 3: Solving the System of Inequalities

Proof of Theorem 4.11 According to Theorem C.1 and Theorem C.2, we have

$$(1-\gamma)G_T \leq 2L\sqrt{G_T W_T} + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon), \quad (25)$$

$$\frac{1}{1-\gamma}W_T \leq 2S_\omega \frac{c_\theta}{c_\omega} \sqrt{G_T W_T} + O\left(\sqrt{T}F_T\right) + O\left(\frac{F_T}{T}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon). \quad (26)$$

Note that

$$2\sqrt{G_T W_T} = 2\sqrt{\frac{1-\gamma}{2L}G_T \cdot \frac{2L}{1-\gamma}W_T} \leq \frac{1-\gamma}{2L}G_T + \frac{2L}{1-\gamma}W_T. \quad (27)$$

Plug (27) into (25), we have

$$\frac{1-\gamma}{2L}G_T \leq \frac{2L}{1-\gamma}W_T + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon). \quad (28)$$

Combining (27) and (28), we have

$$2\sqrt{G_T W_T} \leq \frac{4L}{1-\gamma}W_T + O\left(\sqrt{\frac{F_T}{T}}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon). \quad (29)$$

Plug (29) into (26), we have

$$\frac{1-8LS_\omega \frac{c_\theta}{c_\omega}}{1-\gamma}W_T \leq O\left(\sqrt{T}F_T\right) + O\left(\frac{F_T}{T}\right) + O\left(\frac{1}{\sqrt{T}}\right) + O(\epsilon).$$

Therefore, when $\frac{c_\theta}{c_\omega} \leq \frac{1}{16LS_\omega}$,

$$W_T = O\left(\frac{1}{\sqrt{T}}\right) + O\left(\sqrt{T}F_T\right) + O\left(\frac{F_T}{T}\right) + O(\epsilon).$$

Combined with (28), we have

$$G_T = O\left(\frac{1}{\sqrt{T}}\right) + O\left(\sqrt{T}F_T\right) + O\left(\frac{F_T}{T}\right) + O(\epsilon).$$

D. Proof of Propositions, Preliminary Lemmas and Corollaries

Proof of Proposition 4.2 For any vector x ,

$$\begin{aligned} x^\top A_\theta x &= x^\top \mathbb{E}_{\nu_\rho^{\pi_\theta}, \pi_\theta, \mathcal{P}} \left[\phi(s) (\phi(s) - \gamma \phi(s'))^\top \right] x \\ &= x^\top \left(\mathbb{E}_s [\phi(s) \phi(s)^\top] - \gamma \mathbb{E}_{s,s'} [\phi(s) \phi(s')^\top] \right) x \\ &= \mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x] - \gamma \mathbb{E}_{s,s'} [x^\top \phi(s) \phi(s')^\top x] \end{aligned}$$

According to the Cauchy-Schwartz inequality,

$$\mathbb{E}_{s,s'} [x^\top \phi(s) \phi(s')^\top x] \leq \sqrt{\mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x]} \sqrt{\mathbb{E}_{s'} [x^\top \phi(s') \phi(s')^\top x]}.$$

Note that

$$\Pr(s' = x) = \frac{\nu_\rho^{\pi_\theta}(x) - (1 - \gamma)\rho(x)}{\gamma} \leq \frac{\nu_\rho^{\pi_\theta}(x)}{\gamma},$$

so

$$\mathbb{E}_{s'} [x^\top \phi(s') \phi(s')^\top x] \leq \frac{1}{\gamma} \mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x],$$

and hence

$$\begin{aligned} x^\top A_\theta x &\geq \mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x] - \gamma \sqrt{\mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x]} \sqrt{\frac{1}{\gamma} \mathbb{E}_s [x^\top \phi(s) \phi(s)^\top x]} \\ &= (1 - \sqrt{\gamma}) x^\top \Sigma_\theta x \\ &\geq (1 - \sqrt{\gamma}) \lambda_\Sigma \|x\|_2^2 \end{aligned}$$

Therefore, A_θ is positive definite with singular values lower-bounded by $\lambda = (1 - \sqrt{\gamma}) \lambda_\Sigma$.

