
Near-Optimal Reinforcement Learning for Linear Distributionally Robust Markov Decision Processes

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 We study off-dynamics reinforcement learning (RL), where the policy training and
2 deployment environments are different. To deal with this environmental perturba-
3 tion, we focus on learning policies robust to uncertainties in transition dynamics
4 under the framework of distributionally robust Markov decision processes (DR-
5 MDPs), where the nominal and perturbed dynamics are linear Markov Decision
6 Processes. We propose a novel algorithm We-DRIVE-U that enjoys an average sub-
7 optimality $\tilde{\mathcal{O}}(dH \cdot \min\{1/\rho, H\}/\sqrt{K})$, where K is the number of episodes, H is
8 the horizon length, d is the feature dimension and ρ is the uncertainty level. This
9 result improves the state-of-the-art by $\mathcal{O}(dH/\min\{1/\rho, H\})$. We also construct a
10 novel hard instance and derive the first information-theoretic lower bound in this
11 setting. In stark contrast with standard linear MDPs, our lower bound depends on
12 the uncertainty level ρ , revealing the unique feature of DRMDPs. Our algorithm
13 also enjoys a ‘rare-switching’ design, and thus only requires $\mathcal{O}(dH \log(1 + H^2 K))$
14 policy switches and $\mathcal{O}(d^2 H \log(1 + H^2 K))$ calls for oracle to solve dual opti-
15 mization problems, which significantly improves the computational efficiency of
16 existing algorithms, whose policy switch and oracle complexities are both $\mathcal{O}(K)$.

17 1 Introduction

18 In dynamic decision-making and reinforcement learning (RL), Markov decision processes (MDPs)
19 offer a well-established framework for understanding complex systems and guiding agent behavior
20 [37]. However, MDPs encounter significant challenges in practical applications due to incomplete
21 knowledge of model parameters, especially transition probabilities. This sim-to-real gap, representing
22 the difference between training and testing environments, can lead to failures in fields like infectious
23 disease control and robotics [8, 60, 21, 22, 33]. To address these challenges, off-dynamics RL
24 provides a framework where policies are trained on a source domain and deployed to a distinct
25 target domain, promoting robust performance across varying environments [7, 17, 42]. Within
26 this framework, distributionally robust Markov decision processes (DRMDPs) have emerged as a
27 promising way to model transition uncertainty. DRMDPs focus on learning robust policies that
28 perform well under worst-case scenarios [29, 16]. Prior works [56, 52, 32, 35, 53, 34] have proposed
29 algorithms mainly for tabular DRMDPs with finite number of states and actions, which are infeasible
30 when facing large state and action spaces.

31 In environments characterized by large state and action spaces, function approximation techniques
32 are crucial to overcome the computational burden posed by high dimensionality. Linear function
33 approximation methods, based on relatively simple function classes, have shown significant theoretical
34 and practical successes in standard MDP environments [19, 11, 10, 50, 12]. However, their application
35 in DRMDPs introduces additional complexities. These complexities arise from the nonlinearity
36 caused by the dual formulation in the worst-case analysis, even when the transition dynamics in the

source domain are modeled as linear. Recently, [23] provided the first theoretical results in the *online* setting of d -rectangular linear DRMDPs, a specific type of DRMDPs where the nominal model is a linear MDP [19] and the uncertainty set is defined based on the linear structure of the nominal transition kernel. Apart from this, online DRMDP with linear function approximation is largely underexplored and it is not clear how far existing algorithms are from optimal. Consequently, two natural questions arise:

*Can we improve the current results for online DRMDPs with linear function approximation?
What is the fundamental limit in this setting?*

In this paper, we provide an affirmative answer to the first question and answer the second question by providing an information theoretic lower bound for d -rectangular linear DRMDPs. In particular, motivated by the adoption of variance-weighted ridge regression to achieve nearly optimal result in standard linear MDPs [62, 61, 58, 20, 59, 10, 13], we propose a variance-aware distributionally robust algorithm to solve the off-dynamics RL problem. Due to the nonlinearity caused by the dual optimization of DRMDPs, the adoption of variance information in linear DRMDPs is highly nontrivial. Existing algorithms that incorporate variance information in learning linear DRMDPs requires coverage assumptions on the offline dataset [24, 41], which is infeasible in the online setting where the algorithm needs to interact with the environment to collect data. To be specific, for online DRMDPs, the adoption of variance-weighted ridge regression causes the following **unique challenges** for both algorithm design and theoretical analysis:

- **(Fundamental non-linearity induced by the uncertainty)** The consideration of uncertainty set renders that the (robust) Bellman equations are not linear with respect to the nominal kernel, a key feature in standard MDP. A direct consequence is that though the Q-function remains a linear representation, its parameters must be estimated element-wisely from d (i.e., parameter dimension) variance-weighted ridge regressions, instead of one in standard linear MDP. This poses challenges for algorithm design, as it requires properly incorporating variance information into the estimation process and quantitatively controlling the estimation uncertainty.
- **(Precise control on variance estimation.)** Existing theoretical analyses of online linear MDPs rely heavily on the *Elliptical Potential Lemma*, showing that the estimation error shrinks rapidly enough to guarantee the near-optimality of the learned policy within a small number of rounds. However, this lemma is not applicable in our setting due to the element-wise parameter estimation procedure mentioned above. Instead, we adopt a large- k regime to control the estimation error, based upon the intuition that when the sample size is large, the variance estimation should be close to the true variance (see Lemma D.7). Finally, we leverage the ‘Range Shrinkage’ property (see Lemma H.10) for linear DRMDPs to bound the true variance, and thus obtain an improved bound.

Our work poses a distinct algorithm design and calls for different theoretical analysis techniques. Our **main contributions** are summarized as follows:

- We propose a novel algorithm, We-DRIVE-U, for d -rectangular linear DRMDPs with total-variation (TV) divergence uncertainty sets. We-DRIVE-U is designed based on the optimistic principle [18, 19, 10] to trade off the exploration and exploitation during interacting with the source environment to learn a robust policy. The key novelty of We-DRIVE-U lies in incorporating the variance information into the policy learning, by a carefully designed optimistic estimator of the variance of the optimal robust value function.
- We prove that We-DRIVE-U achieves an average suboptimality of $\tilde{\mathcal{O}}(dH \cdot \min\{1/\rho, H\}/\sqrt{K})$ when the number of episode K is large, which improves the state-of-the-art result [23] by $\tilde{\mathcal{O}}(dH/\min\{1/\rho, H\})$. We highlight that the average suboptimality of We-DRIVE-U demonstrates the ‘Range Shrinkage’ property (refer to Lemma H.10) through the term $\min\{1/\rho, H\}$. We further established an information-theoretic lower bound $\Omega(dH^{1/2} \cdot \min\{1/\rho, H\}/\sqrt{K})$, which shows that We-DRIVE-U is *near-optimal* up to $\mathcal{O}(\sqrt{H})$ for any uncertainty level $\rho \in (0, 1]$.
- We-DRIVE-U is favorable in applications where policy switching is risky or costly, since We-DRIVE-U achieves $\mathcal{O}(dH \log(1 + H^2 K))$ global policy switch (refer to Definition 4.5). Moreover, we note that calls for oracle to solve dual optimizations (3.3) are one of the main sources of computation complexity in DRMDP with linear function approximation. Thanks to the specifically designed ‘rare-switching’ regime, We-DRIVE-U achieves $\mathcal{O}(d^2 H \log(1 + H^2 K))$ oracle complexity (refer to Definition 4.6). Both results improve exiting online DRMDP algorithms by a factor of K . Thus, We-DRIVE-U enjoys low switching cost and low computation cost.

Notations. For any positive integer $H \in \mathbb{Z}_+$, we denote $[H] = \{1, 2, \dots, H\}$. For any set $\mathcal{S}$, define $\Delta(\mathcal{S})$ as the set of probability distributions over $\mathcal{S}$. For any function $V : \mathcal{S} \rightarrow \mathbb{R}$, define $[\mathbb{P}_h V](s, a) = \mathbb{E}_{s' \sim P_h(\cdot|s, a)}[V(s')]$, and $[V(s)]_\alpha = \min\{V(s), \alpha\}$, where $\alpha > 0$ is a constant. For a vector $\mathbf{x}$, define x_j as its j -th entry. Moreover, denote $[x_i]_{i \in [d]}$ as a vector with the i -th entry being x_i . For a matrix A , denote $\lambda_i(A)$ as the i -th eigenvalue of A . For two matrices A and B , denote $A \preceq B$ as the fact that $B - A$ is a positive semi-definite matrix. For any $P, Q \in \Delta(\mathcal{S})$, the total variation divergence of P and Q is defined as $D(P||Q) = 1/2 \int_{\mathcal{S}} |P(s) - Q(s)| ds$.

2 Preliminary

In this section, we introduce the mathematical framework of our setting. We use a tuple $\text{DRMDP}(\mathcal{S}, \mathcal{A}, H, \mathcal{U}^\rho(P^0), r)$ to denote a finite horizon DRMDP, where $\mathcal{S}$ and $\mathcal{A}$ are the state and action spaces, $H \in \mathbb{Z}_+$ is the horizon length, $P^0 = \{P_h^0\}_{h=1}^H$ is the nominal transition kernel, $\mathcal{U}^\rho(P^0) = \bigotimes_{h \in [H]} \mathcal{U}_h^\rho(P_h^0)$ denotes an uncertainty set centered around the nominal transition kernel with an uncertainty level $\rho \geq 0$, $r = \{r_h\}_{h=1}^H$ is the reward function. A policy $\pi = \{\pi_h\}_{h=1}^H$ is a sequence of decision rules. For a policy π , we define the robust value function and Q-function for any $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$ as

$$V_h^{\pi, \rho}(s) = \inf_{P \in \mathcal{U}^\rho(P^0)} \mathbb{E}^P \left[\sum_{t=h}^H r_t(s_t, a_t) \middle| s_h = s, \pi \right],$$

$$Q_h^{\pi, \rho}(s, a) = \inf_{P \in \mathcal{U}^\rho(P^0)} \mathbb{E}^P \left[\sum_{t=h}^H r_t(s_t, a_t) \middle| s_h = s, a_h = a, \pi \right].$$

Moreover, we define the optimal robust value function and optimal robust state-action value function: for any $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$, $V_h^{*, \rho}(s) = \sup_{\pi \in \Pi} V_h^{\pi, \rho}(s)$, $Q_h^{*, \rho}(s, a) = \sup_{\pi \in \Pi} Q_h^{\pi, \rho}(s, a)$, where Π is the set of all policies. Correspondingly, the optimal robust policy is the policy that achieves the optimal robust value function $\pi^* = \text{argsup}_{\pi \in \Pi} V_h^{\pi, \rho}(s)$.

In this paper, we focus on the d -rectangular linear DRMDP [26, 4, 23, 24], where the nominal environment is a linear MDP [19] with a simplex state space, defined as follows.

Assumption 2.1. Given a known feature mapping $\phi : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}^d$ satisfying $\sum_{i=1}^d \phi_i(s, a) = 1$, $\phi_i(s, a) \geq 0$, for any $(i, s, a) \in [d] \times \mathcal{S} \times \mathcal{A}$, we assume the reward functions $\{r_h\}_{h=1}^H$ and nominal transition kernels $\{P_h^0\}_{h=1}^H$ are linearly parameterized. Specifically, for any $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$, $r_h(s, a) = \langle \phi(s, a), \theta_h \rangle$, $P_h^0(\cdot|s, a) = \langle \phi(s, a), \mu_h^0(\cdot) \rangle$, where $\{\theta_h\}_{h=1}^H$ are known vectors with bounded norm $\|\theta_h\|_2 \leq \sqrt{d}$ and $\{\mu_h\}_{h=1}^H$ are unknown probability measures over $\mathcal{S}$.

In d -rectangular linear DRMDPs, an uncertainty set of transition dynamics $\mathcal{U}_h^\rho(P_h^0)$ is defined based on the linear structure of P_h^0 satisfying Assumption 2.1. In particular, for any $(h, i) \in [H] \times [d]$, we first define the factor uncertainty set $\mathcal{U}_{h,i}^\rho(\mu_{h,i}^0) = \{\mu : \mu \in \Delta(\mathcal{S}), D(\mu||\mu_{h,i}^0) \leq \rho\}$, where $D(\cdot||\cdot)$ is a probability divergence which we choose as the total variation (TV) divergence in this paper. Then the uncertainty set of transitions for state s and action a is defined as $\mathcal{U}_h^\rho(s, a; \mu_h^0) = \{\sum_{i=1}^d \phi_i(s, a) \mu_{h,i}(\cdot) : \mu_{h,i}(\cdot) \in \mathcal{U}_{h,i}^\rho(\mu_{h,i}^0), \forall i \in [d]\}$. We also denote $\mathcal{U}_h^\rho(P_h^0) = \bigotimes_{(s,a) \in \mathcal{S} \times \mathcal{A}} \mathcal{U}_h^\rho(s, a; \mu_h^0)$ as the collection of uncertainty sets on the whole state and action spaces. Built on these definitions, [23] showed that the following robust Bellman equations hold for any policy π

$$Q_h^{\pi, \rho}(s, a) = r_h(s, a) + \inf_{P_h(\cdot|s, a) \in \mathcal{U}_h^\rho(s, a; \mu_h^0)} [\mathbb{P}_h V_{h+1}^{\pi, \rho}](s, a), \quad (2.1a)$$

$$V_h^{\pi, \rho}(s) = \mathbb{E}_{a \sim \pi_h(\cdot|s)} [Q_h^{\pi, \rho}(s, a)], \quad (2.1b)$$

Similarly, we have the robust Bellman optimality equations

$$Q_h^{*, \rho}(s, a) = r_h(s, a) + \inf_{P_h(\cdot|s, a) \in \mathcal{U}_h^\rho(s, a; \mu_h^0)} [\mathbb{P}_h V_{h+1}^{*, \rho}](s, a), \quad (2.2a)$$

$$V_h^{*, \rho}(s) = \max_{a \in \mathcal{A}} Q_h^{*, \rho}(s, a). \quad (2.2b)$$

In the context of online DRMDPs, an agent actively interacts with the nominal environment within K episodes to learn the optimal robust policy. Specifically, at the start of episode k , an agent chooses a policy π^k based on the history information and receives the initial state s_1^k . Then the agent interacts with the nominal environment by executing π^k until the end of episode k , and collects a new trajectory.

132 The goal of the agent is to minimize the average suboptimality¹ after K episodes, which is defined as
 133 $\text{AveSubopt}(K) = 1/K \sum_{k=1}^K [V_1^{\star, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k)]$.

134 [25] recently show that sample efficient learning in online tabular DRMDPs is impossible in the
 135 presence of support shift, i.e., the nominal kernel and target kernel do not share the same support.
 136 Built on the hard instance constructed in their work, we carefully design feature mappings for the
 137 transition kernel to extend their lower bound to the following one for online linear DRMDPs.

138 **Proposition 2.2.** (Hardness result) There exists two d -rectangular linear DRMDPs $\{\mathcal{M}_\theta\}_{\theta \in \{0,1\}}$,
 139 such that $\inf_{\mathcal{ALG}} \sup_{\theta \in \{0,1\}} \mathbb{E}[\text{AveSubopt}^{\mathcal{M}_\theta, \mathcal{ALG}}(K)] \geq \Omega(\rho \cdot H)$, where $\text{AveSubopt}^{\mathcal{M}_\theta, \mathcal{ALG}}(K)$
 140 is the average suboptimality of algorithm $\mathcal{ALG}$ in the d -rectangular linear DRMDP $\mathcal{M}_\theta$.

141 Note that the lower bound in Proposition 2.2 does not converge to zero as K increases, which means
 142 that in general no algorithm can guarantee to learn the optimal robust policy approximately. To
 143 circumvent this problem, in the rest of paper we focus on a tractable subclass of d -rectangular linear
 144 DRMDP following [23, 25], which is formally defined in the following assumption.

145 **Assumption 2.3** (Fail-state). Assume there exists a ‘fail state’ s_f in the d -rectangular linear DRMDP,
 146 such that for all $(h, a) \in [H] \times \mathcal{A}$, $r_h(s_f, a) = 0$, $\mathbb{P}_h^0(s_f | s_f, a) = 1$.

147 With Assumption 2.3, we follow the framework in [23], where we have the following results on
 148 robust value functions that are helpful in solving the optimization in (2.2).

149 **Proposition 2.4** (Remark 4.2 of [23]). Under Assumption 2.3, we have $Q_h^{\pi, \rho}(s_f, a) = V_h^{\pi, \rho}(s_f) = 0$,
 150 $\forall (\pi, h, a) \in \Pi \times [H] \times \mathcal{A}$. Moreover, for any function $V : \mathcal{S} \rightarrow [0, H]$ that satisfies $\min_{s \in \mathcal{S}} V(s) =$
 151 $V(s_f) = 0$, we have $\inf_{\mu \in \mathcal{U}^\rho(\mu^0)} \mathbb{E}_{s \sim \mu} V(s) = \max_{\alpha \in [0, H]} \{\mathbb{E}_{s \sim \mu^0} [V(s)]_\alpha - \rho \alpha\}$.

152 3 Algorithm design

153 One prominent property of the d -rectangular DRMDP is that the robust Q-functions possess linear rep-
 154 resentations with respect to the feature mapping ϕ . In particular, under Assumptions 2.1 and 2.3, [23]
 155 show that for any $(\pi, s, a, h) \in \Pi \times \mathcal{S} \times \mathcal{A} \times [H]$, the robust Q-function $Q_h^{\pi, \rho}(s, a)$ has a linear form
 156 as follows $Q_h^{\pi, \rho}(s, a) = (r_h(s, a) + \phi(s, a)^\top \nu_h^{\pi, \rho}) \mathbb{1}\{s \neq s_f\}$, where $\nu_h^{\pi, \rho} = (\nu_{h,1}^{\pi, \rho}, \dots, \nu_{h,d}^{\pi, \rho})^\top$,
 157 $\nu_{h,i}^{\pi, \rho} = \max_{\alpha \in [0, H]} \{z_{h,i}^{\pi}(\alpha) - \rho \alpha\}$, $z_{h,i}^{\pi}(\alpha) = \mathbb{E}^{\mu_{h,i}^0} [V_{h+1}^{\pi, \rho}(s')]_\alpha$ and $\alpha \in [0, H]$ is the dual vari-
 158 able derived from the dual formulation (see Proposition H.1 for more details). Moreover, the robust
 159 Bellman optimality equation (2.2) shows that the greedy policy with respect to the optimal robust
 160 Q-function is exactly the optimal robust policy π^* . Therefore, the core idea behind the algorithm
 161 design is to estimate the optimal robust Q-function using linear function approximation, and then
 162 find π^* by the greedy policy derived from the estimated optimal robust Q-function. We present our
 163 algorithm in Algorithm 1.

164 3.1 Variance-weighted ridge regression for online DRMDPs

165 Algorithm 1 is a value-iteration based algorithm that iteratively estimates the robust Q-function
 166 through variance-weighted ridge regression. Different from the Q-function estimation for standard
 167 linear MDP, we element-wisely estimate the parameters of robust Q-functions. This is a distinct
 168 feature for linear DRMDPs. We next interpret the details of our algorithm design.

169 From Line 6 to 13 of Algorithm 1, we adopt the backward induction procedure to update the robust
 170 Q-function estimation. In particular, for any $(k, h) \in [K] \times [H]$, suppose we have an estimated
 171 robust value function $\hat{V}_{k,h+1}^\rho$. By the robust Bellman optimality equation (2.2) and Proposition 2.4,
 172 conducting one step backward induction on $\hat{V}_{k,h+1}^\rho$ leads to the following linear form [23]:

$$r_h(s, a) + \inf_{P_h \in \mathcal{U}_h^\rho(s, a; \mu^0)} \mathbb{P}_h[\hat{V}_{k,h+1}](s, a) = \phi(s, a)^\top (\theta_h + \nu_h^{\rho, k}) \mathbb{1}\{s \neq s_f\}, \quad (3.1)$$

173 where $\nu_{h,i}^{\rho, k} := \max_{\alpha \in [0, H]} \{z_{h,i}^k(\alpha) - \rho \alpha\}$ and $z_{h,i}^k(\alpha) := \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}(s')]_\alpha$, for any $i \in [d]$. Note
 174 that under Assumption 2.1, for any $\alpha \in [0, H]$, $z_{h,i}^k(\alpha)$ is the i -th element of the parameter of the

¹Our ‘average sub-optimality’ differs from standard ‘regret’ as it measures the gap to the optimal policy in the worst-case target environment, not the nominal one.

175 following linear formulation, $[\mathbb{P}_h^0[\hat{V}_{k,h+1}]_\alpha](s, a) = \langle \phi(s, a), \mathbf{z}_h^k(\alpha) \rangle$. Thus, we can estimate $\mathbf{z}_h^k(\alpha)$
 176 from data to get estimations of $\mathbf{z}_{h,i}^k(\alpha), \forall i \in [d]$. To this end, we introduce the variance-weighted
 177 ridge regression regime to estimate $\mathbf{z}_h^k(\alpha)$ as follows

$$\min_{\mathbf{z} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \left(\mathbf{z}^\top \phi(s_h^\tau, a_h^\tau) - [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_\alpha \right)^2 + \lambda \|\mathbf{z}\|_2^2,$$

178 which leads to the following closed-form estimation

$$\hat{\mathbf{z}}_h^k(\alpha) = \Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_\alpha, \quad (3.2)$$

179 where $\Sigma_{k,h} = \lambda \mathbf{I} + \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) \phi(s_h^\tau, a_h^\tau)^\top$ and $\bar{\sigma}_{\tau,h}$ is a variance estimator that will be
 180 formally introduced in Section 3.2. We then approximate $\nu_h^{\rho,k}$ by solving the optimization problem
 181 element-wisely

$$\hat{\nu}_{h,i}^{\rho,k} = \max_{\alpha \in [0,H]} \{ \hat{z}_{h,i}^k(\alpha) - \rho \alpha \}, \quad i \in [d]. \quad (3.3)$$

182 Then we can estimate the Q-function via (3.1). Due to the nature of online RL, the estimation might
 183 be highly uncertain due to the lack of exploration, and thus we incorporate a bonus term $\hat{\Gamma}_{k,h}(s, a) =$
 184 $\beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$ into the robust Q-function estimation, where $\beta = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d})$.
 185 The final estimator is given on Line 8 of Algorithm 1. We will show in later analysis that the estimated
 186 Q-function is an optimistic estimator for the optimal robust Q-function.

187 Inspired by [10], we also establish estimators for the lower bound of robust Q-functions by the same
 188 backward induction procedure, which will be helpful in constructing the variance estimator $\bar{\sigma}_{\tau,h}$
 189 as shown in the next section. In particular, given $\hat{V}_{k,h}^\rho$, We obtain the variance weighted regression
 190 estimator $\check{\mathbf{z}}_h^k(\alpha) = \Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) [\check{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_\alpha$. Then we get the estimation

$$\check{\nu}_{h,i}^{\rho,k} = \max_{\alpha \in [0,H]} \{ \check{z}_{h,i}^k(\alpha) - \rho \alpha \}, \quad i \in [d]. \quad (3.4)$$

191 By (3.1) and (3.4), we get another estimation of the robust Q-function, which we aim to show is a
 192 pessimistic estimation of the optimal robust Q-function. Similarly, to quantify the uncertainty caused
 193 by online exploration, we introduce a penalty term $\check{\Gamma}_{k,h}(s, a) = \bar{\beta} \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$, where
 194 $\bar{\beta} = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d^3 H^3})$. The final pessimistic estimator $\check{Q}_{k,h}^\rho$ is shown on Line 9 of Algorithm 1.
 195 Though [24] also constructed pessimistic robust Q-function estimations for DRMDPs, our methods
 196 are very different due to 1) they do not update the estimation episodically, 2) their estimators are used
 197 to get the optimal robust policy estimation, while ours are used to construct the variance estimator, as
 198 shown in the next section.

199 3.2 Estimating the variance of the value function

200 In this section, we construct the variance weight $\bar{\sigma}_{\tau,h}$ used in (3.2) and aim to get an optimistic esti-
 201 mator for the variance of the optimal robust value function, $\mathbb{V}_h V_{h+1}^{*,\rho}$. Due to the distinct element-wise
 202 estimation procedure introduced in the previous section, the coarse variance estimation design in [10]
 203 for standard linear MDPs does not apply. Instead, we need to carefully design the variance estimators
 204 used in weighted ridge-regressions based upon the unique characteristics of linear DRMDPs.

205 We desire to design the variance estimator at episode k to be a uniform variance upper bound for
 206 all subsequent episodes. To obtain the optimistic estimator for $\mathbb{V}_h V_{h+1}^{*,\rho}$, we first solve regression
 207 problems to obtain the estimator for $\mathbb{V}_h \hat{V}_{k,h+1}^\rho$, which is denoted as $\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho$. Then we analyze
 208 the error between $\mathbb{V}_h V_{h+1}^{*,\rho}$ and $\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho$ to finish the construction. Different from (5.2) in [24], the
 209 variance estimator here is not trivially constructed from subtracting a specific penalty term because
 210 we should guarantee the monotonicity of estimated variance for the online exploration. The variance
 211 of estimated optimistic value function $\hat{V}_{k,h+1}^\rho$ can be decomposed into

$$[\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s, a) = [\mathbb{P}_h^0(\hat{V}_{k,h+1}^\rho)^2](s, a) - ([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s, a))^2. \quad (3.5)$$

Algorithm 1 Weighted Distributionally Robust Iterative Value Estimation with UCB (We-DRIVE-U)

```

1: Initialization: confidence parameters  $\beta, \bar{\beta}, \tilde{\beta} > 0$  and regularization  $\lambda > 0$ .  $k_{\text{last}} = 0$ . For
   each stage  $h \in [H]$ , initialize  $\Sigma_{0,h} = \Sigma_{1,h} = \Lambda_{1,h} = \lambda \mathbf{I}$  and the upper and lower estimation
    $\hat{Q}_{0,h}^\rho(\cdot, \cdot) = H, \check{Q}_{0,h}^\rho(\cdot, \cdot) = 0$ 
2: for episode  $k = 1, \dots, K$  do
3:   Receive the initial state  $s_1^k$ 
4:   Set  $\hat{V}_{k,H+1}^\rho(\cdot) \leftarrow 0, \check{V}_{k,H+1}^\rho(\cdot) \leftarrow 0$ 
5:   if there exists a stage  $h' \in [H]$  such that  $\det(\Sigma_{k,h'}) \geq 2\det(\Sigma_{k_{\text{last}},h'})$  then
6:     for stage  $h = H, \dots, 1$  do
7:       For  $h = H, \check{\nu}_h^{\rho,k} \leftarrow 0, \check{\nu}_h^{\rho,k} \leftarrow 0$ ; otherwise compute  $\hat{\nu}_{h,i}^{\rho,k}, \forall i \in [d]$  according to (3.3)
         and  $\check{\nu}_{h,i}^{\rho,k}, \forall i \in [d]$  according to (3.4).
8:        $\hat{Q}_{k,h}^\rho(s, a) \leftarrow \min \{r_h(s, a) + \phi(s, a)^\top \hat{\nu}_h^{\rho,k} + \hat{\Gamma}_{k,h}(s, a), \hat{Q}_{k-1,h}^\rho(s, a), H-h+1\} \mathbb{1}\{s \neq s_f\}$ 
9:        $\check{Q}_{k,h}^\rho(s, a) \leftarrow \max \{r_h(s, a) + \phi(s, a)^\top \check{\nu}_h^{\rho,k} - \check{\Gamma}_{k,h}(s, a), \check{Q}_{k-1,h}^\rho(s, a), 0\} \mathbb{1}\{s \neq s_f\}$ 
10:      Set the last updating episode  $k_{\text{last}} \leftarrow k$ 
11:       $\hat{V}_{k,h}^\rho(s) \leftarrow \max_a \hat{Q}_{k,h}^\rho(s, a), \check{V}_{k,h}^\rho(s) \leftarrow \max_a \check{Q}_{k,h}^\rho(s, a)$ 
12:       $\pi_h^k(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} \hat{Q}_{k,h}^\rho(s, a)$ 
13:    end for
14:  else
15:     $\hat{V}_{k,h}^\rho, \check{V}_{k,h}^\rho, \pi_h^k \leftarrow \hat{V}_{k-1,h}^\rho, \check{V}_{k-1,h}^\rho, \pi_h^{k-1}, \forall h \in [H]$ 
16:  end if
17:  for stage  $h = 1, \dots, H$  do
18:    Take  $a_h^k \leftarrow \pi_h^k(s_h^k)$  and receive  $s_{h+1}^k$ 
19:    Calculate the estimated variance  $\sigma_{k,h}$  according to (3.6) and  $\bar{\sigma}_{k,h}$  according to (3.7)
20:     $\Sigma_{k+1,h} \leftarrow \Sigma_{k,h} + \bar{\sigma}_{k,h}^{-2} \phi(s_h^k, a_h^k) \phi(s_h^k, a_h^k)^\top, \Lambda_{k+1,h} \leftarrow \Lambda_{k,h} + \phi(s_h^k, a_h^k) \phi(s_h^k, a_h^k)^\top$ 
21:  end for
22: end for

```

212 Under Assumption 2.1, $\mathbb{P}_h^0(\hat{V}_{k,h+1}^\rho)^2$ and $\mathbb{P}_h^0\hat{V}_{k,h+1}^\rho$ on the RHS of (3.5) are linear in $\phi(s, a)$. Thus we
 213 can approximate the variance as $[\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s, a) \approx [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s, a) = [\phi(s, a)^\top \tilde{\mathbf{w}}_{h,2}^k]_{[0,H^2]} -$
 214 $[\phi(s, a)^\top \tilde{\mathbf{w}}_{h,1}^k]_{[0,H^2]}^2$, where $\tilde{\mathbf{w}}_{h,1}^k = \min_{\mathbf{w} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} (\mathbf{w}^\top \phi(s_h^\tau, a_h^\tau) - \hat{V}_{k,h+1}^\rho(s_h^\tau))^2 + \lambda \|\mathbf{w}\|_2^2$
 215 and $\tilde{\mathbf{w}}_{h,2}^k = \min_{\mathbf{w} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} (\mathbf{w}^\top \phi(s_h^\tau, a_h^\tau) - (\hat{V}_{k,h+1}^\rho(s_h^\tau))^2)^2 + \lambda \|\mathbf{w}\|_2^2$. Different from the
 216 variance estimation in standard MDPs [10], we construct both $\tilde{\mathbf{w}}_{h,2}^k$ and $\tilde{\mathbf{w}}_{h,1}^k$ by solving vanilla
 217 ridge regressions, instead of variance-weighted ridge regressions. This specific choice of parameter
 218 estimation will simplify our analysis of the variance estimation error, while fully capture the variance
 219 information. Now we can construct $\sigma_{k,h}$, which is the estimated variance of the optimal robust value
 220 function $V_h^{*,\rho}$ in episode k , as follows

$$\sigma_{k,h} = \sqrt{[\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) + E_{k,h} + d^3 H D_{k,h} + 1/2}, \quad (3.6)$$

221 where $E_{k,h}$ represents the error between the estimated variance and the true variance of $\hat{V}_{k,h+1}^\rho$,
 222 and $D_{k,h}$ represents the error between the true variance of $\hat{V}_{k,h+1}^\rho$ and the true variance of $V_{h+1}^{*,\rho}$.
 223 Formally, we define

$$\begin{aligned}
 E_{k,h} &= \min \{ \tilde{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \} + \min \{ 2H \bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \}, \\
 D_{k,h} &= \min \{ 4H (\phi(s_h^k, a_h^k)^\top \tilde{\mathbf{w}}_{h,1}^k - \phi(s_h^k, a_h^k)^\top \tilde{\mathbf{w}}_{h,2}^k + 2\bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}), H^2 \},
 \end{aligned}$$

224 where $\Lambda_{k,h} = \lambda \mathbf{I} + \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) \phi(s_h^\tau, a_h^\tau)^\top$, $\bar{\beta} = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d^3 H^3})$, $\tilde{\beta} = \tilde{\mathcal{O}}(H^2\sqrt{d\lambda} +$
 225 $\sqrt{d^3 H^6})$, and $\tilde{\mathbf{w}}_{h,1}^k = \min_{\mathbf{w} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} (\mathbf{w}^\top \phi(s_h^\tau, a_h^\tau) - \hat{V}_{k,h+1}^\rho(s_h^\tau))^2 + \lambda \|\mathbf{w}\|_2^2$. Finally, we
 226 construct weights for the variance-weighted ridge regression problem (3.2): $\forall k, h \in [K] \times [H]$,

$$\bar{\sigma}_{k,h} = \max \{ \sigma_{k,h}, 1, \sqrt{2d^3 H^2} \|\phi(s_h^k, a_h^k)\|_{\Sigma_{k,h}^{-1}}^{1/2} \}. \quad (3.7)$$

227 Compared with the variance estimation for the value function in standard MDPs [10] which is
 228 $\bar{\sigma}_{k,h} = \max \{ \sigma_{k,h}, H, 2d^3 H^2 \|\phi(s_h^k, a_h^k)\|_{\Sigma_{k,h}^{-1}}^{1/2} \}$, our variance estimation is tighter in both the second
 229 and the third terms. The second term in (3.7) is 1, instead of H . This is important in achieving a
 230 tighter dependence on H . The intuition is that, when k is large, $\bar{\sigma}_{k,h}$ should be close to the variance
 231 of the optimal robust value function. This design is motivated by the ‘Range Shrinkage’ phenomenon
 232 unique to DRMDPs (see the discussion after Theorem 4.1 for details), which observes that the true
 233 variance is in the order of $\mathcal{O}(1)$ when $\rho = \mathcal{O}(1)$. To get a precise variance estimation, $\bar{\sigma}_{k,h}$ should
 234 be in the same order of the true variance. Moreover, a constant order lower bound on $\bar{\sigma}_{k,h}$ will also
 235 ensure the weight will not cause any inflation in the weighted regression (3.2). The third term in (3.7)
 236 to be also tighter than that of [10], while maintaining the same theoretical property.

237 3.3 Algorithm interpretation

238 Now we provide some discussions to interpret Algorithm 1.

239 **Remark 3.1.** We highlight that Algorithm 1 is the first algorithm adopting a ‘rare-switching’ update
 240 strategy for distributionally robust RL. Different from [10], the ‘rare-switching’ condition on Line 5
 241 is set at the beginning of each episode. This is achieved by our variance estimator design, which is
 242 independent of the parameter update for $z_h^k(\alpha)$. The update rule on Line 5 determines whether to
 243 update robust Q-function estimations and switch to a new policy for the current episode, and leads to
 244 two advantages, 1) the number of times solving the ridge regression (3.2) and dual optimization (3.3)
 245 significantly decreases, which constitute the main computation cost of Algorithm 1, and 2) in real
 246 application scenarios where policy switching is costly or risky, Algorithm 1 possesses low policy
 247 switching property. We refer the readers to Proposition 4.7 and Remark 4.8 for more details.

248 **Remark 3.2.** On Line 7, we estimate $\nu_h^{\rho,k}$ element-wisely, and thus the estimator $\hat{\nu}_h^{\rho,k}$ is derived
 249 from d separate variance-weighted ridge regressions (3.2) and dual optimizations (3.3). This leads
 250 to the specific form of bonus term $\hat{\Gamma}_{k,h}(s, a) = \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$, which is actually an
 251 upper bound of the robust estimation error (see Lemma D.4 and its proof) at episode k . Though
 252 the bonus term resembles that in [23], we highlight that the sampling covariance matrix $\Sigma_{k,h}$ in
 253 $\hat{\Gamma}_{k,h}(s, a)$ is indeed a variance-weighted one. The specific form of the bonus term leads to the new
 254 variance-weighted d -rectangular robust estimation error defined in (4.1).

255 **Remark 3.3.** On Line 8 and 9, we adopt a monotonic Q-function update strategy, such that the
 256 estimated optimistic (pessimistic) robust value function is monotonically decreasing (increasing) to
 257 the optimal robust value function. This strategy is to make sure that the variance estimator $\sigma_{k,h}$ at
 258 any episode $k \in [H]$ is a uniform upper bound for those in the subsequent episodes, which would be
 259 helpful in bounding the estimation error arising from the variance-weighted ridge regression (3.2).
 260 This idea is first introduced by [2] for standard tabular MDPs and then utilized by [10] for standard
 261 linear MDPs. This is the first time it is utilized in the online linear DRMDP setting, where the
 262 episodic estimation regime proposes additional requirement on the variance estimator construction
 263 compared to the offline setting studied by [24].

264 4 Suboptimality upper bound analysis

265 We now provide theoretical results on the upper bound on the suboptimality of Algorithm 1.

266 **Theorem 4.1.** Under Assumptions 2.1 and 2.3, set $\lambda = 1/H^2$, then for any fixed $\delta \in (0, 1)$ and
 267 $\rho \in (0, 1]$, with probability at least $1 - \delta$, the average suboptimality of We-DRIVE-U satisfies

$$\text{AveSubopt}(K) \leq 2\sqrt{2H^3 \log(6/\delta)/K} + \underbrace{\frac{4\beta}{K} \sum_{k=1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_{h,i}^k \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}}_{\text{variance-weighted } d\text{-rectangular estimation error}}, \quad (4.1)$$

268 where $\beta = \tilde{\mathcal{O}}(\sqrt{d})$, $\phi_{h,i}^k$ is the i -th element of $\phi(s_h^k, a_h^k)$ and $\mathbf{1}_i$ is the i -th standard basis vector.

269 Recall from Remark 3.2, the quantity $\sum_{i=1}^d \phi_{h,i}^k \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$ in (4.1) comes from solving d separate
 270 variance-weighted ridge regressions at step h in episode k . A similar term also appears in the
 271 Theorem 5.1 of [23]. Differently, the quantity $\sum_{i=1}^d \phi_{h,i}^k \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$ is based on the variance-
 272 weighted sampling covariance matrix $\Sigma_{k,h}$, rather than the vanilla sampling covariance matrix

273 $\Lambda_{k,h}$ as in [23]. In order to further bound (4.1), we need to take a closer examination of the
 274 variance estimator. Intuitively, when episode k is large, the variance estimator should be close to
 275 the variance of the optimal robust value function. Recent study [24] shows a ‘Range shrinkage’
 276 phenomenon in the d -rectangular linear DRMDP (refer to Lemma H.10), stating that the range of any
 277 robust value function satisfies $\max_{s \in \mathcal{S}} V_h^{\pi, \rho}(s) - \min_{s \in \mathcal{S}} V_h^{\pi, \rho}(s) \leq \min\{1/\rho, H\}$, $\forall (\pi, h, \rho) \in$
 278 $\Pi \times [H] \times (0, 1]$. This implies that the variance of the optimal robust value function is upper
 279 bounded by $\min\{1/\rho, H\}$. Thus, when k is large, we can expect $\bar{\sigma}_{k,h} \lesssim \tilde{\mathcal{O}}(\min\{1/\rho, H\})$ and
 280 hence $\Sigma_{k,h}^{-1} \preceq \tilde{\mathcal{O}}(\min\{1/\rho^2, H^2\})\Lambda_{k,h}^{-1}$. To this end, next we rigorously bound (4.1) under the same
 281 setting as the Corollary 5.3 of [23], and formally show that the variance information leads to a tighter
 282 dependence on H compared to [23].

283 **Theorem 4.2.** Assume that there exists an absolute constant $c > 0$, such that for all $(\pi, h) \in \Pi \times [H]$

$$\mathbb{E}_{\pi}^{P^0} [\phi(s_h, a_h) \phi(s_h, a_h)^{\top}] \geq c/d \cdot \mathbf{I}. \quad (4.2)$$

284 Then under the same setting in Theorem 4.1 and the assumption in (4.2), for any fixed $\delta \in (0, 1)$,
 285 with probability at least $1 - \delta$, the average suboptimality of We-DRIVE-U is upper bounded by
 286 $\tilde{\mathcal{O}}((dH \cdot \min\{1/\rho, H\} + H^{3/2})/\sqrt{K} + d^{15}H^{13}/K)$.

287 **Remark 4.3.** When $d \geq H$ and the total number of episodes K is sufficiently large, the average
 288 suboptimality can be simplified as $\tilde{\mathcal{O}}\{dH \min\{1/\rho, H\}/\sqrt{K}\}$. Note that under the same assumption
 289 in (4.2), [23] prove that the average suboptimality of their algorithm DR-LSVI-UCB is of the order
 290 $\tilde{\mathcal{O}}(d^2H^2/\sqrt{K})$. Thus, We-DRIVE-U improves the state-of-the-art result by $\mathcal{O}(dH/\min\{1/\rho, H\})$.
 291 Moreover, we highlight that the upper bound in Theorem 4.2 depends on the uncertainty level ρ , which
 292 arises from the ‘Range Shrinkage’ phenomenon. When ρ increases from 0 to 1, the suboptimality
 293 decreases up to a factor of $\mathcal{O}(H)$.

294 **Remark 4.4.** The assumption (4.2) is actually imposed on the DRMDP, requiring that the environ-
 295 ment we encounter is exploratory enough. We would like to note that this assumption is necessary
 296 in deriving our upper bound, since the elliptical potential lemma [1, Lemma 11], which is critical
 297 in deriving upper bounds in linear bandits and linear MDPs, does not apply in the analysis of linear
 298 DRMDPs. We note that the previous work [23] also used this assumption to get the final upper
 299 bound for their algorithm. Moreover, the assumption (4.2) can be deemed as an online version of
 300 the well-known full-type coverage assumption on the offline dataset in offline (non-) robust RL.
 301 Specifically, in the context of standard offline RL, [5, 44, 47] assume the offline dataset should cover
 302 the distribution measure induced by any policy under the nominal environment. In the context of
 303 offline robust RL, [32, 31, 57] assume that the offline dataset should cover the distribution measure
 304 induced by any policy under any transition kernel in the uncertainty set. It would be an interesting
 305 future research direction to study if assumption (4.2) can be relaxed.

306 Next, we study the deployment complexity of Algorithm 1, which constitutes two sources of cost.
 307 The first source is the policy switching cost, say, the total number of changes in the exploration policy.
 308 This might be the main bottleneck in applications where changing the exploration policy is costly
 309 or risky [3, 45]. The second source is the computation cost in solving the dual optimization in (3.3).
 310 Recall in Remark 4.8 we discuss that Algorithm 1 adopts the ‘rare-switching’ update strategy, which
 311 significantly reduces the two sources of cost. Next, we formally define them as follows.

312 **Definition 4.5** (Global Switching Cost). We define the *global switching cost* of an algorithm that
 313 runs for K episodes as $N_{\text{switch}}^{\text{gl}} := \sum_{k=1}^K \mathbb{1}\{\pi_k \neq \pi_{k+1}\}$.

314 **Definition 4.6** (Dual Oracle). We assume access to a maximization oracle, which takes a function
 315 $z : [0, H] \rightarrow \mathbb{R}$ and a fixed constant $\rho > 0$ as input, and outputs the maximum value $z_{\max}$ and the
 316 maximizer $\alpha_{\max}$ defined as $z_{\max} = \max_{\alpha \in [0, H]} \{z(\alpha) - \rho\alpha\}$ and $\alpha_{\max} = \arg\max_{\alpha \in [0, H]} \{z(\alpha) -$
 317 $\rho\alpha\}$. For an algorithm, we define the *oracle complexity* as the number of calls of the dual oracle.
 318 Finally, we show that We-DRIVE-U admits low switching cost and low oracle complexity.

319 Next, we formally present theoretical results on the deployment complexity of Algorithm 1.

320 **Proposition 4.7.** Under the same setting as Theorem 4.1, the switching cost of We-DRIVE-U is
 321 upper bounded by $dH \log(1 + H^2K)$, and the oracle complexity of We-DRIVE-U is upper bounded
 322 by $2d^2H \log(1 + H^2K)$.

323 **Remark 4.8.** The switching cost of the state-of-the-art algorithm DR-LSVI-UCB [23] is K and the
 324 oracle complexity is dK . Thus, We-DRIVE-U improves both the switching and oracle cost by a factor

of K . Different from standard linear MDPs, where the main computation complexity only comes from the policy update [10], in the linear DRMDP setting, the calls of dual oracle, besides policy updates, are also a main source of computational burden. The update rule on Line 5 guarantees that We-DRIVE-U calls the dual oracle and updates the policy only when the criterion is met. Actually, Algorithm 1 is the first DRMDP algorithm that admits low deployment complexity.

5 Information-theoretic lower bound analysis

According to Theorem 4.2, when $\rho = \mathcal{O}(1)$, the suboptimality of We-DRIVE-U is of order $\mathcal{O}(dH/\sqrt{K})$. After multiplying K to recover the cumulative suboptimality, it is smaller than the minimax lower bound for standard linear MDP, $\Omega(d\sqrt{H^3K})$ [62]. To assess the optimality of We-DRIVE-U, next we present an information-theoretic lower bound for online linear DRMDPs.

Theorem 5.1. Let uncertainty level $\rho \in (0, 3/4]$, $H \geq 6$, and $K \geq 9d^2H/32$. Then for any algorithm, there exists a d -rectangular linear DRMDP parameterized by $\xi = (\xi_1, \dots, \xi_{H-1})$ such that the expected average suboptimality is lower bounded as follows:

$$\mathbb{E}_\xi[\text{AveSubopt}(M_\xi, K)] \geq \Omega((dH^{1/2} \cdot \min\{1/\rho, H\})/\sqrt{K}), \quad (5.1)$$

where $\mathbb{E}_\xi$ denotes the expectation over the randomness of the algorithm and the nominal environment.

Remark 5.2. We highlight that the lower bound (5.1) depends on the uncertainty level ρ , which is a distinctive characteristic for DRMDPs. Theorem 5.1 implies We-DRIVE-U is near-optimal up to a factor of $\tilde{\mathcal{O}}(\sqrt{H})$. Moreover, when $\rho \rightarrow 0$, the linear DRMDP degrades to a standard linear MDP, and (5.1) matches the information-theoretic lower bound, $\Omega(d\sqrt{H^3K})$, for standard linear MDPs [62] after multiplying K to recover the cumulative regret. When $\rho = \mathcal{O}(1)$, (5.1) reduces to $\Omega(dH^{1/2}/\sqrt{K})$, which is $\mathcal{O}(H)$ smaller than the lower bound for standard linear MDPs.

Next, we investigate the $\tilde{\mathcal{O}}(\sqrt{H})$ gap between the upper and lower bounds and propose a conjecture on its origin. In the analysis of non-robust MDPs [2, 18, 10] and tabular DRMDPs [25], a tight dependence on H is often achieved by exploiting the total variance law of the value function at each episode. Currently, we bound each term in the variance-weighted d -rectangular estimation error in (4.1) separately. A tight upper bound might be achieved by first bounding the variance-weighted d -rectangular estimation error as a whole by the square root of the total variance and then invoking the total variance law. In particular, inspired by the total variance law in Lemma C.6 of [25], the total variance should be in the order of $\mathcal{O}(H \min\{1/\rho, H\})$. Together with an additional $\sqrt{H}$ arising in the suboptimality analysis, we conjecture the dependence of the upper bound on H could be improved to $\mathcal{O}(\sqrt{H^2 \min\{1/\rho, H\}})$.

When the uncertainty level is small, i.e., $\rho = \mathcal{O}(1/H)$, the conjectured result leads to an upper bound on the suboptimality that depends on $\mathcal{O}(H^{3/2})$, matching the current lower bound we present in (5.1). This suggests that our current lower bound is tight and the total variance analysis could improve our upper bound. When the uncertainty level is relatively large, i.e., $\rho = \mathcal{O}(1)$, the conjectured upper bound is $\mathcal{O}(H)$, which matches the current upper bound in Theorem 4.2. This means the total variance analysis does not further improve the upper bound, and we suspect that a tighter lower bound is instead necessary. This leaves an interesting open problem for future study.

6 Conclusion

This paper advanced the study of online d -rectangular linear DRMDPs by establishing a tighter regret upper bound and the first lower regret bound under this setting. We introduced We-DRIVE-U, a novel variance-aware algorithm that leverages variance-weighted ridge regression and low policy-switching techniques. Under standard MDP structure assumptions, we proved We-DRIVE-U achieves an average suboptimality of $\tilde{\mathcal{O}}(dH \min\{1/\rho, H\}/\sqrt{K})$, improving the state-of-the-art by $\tilde{\mathcal{O}}(dH/\min\{1/\rho, H\})$. We also established an information-theoretic lower bound of $\Omega(dH^{1/2} \min\{1/\rho, H\}/\sqrt{K})$, which implies We-DRIVE-U's near-optimality up to $\mathcal{O}(\sqrt{H})$. Furthermore, We-DRIVE-U reduces computational complexity with $\mathcal{O}(dH \log(1 + H^2K))$ policy switches and $\mathcal{O}(d^2H \log(1 + H^2K))$ oracle complexity, which outperforms existing methods by a factor of K . We also conducted numerical experiments to validate the robustness and improved performance over existing algorithms, which is presented in Appendix B.