Proof of Proposition 4.3

$$\begin{aligned} LHS &= \sqrt{\mathbb{E}_{\nu_\rho^{\pi_\theta}, \pi_\theta, \mathcal{P}} \left[\left(\left(\gamma \hat{V}_{\omega^*}(s') - \hat{V}_{\omega^*}(s) \right) - \left(\gamma \tilde{V}^{\pi_\theta}(s') - \tilde{V}^{\pi_\theta}(s) \right) \right)^2 \right]} \\ &\leq \sqrt{\mathbb{E}_{\nu_\rho^{\pi_\theta}, \pi_\theta, \mathcal{P}} \left[2 \left(\gamma \left(\hat{V}_{\omega^*}(s') - \tilde{V}^{\pi_\theta}(s') \right) \right)^2 + 2 \left(\hat{V}_{\omega^*}(s) - \tilde{V}^{\pi_\theta}(s) \right)^2 \right]} \\ &\leq \sqrt{2 \mathbb{E}_s \left[\left(\hat{V}_{\omega^*}(s) - \tilde{V}^{\pi_\theta}(s) \right)^2 \right] + 2 \gamma^2 \mathbb{E}_{s'} \left[\left(\hat{V}_{\omega^*}(s') - \tilde{V}^{\pi_\theta}(s') \right)^2 \right]} \\ &\leq \sqrt{2} \left(\underbrace{\sqrt{\mathbb{E}_s \left[\left(\hat{V}_{\omega^*}(s) - \tilde{V}^{\pi_\theta}(s) \right)^2 \right]}}_{I_1} + \gamma \underbrace{\sqrt{\mathbb{E}_{s'} \left[\left(\hat{V}_{\omega^*}(s') - \tilde{V}^{\pi_\theta}(s') \right)^2 \right]}}_{I_2} \right) \end{aligned}$$

According to the definition of ϵ (14), $I_1 \leq \epsilon$.

For I_2 , note that

$$\Pr(s' = x) = \frac{\nu_\rho^{\pi_\theta}(x) - (1 - \gamma)\rho(x)}{\gamma} \leq \frac{\nu_\rho^{\pi_\theta}(x)}{\gamma},$$

so

$$I_2 \leq \gamma \sqrt{\frac{1}{\gamma} \mathbb{E}_s \left[\left(\hat{V}_{\omega^*}(s) - \tilde{V}^{\pi_\theta}(s) \right)^2 \right]} \leq \sqrt{\gamma} \epsilon \leq \epsilon.$$

Therefore, $LHS \leq 2\sqrt{2}\epsilon$.

Proof of Proposition 4.5 For any $\theta_1, \theta_2 \in \Omega(\theta)$, let

$$f(a) = \begin{cases} 1 & , \pi_{\theta_1}(a|s) \geq \pi_{\theta_2}(a|s) \\ -1 & , \text{otherwise} \end{cases},$$

then

$$\|\pi_{\theta_1}(\cdot|s) - \pi_{\theta_2}(\cdot|s)\|_1 = \mathbb{E}_{a \sim \pi_{\theta_1}(\cdot|s)}[f(a)] - \mathbb{E}_{a \sim \pi_{\theta_2}(\cdot|s)}[f(a)].$$

Note that

$$\begin{aligned} \|\nabla_{\theta} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[f(a)]\|_2 &= \left\| \nabla_{\theta} \int_{\mathcal{A}} \pi_{\theta}(a|s) f(a) da \right\|_2 \\ &= \left\| \int_{\mathcal{A}} \nabla_{\theta} \pi_{\theta}(a|s) f(a) da \right\|_2 \\ &= \left\| \int_{\mathcal{A}} \pi_{\theta}(a|s) \nabla_{\theta} \log \pi_{\theta}(a|s) f(a) da \right\|_2 \\ &= \left\| \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[\nabla_{\theta} \log \pi_{\theta}(a|s) f(a)] \right\|_2 \\ &\leq \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[\|\nabla_{\theta} \log \pi_{\theta}(a|s)\|_2 |f(a)|] \\ &\leq L. \end{aligned}$$

Therefore,

$$\|\pi_{\theta_1}(\cdot|s) - \pi_{\theta_2}(\cdot|s)\|_1 \leq L \|\theta_1 - \theta_2\|.$$