References

- [1] Yasin Abbasi-Yadkori, Dávid Pál, and Csaba Szepesvári. Improved algorithms for linear stochastic bandits. *Advances in neural information processing systems*, 24, 2011.
- [2] Mohammad Gheshlaghi Azar, Ian Osband, and Rémi Munos. Minimax regret bounds for reinforcement learning. In *International Conference on Machine Learning*, pages 263–272. PMLR, 2017.
- [3] Yu Bai, Tengyang Xie, Nan Jiang, and Yu-Xiang Wang. Provably efficient q-learning with low switching cost. *Advances in Neural Information Processing Systems*, 32, 2019.
- [4] Jose Blanchet, Miao Lu, Tong Zhang, and Han Zhong. Double pessimism is provably efficient for distributionally robust offline reinforcement learning: Generic algorithm and robust partial coverage. *arXiv preprint arXiv:2305.09659*, 2023.
- [5] Jinglin Chen and Nan Jiang. Information-theoretic considerations in batch reinforcement learning. In *International Conference on Machine Learning*, pages 1042–1051. PMLR, 2019.
- [6] Jing Dong, Jingwei Li, Baoxiang Wang, and Jingzhao Zhang. Online policy optimization for robust mdp. *arXiv preprint arXiv:2209.13841*, 2022.
- [7] Benjamin Eysenbach, Swapnil Asawa, Shreyas Chaudhari, Sergey Levine, and Ruslan Salakhutdinov. Off-dynamics reinforcement learning: Training for transfer with domain classifiers. *arXiv preprint arXiv:2006.13916*, 2020.
- [8] Jesse Farebrother, Marlos C Machado, and Michael Bowling. Generalization and regularization in dqn. *arXiv preprint arXiv:1810.00123*, 2018.
- [9] Vineet Goyal and Julien Grand-Clement. Robust markov decision processes: Beyond rectangularity. *Mathematics of Operations Research*, 48(1):203–226, 2023.
- [10] Jiafan He, Heyang Zhao, Dongruo Zhou, and Quanquan Gu. Nearly minimax optimal reinforcement learning for linear markov decision processes. In *International Conference on Machine Learning*, pages 12790–12822. PMLR, 2023.
- [11] Jiafan He, Dongruo Zhou, and Quanquan Gu. Logarithmic regret for reinforcement learning with linear function approximation. In *International Conference on Machine Learning*, pages 4171–4180. PMLR, 2021.
- [12] Hao-Lun Hsu, Weixin Wang, Miroslav Pajic, and Pan Xu. Randomized exploration in cooperative multi-agent reinforcement learning. *arXiv preprint arXiv:2404.10728*, 2024.
- [13] Pihe Hu, Yu Chen, and Longbo Huang. Nearly minimax optimal reinforcement learning with linear function approximation, 2023.
- [14] Haque Ishfaq, Qiwen Cui, Viet Nguyen, Alex Ayoub, Zhuoran Yang, Zhaoran Wang, Doina Precup, and Lin Yang. Randomized exploration in reinforcement learning with general value function approximation. In *International Conference on Machine Learning*, pages 4607–4616. PMLR, 2021.
- [15] Haque Ishfaq, Qingfeng Lan, Pan Xu, A Rupam Mahmood, Doina Precup, Anima Anandkumar, and Kamyar Azizzadenesheli. Provable and practical: Efficient exploration in reinforcement learning via langevin monte carlo. *arXiv preprint arXiv:2305.18246*, 2023.
- [16] Garud N Iyengar. Robust dynamic programming. *Mathematics of Operations Research*, 30(2):257–280, 2005.
- [17] Yifeng Jiang, Tingnan Zhang, Daniel Ho, Yunfei Bai, C Karen Liu, Sergey Levine, and Jie Tan. Simgan: Hybrid simulator identification for domain adaptation via adversarial reinforcement learning. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 2884–2890. IEEE, 2021.
- [18] Chi Jin, Zeyuan Allen-Zhu, Sebastien Bubeck, and Michael I Jordan. Is q-learning provably efficient? *Advances in neural information processing systems*, 31, 2018.

- [19] Chi Jin, Zhuoran Yang, Zhaoran Wang, and Michael I Jordan. Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pages 2137–2143. PMLR, 2020.
- [20] Yeoneung Kim, Insoon Yang, and Kwang-Sung Jun. Improved regret analysis for variance-adaptive linear bandits and horizon-free linear mixture mdps. *Advances in Neural Information Processing Systems*, 35:1060–1072, 2022.
- [21] Eric B Laber, Nick J Meyer, Brian J Reich, Krishna Pacifici, Jaime A Collazo, and John M Drake. Optimal treatment allocations in space and time for on-line control of an emerging infectious disease. *Journal of the Royal Statistical Society Series C: Applied Statistics*, 67(4):743–789, 2018.
- [22] Zhishuai Liu, Jesse Clifton, Eric B Laber, John Drake, and Ethan X Fang. Deep spatial q-learning for infectious disease control. *Journal of Agricultural, Biological and Environmental Statistics*, pages 1–25, 2023.
- [23] Zhishuai Liu and Pan Xu. Distributionally robust off-dynamics reinforcement learning: Provable efficiency with linear function approximation. *arXiv preprint arXiv:2402.15399*, 2024.
- [24] Zhishuai Liu and Pan Xu. Minimax optimal and computationally efficient algorithms for distributionally robust offline reinforcement learning. *arXiv preprint arXiv:2403.09621*, 2024.
- [25] Miao Lu, Han Zhong, Tong Zhang, and Jose Blanchet. Distributionally robust reinforcement learning with interactive data collection: Fundamental hardness and near-optimal algorithm. *arXiv preprint arXiv:2404.03578*, 2024.
- [26] Xiaoteng Ma, Zhipeng Liang, Li Xia, Jiheng Zhang, Jose Blanchet, Mingwen Liu, Qianchuan Zhao, and Zhengyuan Zhou. Distributionally robust offline reinforcement learning with linear function approximation. *arXiv preprint arXiv:2209.06620*, 2022.
- [27] Shie Mannor, Ofir Mebel, and Huan Xu. Robust mdps with k-rectangular uncertainty. *Mathematics of Operations Research*, 41(4):1484–1509, 2016.
- [28] Aditya Modi, Nan Jiang, Ambuj Tewari, and Satinder Singh. Sample complexity of reinforcement learning using linearly combined model ensembles. In *International Conference on Artificial Intelligence and Statistics*, pages 2010–2020. PMLR, 2020.
- [29] Arnab Nilim and Laurent El Ghaoui. Robust control of markov decision processes with uncertain transition matrices. *Operations Research*, 53(5):780–798, 2005.
- [30] Kishan Panaganti and Dileep Kalathil. Sample complexity of robust reinforcement learning with a generative model. In *International Conference on Artificial Intelligence and Statistics*, pages 9582–9602. PMLR, 2022.
- [31] Kishan Panaganti, Adam Wierman, and Eric Mazumdar. Model-free robust ϕ -divergence reinforcement learning using both offline and online data. *arXiv preprint arXiv:2405.05468*, 2024.
- [32] Kishan Panaganti, Zaiyan Xu, Dileep Kalathil, and Mohammad Ghavamzadeh. Robust reinforcement learning using offline data. *Advances in neural information processing systems*, 35:32211–32224, 2022.
- [33] Xue Bin Peng, Marcin Andrychowicz, Wojciech Zaremba, and Pieter Abbeel. Sim-to-real transfer of robotic control with dynamics randomization. In *2018 IEEE international conference on robotics and automation (ICRA)*, pages 3803–3810. IEEE, 2018.
- [34] Yi Shen, Pan Xu, and Michael Zavlanos. Wasserstein distributionally robust policy evaluation and learning for contextual bandits. *Transactions on Machine Learning Research*, 2024. Featured Certification.
- [35] Laixi Shi and Yuejie Chi. Distributionally robust model-based offline reinforcement learning with near-optimal sample complexity. *arXiv preprint arXiv:2208.05767*, 2022.

- [36] Laixi Shi, Gen Li, Yuting Wei, Yuxin Chen, Matthieu Geist, and Yuejie Chi. The curious price of distributional robustness in reinforcement learning with a generative model. *arXiv preprint arXiv:2305.16589*, 2023.
- [37] Richard S Sutton and Andrew G Barto. *Reinforcement learning: An introduction*. MIT press, 2018.
- [38] Cheng Tang, Zhishuai Liu, and Pan Xu. Robust offline reinforcement learning with linearly structured f -divergence regularization. *arXiv preprint arXiv:2411.18612*, 2024.
- [39] Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge university press, 2018.
- [40] Andrew J Wagenmaker, Yifang Chen, Max Simchowitz, Simon Du, and Kevin Jamieson. Reward-free rl is no harder than reward-aware rl in linear markov decision processes. In *International Conference on Machine Learning*, pages 22430–22456. PMLR, 2022.
- [41] He Wang, Laixi Shi, and Yuejie Chi. Sample complexity of offline distributionally robust linear markov decision processes. *arXiv preprint arXiv:2403.12946*, 2024.
- [42] Ruhan Wang, Yu Yang, Zhishuai Liu, Dongruo Zhou, and Pan Xu. Return augmented decision transformer for off-dynamics reinforcement learning. *arXiv preprint arXiv:2410.23450*, 2024.
- [43] Ruosong Wang, Simon S Du, Lin Yang, and Russ R Salakhutdinov. On reward-free reinforcement learning with linear function approximation. *Advances in neural information processing systems*, 33:17816–17826, 2020.
- [44] Ruosong Wang, Dean P Foster, and Sham M Kakade. What are the statistical limits of offline rl with linear function approximation? *arXiv preprint arXiv:2010.11895*, 2020.
- [45] Tianhao Wang, Dongruo Zhou, and Quanquan Gu. Provably efficient reinforcement learning with linear function approximation under adaptivity constraints. *Advances in Neural Information Processing Systems*, 34:13524–13536, 2021.
- [46] Wolfram Wiesemann, Daniel Kuhn, and Berç Rustem. Robust markov decision processes. *Mathematics of Operations Research*, 38(1):153–183, 2013.
- [47] Tengyang Xie, Ching-An Cheng, Nan Jiang, Paul Mineiro, and Alekh Agarwal. Bellman-consistent pessimism for offline reinforcement learning. *Advances in neural information processing systems*, 34:6683–6694, 2021.
- [48] Huan Xu and Shie Mannor. The robustness-performance tradeoff in markov decision processes. *Advances in Neural Information Processing Systems*, 19, 2006.
- [49] Zaiyan Xu, Kishan Panaganti, and Dileep Kalathil. Improved sample complexity bounds for distributionally robust reinforcement learning. In *International Conference on Artificial Intelligence and Statistics*, pages 9728–9754. PMLR, 2023.
- [50] Lin Yang and Mengdi Wang. Reinforcement learning in feature space: Matrix bandit, kernels, and regret bound. In *International Conference on Machine Learning*, pages 10746–10756. PMLR, 2020.
- [51] Wenhao Yang, Han Wang, Tadashi Kozuno, Scott M Jordan, and Zhihua Zhang. Avoiding model estimation in robust markov decision processes with a generative model. *arXiv preprint arXiv:2302.01248*, 2023.
- [52] Wenhao Yang, Liangyu Zhang, and Zhihua Zhang. Toward theoretical understandings of robust markov decision processes: Sample complexity and asymptotics. *The Annals of Statistics*, 50(6):3223–3248, 2022.
- [53] Zhouhao Yang, Yihong Guo, Pan Xu, Anqi Liu, and Animashree Anandkumar. Distributionally robust policy gradient for offline contextual bandits. In *International Conference on Artificial Intelligence and Statistics*, pages 6443–6462. PMLR, 2023.

- 514 [54] Pengqian Yu and Huan Xu. Distributionally robust counterpart in markov decision processes.
515 *IEEE Transactions on Automatic Control*, 61(9):2538–2543, 2015.
- 516 [55] Andrea Zanette, David Brandfonbrener, Emma Brunskill, Matteo Pirotta, and Alessandro
517 Lazaric. Frequentist regret bounds for randomized least-squares value iteration. In *International*
518 *Conference on Artificial Intelligence and Statistics*, pages 1954–1964. PMLR, 2020.
- 519 [56] Huan Zhang, Hongge Chen, Duane Boning, and Cho-Jui Hsieh. Robust reinforcement learning
520 on state observations with learned optimal adversary. *arXiv preprint arXiv:2101.08452*, 2021.
- 521 [57] Runyu Zhang, Yang Hu, and Na Li. Soft robust mdps and risk-sensitive mdps: Equivalence,
522 policy gradient, and sample complexity. In *The Twelfth International Conference on Learning*
523 *Representations*.
- 524 [58] Zihan Zhang, Jiaqi Yang, Xiangyang Ji, and Simon S Du. Improved variance-aware confidence
525 sets for linear bandits and linear mixture mdp. *Advances in Neural Information Processing*
526 *Systems*, 34:4342–4355, 2021.
- 527 [59] Heyang Zhao, Jiafan He, Dongruo Zhou, Tong Zhang, and Quanquan Gu. Variance-dependent
528 regret bounds for linear bandits and reinforcement learning: Adaptivity and computational
529 efficiency. *arXiv preprint arXiv:2302.10371*, 2023.
- 530 [60] Wenshuai Zhao, Jorge Peña Queralta, and Tomi Westerlund. Sim-to-real transfer in deep
531 reinforcement learning for robotics: a survey. In *2020 IEEE symposium series on computational*
532 *intelligence (SSCI)*, pages 737–744. IEEE, 2020.
- 533 [61] Dongruo Zhou and Quanquan Gu. Computationally efficient horizon-free reinforcement learning
534 for linear mixture mdps, 2022.
- 535 [62] Dongruo Zhou, Quanquan Gu, and Csaba Szepesvari. Nearly minimax optimal reinforcement
536 learning for linear mixture markov decision processes. In *Conference on Learning Theory*,
537 pages 4532–4576. PMLR, 2021.
- 538 [63] Zhengqing Zhou, Zhengyuan Zhou, Qinxun Bai, Linhai Qiu, Jose Blanchet, and Peter Glynn.
539 Finite-sample regret bound for distributionally robust offline tabular reinforcement learning. In
540 *International Conference on Artificial Intelligence and Statistics*, pages 3331–3339. PMLR,
541 2021.

A Related work

Distributionally Robust MDPs There has been a large body of works studying DRMDPs under various settings, for instance, the setting of planning and control [48, 46, 54, 27, 9] where the exact transition model is known, the setting with a generative model [63, 52, 30, 49, 36, 51], the offline setting [32, 35, 4, 38] and the online setting [6, 23, 25]. Among tabular DRMDPs, the most relevant studies to ours are [36, 25]. In particular, [36] studies tabular DRMDPs with TV uncertainty sets. They provide an information-theoretic lower bound, as well as a matching upper bound on the sample complexity. The key message is that the sample complexity bounds depend on the uncertainty level, and when the uncertainty level is of constant order, policy learning in a DRMDP requires less samples than in a standard MDP. Further, [25] studies the online tabular DRMDPs with TV uncertainty sets, they provide an algorithm that achieves the near-optimal sample complexity under a vanishing minimal value assumption to circumvent the curse of support shift.

Online Linear MDPs and Linear DRMDPs The nominal model studied in our paper is assumed to be a linear MDP with a simplex feature space. There is a line of works studying online linear MDPs [50, 19, 28, 55, 43, 11, 40, 15], and the minimax optimality of this setting is studied in the recent work of [10]. In particular, they adopt the variance-weighted ridge regression scheme and the ‘rare-switching’ policy update strategy in their algorithm design. The setting of online linear DRMDP is relatively understudied, with both the lower bound and the near-optimal upper bound remain elusive. Specifically, the only work studies the online linear DRMDP setting is [23]. Under the TV uncertainty set, their algorithm, DR-LSVI-UCB, achieves an average suboptimality of the order $\tilde{O}(d^2 H^2 / \sqrt{K})$. However, recent evidence from studies [24, 41] on offline linear DRMDPs suggests that this rate is far from optimality. In particular, [24] proves that their algorithm, VA-DRPVI, achieves an upper bound on the suboptimality in the order of $\tilde{O}(dH \min\{1/\rho, H\} / \sqrt{K})$. Nonetheless, their algorithm and analysis are based on a pre-collected offline dataset which satisfies some coverage assumption, and thus cannot be utilized in the online setting, where a strategy on data collection is required to deal with the challenge of exploration and exploitation trade-off.

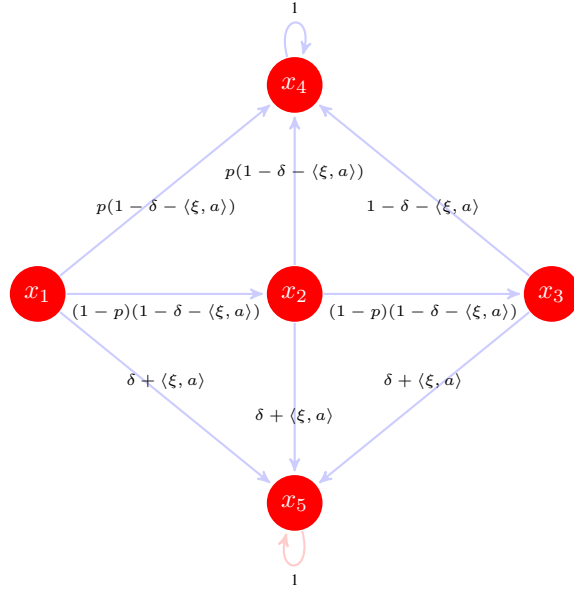
B Experiments on simulated linear DRMDPs

B.1 Simulated linear DRMDPs

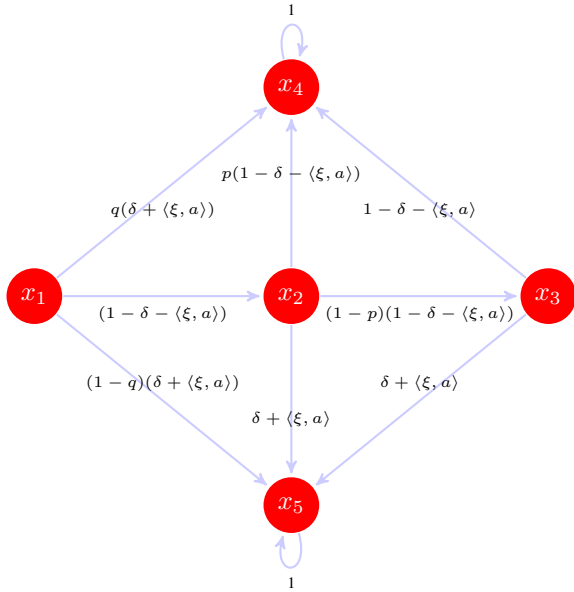
We conduct numerical experiments to illustrate the performances of our proposed algorithm, We-DRIVE-U, and compare it with the state-of-the-art algorithm for d -rectangular linear DRMDPs, DR-LSVI-UCB [23], as well as their non-robust counterpart, LSVI-UCB [19]. All numerical experiments were conducted on a MacBook Pro with a 2.6 GHz 6-Core Intel CPU.

We leverage the simulated linear MDP setting proposed by [23]. For completeness, we recall the experiment setting as follows. The source and target linear MDP environment are shown in Figure 1(a) and Figure 1(b). The state space is $\mathcal{S} = \{x_1, \dots, x_5\}$ and action space $\mathcal{A} = \{-1, 1\}^4 \subset \mathbb{R}^4$. At each episode, the initial state is always x_1 , and it can transit to x_2, x_4, x_5 with probability defined in the figures. x_2 is an intermediate state from which the next state can be x_3, x_4, x_5 . x_4 is the fail state with reward 0 and x_5 is an absorbing state with reward 1. For the reward functions and transition probabilities, they are designed to depend on $\langle \xi, a \rangle$, where $\xi \in \mathbb{R}^4$ is a hyperparameter controls the MDP instances. The target environment is constructed by only perturbing the transition probability at x_1 of the source domain, and the extend of perturbation is controlled by a hyperparameter $q \in (0, 1)$. We refer more details on the construction of the linear DRMDP to the Supplementary A.1 of [23].

We set $\xi = (1/\|\xi\|_1, 1/\|\xi\|_1, 1/\|\xi\|_1, 1/\|\xi\|_1)^\top$ and consider different choices of $\|\xi\|_1$ from the set $\{0.1, 0.2, 0.3\}$. Following the implementation in [23], we use heterogeneous uncertainty level and set $\rho_{1,4} = 0.5$ and $\rho_{h,i} = 0$ for all other cases. We set the number of interactions with the nominal environment to 200. We evaluate policies learned by We-DRIVE-U, DR-LSVI-UCB [23] and LSVI-UCB [19] by the accumulative rewards achieved in the target domain, which are illustrated in Figure 2. Figure 2 shows that: 1) policies learned by We-DRIVE-U are robust to environmental perturbation, and the extent of the robustness depends on the pre-specified parameter ρ ; 2) In most cases, We-DRIVE-U outperforms DR-LSVI-UCB, meaning it being more robust to environment perturbation. Moreover, Table 1 demonstrates the low-switching property of We-DRIVE-U. During 200 interactions of the training process, We-DRIVE-U switches policies only around 24 times,



(a) The source environment.



(b) The target environment.

Figure 1: The source and the target linear MDP environments. The value on each arrow represents the transition probability. For the source MDP, there are five states and three steps, with the initial state being x_1 , the fail state being x_4 , and x_5 being an absorbing state with reward 1. The target MDP on the right is obtained by perturbing the transition probability at the first step of the source MDP, with others remaining the same.

594 which stands in stark contrast to the 200 policy switches by LSVI-UCB and DR-LSVI-UCB. These
 595 numerical results prove the superiority of our proposed algorithm We-DRIVE-U and align well with
 596 our theoretical findings. All numerical experiments were conducted on a MacBook Pro with a 2.6
 597 GHz 6-Core Intel CPU.

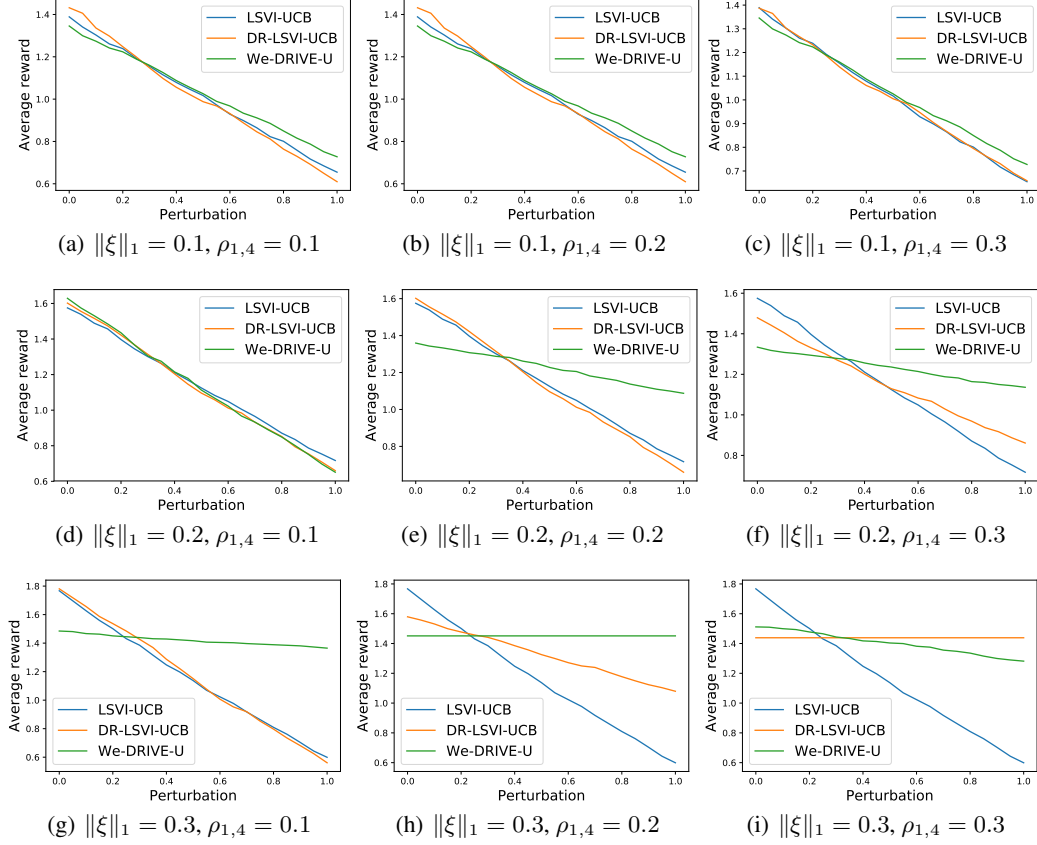


Figure 2: Simulation results under different source domains. The x -axis represents the perturbation level corresponding to different target environments. $\rho_{1,4}$ is the input uncertainty level for our We-DRIVE-U algorithm. $\|\xi\|_1$ is the hyperparameter of the linear DRMDP environment.