Proof of Proposition 4.10

$$\begin{aligned} |\mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)]| &= |\mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[r(s, a; \varphi) - \alpha(\varphi) \log \pi_{\theta}(a|s)]| \\ &= \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}|r(s, a; \varphi)| + \alpha(\varphi) \mathcal{H}(\pi_{\theta}(\cdot|s)) \\ &\leq C_R + C_{\alpha} C_{H,1} \end{aligned}$$

$$\begin{aligned} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[\tilde{r}_{\varphi, \theta}(s, a)^2] &= \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} \left[(r(s, a; \varphi) - \alpha(\varphi) \log \pi_{\theta}(a|s))^2 \right] \\ &= \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [r(s, a; \varphi)^2] - 2\alpha(\varphi) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [r(s, a; \varphi) \log \pi_{\theta}(a|s)] \\ &\quad + \alpha(\varphi)^2 \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [(\log \pi_{\theta}(a|s))^2] \\ &\leq C_R^2 + 2C_R C_{\alpha} C_{H,1} + C_{\alpha}^2 C_{H,2} \\ &= \left(C_R + C_{\alpha} (C_{H,1} \vee \sqrt{C_{H,2}}) \right)^2 \end{aligned}$$

$$\|\nabla_{\theta} \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 = \|\mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a) \nabla_{\theta} \log \pi_{\theta}(a|s)]\|_2 \leq (C_R + C_{\alpha} C_{H,1})L$$

$$\begin{aligned} \|\nabla_{\theta}^2 \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 &= \|\mathbb{E}[(\tilde{r}_{\varphi, \theta}(s, a) - \alpha) \nabla_{\theta} \log \pi_{\theta}(a|s) \nabla_{\theta} \log \pi_{\theta}(a|s)^{\top}] \\ &\quad + \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a) \nabla_{\theta}^2 \log \pi_{\theta}(a|s)]\|_2 \\ &\leq (C_R + C_{\alpha} C_{H,1} + C_{\alpha})L^2 + (C_R + C_{\alpha} C_{H,1})S \\ &\leq (C_R + C_{\alpha}(1 + C_{H,1}))(L^2 + S) \end{aligned}$$

Hence, $C = C_R + C_{\alpha} (1 + C_{H,1} \vee \sqrt{C_{H,2}})$.

$$\begin{aligned} \|\nabla_{\varphi} \mathbb{E}[\tilde{r}_{\varphi, \theta}(s, a)]\|_2 &= \|\mathbb{E}[\nabla_{\varphi} r(s, a; \varphi)] + \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)}[-\log \pi_{\theta}(a|s)] \nabla_{\varphi} \alpha(\varphi)\|_2 \\ &\leq D_R + C_{H,1} D_{\alpha} \end{aligned}$$

Hence, $D = D_R + C_{H,1} D_{\alpha}$.

Proof of Corollary 4.12 Assume that $\mathbb{E}\|h_\varphi(t)\|_2^2 \leq C_\varphi^2$ and $\eta_t^\varphi = \frac{c_\varphi}{\sqrt{t}}$, we have

$$\begin{aligned} F_T &= \frac{1}{T/2} \sum_{t=T/2}^{T-1} \mathbb{E}\|\varphi_{t+1} - \varphi_t\|_2^2 \\ &= \frac{1}{T/2} \sum_{t=T/2}^{T-1} \eta_t^{\varphi^2} \mathbb{E}\|h_\varphi(t)\|_2^2 \\ &\leq \frac{C_\varphi^2}{T/2} \sum_{t=T/2}^{T-1} \frac{c_\varphi^2}{t} \\ &= O(1/T) \end{aligned}$$

Hence, the terms $O(F_T\sqrt{T})$ and $O(\sqrt{F_T/T})$ are both dominated by $O(1/\sqrt{T})$, leading to an overall $O(1/\sqrt{T}) + O(\epsilon)$ bound.

Proof of Proposition 4.13 We abuse the notation $\widehat{\mathcal{P}}_\theta : \Delta(\mathcal{S}) \rightarrow \Delta(\mathcal{S})$ to denote an operator that acts on a state distribution ν , defined by

$$\begin{aligned} (\widehat{\mathcal{P}}_\theta \nu)(s') &= \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \nu(s) \pi_\theta(a|s) \widehat{\mathcal{P}}(s'|s, a) \\ &= \gamma \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \nu(s) \pi_\theta(a|s) \mathcal{P}(s'|s, a) + (1 - \gamma) \rho(s') \end{aligned}$$

Then, $\widehat{\mathcal{P}}_\theta$ is a contraction mapping and $\nu_\rho^{\pi_\theta}$ is the unique fix point of it. Formally, $\forall \nu_1, \nu_2 \in \Delta(\mathcal{S})$, we have

$$\begin{aligned} \|\widehat{\mathcal{P}}_\theta \nu_1 - \widehat{\mathcal{P}}_\theta \nu_2\|_1 &= \int_{\mathcal{S}} ds' |(\widehat{\mathcal{P}}_\theta \nu_1)(s') - (\widehat{\mathcal{P}}_\theta \nu_2)(s')| \\ &= \gamma \int_{\mathcal{S}} ds' \left| \int_{\mathcal{S}} ds \int_{\mathcal{A}} da (\nu_1(s) - \nu_2(s)) \pi_\theta(a|s) \mathcal{P}(s'|s, a) \right| \\ &\leq \gamma \int_{\mathcal{S}} ds |\nu_1(s) - \nu_2(s)| \int_{\mathcal{S}} ds' \int_{\mathcal{A}} da \pi_\theta(a|s) \mathcal{P}(s'|s, a) \\ &= \gamma \|\nu_1 - \nu_2\|_1, \end{aligned}$$

and

$$(\widehat{\mathcal{P}}_\theta \nu_\rho^{\pi_\theta})(s) = \nu_\rho^{\pi_\theta}(s), \forall s \in \mathcal{S}.$$

Therefore,

$$\begin{aligned} \mathbb{E}\|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_t}}\|_1 &\leq \mathbb{E}\|\hat{\nu}_t - \nu_\rho^{\pi_{\theta_{t-1}}}\|_1 + \mathbb{E}\|\nu_\rho^{\pi_{\theta_{t-1}}} - \nu_\rho^{\pi_{\theta_t}}\|_1 \\ &\leq \mathbb{E}\left\|\widehat{\mathcal{P}}_{\theta_{t-1}} \hat{\nu}_{t-1} - \widehat{\mathcal{P}}_{\theta_{t-1}} \nu_\rho^{\pi_{\theta_{t-1}}}\right\|_1 + L_\nu \mathbb{E}\|\theta_{t-1} - \theta_t\|_2 \\ &\leq \gamma \mathbb{E}\|\hat{\nu}_{t-1} - \nu_\rho^{\pi_{\theta_{t-1}}}\|_1 + LC_\delta L_\nu \eta_{t-1}^\theta \\ &\leq LC_\delta L_\nu \sum_{k=0}^{t-1} \gamma^{t-1-k} \eta_k^\theta + \gamma^t \|\rho - \nu_\rho^{\pi_{\theta_0}}\|_1. \end{aligned}$$

Proof of Lemma B.1

$$\begin{aligned} |\tilde{V}_\varphi^{\pi_\theta}(s)| &= \left| \frac{1}{1-\gamma} \mathbb{E}_{s \sim \nu_\rho^{\pi_\theta}(\cdot)} [\mathbb{E}_{a \sim \pi_\theta(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]] \right| \\ &\leq \frac{1}{1-\gamma} \mathbb{E}_{s \sim \nu_\rho^{\pi_\theta}(\cdot)} [\mathbb{E}_{a \sim \pi_\theta(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]] \\ &\leq \frac{C}{1-\gamma} \end{aligned}$$

Hence, $C_J = O((1 - \gamma)^{-1})$. Note that by letting $\rho(s) = \mathbb{I}[s = s_0]$, we have $\left| \tilde{V}_{\varphi}^{\pi_{\theta}}(s_0) \right| \leq C_J$ for any $s_0 \in \mathcal{S}$.

$$\begin{aligned} \|\nabla_{\theta} J_{\varphi}(\theta)\|_2 &= \left\| \frac{1}{1 - \gamma} \mathbb{E}_{s \sim \nu_{\rho}^{\pi_{\theta}}(\cdot)} \left[\mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} \left[\tilde{Q}_{\varphi}^{\pi_{\theta}}(s, a) \nabla_{\theta} \log \pi_{\theta}(a|s) \right] \right] \right\|_2 \\ &\leq \frac{L}{1 - \gamma} \mathbb{E}_{s \sim \nu_{\rho}^{\pi_{\theta}}(\cdot), a \sim \pi_{\theta}(\cdot|s)} \left| \tilde{r}(s, a) + \gamma \mathbb{E}_{s' \sim \mathcal{P}(\cdot|s, a)} \left[\tilde{V}_{\varphi}^{\pi_{\theta}}(s') \right] \right| \\ &\leq \frac{CL}{(1 - \gamma)^2} \end{aligned}$$

Hence, $L_J = O((1 - \gamma)^{-2})$. Similarly, by letting $\rho(s) = \mathbb{I}[s = s_0]$, we have $\left| \nabla_{\theta} \tilde{V}_{\varphi}^{\pi_{\theta}}(s_0) \right| \leq L_J$ for any $s_0 \in \mathcal{S}$.