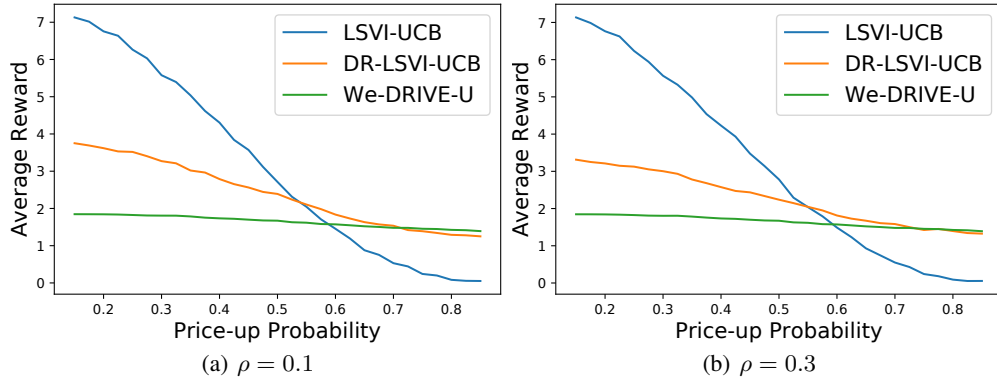


Figure 3: Results for the simulated American put option problem. ρ is the uncertainty level in We-DRIVE-U.

Table 1: Simulation results of the switch complexity of We-DRIVE-U. We present the average policy switch times of We-DRIVE-U during 200 interactions with the nominal environment, averaged over 10 replications. As a comparison, the policy switch times for LSVI-UCB and DR-LSVI-UCB are both **200** under each setting.

	$\rho=0.1$	$\rho=0.2$	$\rho=0.3$
$\ \xi\ _1=0.1$	23.8	24.0	23.8
$\ \xi\ _1=0.2$	24.2	24.4	24.0
$\ \xi\ _1=0.3$	24.3	23.6	24.8

598 B.2 Simulated American put option

599 We additionally conduct a simulation study in the American put option environment (see more details
600 in Section 6.2 of [1]), which under the hood is not a linear MDP. By manually constructing ϕ , we
601 show in Figure 3 that our algorithm still achieves some degree of robustness.

602 C Proof of Proposition 2.2

603 *Proof.* We instantiate the hard example in Example 3.1 of [25] in terms of the formulation of d -
604 rectangular linear DRMDP satisfying Assumption 2.1. Consider two d -rectangular linear DRMDPs ,
605 $\mathcal{M}_0$ and $\mathcal{M}_1$. The state space $\mathcal{S} = \{s_{\text{good}}, s_{\text{bad}}\}$, and the action space is $\mathcal{A} = \{0, 1\}$. We define the
606 feature mapping as

$$\phi^\varrho(s_{\text{good}}, a) = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \forall a \in \mathcal{A}, \phi^\varrho(s_{\text{bad}}, 0) = \begin{pmatrix} 0 \\ p(1-\varrho) \\ q\varrho \\ (1-p)(1-\varrho) \\ (1-q)\varrho \end{pmatrix}, \phi^\varrho(s_{\text{bad}}, 1) = \begin{pmatrix} 0 \\ p\varrho \\ q(1-\varrho) \\ (1-p)\varrho \\ (1-q)(1-\varrho) \end{pmatrix},$$

607 where $\varrho \in \{0, 1\}$ is the index of the d -rectangular linear DRMDP instance. Define the factor
608 distributions $\mu = (\delta_{s_{\text{good}}}, \delta_{s_{\text{good}}}, \delta_{s_{\text{good}}}, \delta_{s_{\text{bad}}}, \delta_{s_{\text{bad}}})^\top$ and the reward parameter $\theta = (1, 0, 0, 0, 0)^\top$.
609 Then it is trivial to check that equipped with the d -rectangular TV divergence uncertainty set, this
610 example recover the hard example in Example 3.1 of [25]. $\square$

611 D Proof of the Upper Bound on the Suboptimality of We-DRIVE-U

612 In this section, we present the proofs of our main theoretical results Theorems 4.1 and 4.2. We start
613 with presenting the technical lemmas in Appendix D.1, and then we derive the upper bound on the
614 suboptimality of We-DRIVE-U in Appendices D.2 and D.3.

615 D.1 Technical Lemmas

616 **Definition D.1** (Good event). Under Assumptions 2.1 and 2.3, then for any fixed $\delta \in (0, 1)$, $\alpha' \in$
617 $[0, H]$ and $\rho \in (0, 1]$, we define $\mathcal{E}_h$ be the event that for all episode $k \in [K]$, stage $h \leq h' \leq H$,

$$\left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau, h'}^{-2} \phi_{h'}^\tau \left[[\hat{V}_{k, h'+1}^\rho(s_{h'+1}^\tau)]_{\alpha'} - [\mathbb{P}_{h'}^0[\hat{V}_{k, h'+1}^\rho]_{\alpha'}](s_{h'}^\tau, a_{h'}^\tau) \right] \right\|_{\Sigma_{k, h'}^{-1}} \leq \gamma, \quad (\text{D.1})$$

618 where $\gamma = \tilde{O}(\sqrt{d})$.

619 **Lemma D.2.** We define $\bar{\mathcal{E}}$ as the event that the following inequalities hold for all $(s, a) \in \mathcal{S} \times \mathcal{A}$,
620 $k \in [K]$, $h \in [H]$,

$$\begin{aligned} |\phi(s, a)^\top \hat{w}_{h,1}^k - [\mathbb{P}_h^0 \hat{V}_{k, h+1}^\rho](s, a)| &\leq \bar{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k, h}^{-1} \phi(s, a)}, \\ |\phi(s, a)^\top \tilde{w}_{h,1}^k - [\mathbb{P}_h^0 \tilde{V}_{k, h+1}^\rho](s, a)| &\leq \bar{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k, h}^{-1} \phi(s, a)}, \end{aligned}$$

$$|\phi(s, a)^\top \tilde{\mathbf{w}}_{h,2}^k - [\mathbb{P}_h^0(\hat{V}_{k,h+1}^\rho)^2](s, a)| \leq \tilde{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k,h}^{-1} \phi(s, a)},$$

621 where $\bar{\beta} = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d^3H^3})$ and $\tilde{\beta} = \tilde{\mathcal{O}}(H^2\sqrt{d\lambda} + \sqrt{d^3H^6})$. Then event $\bar{\mathcal{E}}$ holds with
 622 probability at least $1 - \delta$.

623 **Lemma D.3** (Variance error). On the event $\mathcal{E}_{h+1}$ and $\bar{\mathcal{E}}$, for all episode $k \in [K]$, the estimated
 624 variance satisfies

$$\begin{aligned} |[\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k)| &\leq E_{k,h}, \\ |[\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k)| &\leq E_{k,h} + D_{k,h}. \end{aligned}$$

625 Thus we also have

$$\bar{\sigma}_{k,h}^2 \geq [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) + E_{k,h} + D_{k,h} \geq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k).$$

626 **Lemma D.4.** For any fixed policy π , on the event $\mathcal{E}_h$ and $\bar{\mathcal{E}}$, for all $(s, a, k) \in \mathcal{S}/\{s_f\} \times \mathcal{A} \times [K]$,
 627 for stage $h \leq h' \leq H$, we have

$$\begin{aligned} (r_{h'}(s, a) + \phi(s, a)^\top \hat{\nu}_{h'}^{\rho,k}) - Q_{h'}^{\pi,\rho}(s, a) &= \inf_{P_{h'}(\cdot|s,a) \in \mathcal{U}_{h'}^\rho(s,a;\mu_{h'}^0)} [\mathbb{P}_{h'} \hat{V}_{k,h'+1}^\rho](s, a) \\ &\quad - \inf_{P_{h'}(\cdot|s,a) \in \mathcal{U}_{h'}^\rho(s,a;\mu_{h'}^0)} [\mathbb{P}_{h'} V_{h'+1}^{\pi,\rho}](s, a) + \Delta_{h'}^k(s, a), \end{aligned}$$

628 where $\Delta_{h'}^k(s, a)$ that satisfies $|\Delta_{h'}^k(s, a)| \leq \hat{\Gamma}_{k,h'}(s, a) = \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h'}^{-1} \mathbf{1}_i}$ where
 629 $\beta = \tilde{\mathcal{O}}(\sqrt{\lambda d}H + \sqrt{d})$.

630 **Lemma D.5** (Optimism and pessimism). On the event $\mathcal{E}_h$ and $\bar{\mathcal{E}}$, for all episode $k \in [K]$ and stage
 631 $h \leq h' \leq H$, for all $(s, a) \in \mathcal{S} \times \mathcal{A}$, we have $\hat{Q}_{k,h'}^\rho(s, a) \geq Q_{h'}^{\pi,\rho}(s, a) \geq \check{Q}_{k,h'}^\rho(s, a)$. In addition,
 632 we have $\hat{V}_{k,h'}^\rho(s) \geq V_{h'}^{\pi,\rho}(s) \geq \check{V}_{k,h'}^\rho(s)$.

633 **Lemma D.6.** On the event $\bar{\mathcal{E}}$, event $\mathcal{E} = \mathcal{E}_1$ holds with probability at least $1 - \delta$.

634 **Lemma D.7.** Under the assumption (4.2) and events $\mathcal{E}$ and $\bar{\mathcal{E}}$, for any $h \in [H]$, set $\lambda = 1/H^2$ and
 635 $\rho \in (0, 1]$. Then when $k \geq \tilde{K}$ where $\tilde{K} = \tilde{\mathcal{O}}(d^{15}H^{12})$, with probability at least $1 - \delta$, then we have

$$\bar{\sigma}_{k,h}^2 \leq \mathcal{O}\left(\min\left\{\frac{1}{\rho^2}, H^2\right\}\right).$$

636 D.2 Proof of Theorem 4.1

637 *Proof of Theorem 4.1.* Conditioned on the event $\mathcal{E}$ and $\bar{\mathcal{E}}$, we first do the following decomposition

$$\begin{aligned} &\hat{V}_{k,h}^\rho(s_h^k) - V_h^{\pi^k,\rho}(s_h^k) \\ &= \hat{Q}_{k,h}^\rho(s_h^k, a_h^k) - Q_h^{\pi^k,\rho}(s_h^k, a_h^k) \\ &\leq r_h(s_h^k, a_h^k) + \phi(s_h^k, a_h^k)^\top \hat{\nu}_h^{\rho,k_{\text{last}}} + \hat{\Gamma}_{k_{\text{last}},h}(s_h^k, a_h^k) - Q_h^{\pi^k,\rho}(s_h^k, a_h^k) \\ &\leq \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h V_{h+1}^{\pi^k,\rho}](s_h^k, a_h^k) + 2\hat{\Gamma}_{k_{\text{last}},h}(s_h^k, a_h^k) \\ &\leq \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h V_{h+1}^{\pi^k,\rho}](s_h^k, a_h^k) + 4\hat{\Gamma}_{k,h}(s_h^k, a_h^k), \end{aligned}$$

638 where the first equality holds due to the selection of π_h^k , the first inequality holds due to the definition
 639 of $\hat{Q}_{k,h}^\rho$, the second inequality hold from Lemma D.4, the third inequality holds from Lemma H.2.
 640 Note that

$$\begin{aligned} &\inf_{P_h(\cdot|s_h^k, a_h^k) \in \mathcal{U}_h^\rho(s_h^k, a_h^k;\mu_h^0)} [\mathbb{P}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - \inf_{P_h(\cdot|s_h^k, a_h^k) \in \mathcal{U}_h^\rho(s_h^k, a_h^k;\mu_h^0)} [\mathbb{P}_h V_{h+1}^{\pi^k,\rho}](s_h^k, a_h^k) \\ &= \left\langle \phi(s_h^k, a_h^k), \left[\max_{\alpha_i \in [0, H]} \left\{ \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i} - \rho \alpha_i \right\} \right]_{i \in [d]} \right\rangle \\ &\quad - \left\langle \phi(s_h^k, a_h^k), \left[\max_{\alpha_i \in [0, H]} \left\{ \mathbb{E}^{\mu_{h,i}^0} [V_{h+1}^{\pi^k,\rho}(s)]_{\alpha_i} - \rho \alpha_i \right\} \right]_{i \in [d]} \right\rangle \end{aligned}$$

$$\begin{aligned}
&\leq \left\langle \phi(s_h^k, a_h^k), \left[\max_{\alpha_i \in [0, H]} \left\{ \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i} - \mathbb{E}^{\mu_{h,i}^0} [V_{h+1}^{\pi^k, \rho}(s)]_{\alpha_i} \right\} \right]_{i \in [d]} \right\rangle \\
&\leq \langle \phi(s_h^k, a_h^k), \mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s) - V_{h+1}^{\pi^k, \rho}(s)] \rangle \\
&= \mathbb{P}_h^0 [\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k, \rho}] (s_h^k, a_h^k) \\
&= [\mathbb{P}_h^0 [\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k, \rho}]] (s_h^k, a_h^k) - [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k, \rho}(s_{h+1}^k)] + [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k, \rho}(s_{h+1}^k)],
\end{aligned}$$

641 where the second inequality holds from Lemma D.5. Then we have

$$\begin{aligned}
&\hat{V}_{k,h}^\rho(s_h^k) - V_h^{\pi^k, \rho}(s_h^k) \\
&\leq [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k, \rho}(s_{h+1}^k)] + [\mathbb{P}_h^0 [\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k, \rho}]] (s_h^k, a_h^k) - [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k, \rho}(s_{h+1}^k)] \\
&\quad + 4\hat{\Gamma}_{k,h}(s_h^k, a_h^k), \tag{D.2}
\end{aligned}$$

642 Then by applying (D.2) iteratively and applying Azuma-Hoeffding inequality, with probability at
643 least $1 - \delta/3$, we have

$$\begin{aligned}
K \times \text{AveSubopt}(K) &= \sum_{k=1}^K \left(V_1^{*, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k) \right) \\
&\leq \sum_{k=1}^K \left(\hat{V}_{k,1}^\rho(s_1^k) - V_1^{\pi^k, \rho}(s_1^k) \right) \\
&\leq \sum_{k=1}^K \sum_{h=1}^H \left([\mathbb{P}_h^0 [\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k, \rho}]] (s_h^k, a_h^k) - [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k, \rho}(s_{h+1}^k)] \right) \\
&\quad + \sum_{k=1}^K \sum_{h=1}^H 4\hat{\Gamma}_{k,h}(s_h^k, a_h^k) \\
&\leq 2\sqrt{2H^3 K \log(6/\delta)} + 4\beta \sum_{k=1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i},
\end{aligned}$$

644 where the first inequality holds from Lemma D.5, the second inequality holds from (D.2), the third
645 inequality holds from Azuma-Hoeffding inequality and the definition of $\hat{\Gamma}_{k,h}(s_h^k, a_h^k)$. Finally, by
646 taking probability union bound over $\mathcal{E}$ and $\bar{\mathcal{E}}$, with probability at least $1 - \delta$, we can get the result of
647 Theorem 4.2,

$$\text{AveSubopt}(K) \leq 2\sqrt{2H^3 \log(6/\delta)/K} + 4\beta/K \sum_{k=1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_{h,i}^k \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}.$$

648 This completes the proof. $\square$

649 D.3 Proof of Theorem 4.2

650 *Proof of Theorem 4.2.* Conditioned on the event $\mathcal{E}$ and $\bar{\mathcal{E}}$, we first do the decomposition as follows

$$\begin{aligned}
K \times \text{AveSubopt}(K) &= \sum_{k=1}^K (V_1^{*, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k)) \\
&= \sum_{k=1}^{\tilde{K}} (V_1^{*, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k)) + \sum_{k=\tilde{K}+1}^K (V_1^{*, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k)) \\
&\leq H\tilde{K} + \sum_{k=\tilde{K}+1}^K (V_1^{*, \rho}(s_1^k) - V_1^{\pi^k, \rho}(s_1^k)).
\end{aligned}$$

651 Recall from (D.2) in the proof of Theorem 4.1, we have

$$\hat{V}_{k,h}^\rho(s_h^k) - V_h^{\pi^k, \rho}(s_h^k)$$

$$\leq [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k,\rho}(s_{h+1}^k)] + [\mathbb{P}_h[\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k,\rho}]](s_h^k, a_h^k) - [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k,\rho}(s_{h+1}^k)] + 4\hat{\Gamma}_{k,h}(s_h^k, a_h^k).$$

Then by applying (D.2) iteratively and applying Azuma-Hoeffding inequality, with probability at least $1 - \delta/4$, we have

$$\begin{aligned} K \times \text{AveSubopt}(K) &\leq H\tilde{K} + \sum_{k=\tilde{K}+1}^K (V_1^{*,\rho}(s_1^k) - V_1^{\pi^k,\rho}(s_1^k)) \\ &\leq H\tilde{K} + \sum_{k=\tilde{K}+1}^K (\hat{V}_{k,1}^\rho(s_1^k) - V_1^{\pi^k,\rho}(s_1^k)) \\ &\leq H\tilde{K} + \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \left([\mathbb{P}_h^0[\hat{V}_{k,h+1}^\rho - V_{h+1}^{\pi^k,\rho}]](s_h^k, a_h^k) \right. \\ &\quad \left. - [\hat{V}_{k,h+1}^\rho(s_{h+1}^k) - V_{h+1}^{\pi^k,\rho}(s_{h+1}^k)] \right) + \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H 4\hat{\Gamma}_{k,h}(s_h^k, a_h^k) \\ &\leq H\tilde{K} + 2\sqrt{2H^3K \log(8/\delta)} + 4\beta \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}, \end{aligned}$$

where the second inequality holds from Lemma D.5, the third inequality holds from (D.2) and the last inequality holds from Azuma-Hoeffding inequality and the definition of $\hat{\Gamma}_{k,h}(s_h^k, a_h^k)$. Based on (4.2) and Lemma D.7, with probability at least $1 - \delta/4$, we can further have

$$\begin{aligned} &4\beta \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i} \\ &\leq 4c_1\beta \min\left\{\frac{1}{\rho}, H\right\} \cdot \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Lambda_{k,h}^{-1} \mathbf{1}_i} \\ &\leq 4c_1\beta \min\left\{\frac{1}{\rho}, H\right\} \cdot \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\lambda_{\max}(\Lambda_{k,h}^{-1})} \\ &\leq 4c_1\beta \min\left\{\frac{1}{\rho}, H\right\} \cdot \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sqrt{\frac{1}{\lambda_{\min}(\Lambda_{k,h})}} \\ &\leq 4c_1\beta \min\left\{\frac{1}{\rho}, H\right\} \cdot \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sqrt{\frac{2d}{k \cdot c}} \\ &\leq 4c_1\sqrt{2d}\beta \frac{H}{\sqrt{c}} \cdot \min\left\{\frac{1}{\rho}, H\right\} \cdot \int_{\tilde{K}+1}^K \frac{1}{\sqrt{k-1}} dk \\ &\leq 4c_1\sqrt{2d}\beta \frac{H}{\sqrt{c}} \cdot \min\left\{\frac{1}{\rho}, H\right\} \cdot 2\sqrt{\tilde{K}} \\ &\leq \tilde{\mathcal{O}}\left(dH\sqrt{\tilde{K}} \cdot \min\left\{\frac{1}{\rho}, H\right\}\right), \end{aligned}$$

where $c_1 > 0$ is an absolute constant. The first inequality holds from Lemma D.7, the third inequality holds because $\sum_{i=1}^d \phi_i(s, a) = 1$ and the fourth inequality holds due to (E.10) with $\tilde{K} > 512/\eta^2 \log(dKH/\delta)$. Therefore, we can further bound the regret that

$$K \times \text{AveSubopt}(K) \leq H\tilde{K} + 2\sqrt{2H^3K \log(8/\delta)} + 4\beta \sum_{k=\tilde{K}+1}^K \sum_{h=1}^H \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$$

$$\leq \tilde{O}\left(dH\sqrt{K} \cdot \min\left\{\frac{1}{\rho}, H\right\} + H^{\frac{3}{2}}\sqrt{K} + d^{15}H^{13}\right).$$

660 Finally, by taking probability union bound over $\mathcal{E}$ and $\bar{\mathcal{E}}$, with probability at least $1 - \delta$, we can bound
 661 the average suboptimality of We-DRIVE-U as follows

$$\text{AveSubopt}(K) \leq \tilde{O}\left(\frac{dH \cdot \min\left\{\frac{1}{\rho}, H\right\} + H^{\frac{3}{2}}}{\sqrt{K}} + \frac{d^{15}H^{13}}{K}\right). \quad (\text{D.3})$$

662 We complete the proof by substituting $\eta = O(1/d)$ into (D.3). $\square$

663 E Proof of the Technical Lemmas

664 E.1 Proof of Lemma D.2

665 Before the proof of Lemma D.2, we first present a lemma that defines the optimistic value function
 666 class and gives a upper bound for its covering number.

667 **Lemma E.1** (Function class covering number). In Algorithm 1, for each episode $k \in [K]$ and
 668 $h \in [H]$, the optimistic value function $\hat{V}_{k,h}^\rho$ belongs to the following function class

$$\mathcal{V}_h = \left\{ V \mid V(\cdot) = \max_a \max_{1 \leq j \leq \ell} \min \left\{ r_h(\cdot, a) + \phi(\cdot, a)^\top \mathbf{w}_j + \beta \sum_{i=1}^d \phi_i(\cdot, a) \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_j \mathbf{1}_i}, H \right\}, \right. \\ \left. \|\mathbf{w}_j\| \leq L, \|\mathbf{\Gamma}_j\|_F \leq \lambda^{-1} \sqrt{d} \right\},$$

669 where $\ell \leq dH \log(1 + K/\lambda)$ is the number of value function updates from Lemma F.1 and $L =$
 670 $2H\sqrt{dK/\lambda}$ from Lemma F.2. Define $\mathcal{N}_\epsilon$ be the ϵ -covering number of $\mathcal{V}_h$ with respect to the distance
 671 $\text{dist}(V_1, V_2) = \sup_s |V_1(s) - V_2(s)|$. Then the covering entropy can be bounded by

$$\log \mathcal{N}_\epsilon \leq d\ell \log(1 + 4L/\epsilon) + d^2 \ell \log(1 + 8\sqrt{d}\beta^2/\lambda\epsilon^2).$$

672 *Proof of Lemma E.1.* For any two function $V_1, V_2 \in \mathcal{V}_h$, we can write V_1, V_2 as follows

$$V_1(\cdot) = \max_a \max_{1 \leq j \leq \ell} \min \left\{ r_h(\cdot, a) + \phi(\cdot, a)^\top \mathbf{w}_{1,j} + \beta \sum_{i=1}^d \phi_i(\cdot, a) \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{1,j} \mathbf{1}_i}, H \right\}, \\ V_2(\cdot) = \max_a \max_{1 \leq j \leq \ell} \min \left\{ r_h(\cdot, a) + \phi(\cdot, a)^\top \mathbf{w}_{2,j} + \beta \sum_{i=1}^d \phi_i(\cdot, a) \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{2,j} \mathbf{1}_i}, H \right\},$$

673 where $\|\mathbf{w}_{1,j}\|, \|\mathbf{w}_{2,j}\| \leq L$, $\mathbf{\Gamma}_{1,j}, \mathbf{\Gamma}_{2,j} \preceq \lambda^{-1} \mathbf{I}$ and $\|\mathbf{\Gamma}_{1,j}\|_F, \|\mathbf{\Gamma}_{2,j}\|_F \leq \lambda^{-1} \sqrt{d}$. Then we have

$$\begin{aligned} \text{dist}(V_1, V_2) &= \sup_s |V_1(s) - V_2(s)| \\ &\leq \sup_{1 \leq j \leq \ell, s \in \mathcal{S}, a \in \mathcal{A}} \left| \phi(s, a)^\top \mathbf{w}_{1,j} + \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{1,j} \mathbf{1}_i} \right. \\ &\quad \left. - \phi(s, a)^\top \mathbf{w}_{2,j} - \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{2,j} \mathbf{1}_i} \right| \\ &\leq \beta \sup_{1 \leq j \leq \ell, s \in \mathcal{S}, a \in \mathcal{A}} \left| \sum_{i=1}^d \phi_i(s, a) \left(\sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{1,j} \mathbf{1}_i} - \sqrt{\mathbf{1}_i^\top \mathbf{\Gamma}_{2,j} \mathbf{1}_i} \right) \right| \\ &\quad + \sup_{1 \leq j \leq \ell, s \in \mathcal{S}, a \in \mathcal{A}} |\phi(s, a)^\top (\mathbf{w}_{1,j} - \mathbf{w}_{2,j})| \\ &\leq \beta \sup_{1 \leq j \leq \ell, s \in \mathcal{S}, a \in \mathcal{A}} \left| \sum_{i=1}^d \sqrt{\phi_i(s, a) \mathbf{1}_i^\top (\mathbf{\Gamma}_{1,j} - \mathbf{\Gamma}_{2,j}) \phi_i(s, a) \mathbf{1}_i} \right| \end{aligned}$$

$$\begin{aligned}
& + \sup_{1 \leq j \leq \ell, s \in \mathcal{S}, a \in \mathcal{A}} |\phi(s, a)^\top (\mathbf{w}_{1,j} - \mathbf{w}_{2,j})| \\
& \leq \beta \sup_{1 \leq j \leq \ell} \sqrt{\|\Gamma_{1,j} - \Gamma_{2,j}\|_F} + \sup_{1 \leq j \leq \ell} \|\mathbf{w}_{1,j} - \mathbf{w}_{2,j}\|_2
\end{aligned} \tag{E.1}$$

where the third inequality holds because $|\sqrt{x-y}| \geq |\sqrt{x} - \sqrt{y}|$, the fourth inequality holds because Cauchy-Schwarz inequality, $\|\phi(s, a)\|_2 \leq 1$ and $\sum_{i=1}^d \phi_i = 1$. Moreover, $\|\cdot\|_F$ is the Frobenius norm.