$$\begin{aligned} \nabla_{\theta}^2 J_{\varphi}(\theta) &= \frac{1}{1 - \gamma} \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}} \left[\tilde{Q}^{\pi_{\theta}}(s, a) (\nabla_{\theta} \log \pi_{\theta}(a|s) \nabla_{\theta} \log \pi_{\theta}(a|s)^{\top} + \nabla_{\theta}^2 \log \pi_{\theta}(a|s)) \right] \\ &\quad + \frac{\gamma}{1 - \gamma} \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}, \mathcal{P}} \left[\nabla_{\theta} \log \pi_{\theta}(a|s) \nabla_{\theta} \tilde{V}_{\varphi}^{\pi_{\theta}}(s')^{\top} + \nabla_{\theta} \tilde{V}_{\varphi}^{\pi_{\theta}}(s') \nabla_{\theta} \log \pi_{\theta}(a|s)^{\top} \right] \\ \|\nabla_{\theta}^2 J_{\varphi}(\theta)\|_2 &\leq \frac{C_J(L^2 + S)}{1 - \gamma} + \frac{2\gamma LL_J}{1 - \gamma} = O((1 - \gamma)^{-3}). \end{aligned}$$

Hence, $S_J = O((1 - \gamma)^{-3})$.

$$\begin{aligned} \|\nabla_{\varphi} J_{\varphi}(\theta)\|_2 &= \left\| \nabla_{\varphi} \left(\frac{1}{1 - \gamma} \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}} [\tilde{r}_{\varphi, \theta}(s, a)] \right) \right\|_2 \\ &= \left\| \frac{1}{1 - \gamma} \mathbb{E}_{s \sim \nu_{\rho}^{\pi_{\theta}}(\cdot)} \left[\nabla_{\varphi} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)] \right] \right\|_2 \\ &\leq \frac{D}{1 - \gamma} \end{aligned}$$

Hence, $D_J = O((1 - \gamma)^{-1})$.

Proof of Corollary B.2 For any state $\theta_1, \theta_2 \in \Omega(\theta)$, consider the MDP $(\mathcal{S}, \mathcal{A}, \mathcal{P}, r', \gamma)$ where for any $s \in \mathcal{S}, a \in \mathcal{A}$,

$$r'(s, a) = f(s) := \begin{cases} 1 & , \nu_{\rho}^{\pi_{\theta_1}}(s) > \nu_{\rho}^{\pi_{\theta_2}}(s) \\ -1 & , \text{otherwise} \end{cases}.$$

Assume the regularization factor $\alpha = 0$, then for this RL problem, $\tilde{r}(s, a) = r'(s, a)$, and

$$\begin{aligned} \|\nu_{\rho}^{\pi_{\theta_1}} - \nu_{\rho}^{\pi_{\theta_2}}\|_1 &= \int_{\mathcal{S}} ds |\nu_{\rho}^{\pi_{\theta_1}}(s) - \nu_{\rho}^{\pi_{\theta_2}}(s)| \\ &= \int_{\mathcal{S}} ds (\nu_{\rho}^{\pi_{\theta_1}}(s) - \nu_{\rho}^{\pi_{\theta_2}}(s)) f(s) \\ &= \int_{\mathcal{S}} ds \nu_{\rho}^{\pi_{\theta_1}}(s) f(s) - \int_{\mathcal{S}} ds \nu_{\rho}^{\pi_{\theta_2}}(s) f(s) \\ &= \mathbb{E}_{s \sim \nu_{\rho}^{\pi_{\theta_1}}(\cdot)} \left[\mathbb{E}_{a \sim \pi_{\theta_1}(\cdot|s)} [\tilde{r}(s, a)] \right] - \mathbb{E}_{s \sim \nu_{\rho}^{\pi_{\theta_2}}(\cdot)} \left[\mathbb{E}_{a \sim \pi_{\theta_2}(\cdot|s)} [\tilde{r}(s, a)] \right] \\ &= (1 - \gamma)(J(\theta_1) - J(\theta_2)). \end{aligned}$$

Then we can apply Lemma B.1 with $C = 1$ to obtain $L_{\nu} = (1 - \gamma)L_J$ and $S_{\nu} = (1 - \gamma)S_J$.