Now, we denote $\mathcal{C}_{\mathbf{w}}$ as a $\epsilon/2$ -cover of the set $\{\mathbf{w} \in \mathbb{R}^d \mid \|\mathbf{w}\|_2 \leq L\}$ and $\mathcal{C}_{\Gamma}$ as a $\epsilon^2/4\beta^2$ -cover of the set $\{\Gamma \in \mathbb{R}^{d \times d} \mid \|\Gamma\|_F \leq \lambda^{-1}\sqrt{d}\}$ with respect to the Frobenius norm. Then according to Lemma H.6, we have

$$|\mathcal{C}_{\mathbf{w}}| \leq (1 + 4L/\epsilon)^d, |\mathcal{C}_{\Gamma}| \leq (1 + 8\sqrt{d}\beta^2/\lambda\epsilon^2)^{d^2}.$$

Then for any function $V_1 \in \mathcal{V}_h$ with parameters $\mathbf{w}_{1,j}, \Gamma_{1,j}, 1 \leq j \leq \ell$, we can find parameters $\mathbf{w}_{2,j} \in \mathcal{C}_{\mathbf{w}}, \Gamma_{2,j} \in \mathcal{C}_{\Gamma}, 1 \leq j \leq \ell$, such that $\|\mathbf{w}_{2,j} - \mathbf{w}_{1,j}\|_2 \leq \epsilon/2, \|\Gamma_{2,j} - \Gamma_{1,j}\|_F \leq \epsilon^2/4\beta^2$. Thus we have

$$\text{dist}(V_1, V_2) \leq \beta \sup_{1 \leq j \leq \ell} \sqrt{\|\Gamma_{1,j} - \Gamma_{2,j}\|_F} + \sup_{1 \leq j \leq \ell} \|\mathbf{w}_{1,j} - \mathbf{w}_{2,j}\|_2 \leq \epsilon,$$

where the inequality holds from (E.1). Therefore, the ϵ -covering number of optimistic function class $\mathcal{V}_h$ is bounded by $\mathcal{N}_\epsilon \leq |\mathcal{C}_{\mathbf{w}}|^\ell \cdot |\mathcal{C}_{\Gamma}|^\ell$, thus we have

$$\log \mathcal{N}_\epsilon \leq d\ell \log(1 + 4L/\epsilon) + d^2\ell \log(1 + 8\sqrt{d}\beta^2/\lambda\epsilon^2),$$

which completes the proof. $\square$

Now we are ready to prove Lemma D.2.

Proof of Lemma D.2. For any stage $h \in [H]$ and the optimistic value function $\hat{V}_{k,h+1}^\rho$, according to Lemma F.3, there exists a vector $\mathbf{w}_h^k$ such that $\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho(s, a)$ can be represented by $\phi(s, a)^\top \mathbf{w}_h^k$ and $\|\mathbf{w}_h^k\|_2 \leq H\sqrt{d}$. Therefore, the parameter estimation error can be decomposed as

$$\begin{aligned}
& \|\hat{\mathbf{w}}_{h,1}^k - \mathbf{w}_h^k\|_{\Lambda_{k,h}} \\
& \leq \left\| \Lambda_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) \hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - \Lambda_{k,h}^{-1} \left(\lambda \mathbf{I} + \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) \phi(s_h^\tau, a_h^\tau)^\top \right) \mathbf{w}_h^k \right\|_{\Lambda_{k,h}} \\
& \leq \left\| \Lambda_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) (\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - \mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho(s_h^\tau, a_h^\tau)) - \lambda \Lambda_{k,h}^{-1} \mathbf{w}_h^k \right\|_{\Lambda_{k,h}} \\
& \leq \underbrace{\left\| \lambda \Lambda_{k,h}^{-1} \mathbf{w}_h^k \right\|_{\Lambda_{k,h}}}_{I_1} + \underbrace{\left\| \Lambda_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) (\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - \mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho(s_h^\tau, a_h^\tau)) \right\|_{\Lambda_{k,h}}}_{I_2}.
\end{aligned}$$

Bound term I_1 :

$$I_1 = \|\lambda \Lambda_{k,h}^{-1} \mathbf{w}_h^k\|_{\Lambda_{k,h}} = \lambda \|\mathbf{w}_h^k\|_{\Lambda_{k,h}^{-1}} \leq \sqrt{\lambda} \|\mathbf{w}_h^k\|_2 \leq H\sqrt{d\lambda},$$

where we have $\Lambda_{k,h} \succcurlyeq \lambda \mathbf{I}$ and $\|\mathbf{w}_h^k\|_2 \leq H\sqrt{d}$.

Bound term I_2 : we apply Lemma H.9 with the optimistic value function class $\mathcal{V}_h$ and $\epsilon = H\sqrt{\lambda}/K$, then for any fixed $h \in [H]$, with probability at least $1 - \delta/3H$, for all episode $k \in [K]$, we have

$$\begin{aligned}
I_2 & = \left\| \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) (\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - \mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho(s_h^\tau, a_h^\tau)) \right\|_{\Lambda_{k,h}^{-1}} \\
& \leq \sqrt{4H^2 \left[\frac{d}{2} \log \left(\frac{k+\lambda}{\lambda} \right) + \log \frac{\mathcal{N}_\epsilon}{\delta} \right] + \frac{8k^2\epsilon^2}{\lambda}}
\end{aligned}$$

$$\leq \tilde{\mathcal{O}}(\sqrt{d^3 H^3}),$$

694 where the first inequality holds because of Lemma H.9, the second inequality holds from Lemma E.1.
 695 Thus we have

$$\|\hat{\mathbf{w}}_{h,1}^k - \mathbf{w}_h^k\|_{\Lambda_{k,h}} \leq I_1 + I_2 = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d^3 H^3}) = \bar{\beta}.$$

696 Therefore, the estimation error can be bounded by

$$\begin{aligned} |\phi(s, a)^\top \hat{\mathbf{w}}_{h,1}^k - [\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s, a)| &= |\phi(s, a)^\top \hat{\mathbf{w}}_{h,1}^k - \phi(s, a)^\top \mathbf{w}_h^k| \\ &\leq \|\hat{\mathbf{w}}_{h,1}^k - \mathbf{w}_h^k\|_{\Lambda_{k,h}} \cdot \|\phi(s, a)\|_{\Lambda_{k,h}^{-1}} \\ &\leq \bar{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k,h}^{-1} \phi(s, a)}, \end{aligned}$$

697 where the first inequality holds from Cauchy-Schwarz inequality. Similarly, for the pessimistic
 698 function class $\mathcal{V}_h$ (or squared value function class $\mathcal{V}_h^2$), we have the similar result as follows

$$\begin{aligned} |\phi(s, a)^\top \hat{\mathbf{w}}_{h,1}^k - [\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s, a)| &\leq \bar{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k,h}^{-1} \phi(s, a)}, \\ |\phi(s, a)^\top \tilde{\mathbf{w}}_{h,2}^k - [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s, a)| &\leq \tilde{\beta} \sqrt{\phi(s, a)^\top \Lambda_{k,h}^{-1} \phi(s, a)}, \end{aligned}$$

699 where $\bar{\beta} = \tilde{\mathcal{O}}(H\sqrt{d\lambda} + \sqrt{d^3 H^3})$ and $\tilde{\beta} = \tilde{\mathcal{O}}(H^2\sqrt{d\lambda} + \sqrt{d^3 H^6})$. By taking union bound over
 700 $h \in [H]$ and three function classes, we have that the event $\bar{\mathcal{E}}$ holds with probability at least $1 - \delta$.
 701 This completes the proof. $\square$

702 E.2 Proof of Lemma D.3

703 *Proof of Lemma D.3.* First, recall from (3.6), we have

$$[\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s, a) \approx [\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s, a) = [\phi(s, a)^\top \tilde{\mathbf{w}}_{h,2}^k]_{[0, H^2]} - [\phi(s, a)^\top \hat{\mathbf{w}}_{h,1}^k]_{[0, H]}^2,$$

704 where $\hat{\mathbf{w}}_{h,1}^k$ and $\tilde{\mathbf{w}}_{h,2}^k$ is the solution of the following ridge regression problems

$$\begin{aligned} \tilde{\mathbf{w}}_{h,2}^k &= \operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} \left(\mathbf{w}^\top \phi(s_h^\tau, a_h^\tau) - (\hat{V}_{k,h+1}^\rho(s_h^\tau)) \right)^2 + \lambda \|\mathbf{w}\|_2^2, \\ \hat{\mathbf{w}}_{h,1}^k &= \operatorname{argmin}_{\mathbf{w} \in \mathbb{R}^d} \sum_{\tau=1}^{k-1} \left(\mathbf{w}^\top \phi(s_h^\tau, a_h^\tau) - \hat{V}_{k,h+1}^\rho(s_h^\tau) \right)^2 + \lambda \|\mathbf{w}\|_2^2. \end{aligned}$$

705 Then we have

$$\begin{aligned} &|[\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k)| \\ &\leq \left| [\phi(s_h^k, a_h^k)^\top \tilde{\mathbf{w}}_{h,2}^k]_{[0, H^2]} - [\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k]_{[0, H]}^2 - [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s_h^k, a_h^k) + ([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k))^2 \right| \\ &\leq \left| [\phi(s_h^k, a_h^k)^\top \tilde{\mathbf{w}}_{h,2}^k]_{[0, H^2]} - [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s_h^k, a_h^k) \right| + \left| [\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k]_{[0, H]}^2 - ([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k))^2 \right| \\ &\leq \left| [\phi(s_h^k, a_h^k)^\top \tilde{\mathbf{w}}_{h,2}^k]_{[0, H^2]} - [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s_h^k, a_h^k) \right| + 2H \left| [\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k]_{[0, H]} - [\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) \right| \\ &\leq \min \left\{ \tilde{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \right\} + \min \left\{ 2H\bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \right\} \\ &= E_{k,h}, \end{aligned}$$

706 where the last inequality holds from Lemma D.2. For the second result, we have

$$\begin{aligned} &|[\mathbb{V}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k)| \\ &= \left| [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s_h^k, a_h^k) - ([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k))^2 - [\mathbb{P}_h^0 (V_{h+1}^{*,\rho})^2](s_h^k, a_h^k) + ([\mathbb{P}_h^0 V_{h+1}^{*,\rho}](s_h^k, a_h^k))^2 \right| \\ &\leq \left| [\mathbb{P}_h^0 (\hat{V}_{k,h+1}^\rho)^2](s_h^k, a_h^k) - [\mathbb{P}_h^0 (V_{h+1}^{*,\rho})^2](s_h^k, a_h^k) \right| + \left| ([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k))^2 - ([\mathbb{P}_h^0 V_{h+1}^{*,\rho}](s_h^k, a_h^k))^2 \right| \\ &\leq 4H \left| [\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{P}_h^0 V_{h+1}^{*,\rho}](s_h^k, a_h^k) \right| \end{aligned}$$

$$\begin{aligned}
&\leq 4H \left([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{P}_h^0 V_{h+1}^{*,\rho}](s_h^k, a_h^k) \right) \\
&\leq 4H \left([\mathbb{P}_h^0 \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{P}_h^0 \check{V}_{k,h+1}^\rho](s_h^k, a_h^k) \right) \\
&\leq \min \left\{ 4H \left(\phi(s_h^k, a_h^k)^\top \check{\mathbf{w}}_{h,1}^k - \phi(s_h^k, a_h^k)^\top \check{\mathbf{w}}_{h,1}^k + 2\bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}} \right), H^2 \right\} \\
&= D_{k,h}.
\end{aligned}$$

707 where the second inequality holds because $0 \leq V_{h+1}^{*,\rho}, \hat{V}_{k,h+1}^\rho \leq H$, the third and fourth inequality
708 holds because of Lemma D.5, the fifth inequality holds due to Lemma D.2 and the last inequality
709 holds because the trivial result $0 \leq [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k), [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) \leq H^2$. Thus we have

$$|[\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) - [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k)| \leq E_{k,h} + D_{k,h}.$$

710 Then we also have

$$\bar{\sigma}_{k,h}^2 \geq [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) + E_{k,h} + D_{k,h} \geq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k),$$

711 where we use the definition of $\bar{\sigma}_{k,h}$ in (3.6). This completes the proof. $\square$

712 E.3 Proof of Lemma D.4

713 *Proof of Lemma D.4.* For all $(s, a) \in \mathcal{S}/\{s_f\} \times \mathcal{A}$, for stage $h \leq h' \leq H$ (we use h to replace h' in
714 this part for simplicity), we have

$$Q_h^{\pi,\rho}(s, a) = r_h(s, a) + \phi(s, a)^\top \boldsymbol{\nu}_h^{\pi,\rho} = r_h(s, a) + \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s, a; \boldsymbol{\mu}_h^0)} [\mathbb{P}_h V_{h+1}^{\pi,\rho}](s, a).$$

715 We first decompose the gap $\hat{\nu}_h^{\rho,k} - \nu_h^{\pi,\rho}$ into two terms

$$\hat{\nu}_h^{\rho,k} - \nu_h^{\pi,\rho} = \underbrace{\hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{\rho,k}}_{\text{I}} + \underbrace{\tilde{\nu}_h^{\rho,k} - \nu_h^{\pi,\rho}}_{\text{II}}, \quad (\text{E.2})$$

716 where $\tilde{\nu}_h^{\rho,k} = [\tilde{\nu}_{h,i}^{\rho,k}]_{i \in [d]}$, and $\tilde{\nu}_{h,i}^{\rho,k} = \max_{\alpha \in [0, H]} \{ \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}^\rho(s)]_\alpha - \rho \alpha \}$. Then we will bound
717 these two terms separately.

718 **Bound term I in (E.2):** we have

$$\hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{\rho,k} \leq \left[\max_{\alpha \in [0, H]} \left\{ \hat{z}_{h,i}^k(\alpha) - \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}^\rho(s)]_\alpha \right\} \right]_{i \in [d]}.$$

719 Denote $\alpha_i^k = \operatorname{argmax}_{\alpha \in [0, H]} \{ \hat{z}_{h,i}^k(\alpha) - \mathbb{E}^{\mu_{h,i}^0} [\hat{V}_{k,h+1}^\rho(s)]_\alpha \}$, $i = 1, \dots, d$. Then we have

$$\begin{aligned}
\hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{\rho,k} &\leq \left[\left(\boldsymbol{\Sigma}_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha_i^k} \right)_i - \left(\mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i^k} \right)_i \right]_{i \in [d]} \\
&= \left[\left(-\lambda \boldsymbol{\Sigma}_{k,h}^{-1} \mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i^k} \right)_i + \left(\boldsymbol{\Sigma}_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha_i^k} \right. \right. \\
&\quad \left. \left. - [\mathbb{P}_h^0 [\hat{V}_{k,h+1}^\rho]_{\alpha_i^k}](s_h^\tau, a_h^\tau) \right) \right]_{i \in [d]}. \quad (\text{E.3})
\end{aligned}$$

720 For the first term on the RHS of (E.3),

$$\begin{aligned}
&\left| \left\langle \phi(s, a), \left[\left(-\lambda \boldsymbol{\Sigma}_{k,h}^{-1} \mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i^k} \right)_i \right]_{i \in [d]} \right\rangle \right| \\
&= \left| \sum_{i=1}^d \phi_i(s, a) \mathbf{1}_i^\top (-\lambda) \boldsymbol{\Sigma}_{k,h}^{-1} \mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i^k} \right| \\
&\leq \lambda \sum_{i=1}^d \sqrt{\phi_i(s, a) \mathbf{1}_i^\top \boldsymbol{\Sigma}_{k,h}^{-1} \phi_i(s, a) \mathbf{1}_i} \cdot \left\| \mathbb{E}^{\mu_h^0} [\hat{V}_{k,h+1}^\rho(s)]_{\alpha_i^k} \right\|_{\boldsymbol{\Sigma}_{k,h}^{-1}}
\end{aligned}$$

$$\leq \sqrt{\lambda d H} \sum_{i=1}^d \sqrt{\phi_i(s, a) \mathbf{1}_i^\top \Sigma_{k,h}^{-1} \phi_i(s, a) \mathbf{1}_i}, \quad (\text{E.4})$$

where $\mathbf{1}_i$ is the vector with the i -th entry being 1 and else being 0. The first inequality holds due to the Cauchy-Schwarz inequality. For the second term on the RHS of (E.3), given the event $\mathcal{E}_h$ defined in Definition D.1, we have

$$\begin{aligned} & \left| \left\langle \phi(s, a), \left[\left(\Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau \left[\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) \right]_{\alpha_i^k} - \left[\mathbb{P}_h^0 \left[\hat{V}_{k,h+1}^\rho \right]_{\alpha_i^k} \right] (s_h^\tau, a_h^\tau) \right) \right]_{i \in [d]} \right\rangle \right| \\ &= \left| \sum_{i=1}^d \phi_i(s, a) \mathbf{1}_i^\top \Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau \left[\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) \right]_{\alpha_i^k} - \left[\mathbb{P}_h^0 \left[\hat{V}_{k,h+1}^\rho \right]_{\alpha_i^k} \right] (s_h^\tau, a_h^\tau) \right| \\ &\leq \sum_{i=1}^d \sqrt{\phi_i(s, a) \mathbf{1}_i^\top \Sigma_{k,h}^{-1} \phi_i(s, a) \mathbf{1}_i} \cdot \left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau \left[\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) \right]_{\alpha_i^k} - \left[\mathbb{P}_h^0 \left[\hat{V}_{k,h+1}^\rho \right]_{\alpha_i^k} \right] (s_h^\tau, a_h^\tau) \right\|_{\Sigma_{k,h}^{-1}} \\ &\leq \gamma \sum_{i=1}^d \sqrt{\phi_i(s, a) \mathbf{1}_i^\top \Sigma_{k,h}^{-1} \phi_i(s, a) \mathbf{1}_i}, \quad (\text{E.5}) \end{aligned}$$

where the first inequality follows from Cauchy-Schwarz inequality and the second inequality holds from the event $\mathcal{E}_h$. Combining (E.3), (E.4) and (E.5), we have

$$\langle \phi(s, a), \hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{\rho,k} \rangle \leq \left(\sqrt{\lambda d H} + \gamma \right) \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i},$$

On the other hand, we can similarly do analysis for $\langle \phi(s, a), \tilde{\nu}_h^{\rho,k} - \hat{\nu}_h^{\rho,k} \rangle$. Then we have

$$|\langle \phi(s, a), \hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{\rho,k} \rangle| \leq \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}, \quad (\text{E.6})$$

where $\beta = \left(\sqrt{\lambda d H} + \gamma \right) = \tilde{\mathcal{O}}\left(\sqrt{\lambda d H} + \sqrt{d} \right)$.

Bound term II in (E.2): we have

$$\langle \phi(s, a), \tilde{\nu}_h^{\rho,k} - \nu_h^{\pi,\rho} \rangle = \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h \hat{V}_{k,h+1}^\rho \right](s, a) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h V_{h+1}^{\pi,\rho} \right](s, a).$$

Finally we have

$$\begin{aligned} & (r_h(s, a) + \phi(s, a)^\top \hat{\nu}_h^{\rho,k}) - Q_h^{\pi,\rho}(s, a) \\ &= \langle \phi(s, a), \hat{\nu}_h^{\rho,k} - \tilde{\nu}_h^{k,\rho} + \tilde{\nu}_h^{k,\rho} - \nu_h^{\pi,\rho} \rangle \\ &= \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h \hat{V}_{k,h+1}^\rho \right](s, a) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h V_{h+1}^{\pi,\rho} \right](s, a) + \Delta_h^k(s, a), \end{aligned}$$

where $|\Delta_h^k(s, a)| \leq \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i}$. This completes the proof. $\square$

E.4 Proof of Lemma D.5

Proof of Lemma D.5. We prove this lemma by induction. For last stage $H+1$, it is trivial because for all $(s, a) \in \mathcal{S} \times \mathcal{A}$, we have $\hat{Q}_{k,H+1}^\rho(s, a) = Q_{H+1}^{*,\rho}(s, a) = \check{Q}_{k,H+1}^\rho(s, a) = 0$.

Assume that the lemma holds at stage $h'+1$, now consider the situation at stage h' (we use h to replace h' in this part for simplicity). For all episode $k \in [K]$, we have

$$\begin{aligned} & r_h(s, a) + \phi(s, a)^\top \hat{\nu}_h^{\rho,k} + \hat{\Gamma}_{k,h}(s, a) - Q_h^{*,\rho}(s, a) \\ &\geq \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h \hat{V}_{k,h+1}^\rho \right](s, a) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h V_{h+1}^{*,\rho} \right](s, a) + \Delta_h^k(s, a) + \hat{\Gamma}_{k,h}(s, a) \\ &\geq \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} \left[\mathbb{P}_h (\hat{V}_{k,h+1}^\rho - V_{h+1}^{*,\rho}) \right](s, a) \end{aligned}$$

≥ 0 ,

where the first inequality holds from Lemma D.4, the second inequality holds because $|\Delta_h^k(s, a)| \leq \hat{\Gamma}_{k,h}(s, a)$, the third inequality holds from induction assumption. Thus we have

$$Q_h^{*,\rho}(s, a) \leq \min_{i \in [k]} \left\{ \min_{i \in [k]} r_h(s, a) + \phi(s, a)^\top \hat{\nu}_h^{\rho,i} + \hat{\Gamma}_{i,h}(s, a), H - h + 1 \right\} \leq \hat{Q}_{k,h}^\rho(s, a).$$

Thus for value function V , we have

$$\hat{V}_{k,h}^\rho(s) = \max_a \hat{Q}_{k,h}^\rho(s, a) \geq \max_a Q_h^{*,\rho}(s, a) = V_h^{*,\rho}(s).$$

For the pessimistic value function $\check{Q}_{k,h}^\rho(s, a)$, we can do the similar analysis. Finally, by induction, we finish the proof. $\square$

E.5 Proof of Lemma D.6

Proof of Lemma D.6. We use backward induction to prove this lemma. For the base case, the stage H , it is trivial to obtain (D.1) because $\hat{V}_{k,H+1}^\rho = 0$. Assume (D.1) hold for the stage $h + 1$, then we consider the stage h .

For all episode $k \in [K]$, we first do the following decomposition

$$\begin{aligned} & \left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_\tau^\top \left[[\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha'} - [\mathbb{P}_h^0[\hat{V}_{k,h+1}^\rho]]_{\alpha'}(s_h^\tau, a_h^\tau) \right] \right\|_{\Sigma_{k,h}^{-1}} \\ & \leq \underbrace{\left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_\tau^\top \left[[V_{h+1}^{*,\rho}(s_{h+1}^\tau)]_{\alpha'} - [\mathbb{P}_h^0[V_{h+1}^{*,\rho}]]_{\alpha'}(s_h^\tau, a_h^\tau) \right] \right\|_{\Sigma_{k,h}^{-1}}}_{J_1} \\ & \quad + \underbrace{\left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_\tau^\top \left[\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - [\mathbb{P}_h^0(\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho)](s_h^\tau, a_h^\tau) \right] \right\|_{\Sigma_{k,h}^{-1}}}_{J_2}, \end{aligned} \quad (\text{E.7})$$

where $\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) = [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha'} - [V_{h+1}^{*,\rho}(s_{h+1}^\tau)]_{\alpha'}$.

Bound term J_1 in (E.7): For term J_1 , we apply Lemma H.8 with $\mathbf{x}_i = \bar{\sigma}_{i,h}^{-1} \phi(s_h^i, a_h^i)$ and $\eta_i = \bar{\sigma}_{i,h}^{-1} ([V_{h+1}^{*,\rho}(s_{h+1}^i)]_{\alpha'} - [\mathbb{P}_h^0[V_{h+1}^{*,\rho}]]_{\alpha'}(s_h^i, a_h^i))$. Note that based on Lemma D.3, we have $\bar{\sigma}_{k,h}^2 \geq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) \geq [\mathbb{V}_h [V_{h+1}^{*,\rho}]]_{\alpha'}(s_h^k, a_h^k)$. Then for $\mathbf{x}_i$ and η_i , we have

$$\begin{aligned} \|\mathbf{x}_i\|_2 &= \|\phi(s_h^i, a_h^i)\|_2 / \bar{\sigma}_{i,h} \leq 1, \\ \mathbb{E}[\eta_i | \mathcal{F}_i] &= 0, |\eta_i| \leq |\bar{\sigma}_{i,h}^{-1} ([V_{h+1}^{*,\rho}(s_{h+1}^i)]_{\alpha'} - [\mathbb{P}_h^0[V_{h+1}^{*,\rho}]]_{\alpha'}(s_h^i, a_h^i))| \leq 2H, \\ \mathbb{E}[\eta_i^2 | \mathcal{F}_i] &= \mathbb{E}[\bar{\sigma}_{i,h}^{-2} ([V_{h+1}^{*,\rho}(s_{h+1}^i)]_{\alpha'} - [\mathbb{P}_h^0[V_{h+1}^{*,\rho}]]_{\alpha'}(s_h^i, a_h^i))^2] \leq 1 = \sigma^2, \\ \max_{1 \leq i \leq k} \left\{ |\eta_i| \cdot \min \{1, \|\mathbf{x}_i\|_{\Sigma_{i,h}^{-1}}\} \right\} &\leq \max_{1 \leq i \leq k} \left\{ 2H \bar{\sigma}_{i,h}^{-1} \|\mathbf{x}_i\|_{\Sigma_{i,h}^{-1}} \right\} \leq \sqrt{d}, \end{aligned}$$

where we use the definition of $\bar{\sigma}_{i,h}$ in (3.7). Then for all $k \in [K]$, with probability at least $1 - \delta/2H$, we have

$$J_1 = \left\| \sum_{i=1}^{k-1} \mathbf{x}_i \eta_i \right\|_{\Sigma_{k,h}^{-1}} \leq \tilde{O}(\sigma \sqrt{d} + \max_{1 \leq i \leq k} |\eta_i| \min \{1, \|\mathbf{x}_i\|_{\Sigma_{i,h}^{-1}}\}) = \tilde{O}(\sqrt{d}).$$

Bound term J_2 in (E.7): To bound term J_2 , we need to use ϵ -covering for function class $\tilde{\mathcal{V}}_{h+1} - [V_{h+1}^{*,\rho}]_{\alpha'}$, where $\tilde{\mathcal{V}}_{h+1} = \{[V]_{\alpha'} | V \in \mathcal{V}_{h+1}, \alpha \in [0, H]\}$ is the truncated optimistic value function class. For any two function $\tilde{V}_1, \tilde{V}_2 \in \tilde{\mathcal{V}}_{h+1}$, we can write that

$$\tilde{V}_1 = [V_1]_{\alpha_1}, \tilde{V}_2 = [V_2]_{\alpha_2},$$

756 where $V_1, V_2 \in \mathcal{V}_{h+1}$, $\alpha_1, \alpha_2 \in [0, H]$. Then we have

$$\begin{aligned} \text{dist}(\tilde{V}_1, \tilde{V}_2) &= \sup_s |\tilde{V}_1(s) - \tilde{V}_2(s)| \\ &= \sup_s |[V_1]_{\alpha_1}(s) - [V_2]_{\alpha_2}(s)| \\ &\leq \sup_s |[V_1]_{\alpha_1}(s) - [V_1]_{\alpha_2}(s)| + \sup_s |[V_1]_{\alpha_2}(s) - [V_2]_{\alpha_2}(s)| \\ &\leq |\alpha_1 - \alpha_2| + \text{dist}(V_1, V_2). \end{aligned}$$

757 This indicates that the ϵ -covering number $\tilde{\mathcal{N}}_\epsilon$ for function class $\tilde{\mathcal{V}}_{h+1}$ can be bounded by

$$\tilde{\mathcal{N}}_\epsilon \leq \mathcal{N}_{1, \frac{\epsilon}{2}} \cdot \mathcal{N}_{2, \frac{\epsilon}{2}},$$

758 where $\mathcal{N}_{1, \frac{\epsilon}{2}}$ is the $\frac{\epsilon}{2}$ -covering number for optimistic value function class $\mathcal{V}_{h+1}$ and $\mathcal{N}_{2, \frac{\epsilon}{2}}$ is the
759 $\frac{\epsilon}{2}$ -covering number for closed interval $[0, H]$. Then based on Lemma E.1 and Lemma H.7, we have

$$\log \tilde{\mathcal{N}}_\epsilon \leq d\ell \log(1 + 8L/\epsilon) + d^2\ell \log(1 + 32\sqrt{d}\beta^2/\lambda\epsilon^2) + \log(6H/\epsilon),$$

760 where $\ell = dH \log(1 + K/\lambda)$ and $L = 2H\sqrt{dK/\lambda}$. Here we set $\epsilon = \sqrt{\lambda}/4H^2d^3K$, then the
761 covering entropy can be bounded by

$$\log \tilde{\mathcal{N}}_\epsilon \leq \tilde{\mathcal{O}}(d^3H).$$

762 For simplicity, we denote $\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho$ here as ΔV , then for ΔV , there exists a function $\tilde{V}$ in the ϵ -net
763 satisfies that

$$\text{dist}(\Delta V, \tilde{V}) \leq \epsilon.$$

764 Then the difference of the variance of ΔV and $\tilde{V}$ can be bounded by

$$\begin{aligned} &[\mathbb{V}_h \tilde{V}](s_h^k, a_h^k) - [\mathbb{V}_h \Delta V](s_h^k, a_h^k) \\ &= [\mathbb{P}_h \tilde{V}^2](s_h^k, a_h^k) - [\mathbb{P}_h (\Delta V)^2](s_h^k, a_h^k) - ([\mathbb{P}_h \tilde{V}](s_h^k, a_h^k))^2 + ([\mathbb{P}_h (\Delta V)](s_h^k, a_h^k))^2 \\ &\leq 2 \sup_s |\Delta V(s) - \tilde{V}(s)| \cdot \sup_s |\Delta V(s) + \tilde{V}(s)| \\ &\leq 4H \text{dist}(\Delta V, \tilde{V}) \\ &\leq \frac{1}{2d^3H}. \end{aligned}$$

765 This indicates that

$$\begin{aligned} [\mathbb{V}_h \tilde{V}](s_h^k, a_h^k) &\leq [\mathbb{V}_h \Delta V](s_h^k, a_h^k) + \frac{1}{2d^3H} \\ &\leq [\mathbb{P}_h([\hat{V}_{k,h+1}^\rho]_{\alpha'} - [V_{h+1}^{*,\rho}]_{\alpha'})^2](s_h^k, a_h^k) + \frac{1}{2d^3H} \\ &\leq [\mathbb{P}_h(\hat{V}_{k,h+1}^\rho - V_{h+1}^{*,\rho})^2](s_h^k, a_h^k) + \frac{1}{2d^3H} \\ &\leq 2H [\mathbb{P}_h(\hat{V}_{k,h+1}^\rho - \check{V}_{k,h+1}^\rho)](s_h^k, a_h^k) + \frac{1}{2d^3H} \\ &\leq 2H (\mathbb{P}_h \hat{V}_{k,h+1}^\rho(s_h^k, a_h^k) - \mathbb{P}_h \check{V}_{k,h+1}^\rho(s_h^k, a_h^k)) + \frac{1}{2d^3H} \\ &\leq D_{k,h} + \frac{1}{2d^3H} \\ &\leq \bar{\sigma}_{k,h}^2/d^3H, \end{aligned} \tag{E.8}$$

766 where the fourth inequality holds due to Lemma D.5 with induction assumption $\mathcal{E}_{h+1}$, the sixth
767 inequality holds due to the definition of $D_{k,h}$ and the last inequality holds because of the definition
768 of $\sigma_{k,h}$. Then we apply we apply Lemma H.8 with $\mathbf{x}_i = \bar{\sigma}_{i,h}^{-1} \phi(s_h^i, a_h^i)$ and $\eta_i = \bar{\sigma}_{i,h}^{-1} (\tilde{V}(s_{h+1}^i) -$
769 $\mathbb{P}_h^0 \tilde{V}(s_h^i, a_h^i))$. For $\mathbf{x}_i$ and η_i , we have

$$\|\mathbf{x}_i\|_2 \leq \|\phi(s_h^i, a_h^i)\|_2 / \bar{\sigma}_{i,h} \leq 1,$$

$$\begin{aligned}\mathbb{E}[\eta_i|\mathcal{F}_i] &= 0, |\eta_i| \leq |\bar{\sigma}_{i,h}^{-1}(\tilde{V}(s_{h+1}^i) - \mathbb{P}_h^0 \tilde{V}(s_h^i, a_h^i))| \leq 2H, \\ \mathbb{E}[\eta_i^2|\mathcal{F}_i] &= \bar{\sigma}_{i,h}^{-2}[\mathbb{V}_h \tilde{V}](s_h^i, a_h^i) \leq 1/d^3 H, \\ \max_{1 \leq i \leq k} \left\{ |\eta_i| \cdot \min \{1, \|\mathbf{x}_i\|_{\Sigma_{i,h}^{-1}}\} \right\} &\leq \max_{1 \leq i \leq k} \left\{ 2H \bar{\sigma}_{i,h}^{-1} \|\mathbf{x}_i\|_{\Sigma_{i,h}^{-1}} \right\} \leq 1/d^3 H,\end{aligned}$$

770 where we use the construction of $\bar{\sigma}_{i,h}$ in (3.7) and (E.8). After taking union probability bound over
771 ϵ -covering for function class $\tilde{\mathcal{V}}_{h+1} - [V_{h+1}^{*,\rho}]_{\alpha'}$, we have

$$\left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\tilde{V}(s_{h+1}^\tau) - [\mathbb{P}_h^0 \tilde{V}](s_h^\tau, a_h^\tau)] \right\|_{\Sigma_{k,h}^{-1}} \leq \tilde{\mathcal{O}}(\sqrt{d}).$$

772 For simplicity, we denote that $\bar{V} = \Delta V - \tilde{V} = \Delta_{\alpha'} \hat{V}_{k,h+1}^\rho - \tilde{V}$ and have $\sup_s |\bar{V}(s)| \leq \epsilon$. Then we
773 obtain

$$\begin{aligned}J_2 &= \left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - [\mathbb{P}_h^0(\Delta_{\alpha'} \hat{V}_{k,h+1}^\rho)](s_h^\tau, a_h^\tau)] \right\|_{\Sigma_{k,h}^{-1}} \\ &\leq 2 \left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\tilde{V}(s_{h+1}^\tau) - [\mathbb{P}_h^0 \tilde{V}](s_h^\tau, a_h^\tau)] \right\|_{\Sigma_{k,h}^{-1}} \\ &\quad + 2 \left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\bar{V}(s_{h+1}^\tau) - [\mathbb{P}_h^0 \bar{V}](s_h^\tau, a_h^\tau)] \right\|_{\Sigma_{k,h}^{-1}} \\ &\leq \tilde{\mathcal{O}}(\sqrt{d}) + 4\epsilon k / \sqrt{\lambda} \\ &\leq \tilde{\mathcal{O}}(\sqrt{d}),\end{aligned}$$

774 where we use that $\epsilon = \sqrt{\lambda}/4H^2 d^3 K$. Finally, we have

$$\left\| \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi_h^\tau [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)_{\alpha'} - [\mathbb{P}_h^0[\hat{V}_{k,h+1}^\rho]_{\alpha'}](s_h^\tau, a_h^\tau)] \right\|_{\Sigma_{k,h}^{-1}} = J_1 + J_2 \leq \gamma,$$

775 where $\gamma = \tilde{\mathcal{O}}(\sqrt{d})$. Thus, by induction we complete the proof. $\square$

776 E.6 Proof of Lemma D.7

777 *Proof of Lemma D.7.* Conditioned on the event $\mathcal{E}$ and $\bar{\mathcal{E}}$, to bound the weight $\bar{\sigma}_{k,h}^2$, recall from the
778 definition (3.7), we have

$$\bar{\sigma}_{k,h} = \max \left\{ \sigma_{k,h}, 1, \sqrt{2d^3 H^2} \|\phi(s_h^k, a_h^k)\|_{\Sigma_{k,h}^{-1}}^{\frac{1}{2}} \right\},$$

779 According to (3.6), we have

$$\sigma_{k,h}^2 = [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) + E_{k,h} + d^3 H \cdot D_{k,h} + \frac{1}{2},$$

780 where $E_{k,h}, D_{k,h}$ are defined as follows

$$\begin{aligned}E_{k,h} &= \min \left\{ \tilde{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \right\} + \min \left\{ 2H \bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}, H^2 \right\}, \\ D_{k,h} &= \min \left\{ 4H \left(\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k - \phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k + 2\bar{\beta} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}} \right), H^2 \right\},\end{aligned}$$

781 where $\bar{\beta} = \tilde{\mathcal{O}}(d^{\frac{3}{2}} H^{\frac{3}{2}})$, $\tilde{\beta} = \tilde{\mathcal{O}}(d^{\frac{3}{2}} H^3)$ when we set $\lambda = 1/H^2$. Note that

$$\sqrt{2d^3 H^2} \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}}^{\frac{1}{2}} \leq \sqrt{2d^3 H^2} \|\phi(s_h^k, a_h^k)\|_2^{\frac{1}{2}} / \lambda^{\frac{1}{4}} \leq \sqrt{2d^3 H^3}.$$

782 Also note that

$$\sigma_{k,h}^2 = [\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) + E_{k,h} + d^3 H \cdot D_{k,h} + \frac{1}{2}$$

$$\begin{aligned} &\leq H^2 + 2H^2 + d^3 H \cdot H^2 + \frac{1}{2} \\ &\leq 2d^3 H^3. \end{aligned}$$

783 Then we obtain the trivial upper bound $\bar{\alpha}$ for $\bar{\sigma}_{k,h}$

$$\bar{\sigma}_{k,h} \leq 2\sqrt{d^3 H^3} = \bar{\alpha}, \quad (\text{E.9})$$

784 Based on Lemma D.3, we have

$$[\bar{\mathbb{V}}_h \hat{V}_{k,h+1}^\rho](s_h^k, a_h^k) \leq E_{k,h} + D_{k,h} + [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k).$$

785 Then we have

$$\sigma_{k,h}^2 \leq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) + 2E_{k,h} + 2d^3 H \cdot D_{k,h} + \frac{1}{2}.$$

786 Next, we carefully bound $\sigma_{k,h}^2$ and $\bar{\sigma}_{k,h}^2$. To this end, we bound term $E_{k,h}, D_{k,h}$ when k is large
787 enough. The intuition is that, when the episode k is large enough, all the error terms should be small
788 under the assumption (E.10).

789 **Bound term $E_{k,h}$:** Note that based on (4.2), with the same analysis as the proof of Corollary 5.3 in
790 [23], with probability at least $1 - \delta$, we have

$$\lambda_{\min}(\mathbf{\Lambda}_{k,h}) \geq \max\{c(k-1)/d + \lambda - \sqrt{32k \log(dKH/\delta)}, \lambda\}.$$

791 Then when we choose $k > 512d^2 \log(dKH/\delta)/c^2$ and note that $\lambda = 1/H^2$, we have

$$c(k-1)/d + \lambda - \sqrt{32k \log(dKH/\delta)} \geq \frac{c}{2d}k,$$

792 which indicates that

$$\lambda_{\min}(\mathbf{\Lambda}_{k,h}) \geq \frac{c}{2d}k. \quad (\text{E.10})$$

793 Then when $k > 512d^2 \log(dKH/\delta)/c^2$, we can calculate that

$$\|\phi(s_h^k, a_h^k)\|_{\mathbf{\Lambda}_{k,h}^{-1}} = \left\| \mathbf{\Lambda}_{k,h}^{-\frac{1}{2}} \phi(s_h^k, a_h^k) \right\|_2 \leq \sqrt{\lambda_{\max}(\mathbf{\Lambda}_{k,h}^{-1})} \leq \sqrt{\frac{2d}{kc}}, \quad (\text{E.11})$$

794 where in the first inequality we use the fact that $\|\phi(s, a)\|_2 \leq 1$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$. Then when k
795 is large enough and also at least $k > 512d^2 \log(dKH/\delta)/c^2$, we can have that

$$E_{k,h} \leq \tilde{\mathcal{O}}\left(\frac{1}{\sqrt{kc}} d^2 H^3\right).$$

796 This indicates that there exists an absolute constant $c_E > 0$ such that

$$E_{k,h} \leq c_E \frac{d^2 H^3}{\sqrt{k}}.$$

797 **Bound term $D_{k,h}$:** note that $\hat{\mathbf{w}}_{h,1}^k$ and $\check{\mathbf{w}}_{h,1}^k$ have the closed-form expression as follows

$$\begin{aligned} \hat{\mathbf{w}}_{h,1}^k &= \mathbf{\Lambda}_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) \hat{V}_{k,h+1}^\rho(s_{h+1}^\tau), \\ \check{\mathbf{w}}_{h,1}^k &= \mathbf{\Lambda}_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) \check{V}_{k,h+1}^\rho(s_{h+1}^\tau). \end{aligned}$$

798 Thus we can calculate that

$$\begin{aligned} &\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k - \phi(s_h^k, a_h^k)^\top \check{\mathbf{w}}_{h,1}^k \\ &= \phi(s_h^k, a_h^k)^\top \mathbf{\Lambda}_{k,h}^{-1} \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) (\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau) - \check{V}_{k,h+1}^\rho(s_{h+1}^\tau)), \end{aligned}$$

$$\begin{aligned}
&\leq \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}} \cdot \left\| \sum_{\tau=1}^{k-1} \phi(s_h^\tau, a_h^\tau) (\hat{V}_{k,h+1}^\rho(s_h^\tau) - \check{V}_{k,h+1}^\rho(s_h^\tau)) \right\|_{\Lambda_{k,h}^{-1}} \\
&\leq \|\phi(s_h^k, a_h^k)\|_{\Lambda_{k,h}^{-1}} \cdot \sum_{\tau=1}^{k-1} \|\phi(s_h^\tau, a_h^\tau)\|_{\Lambda_{k,h}^{-1}} \cdot (\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k)) \\
&\leq \sqrt{\frac{2d}{kc}} \sum_{\tau=1}^{k-1} \sqrt{\frac{2d}{kc}} \cdot (\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k)), \\
&\leq 2d/c \cdot (\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k)), \tag{E.12}
\end{aligned}$$

799 where $\tilde{s}_{h+1}^k = \operatorname{argmax}_{s \in \mathcal{S}} \{ \hat{V}_{k,h+1}^\rho(s) - \check{V}_{k,h+1}^\rho(s) \}$, the first inequality holds because of Cauchy-
800 Schwarz inequality, the third inequality holds due to (E.11). Next, we bound $\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) -$
801 $\check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k)$. The intuition is that when k is large, both $\hat{V}_{k,h+1}^\rho$ and $\check{V}_{k,h+1}^\rho$ should be close to the
802 robust optimal value function. Thus, $\hat{V}_{k,h+1}^\rho$ should be close to $\check{V}_{k,h+1}^\rho$, and the closeness could be
803 quantified by the bonus terms, which is of order $\tilde{O}(d/\sqrt{k})$ under the assumption (E.10). In particular,
804 we have

$$\begin{aligned}
&\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) \\
&= \underbrace{\hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - V_{h+1}^{*,\rho}(\tilde{s}_{h+1}^k)}_{\text{I}} + \underbrace{V_{h+1}^{*,\rho}(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k)}_{\text{II}}. \tag{E.13}
\end{aligned}$$

805 **Bound term I in (E.13):** note that

$$\begin{aligned}
&\hat{V}_{k,h}^\rho(s) - V_h^{*,\rho}(s) \\
&= \hat{Q}_{k,h}^\rho(s, \pi_h^k(s)) - Q_h^{*,\rho}(s, \pi_h^*(s)) \\
&\leq \hat{Q}_{k,h}^\rho(s, \pi_h^k(s)) - Q_h^{*,\rho}(s, \pi_h^k(s)) \\
&\leq \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h \hat{V}_{k,h+1}^\rho](s, \pi_h^k(s)) - \inf_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_h^0)} [\mathbb{P}_h V_{h+1}^{*,\rho}](s, \pi_h^k(s)) \\
&\quad + \Delta_h^k(s, \pi_h^k(s)) + \hat{\Gamma}_{k,h}(s, \pi_h^k(s)) \\
&\leq [\hat{\mathbb{P}}_h (\hat{V}_{k,h+1}^\rho - V_{h+1}^{*,\rho})](s, \pi_h^k(s)) + 2\hat{\Gamma}_{k,h}(s, \pi_h^k(s)),
\end{aligned}$$

806 where the second inequality holds due to the definition of $\hat{Q}_{k,h}^\rho$, robust Bellman equation and
807 Lemma D.4, $\hat{P}_h(\cdot|s,a) = \operatorname{arginf}_{P_h(\cdot|s,a) \in \mathcal{U}_h^\rho(s,a;\mu_{h,i}^0)} [\mathbb{P}_h V_{h+1}^{*,\rho}](s,a), \forall (s,a) \in \mathcal{S} \times \mathcal{A}$. By recur-
808 sively applying it, then we have

$$\begin{aligned}
\hat{V}_{k,h}^\rho(s) - V_h^{*,\rho}(s) &\leq [\hat{\mathbb{P}}_h (\hat{V}_{k,h+1}^\rho - V_{h+1}^{*,\rho})](s, \pi_h^k(s)) + 2\hat{\Gamma}_{k,h}(s, \pi_h^k(s)) \\
&\leq 2 \sum_{h'=h}^H \mathbb{E}^{\pi_{h'}, \hat{P}} [\hat{\Gamma}_{k,h'}(s, a) | s_{h'} = s].
\end{aligned}$$

809 Note that by (E.9) $\bar{\sigma}_{k,h}^2 \leq \bar{\alpha}^2$, we have $\Sigma_{k,h} \succcurlyeq \bar{\alpha}^{-2} \Lambda_{k,h}$. Similar to the analysis of (E.11), when
810 $k > 512/c^2 \log(dKH/\delta)$, we have

$$\begin{aligned}
\hat{\Gamma}_{k,h}(s, a) &= \beta \sum_{i=1}^d \phi_i(s, a) \sqrt{\mathbf{1}_i^\top \Sigma_{k,h}^{-1} \mathbf{1}_i} \\
&\leq \beta \bar{\alpha} \sum_{i=1}^d \phi_i(s_h^k, a_h^k) \sqrt{\mathbf{1}_i^\top \Lambda_{k,h}^{-1} \mathbf{1}_i} \\
&\leq \beta \bar{\alpha} \sqrt{\lambda_{\max}(\Lambda_{k,h}^{-1})} \\
&\leq \sqrt{\frac{2d}{kc}} \beta \bar{\alpha}.
\end{aligned}$$

811 Then we have

$$\hat{V}_{k,h}^\rho(s) - V_h^{*,\rho}(s) \leq 2H\sqrt{\frac{2d}{kc}}\beta\bar{\alpha} \leq \frac{4\beta\sqrt{d}\bar{\alpha}H}{\sqrt{kc}}.$$

812 Therefore, we can bound I as follows

$$I = \hat{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) - V_{h+1}^{*,\rho}(\tilde{s}_{h+1}^k) \leq \frac{4\beta\sqrt{d}\bar{\alpha}H}{\sqrt{kc}}.$$

813 **Bound term II in (E.13):** Similar to the analysis above, we can derive the similar result as follows

$$\Pi = V_{h+1}^{*,\rho}(\tilde{s}_{h+1}^k) - \check{V}_{k,h+1}^\rho(\tilde{s}_{h+1}^k) \leq \frac{4\bar{\beta}\sqrt{d}\bar{\alpha}H}{\sqrt{kc}}.$$

814 Now we can bound that

$$\phi(s_h^k, a_h^k)^\top \hat{\mathbf{w}}_{h,1}^k - \phi(s_h^k, a_h^k)^\top \check{\mathbf{w}}_{h,1}^k \leq \frac{2d}{c} \cdot \frac{4(\bar{\beta} + \beta)\sqrt{d}\bar{\alpha}H}{\sqrt{kc}} \leq \frac{16\bar{\beta}d^{3/2}\bar{\alpha}H}{\sqrt{kc^3}}.$$

815 Then when k is large enough, we can have that

$$D_{k,h} \leq \tilde{\mathcal{O}}\left(\frac{\bar{\alpha}}{\sqrt{k}}d^3H^{\frac{7}{2}}\right).$$

816 This indicates that there exists an absolute constant $c_D > 0$ such that

$$D_{k,h} \leq c_D \frac{\bar{\alpha}}{\sqrt{k}}d^3H^{\frac{7}{2}}.$$

817 When k is large enough, we have

$$\begin{aligned} \sigma_{k,h}^2 &\leq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) + (2E_{k,h} + 2d^3H \cdot D_{k,h}) + \frac{1}{2} \\ &\leq [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) + 2c_E \frac{\bar{\alpha}}{\sqrt{k}}d^2H^3 + 2c_D \frac{\bar{\alpha}}{\sqrt{k}}d^6H^{\frac{9}{2}} + \frac{1}{2}. \end{aligned}$$

818 When we choose $\tilde{K} = \tilde{c} \cdot \bar{\alpha}^2 d^{12} H^9$ where $\tilde{c} = \tilde{\mathcal{O}}(1)$. When $k > \tilde{K}$, then we have

$$\begin{aligned} \bar{\sigma}_{k,h}^2 &= \max \left\{ \sigma_{k,h}^2, 1, 2d^3H^2 \|\phi(s_h^k, a_h^k)\|_{\Sigma_{k,h}^{-1}} \right\} \\ &\leq \max \left\{ [\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k) + 1, 1 \right\} \\ &\leq 2[\mathbb{V}_h V_{h+1}^{*,\rho}](s_h^k, a_h^k)_{[1, H^2]}. \end{aligned}$$

819 Based on Lemma H.10, we have

$$[\mathbb{V}_h V_{h+1}^{*,\rho}](s, a) \leq \left(\frac{1 - (1 - \rho)^{H-h+1}}{\rho} \right)^2 \leq \left(\frac{1 - (1 - \rho)^H}{\rho} \right)^2 = \Theta\left(\min\left\{\frac{1}{\rho^2}, H^2\right\}\right).$$

820 Then when $k > \tilde{K}$, we have

$$\bar{\sigma}_{k,h}^2 \leq \mathcal{O}\left(\min\left\{\frac{1}{\rho^2}, H^2\right\}\right).$$

821 Additionally, note that $\bar{\alpha}^2 = \mathcal{O}(d^3H^3)$, we have

$$\tilde{K} = \tilde{\mathcal{O}}(d^{15}H^{12}).$$

822 This completes the proof. □

F Supporting Lemmas

Lemma F.1 (Number of value function updates). The number of episodes where the algorithm updates the value function in Algorithm 1 is upper bounded by $dH \log(1 + K/\lambda)$.