Proof of Lemma B.3

$$\begin{aligned}
 \|\nabla_{\boldsymbol{\theta}} \mathbf{A}_{\boldsymbol{\theta}}\|_2 &= \left\| \nabla_{\boldsymbol{\theta}} \mathbb{E}_{\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}, \pi_{\boldsymbol{\theta}}, \mathcal{P}} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}] \right\|_2 \\
 &= \left\| \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \nabla_{\boldsymbol{\theta}} (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s)) \mathbb{E}_{s' \sim \mathcal{P}(\cdot|s,a)} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}] \right\|_2 \\
 &\leq \left(\int_{\mathcal{S}} ds \int_{\mathcal{A}} da \|\nabla_{\boldsymbol{\theta}} (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s))\|_2 \right) \left(\max_s \|\mathbb{E}_{s' \sim \mathcal{P}(\cdot|s,a)} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}]\|_2 \right) \\
 &\leq (1 + \gamma) \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \|\nabla_{\boldsymbol{\theta}} (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s))\|_2 \\
 &= (1 + \gamma) \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \|\nabla_{\boldsymbol{\theta}} \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s) + \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \nabla_{\boldsymbol{\theta}} \pi_{\boldsymbol{\theta}}(a|s)\|_2 \\
 &\leq (1 + \gamma) \left(\int_{\mathcal{S}} ds \|\nabla_{\boldsymbol{\theta}} \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s)\|_2 \int_{\mathcal{A}} da \pi_{\boldsymbol{\theta}}(a|s) + \int_{\mathcal{S}} ds \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \int_{\mathcal{A}} da \pi_{\boldsymbol{\theta}}(a|s) \|\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\theta}}(a|s)\|_2 \right) \\
 &\leq (1 + \gamma)(L_{\mu} + L) \\
 &= O((1 - \gamma)^{-1})
 \end{aligned}$$

Hence, $L_A = O((1 - \gamma)^{-1})$.

$$\begin{aligned}
 \|\nabla_{\boldsymbol{\theta}}^2 \mathbf{A}_{\boldsymbol{\theta}}\|_2 &= \left\| \nabla_{\boldsymbol{\theta}}^2 \mathbb{E}_{\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}, \pi_{\boldsymbol{\theta}}, \mathcal{P}} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}] \right\|_2 \\
 &= \left\| \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \nabla_{\boldsymbol{\theta}}^2 (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s)) \mathbb{E}_{s' \sim \mathcal{P}(\cdot|s,a)} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}] \right\|_2 \\
 &\leq \left(\int_{\mathcal{S}} ds \int_{\mathcal{A}} da \|\nabla_{\boldsymbol{\theta}}^2 (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s))\|_2 \right) \left(\max_s \|\mathbb{E}_{s' \sim \mathcal{P}(\cdot|s,a)} [\phi(s)(\phi(s) - \gamma\phi(s'))^{\top}]\|_2 \right) \\
 &\leq (1 + \gamma) \int_{\mathcal{S}} ds \int_{\mathcal{A}} da \|\nabla_{\boldsymbol{\theta}}^2 (\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \pi_{\boldsymbol{\theta}}(a|s))\|_2 \\
 &\leq (1 + \gamma) \left[\int_{\mathcal{S}} ds \|\nabla_{\boldsymbol{\theta}}^2 \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s)\|_2 \int_{\mathcal{A}} da \pi_{\boldsymbol{\theta}}(a|s) \right. \\
 &\quad \left. + 2 \int_{\mathcal{S}} ds \|\nabla_{\boldsymbol{\theta}} \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s)\|_2 \int_{\mathcal{A}} da \pi_{\boldsymbol{\theta}}(a|s) \|\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\theta}}(a|s)\|_2 \right. \\
 &\quad \left. + \int_{\mathcal{S}} ds \nu_{\rho}^{\pi_{\boldsymbol{\theta}}}(s) \int_{\mathcal{A}} da \pi_{\boldsymbol{\theta}}(a|s) \|\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\theta}}(a|s) \nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\theta}}(a|s)^{\top} + \nabla_{\boldsymbol{\theta}}^2 \log \pi_{\boldsymbol{\theta}}(a|s)\|_2 \right] \\
 &\leq (1 + \gamma) (S_{\nu} + 2LL_{\nu} + L^2 + S) \\
 &= O((1 - \gamma)^{-2})
 \end{aligned}$$

Hence, $S_A = O((1 - \gamma)^{-2})$.