Proof of Lemma F.1. This proof is the same as [10, Lemma F.1] because of the same rare-switching condition (Line 5 in Algorithm 1). $\square$

Lemma F.2. For any $(k, h) \in [K] \times [H]$, the weight $\hat{\nu}_h^{\rho, k}$ satisfies

$$\|\hat{\nu}_h^{\rho, k}\|_2 \leq 2H\sqrt{dk/\lambda}.$$

Proof of Lemma F.2. Denote $\alpha_i = \operatorname{argmax}_{\alpha \in [0, H]} \{\hat{z}_{h,i}^k(\alpha) - \rho\alpha\}$, $i \in [d]$. Then we have

$$\begin{aligned} \|\hat{\nu}_h^{\rho, k}\|_2 &= \left\| \left[\max_{\alpha \in [0, H]} \{\hat{z}_{h,i}^k(\alpha) - \rho\alpha\} \right]_{i \in [d]} \right\|_2 \\ &\leq \rho\sqrt{d}\alpha + \left\| \left[\left(\Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha_i} \right) \right]_{i \in [d]} \right\|_2 \\ &\leq H\sqrt{d} + H \cdot \left\| \Sigma_{k,h}^{-1} \sum_{\tau=1}^{k-1} \bar{\sigma}_{\tau,h}^{-2} \phi(s_h^\tau, a_h^\tau) \right\|_2 \\ &\leq H\sqrt{d} + H\sqrt{k/\lambda} \cdot \left(\sum_{\tau=1}^{k-1} (\bar{\sigma}_{\tau,h}^{-1} \phi(s_h^\tau, a_h^\tau))^\top \Sigma_{k,h}^{-1} (\bar{\sigma}_{\tau,h}^{-1} \phi(s_h^\tau, a_h^\tau)) \right)^{\frac{1}{2}} \\ &\leq H\sqrt{d} + H\sqrt{dk/\lambda} \\ &\leq 2H\sqrt{dk/\lambda}, \end{aligned}$$

where the first inequality holds due to the triangle inequality, the second inequality holds from the fact that $\rho \leq 1$, $0 \leq \alpha \leq H$ and $0 \leq [\hat{V}_{k,h+1}^\rho(s_{h+1}^\tau)]_{\alpha_i} \leq H$, the third inequality holds because of Lemma H.5 and the fourth inequality holds because $\Sigma_{k,h} \succcurlyeq \lambda \mathbf{I}$ and Lemma H.4. This completes the proof. $\square$

Lemma F.3. Under a linear MDP, for any stage $h \in [H]$ and any bounded function $V : \mathcal{S} \rightarrow [0, H]$, there always exists a vector $\mathbf{z} \in \mathbb{R}^d$ such that for all $(s, a) \in \mathcal{S} \times \mathcal{A}$, we have

$$[\mathbb{P}_h^0 V](s, a) = \mathbf{z}^\top \phi(s, a),$$

where $\mathbf{z}$ satisfies that $\|\mathbf{z}\|_2 \leq H\sqrt{d}$.

Proof of Lemma F.3. Based on Assumption 2.1, we have

$$\begin{aligned} [\mathbb{P}_h^0 V](s, a) &= \int \mathbb{P}_h^0(s'|s, a) V(s') ds' \\ &= \int \phi(s, a)^\top V(s') d\mu_h^0(s') \\ &= \phi(s, a)^\top \int V(s') d\mu_h^0(s') \\ &= \phi(s, a)^\top \mathbf{z}, \end{aligned}$$

where $\mathbf{z} = \int V(s') d\mu_h^0(s')$. Thus we have

$$\|\mathbf{z}\|_2 = \left\| \int V(s') d\mu_h^0(s') \right\|_2 \leq \max_{s'} V(s') \cdot \|\mu_h^0(\mathcal{S})\|_2 \leq H\sqrt{d}.$$

This completes the proof. $\square$

840 G Proof of the Minimax Lower Bound

841 In this section, we prove the minimax lower bound. To this end, we first introduce the construction of
842 hard instances in Appendix G.1, and then we prove Theorem 5.1 in Appendix G.3.

843 G.1 Construction of Hard Instances

844 We construct a family of d -rectangular linear DRMDPs based on the hard-to-learn linear MDP
845 introduced in [62]. Let $\delta = 1/H$, $\Delta = \sqrt{\delta/K}/(4\sqrt{2})$. Each d -rectangular linear DRMDP in
846 this family is parameterized by a Boolean vector $\xi = \{\xi_h\}_{h \in [H-1]}$, where $\xi_h \in \{-\Delta, \Delta\}^d$. For
847 a given ξ and uncertainty level $\rho \in (0, 3/4]$, the corresponding d -rectangular linear DRMDP M_ξ^ρ
848 has the following structure. The state space $\mathcal{S} = \{x_1, x_2, \dots, x_H, x_{H+1}\}$ and the action space
849 $\mathcal{A} = \{-1, 1\}^d$. The first state is always x_1 . The feature mapping $\phi : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}^{2d+2}$ is defined to
850 depend on the state x_h through ξ_h as follows:

$$\begin{aligned} \phi(x_1, a) &= \begin{pmatrix} \frac{1}{2d} - \frac{\delta}{d} - \xi_{11}a_1 \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{12}a_2 \\ \vdots \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{1d}a_d \\ \frac{1}{2} \\ \frac{\delta}{d} + \xi_{11}a_1 \\ \frac{\delta}{d} + \xi_{12}a_2 \\ \vdots \\ \frac{\delta}{d} + \xi_{1d}a_d \\ 0 \end{pmatrix}, \phi(x_2, a) = \begin{pmatrix} \frac{1}{2d} - \frac{\delta}{d} - \xi_{21}a_1 \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{22}a_2 \\ \vdots \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{2d}a_d \\ \frac{1}{2} \\ \frac{\delta}{d} + \xi_{21}a_1 \\ \frac{\delta}{d} + \xi_{22}a_2 \\ \vdots \\ \frac{\delta}{d} + \xi_{2d}a_d \\ 0 \end{pmatrix}, \dots, \\ \phi(x_{H-1}, a) &= \begin{pmatrix} \frac{1}{2d} - \frac{\delta}{d} - \xi_{H-1,1}a_1 \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{H-1,2}a_2 \\ \vdots \\ \frac{1}{2d} - \frac{\delta}{d} - \xi_{H-1,d}a_d \\ \frac{1}{2} \\ \frac{\delta}{d} + \xi_{H-1,1}a_1 \\ \frac{\delta}{d} + \xi_{H-1,2}a_2 \\ \vdots \\ \frac{\delta}{d} + \xi_{H-1,d}a_d \\ 0 \end{pmatrix}, \phi(x_H, a) = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 0 \\ 0 \\ 0 \\ \vdots \\ 0 \\ 1 \end{pmatrix}, \phi(x_{H+1}, a) = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 0 \\ \frac{1}{d} \\ \frac{1}{d} \\ \vdots \\ \frac{1}{d} \\ 0 \end{pmatrix}. \end{aligned}$$

851 We assume that

$$K \geq 9d^2H/32 \text{ and } H \geq 6, \quad (\text{G.1})$$

such that $\frac{1}{2d} - \frac{1}{dH} - \delta \geq 0$. Then it can be easily checked that for any $s \in \mathcal{S}$, we have $\phi_i(s, a) \geq 0$ and $\sum_{i=1}^{2d+2} \phi_i(s, a) = 1$. The factor distribution $\mu_1 : \mathcal{S} \rightarrow \mathbb{R}^{2d+2}$ is defined as follows.

$$\mu_1(\cdot) = (\underbrace{\delta_{x_2}(\cdot), \dots, \delta_{x_2}(\cdot)}_{d \text{ terms}}, \delta_{x_2}(\cdot), \underbrace{\delta_{x_{H+1}}(\cdot), \dots, \delta_{x_{H+1}}(\cdot)}_{d \text{ terms}}, \delta_{x_H}(\cdot))^\top.$$

Similarly, for $h = 2, \dots, H$, we have

$$\mu_2(\cdot) = (\delta_{x_3}(\cdot), \dots, \delta_{x_3}(\cdot), \delta_{x_3}(\cdot), \delta_{x_{H+1}}(\cdot), \dots, \delta_{x_{H+1}}(\cdot), \delta_{x_H}(\cdot))^\top, \\ \dots$$

$$\mu_{H-1}(\cdot) = \mu_H(\cdot) = (\delta_{x_H}(\cdot), \dots, \delta_{x_H}(\cdot), \delta_{x_H}(\cdot), \delta_{x_{H+1}}(\cdot), \dots, \delta_{x_{H+1}}(\cdot), \delta_{x_H}(\cdot))^\top,$$

Note that for each episode k , the initial state s_1^k is always x_1 . In the nominal environment, at step h , the state s_h^k is either x_h or x_{H+1} . State x_H and x_{H+1} are absorbing states. Figure 4(a) illustrates the nominal MDP.

Now we construct the reward parameters $\{\theta_h\}_{h \in [H]}$ as follows.

$$\theta_h = (1, 1, \dots, 1, -1, 1, 1, \dots, 1, 0)^\top, \forall h \in [H].$$

We have $\forall h \in [H]$,

$$r_h(x_H, a) = \phi(x_H, \mathbf{a})^\top \theta_h = 0, \\ r_h(x_h, a) = \phi(x_h, \mathbf{a})^\top \theta_h = 0, \\ r_h(x_{H+1}, a) = \phi(x_{H+1}, \mathbf{a})^\top \theta_h = 1.$$

Thus, only the transition starting from x_{H+1} generates a reward of 1, and transitions starting from any other state generate 0 reward. Next, we consider the model perturbation. An observation is that x_H is the worst state since it is an absorbing state with zero reward. By the definition of the d -rectangular uncertainty set, the worst case kernel is the linear combination of worst case factor distributions. Further, by the definition of the factor uncertainty set, the worst case factor distribution is the one that leads to the highest probability ρ to the worst state x_H . Thus, the worst factor distributions are

$$\check{\mu}_1 = ((1 - \rho)\delta_{x_2} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_2} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_2} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_2} + \rho\delta_{x_H}, \\ (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \delta_{x_H})^\top, \\ \check{\mu}_2 = ((1 - \rho)\delta_{x_3} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_3} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_3} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_3} + \rho\delta_{x_H}, \\ (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \delta_{x_H})^\top, \\ \dots \\ \check{\mu}_{H-1} = ((1 - \rho)\delta_{x_{H-1}} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_{H-1}} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_{H-1}} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_{H-1}} + \rho\delta_{x_H}, \\ (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \dots, (1 - \rho)\delta_{x_{H+1}} + \rho\delta_{x_H}, \delta_{x_H})^\top, \\ \check{\mu}_{H-1} = \mu_H.$$

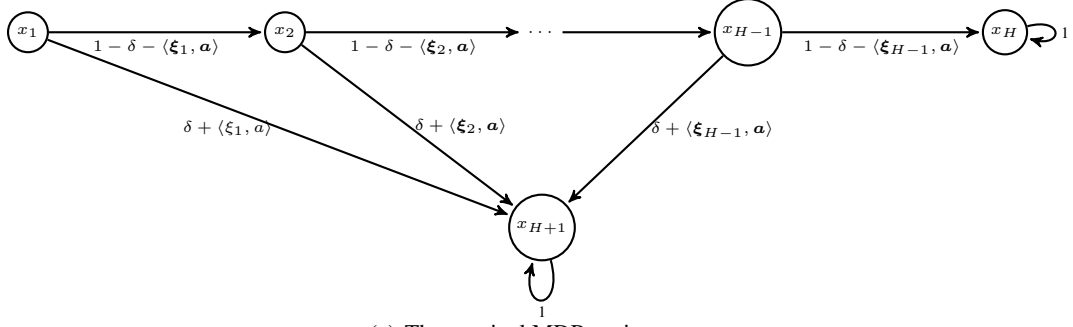
Figure 4(b) illustrates the worst case MDP.

G.2 Reduction from d -Rectangular DRMDP to Linear Bandits

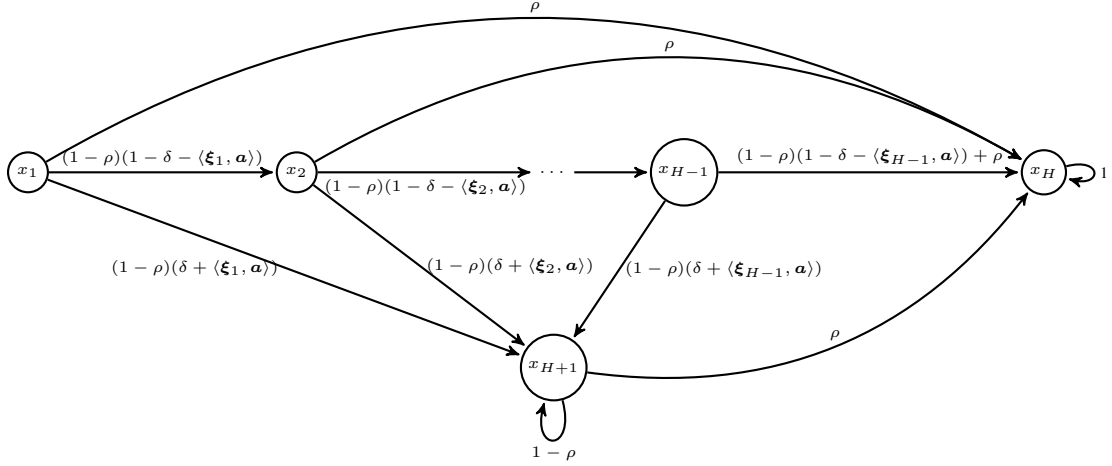
Note that by construction, at steps $h = 1, \dots, H - 2$, the probability of transitioning to the worst case state x_H is independent of the action a . Moreover, since x_{H+1} is the only rewarding state, so the optimal action at step h is the one that leads to the largest probability to x_{H+1} , i.e., $\mathbf{a}_h^* = \arg\max_{\mathbf{a} \in \mathcal{A}} \langle \xi_h, \mathbf{a} \rangle$. Further, in the nominal environment, state x_h can only be reached through states $x_1, x_2, \dots, x_{h-1}$. As discussed by [62], knowing the state x_h is equivalent to knowing the entire history starting from the initial state at current episode. Consequently, policies dictating what actions to take upon reaching a state at the beginning of an episode are equivalent to policies relying on the “within episode” history (we refer to the discussion in E.1 of [62] for more details). In the following lemma, we show that the average suboptimality of the d -rectangular DRMDP can be lower bounded by the regret of $H/2$ bandit instances.

Lemma G.1. With the choice of d, K, H in (G.1), we have $3d\Delta \leq \delta$. Fix $\xi = \{\xi_h\}_{h \in [H-1]}$. Fix a possibly history dependent policy π and define $\bar{\mathbf{a}}_h^\pi = \mathbb{E}_\xi[\mathbf{a}_h | s_h = x_h]$ as the expected action taken by the policy when it visits state x_h in stage h . Then, there exist a constant $c > 0$ such that

$$V_1^{\star, \rho}(x_1) - V_1^{\pi, \rho}(x_1) \geq c \min\left\{\frac{1}{\rho}, H\right\} \sum_{h=1}^{H/2} \left(\max_{\mathbf{a} \in \mathcal{A}} \langle \mu_h, \mathbf{a} \rangle - \langle \mu_h, \bar{\mathbf{a}}_h^\pi \rangle \right).$$



(a) The nominal MDP environment.



(b) The worst-case MDP environment.

Figure 4: Constructions of the nominal MDP and the worst-case MDP environments.

881 *Proof of Lemma G.1.* For the fixed policy π , we first get the ground truth robust value $V_1^{\pi, \rho}(x_1)$ by
 882 induction. Starting from the last step H , we have

$$V_H^{\pi, \rho}(x_H) = 0, \quad V_H^{\pi, \rho}(x_{H+1}) = 1.$$

883 For step $H - 1$, we have

$$V_{H-1}^{\pi, \rho}(x_H) = 0, \quad V_{H-1}^{\pi, \rho}(x_{H-1}) = (1 - \rho)(\delta + \langle \xi_{H-1}, \mathbf{a}_{H-1}^\pi \rangle) \cdot 1, \quad V_{H-1}^{\pi, \rho}(x_{H+1}) = 1 + (1 - \rho) \cdot 1.$$

884 For step $H - 2$, we have $V_{H-2}^{\pi, \rho}(x_H) = 0$ and

$$\begin{aligned} V_{H-2}^{\pi, \rho}(x_{H+1}) &= 1 + (1 - \rho) \cdot V_{H-1}(x_{H+1}) = 1 + (1 - \rho) + (1 - \rho)^2, \\ V_{H-2}^{\pi, \rho}(x_{H-2}) &= (1 - \rho)(\delta + \langle \xi_{H-2}, \bar{\mathbf{a}}_{H-2}^\pi \rangle) \cdot V_{H-1}^{\pi, \rho}(x_{H+1}) \\ &\quad + (1 - \rho)(1 - \delta - \langle \xi_{H-2}, \bar{\mathbf{a}}_{H-2}^\pi \rangle) \cdot V_{H-1}^{\pi, \rho}(x_{H-1}) \\ &= [(1 - \rho) + (1 - \rho)^2](\delta + \langle \xi_{H-2}, \bar{\mathbf{a}}_{H-2}^\pi \rangle) \\ &\quad + (1 - \rho)^2(1 - \delta - \langle \xi_{H-2}, \bar{\mathbf{a}}_{H-2}^\pi \rangle)(\delta + \langle \xi_{H-1}, \bar{\mathbf{a}}_{H-1}^\pi \rangle). \end{aligned}$$

885 For step $H - 3$, we have $V_{H-3}^{\pi, \rho}(x_H) = 0$ and

$$\begin{aligned} V_{H-3}^{\pi, \rho}(x_{H+1}) &= 1 + (1 - \rho) \cdot V_{H-2}(x_{H+1}) = 1 + (1 - \rho) + (1 - \rho)^2 + (1 - \rho)^3, \\ V_{H-3}^{\pi, \rho}(x_{H-3}) &= (1 - \rho)(\delta + \langle \xi_{H-3}, \bar{\mathbf{a}}_{H-3}^\pi \rangle) \cdot V_{H-2}^{\pi, \rho}(x_{H+1}) \\ &\quad + (1 - \rho)(1 - \delta - \langle \xi_{H-3}, \bar{\mathbf{a}}_{H-3}^\pi \rangle) \cdot V_{H-2}^{\pi, \rho}(x_{H-2}) \\ &= [(1 - \rho) + (1 - \rho)^2 + (1 - \rho)^3](\delta + \langle \xi_{H-3}, \bar{\mathbf{a}}_{H-3}^\pi \rangle) \\ &\quad + [(1 - \rho)^2 + (1 - \rho)^3](1 - \delta - \langle \xi_{H-3}, \bar{\mathbf{a}}_{H-3}^\pi \rangle)(\delta + \langle \xi_{H-2}, \bar{\mathbf{a}}_{H-2}^\pi \rangle) \end{aligned}$$

$$+ (1 - \rho)^3 (1 - \delta - \langle \xi_{H-3}, \bar{a}_{H-3}^\pi \rangle) (1 - \delta - \langle \xi_{H-2}, \bar{a}_{H-2}^\pi \rangle) (\delta + \langle \xi_{H-1}, \bar{a}_{H-1}^\pi \rangle).$$

886 Keep performing the backward induction until step $h = 1$, we have

$$\begin{aligned} & V_1^{\pi, \rho}(x_1) \\ &= V_{H-(H-1)}^{\pi, \rho}(x_1) \\ &= [(1 - \rho) + \dots + (1 - \rho)^{H-1}] (\delta + \langle \xi_1, \bar{a}_1^\pi \rangle) + \\ & \quad [(1 - \rho)^2 + \dots + (1 - \rho)^{H-1}] (1 - \delta - \langle \xi_1, \bar{a}_1^\pi \rangle) (\delta + \langle \xi_2, \bar{a}_2^\pi \rangle) + \\ & \quad [(1 - \rho)^3 + \dots + (1 - \rho)^{H-1}] (1 - \delta - \langle \xi_1, \bar{a}_1^\pi \rangle) (1 - \delta - \langle \xi_2, \bar{a}_2^\pi \rangle) (\delta + \langle \xi_3, \bar{a}_3^\pi \rangle) + \\ & \quad \dots + \\ & \quad (1 - \rho)^{H-1} (1 - \delta - \langle \xi_1, \bar{a}_1^\pi \rangle) (1 - \delta - \langle \xi_2, \bar{a}_2^\pi \rangle) \dots (1 - \delta - \langle \xi_{H-2}, \bar{a}_{H-2}^\pi \rangle) (\delta + \langle \xi_{H-1}, \bar{a}_{H-1}^\pi \rangle) \\ &= \sum_{h=1}^{H-1} \left(\sum_{i=h}^{H-1} (1 - \rho)^i \right) (o_h + \delta) \prod_{j=1}^{h-1} (1 - o_j - \delta), \end{aligned} \quad (G.2)$$