Proof of Lemma B.4

$$\begin{aligned}
 \|\mathbf{b}_{\boldsymbol{\varphi}, \boldsymbol{\theta}}\|_2 &= \left\| \mathbb{E}_{\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}, \pi_{\boldsymbol{\theta}}} [\tilde{r}_{\boldsymbol{\varphi}, \boldsymbol{\theta}}(s, a) \phi(s)] \right\|_2 \\
 &\leq \left| \mathbb{E}_{\nu_{\rho}^{\pi_{\boldsymbol{\theta}}}, \pi_{\boldsymbol{\theta}}} [\tilde{r}_{\boldsymbol{\varphi}, \boldsymbol{\theta}}(s, a)] \right| \\
 &\leq C
 \end{aligned}$$

Hence, $C_b = O(1)$.

$$\begin{aligned}
 \|\nabla_{\theta} \mathbf{b}_{\varphi, \theta}\|_2 &= \left\| \nabla_{\theta} \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}} [\tilde{r}_{\varphi, \theta}(s, a) \phi(s)] \right\|_2 \\
 &= \left\| \int_S ds \nabla_{\theta} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \phi(s) \right\|_2 \\
 &\leq \left\| \int_S ds \nabla_{\theta} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \right\|_2 \left(\max_s \|\phi(s)\|_2 \right) \\
 &\leq \int_S ds \|\nabla_{\theta} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)])\|_2 \\
 &\leq \int_S ds \|\nabla_{\theta} \nu_{\rho}^{\pi_{\theta}}(s)\|_2 \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)] \\
 &\quad + \int_S ds \nu_{\rho}^{\pi_{\theta}}(s) \|\nabla_{\theta} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \\
 &\leq CL_{\nu} + CL + 0 \\
 &= O((1 - \gamma)^{-1})
 \end{aligned}$$

Hence, $L_b = O((1 - \gamma)^{-1})$.

$$\begin{aligned}
 \|\nabla_{\theta}^2 \mathbf{b}_{\varphi, \theta}\|_2 &= \left\| \nabla_{\theta}^2 \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}} [\tilde{r}_{\varphi, \theta}(s, a) \phi(s)] \right\|_2 \\
 &= \left\| \int_S ds \nabla_{\theta}^2 (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \phi(s) \right\|_2 \\
 &\leq \left\| \int_S ds \nabla_{\theta}^2 (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \right\|_2 \left(\max_s \|\phi(s)\|_2 \right) \\
 &\leq \int_S ds \|\nabla_{\theta}^2 (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)])\|_2 \\
 &\leq \int_S ds \|\nabla_{\theta}^2 \nu_{\rho}^{\pi_{\theta}}(s)\|_2 \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)] \\
 &\quad + \int_S ds \|\nabla_{\theta} \nu_{\rho}^{\pi_{\theta}}(s)\|_2 \|\nabla_{\theta} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \\
 &\quad + \int_S ds \nu_{\rho}^{\pi_{\theta}}(s) \|\nabla_{\theta}^2 \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \\
 &\leq CS_{\nu} + CLL_{\nu} + C(L^2 + S) \\
 &= O((1 - \gamma)^{-2})
 \end{aligned}$$

Hence, $S_b = O((1 - \gamma)^{-2})$.

$$\begin{aligned}
 \|\nabla_{\varphi} \mathbf{b}_{\varphi, \theta}\|_2 &= \left\| \nabla_{\varphi} \mathbb{E}_{\nu_{\rho}^{\pi_{\theta}}, \pi_{\theta}} [\tilde{r}_{\varphi, \theta}(s, a) \phi(s)] \right\|_2 \\
 &= \left\| \int_S ds \nabla_{\varphi} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \phi(s) \right\|_2 \\
 &\leq \left\| \int_S ds \nabla_{\varphi} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]) \right\|_2 \left(\max_s \|\phi(s)\|_2 \right) \\
 &\leq \int_S ds \|\nabla_{\varphi} (\nu_{\rho}^{\pi_{\theta}}(s) \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)])\|_2 \\
 &= \int_S ds \nu_{\rho}^{\pi_{\theta}}(s) \|\nabla_{\varphi} \mathbb{E}_{a \sim \pi_{\theta}(\cdot|s)} [\tilde{r}_{\varphi, \theta}(s, a)]\|_2 \\
 &\leq D
 \end{aligned}$$

Hence, $D_b = O(1)$.