887 where $o_h = \langle \xi_h, \bar{a}_h^\pi \rangle, \forall h \in [H]$. Recall that the optimal robust action at step h is $\mathbf{a}_h^* =$
888 $\operatorname{argmax}_{\mathbf{a} \in \mathcal{A}} \langle \xi_h, \mathbf{a} \rangle$, and hence $\max_{\mathbf{a} \in \mathcal{A}} \langle \xi_h, \mathbf{a} \rangle = \Delta d$. Thus, we have

$$V_1^{\star, \rho}(x_1) = \sum_{h=1}^{H-1} \left(\sum_{i=h}^{H-1} (1 - \rho)^i \right) (d\Delta + \delta) \prod_{j=1}^{h-1} (1 - d\Delta - \delta). \quad (G.3)$$

889 For $k \in [H - 1]$, we define

$$S_k = \sum_{h=k}^{H-1} \left(\sum_{i=h-k+1}^{H-k} (1 - \rho)^i \right) \prod_{j=k}^{h-1} (1 - o_j - \delta) (o_h + \delta), \quad (G.4)$$

$$T_k = \sum_{h=k}^{H-1} \left(\sum_{i=h-k+1}^{H-k} (1 - \rho)^i \right) \prod_{j=k}^{h-1} (1 - d\Delta - \delta) (d\Delta + \delta). \quad (G.5)$$

890 Then by (G.2), (G.3), (G.4) and (G.5), we know $V_1^{\star, \rho}(x_1) - V_1^{\pi, \rho}(x_1) = T_1 - S_1$. Next, we aim to
891 lower bound $T_1 - S_1$. Inspired by the backward induction process, we have

$$\begin{aligned} S_k &= \left(\sum_{i=1}^{H-k} (1 - \rho)^i \right) (o_k + \delta) + S_{k+1} (1 - o_k - \delta), \\ T_k &= \left(\sum_{i=1}^{H-k} (1 - \rho)^i \right) (d\Delta + \delta) + T_{k+1} (1 - d\Delta - \delta). \end{aligned}$$

892 Then, we have

$$\begin{aligned} T_k - S_k &= \left(\sum_{i=1}^{H-k} (1 - \rho)^i \right) (d\Delta - o_k) - S_{k+1} (1 - o_k - \delta) + T_{k+1} (1 - d\Delta - \delta) \\ &= \left(\sum_{i=1}^{H-k} (1 - \rho)^i - T_{k+1} \right) (d\Delta - o_k) + (1 - o_k - \delta) (T_{k+1} - S_{k+1}). \end{aligned} \quad (G.6)$$

893 Define $T_H = S_H = 0$, then by the recursive formula (G.6), we have

$$T_1 - S_1 = \sum_{h=1}^{H-1} (d\Delta - o_h) \underbrace{\left(\sum_{i=1}^{H-h} (1 - \rho)^i - T_{h+1} \right)}_I \prod_{j=1}^{h-1} (1 - o_j - \delta). \quad (G.7)$$

894 To further bound (G.7), we first study the term I. Next we derive a close form expression of T_k . In
895 specific, we have

$$T_k = \sum_{h=k}^{H-1} \left(\sum_{i=h-k+1}^{H-k} (1 - \rho)^i \right) \prod_{j=k}^{h-1} (1 - d\Delta - \delta) (d\Delta + \delta)$$

$$\begin{aligned}
&= \left(\sum_{i=1}^{H-k} (1-\rho)^i \right) (d\Delta + \delta) + \left(\sum_{i=2}^{H-k} (1-\rho)^i \right) (1-d\Delta - \delta)(d\Delta + \delta) \\
&\quad + \left(\sum_{i=3}^{H-k} (1-d\Delta - \delta)^2 (d\Delta + \delta) \right) + \cdots + (1-\rho)^{H-k} (1-d\Delta - \delta)^{H-k-1} (d\Delta + \delta).
\end{aligned} \tag{G.8}$$

896 Multiply T_k by $(1-d\Delta - \delta)$, we have

$$\begin{aligned}
&(1-d\Delta - \delta)T_k \\
&= \left(\sum_{i=1}^{H-k} (1-\rho)^i \right) (d\Delta + \delta)(1-d\Delta - \delta) + \left(\sum_{i=2}^{H-k} (1-\rho)^i \right) (1-d\Delta - \delta)^2 (d\Delta + \delta) \\
&\quad + \left(\sum_{i=3}^{H-k} (1-d\Delta - \delta)^2 (d\Delta + \delta) \right) + \cdots + (1-\rho)^{H-k} (1-d\Delta - \delta)^{H-k} (d\Delta + \delta).
\end{aligned} \tag{G.9}$$

897 Then we have

$$\begin{aligned}
&(\text{G.8}) - (\text{G.9}) \\
&= (d\Delta + \delta)T_k \\
&= \left(\sum_{i=1}^{H-k} (1-\rho)^i \right) (d\Delta + \delta) - (1-\rho)(1-d\Delta - \delta)(d\Delta + \delta) - (1-\rho)^2(1-d\Delta - \delta)^2(d\Delta + \delta) \\
&\quad - \cdots - (1-\rho)^{H-k}(1-d\Delta - \delta)^{H-k}(d\Delta + \delta).
\end{aligned} \tag{G.10}$$

898 Divide both side of equation (G.10) by $(d\Delta + \delta)$ and then apply the formula for the sum of a geometric
899 series, we know T_k has the following closed form expression

$$T_k = \left(\sum_{i=1}^{H-k} (1-\rho)^i \right) - \frac{(1-\rho)(1-d\Delta - \delta)(1 - (1-\rho)^{H-k}(1-d\Delta - \delta)^{H-k})}{1 - (1-\rho)(1-d\Delta - \delta)}.$$

900 Then, for any $h \leq H/2$, we have the following bound on the term I of (G.7),

$$\begin{aligned}
&\sum_{i=1}^{H-h} (1-\rho)^i - T_{h+1} \\
&= \sum_{i=1}^{H-h} (1-\rho)^i - \sum_{i=1}^{H-h-1} (1-\rho)^i + \frac{(1-\rho)(1-d\Delta - \delta)(1 - (1-\rho)^{H-h-1}(1-d\Delta - \delta)^{H-h-1})}{1 - (1-\rho)(1-d\Delta - \delta)} \\
&= (1-\rho)^{H-h} + \frac{(1-\rho)(1-d\Delta - \delta)(1 - (1-\rho)^{H-h-1}(1-d\Delta - \delta)^{H-h-1})}{1 - (1-\rho)(1-d\Delta - \delta)} \\
&= (1-\rho)^{H-h} + (1-\rho)(1-d\Delta - \delta) + \cdots + (1-\rho)^{H-h-1}(1-d\Delta - \delta)^{H-h-1} \\
&\geq (1-d\Delta - \delta)^H ((1-\rho) + \cdots + (1-\rho)^{H-h-1} + (1-\rho)^{H-h}) \\
&\geq \left(1 - \frac{2}{H}\right)^H ((1-\rho) + \cdots + (1-\rho)^{H-h-1} + (1-\rho)^{H-h})
\end{aligned} \tag{G.11}$$

$$\geq \frac{1}{12} \sum_{i=1}^{H-h} (1-\rho)^i, \tag{G.12}$$

901 where (G.11) holds due to $3d\Delta \leq \delta = 1/H$ and (G.12) holds due to $H \geq 6$. Next, we carefully
902 bound the LHS of (G.12) with respect to ρ . For any $h \leq H/2$ and $\rho \in (0, 3/4]$, we have

$$\frac{1}{12} \sum_{i=1}^{H-h} (1-\rho)^i \geq \frac{1}{12} \frac{(1-\rho)(1 - (1-\rho)^{H/2})}{\rho} \geq \frac{1}{50} \frac{1 - (1-\rho)^{H/2}}{\rho}.$$

903 Given the fact that

$$\frac{1 - (1-\rho)^{H/2}}{\rho} = \Theta\left(\min\left(H, \frac{1}{\rho}\right)\right),$$

904 there exist a constant $c > 0$, such that

$$\frac{1 - (1 - \rho)^{H/2}}{\rho} \geq c \cdot \min\left(H, \frac{1}{\rho}\right).$$

905 Then we have

$$\sum_{i=1}^{H-h} (1 - \rho)^i - T_{h+1} \geq c' \cdot \min\left(H, \frac{1}{\rho}\right), \quad (\text{G.13})$$

906 where $c' = c/50$. Moreover, with the choice of parameter $3d\Delta \leq \delta$, $\delta = 1/H$, and $H \geq 6$, we have

$$\prod_{j=1}^{h-1} (1 - o_j - \delta) \geq (1 - 4\delta/3)^H \geq 1/3. \quad (\text{G.14})$$

907 Therefore, by (G.7), (G.13) and (G.14), we have

$$\begin{aligned} V_1^{\star, \rho}(x_1) - V_1^{\pi, \rho}(x_1) &= T_1 - S_1 \\ &\geq c'' \cdot \min\{H, 1/\rho\} \cdot \sum_{h=1}^{H/2} (d\Delta - o_h) \\ &= c'' \cdot \min\{H, 1/\rho\} \cdot \sum_{h=1}^{H/2} \left(\max_{a \in \mathcal{A}} \langle \boldsymbol{\mu}_h, \mathbf{a} \rangle - \langle \boldsymbol{\mu}_h, \bar{\mathbf{a}}_h^\pi \rangle \right), \end{aligned}$$

908 where $c'' = c'/3$. This completes the proof. $\square$

909 G.3 Proof of Theorem 5.1

910 Next, we present an existing result on lower bounding the regret of linear bandits induced by
911 Lemma G.1. This result is useful in deriving the lower bound in Theorem 5.1.

912 **Lemma G.2.** [62, Lemma 25] Fix a positive real $0 \leq \delta \leq 1/3$, and positive integers K, d and
913 assume that $K \geq d^2/(2\delta)$. Let $\Delta = \sqrt{\delta/K}/(4\sqrt{2})$ and consider the linear bandit problems $\mathcal{L}_\mu$
914 parameterized with a parameter vector $\boldsymbol{\mu} \in \{-\Delta, \Delta\}^d$ and action set $\mathcal{A} = \{-1, 1\}^d$ so that the
915 reward distribution for taking action $\mathbf{a} \in \mathcal{A}$ is a Bernoulli distribution $\text{Bernoulli}(\delta + \langle \boldsymbol{\mu}, \mathbf{a} \rangle)$. Then
916 for any bandit algorithm $\mathcal{B}$, there exists a $\boldsymbol{\mu}^\star \in \{-\Delta, \Delta\}^d$ such that the expected pseudo-regret of $\mathcal{B}$
917 over first K steps on bandit $\mathcal{L}_{\boldsymbol{\mu}^\star}$ is lower bounded as follows:

$$\mathbb{E}_{\boldsymbol{\mu}^\star} \text{Regret}(K) \geq \frac{d\sqrt{K}\delta}{8\sqrt{2}}.$$

918 Note that the expectation is with respect to a distribution that depends both on $\mathcal{B}$ and $\boldsymbol{\mu}^\star$, but since $\mathcal{B}$
919 is fixed, this dependence is hidden.

920 Now we are ready to prove the lower bound in Theorem 5.1.

921 *Proof of Theorem 5.1.* By Lemma G.1, we have

$$\begin{aligned} \mathbb{E}_{\boldsymbol{\xi}} \text{AveSubopt}(M_{\boldsymbol{\xi}}, K) &= \frac{1}{K} \mathbb{E}_{\boldsymbol{\xi}} \left[\sum_{k=1}^K [V_1^{\star, \rho}(x_1) - V_1^{\pi, \rho}(x_1)] \right] \\ &\geq c \cdot \frac{\min\{H, 1/\rho\}}{K} \sum_{h=1}^{H/2} \mathbb{E}_{\boldsymbol{\xi}} \left[\sum_{k=1}^K \left(\max_{a \in \mathcal{A}} \langle \boldsymbol{\xi}_h, \mathbf{a} \rangle - \langle \boldsymbol{\xi}_h, \bar{\mathbf{a}}_h^{\pi_k} \rangle \right) \right]. \end{aligned}$$

922 Note that the learning process is conducted on the nominal environment, which is exactly the MDP in
923 [62], thus the rest proof of Theorem 4.2 follows the argument in the proof of Theorem 8 in [62]. In
924 particular, define $\boldsymbol{\xi}^{-h} = (\boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_{h-1}, \boldsymbol{\xi}_{h+1}, \dots, \boldsymbol{\xi}_H)$, then every MDP policy π induces a bandit
925 algorithm $\mathcal{B}_{\pi, h, \boldsymbol{\xi}^{-h}}$ for the linear bandit of Lemma G.2. Moreover, our choice of parameters in

(G.1) satisfy the requirement of Lemma G.2. Denote the regret of this bandit problem on $\mathcal{L}_\xi$ as $\text{BanditRegret}(\mathcal{B}_{\pi,h,\xi^{-h}}, \xi_h)$, then we have

$$\begin{aligned}
\sup_{\xi} \mathbb{E}_{\xi} \text{AveSubopt}(M_{\xi}, K) &\geq \sup_{\xi} c \cdot \frac{\min\{H, 1/\rho\}}{K} \sum_{h=1}^{H/2} \text{BanditRegret}(\mathcal{B}_{\pi,h,\xi^{-h}}, \xi_h) \\
&\geq \sup_{\xi} c \cdot \frac{\min\{H, 1/\rho\}}{K} \sum_{h=1}^{H/2} \inf_{\xi^{-h}} \text{BanditRegret}(\mathcal{B}_{\pi,h,\xi^{-h}}, \xi_h) \\
&= c \cdot \frac{\min\{H, 1/\rho\}}{K} \sum_{h=1}^{H/2} \sup_{\xi} \inf_{\xi^{-h}} \text{BanditRegret}(\mathcal{B}_{\pi,h,\xi^{-h}}, \xi_h) \\
&\geq c \cdot \frac{\min\{H, 1/\rho\} d H \sqrt{K} \delta}{16\sqrt{2} \cdot K} \\
&= \frac{c}{16\sqrt{2}} \cdot \frac{d\sqrt{H} \cdot \min\{H, 1/\rho\}}{\sqrt{K}}.
\end{aligned}$$

This completes the proof. $\square$

H Auxiliary Lemmas

In this section, we present some standard technical results in the literature that our proofs are built on.

Proposition H.1. (Strong duality for TV [36, Lemma 4]). Given any probability measure μ^0 over $\mathcal{S}$, a fixed uncertainty level ρ , the uncertainty set $\mathcal{U}^\rho(\mu^0) = \{\mu : \mu \in \Delta(\mathcal{S}), D_{TV}(\mu || \mu^0) \leq \rho\}$, and any function $V : \mathcal{S} \rightarrow [0, H]$, we obtain

$$\inf_{\mu \in \mathcal{U}^\rho(\mu^0)} \mathbb{E}_{s \sim \mu} V(s) = \max_{\alpha \in [V_{\min}, V_{\max}]} \left\{ \mathbb{E}_{s \sim \mu^0} [V(s)]_{\alpha} - \rho(\alpha - \min_{s'} [V(s')]_{\alpha}) \right\}, \quad (\text{H.1})$$

where $[V(s)]_{\alpha} = \min\{V(s), \alpha\}$, $V_{\min} = \min_s V(s)$ and $V_{\max} = \max_s V(s)$. Notably, the range of α can be relaxed to $[0, H]$ without impacting the optimization.

Lemma H.2. [1, Lemma 12] Let $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{C}$ be positive semi-definite matrices such that $\mathbf{A} = \mathbf{B} + \mathbf{C}$. Then we have that

$$\sup_{\mathbf{x} \neq 0} \frac{\mathbf{x}^\top \mathbf{A} \mathbf{x}}{\mathbf{x}^\top \mathbf{B} \mathbf{x}} \leq \frac{\det(\mathbf{A})}{\det(\mathbf{B})}.$$

Lemma H.3. [1, Confidence Ellipsoid, Theorem 2] Let $\{\mathcal{G}_k\}_{k=1}^\infty$ be a filtration, and $\{\mathbf{x}_k, \eta_k\}_{k \geq 1}$ be a stochastic process such that $\mathbf{x}_k \in \mathbb{R}^d$ is $\mathcal{G}_k$ -measurable and $\eta_k \in \mathbb{R}$ is $\mathcal{G}_{k+1}$ -measurable. Let $L, \sigma, \Sigma, \epsilon > 0, \mu^* \in \mathbb{R}^d$. For $k \geq 1$, let $y_k = \langle \mu^*, \mathbf{x}_k \rangle + \eta_k$ and suppose that $\eta_k, \mathbf{x}_k$ also satisfy

$$\mathbb{E}[\eta_k | \mathcal{G}_k] = 0, |\eta_k| \leq R, \|\mathbf{x}_k\|_2 \leq L.$$

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

$$\beta_k = R \sqrt{d \log \left(1 + \frac{kL^2}{d\lambda} \right) + 2 \log \frac{1}{\delta}}.$$

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

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

Lemma H.4. [19, Lemma D.1] Let $\mathbf{A}_t = \lambda \mathbf{I} + \sum_{i=1}^t \phi_i \phi_i^\top$, where $\phi_i \in \mathbb{R}^d$ and $\lambda > 0$. Then we have

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

945 **Lemma H.5.** [14, Lemma D.5] Let $\mathbf{A} \in \mathbb{R}^{d \times d}$ be a positive definite matrix where its largest
 946 eigenvalue $\lambda_{\max}(\mathbf{A}) \leq \lambda$. Let $\mathbf{x}_1, \dots, \mathbf{x}_k$ be k vectors in $\mathbb{R}^d$. Then it holds that

$$\left\| \mathbf{A} \sum_{i=1}^k \mathbf{x}_i \right\| \leq \sqrt{\lambda k} \left(\sum_{i=1}^k \|\mathbf{x}_i\|_{\mathbf{A}}^2 \right)^{1/2}.$$

947 **Lemma H.6.** [39, Covering number of Euclidean ball] For any $\varepsilon > 0$, $\mathcal{N}_\varepsilon$, the ε -covering number
 948 of the Euclidean ball of radius $B > 0$ in $\mathbb{R}^d$ satisfies

$$\mathcal{N}_\varepsilon \leq \left(1 + \frac{2B}{\varepsilon} \right)^d \leq \left(\frac{3B}{\varepsilon} \right)^d.$$

949 **Lemma H.7.** [39, Covering number of an interval] Denote the ϵ -covering number of the closed
 950 interval $[a, b]$ for some real number $b > a$ with respect to the distance metric $d(\alpha_1, \alpha_2) = |\alpha_1 - \alpha_2|$
 951 as $\mathcal{N}_\epsilon([a, b])$. Then we have $\mathcal{N}_\epsilon([a, b]) \leq 3(b - a)/\epsilon$.

952 **Lemma H.8.** [61, Theorem 4.3] Let $\{\mathcal{G}_k\}_{k=1}^\infty$ be a filtration, and $\{\mathbf{x}_k, \eta_k\}_{k \geq 1}$ be a stochastic process
 953 such that $\mathbf{x}_k \in \mathbb{R}^d$ is $\mathcal{G}_k$ -measurable and $\eta_k \in \mathbb{R}$ is $\mathcal{G}_{k+1}$ -measurable. Let $L, \sigma > 0, \boldsymbol{\mu}^* \in \mathbb{R}^d$. For
 954 $k \geq 1$, let $y_k = \langle \boldsymbol{\mu}^*, \mathbf{x}_k \rangle + \eta_k$ and suppose that $\eta_k, \mathbf{x}_k$ also satisfy

$$\mathbb{E}[\eta_k \mid \mathcal{G}_k] = 0, \mathbb{E}[\eta_k^2 \mid \mathcal{G}_k] \leq \sigma^2, |\eta_k| \leq R, \|\mathbf{x}_k\|_2 \leq L.$$

955 For $k \geq 1$, let $\beta_k = \tilde{O}(\sigma\sqrt{d} + \max_{1 \leq i \leq k} |\eta_i| \min \{1, \|\mathbf{x}_i\|_{\mathbf{Z}_{i-1}^{-1}}\})$ and $\mathbf{Z}_k = \lambda \mathbf{I} + \sum_{i=1}^k \mathbf{x}_i \mathbf{x}_i^\top$,
 956 $\mathbf{b}_k = \sum_{i=1}^k y_i \mathbf{x}_i$, $\boldsymbol{\mu}_k = \mathbf{Z}_k^{-1} \mathbf{b}_k$. Then, for any $0 < \delta < 1$, with probability at least $1 - \delta$, for all
 957 $k \in [K]$, we have

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

958 **Lemma H.9.** [19, Lemma D.4] Let $\{s_i\}_{i=1}^\infty$ be a stochastic process on state space $\mathcal{S}$ with corre-
 959 sponding filtration $\{\mathcal{F}_i\}_{i=1}^\infty$. Let $\{\phi_i\}_{i=1}^\infty$ be an $\mathbb{R}^d$ -valued stochastic process where $\phi_i \in \mathcal{F}_{i-1}$, and
 960 $\|\phi_i\| \leq 1$. Let $\boldsymbol{\Lambda}_k = \lambda \mathbf{I} + \sum_{i=1}^k \phi_i \phi_i^\top$. Then for any $\delta > 0$, with probability at least $1 - \delta$, for all
 961 $k \geq 0$, and any $V \in \mathcal{V}$ with $\sup_{s \in \mathcal{S}} |V(s)| \leq H$, we have

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

962 where $\mathcal{N}_\varepsilon$ is the ε -covering number of $\mathcal{V}$ with respect to the distance $\text{dist}(V, V') = \sup_{s \in \mathcal{S}} |V(s) -$
 963 $V'(s)|$.

964 **Lemma H.10.** [24, Lemma 5.1 (Range Shrinkage)] For any $(\rho, \pi, h) \in (0, 1] \times \Pi \times [H]$, we have
 965 $\max_{s \in \mathcal{S}} V_h^{\pi, \rho}(s) - \min_{s \in \mathcal{S}} V_h^{\pi, \rho}(s) \leq (1 - (1 - \rho)^{H-h+1})/\rho$.

NeurIPS Paper Checklist

1. Claims

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

Answer: [Yes]

Justification: The main claims made in the abstract and introduction accurately reflect the paper's contributions and scope. This work focus on the theoretical side of online learning of linear DRMDP. A tight upper bound and the first lower bound is provided, both of which depend on non-trivial algorithm and hard instance design, as well as technical analysis.

Guidelines:

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

2. Limitations

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

Answer: [Yes]

Justification: There is a $O(\sqrt{H})$ gap between our upper bound and lower bound. We discuss this limitation in Section 5.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

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

Answer: [Yes]

Justification: We provide detailed discussion and justification of assumptions used in the main context and rigorous proof of theorems in the supplementary material.

Guidelines:

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

4. Experimental result reproducibility

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

Answer: [Yes]

Justification: Details on experiment setup and implementation are provided in Appendix B. They are enough for reproduction of all experiment results.

Guidelines:

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

5. Open access to data and code

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

Answer: [Yes]

Justification: All experiment results can be reproduced by the code in this link: <https://anonymous.4open.science/r/We-Drive-U-9EC3>

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

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

Answer: [Yes]

Justification: Details on experiment setup and implementation are provided in Appendix B. More details can be found in the released code.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [No]

Justification: Not applicable to our experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: All numerical experiments were conducted on a MacBook Pro with a 2.6 GHz 6-Core Intel CPU.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: We have checked that the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics.

Guidelines:

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

10. Broader impacts

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

Answer: [NA]

Justification: There is no societal impact of the work performed. This work focuses on the theoretical side of robust RL, and methods in this paper do not lead to a direct path to any negative applications.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [NA]

Justification: The paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

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

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

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

14. Crowdsourcing and research with human subjects

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

Answer: [NA]

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

Guidelines:

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

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

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

Answer: [NA]

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

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- 1276 • Depending on the country in which research is conducted, IRB approval (or equivalent)
1277 may be required for any human subjects research. If you obtained IRB approval, you
1278 should clearly state this in the paper.
- 1279 • We recognize that the procedures for this may vary significantly between institutions
1280 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
1281 guidelines for their institution.
- 1282 • For initial submissions, do not include any information that would break anonymity (if
1283 applicable), such as the institution conducting the review.

1284 16. **Declaration of LLM usage**

1285 Question: Does the paper describe the usage of LLMs if it is an important, original, or
1286 non-standard component of the core methods in this research? Note that if the LLM is used
1287 only for writing, editing, or formatting purposes and does not impact the core methodology,
1288 scientific rigorousness, or originality of the research, declaration is not required.

1289 Answer: [NA]

1290 Justification: The core method development in this research does not involve LLMs as any
1291 important, original, or non-standard components.

1292 Guidelines:

- 1293 • The answer NA means that the core method development in this research does not
1294 involve LLMs as any important, original, or non-standard components.
- 1295 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>)
1296 for what should or should not be described.