Proof of Lemma B.5

$$\|\omega^*(\varphi, \theta)\|_2 = \|\mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta}\|_2 \leq \|\mathbf{A}_\theta^{-1}\|_2 \|\mathbf{b}_{\varphi, \theta}\|_2 \leq \frac{C_b}{\lambda} = \frac{C}{\lambda}$$

Hence, $C_\omega = O(\lambda^{-1})$.

$$\begin{aligned} \|\nabla_\theta \omega^*(\varphi, \theta)\|_2 &= \|\nabla_\theta (\mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta})\|_2 \\ &= \|\nabla_\theta (\mathbf{A}_\theta^{-1}) \mathbf{b}_{\varphi, \theta} + \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{b}_{\varphi, \theta}\|_2 \\ &= \|\mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta} + \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{A}_\theta\|_2 \|\mathbf{A}_\theta^{-1}\|_2 \|\mathbf{b}_{\varphi, \theta}\|_2 + \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq L_A C_b \lambda^{-2} + L_b \lambda^{-1} \\ &= O((1 - \gamma)^{-1} \lambda^{-2}) \end{aligned}$$

Hence, $L_\omega = O((1 - \gamma)^{-1} \lambda^{-2})$.

$$\begin{aligned} \|\nabla_\theta^2 \omega^*(\varphi, \theta)\|_2 &= \|\nabla_\theta^2 (\mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta})\|_2 \\ &= \|\nabla_\theta (\mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta} + \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{b}_{\varphi, \theta})\|_2 \\ &= \|\mathbf{A}_\theta^{-1} \nabla_\theta^2 \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta} + 2 \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta} \\ &\quad + 2 \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{A}_\theta \mathbf{A}_\theta^{-1} \nabla_\theta \mathbf{b}_{\varphi, \theta} + \mathbf{A}_\theta^{-1} \nabla_\theta^2 \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta^2 \mathbf{A}_\theta\|_2 \|\mathbf{A}_\theta^{-1}\|_2 \|\mathbf{b}_{\varphi, \theta}\|_2 \\ &\quad + 2 \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{A}_\theta\|_2 \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{A}_\theta\|_2 \|\mathbf{A}_\theta^{-1}\|_2 \|\mathbf{b}_{\varphi, \theta}\|_2 \\ &\quad + 2 \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{A}_\theta\|_2 \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta \mathbf{b}_{\varphi, \theta}\|_2 \\ &\quad + \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\theta^2 \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq S_A C_b \lambda^{-2} + 2L_A^2 C_b \lambda^{-3} + 2L_A L_b \lambda^{-2} + S_b \lambda^{-1} \\ &= O((1 - \gamma)^{-2} \lambda^{-3}) \end{aligned}$$

Hence, $S_\omega = O((1 - \gamma)^{-2} \lambda^{-3})$.

$$\begin{aligned} \|\nabla_\varphi \omega^*(\varphi, \theta)\|_2 &= \|\nabla_\varphi (\mathbf{A}_\theta^{-1} \mathbf{b}_{\varphi, \theta})\|_2 \\ &= \|\mathbf{A}_\theta^{-1} \nabla_\varphi \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq \|\mathbf{A}_\theta^{-1}\|_2 \|\nabla_\varphi \mathbf{b}_{\varphi, \theta}\|_2 \\ &\leq \frac{D_b}{\lambda} \\ &= O(\lambda^{-1}) \end{aligned}$$

Hence, $D_\omega = O(\lambda^{-1})$.

Proof of Lemma B.6

$$\begin{aligned} \mathbb{E}_{\nu, \pi_\theta, \mathcal{P}} \|\hat{\delta}(s, a, s')\|_2^2 &= \int_S ds \nu(s) \mathbb{E}_{\pi_\theta, \mathcal{P}} \left[(\tilde{r}(s, a) + (\phi(s') - \phi(s))^\top \omega)^2 \right] \\ &\leq \int_S ds \nu(s) \mathbb{E}_{\pi_\theta} \left[(\tilde{r}(s, a) + 2C_\omega)^2 \right] \\ &\leq \int_S ds \nu(s) (\mathbb{E}_{\pi_\theta} [\tilde{r}(s, a)^2] + 4C_\omega \mathbb{E}_{\pi_\theta} [\tilde{r}(s, a)] + 4C_\omega^2) \\ &\leq (C^2 + 4CC_\omega + 4C_\omega^2) \\ &= (C + 2C_\omega)^2 \end{aligned}$$

Hence, $C_\delta = C + 2C_\omega = O(\lambda^{-1})$.