

Distributionally Robust Model-Based Offline Reinforcement Learning with Near-Optimal Sample Complexity

Laixi Shi

*Computing Mathematical Sciences
California Institute of Technology
Pasadena, CA, 91125, USA*

LAIXIS@CALTECH.EDU

Yuejie Chi

*Department of Electrical and Computer Engineering
Carnegie Mellon University
Pittsburgh, PA, 15213, USA*

YUEJIECHI@CMU.EDU

Editor: Shipra Agrawal

Abstract

This paper concerns the central issues of model robustness and sample efficiency in offline reinforcement learning (RL), which aims to learn to perform decision making from history data without active exploration. Due to uncertainties and variabilities of the environment, it is critical to learn a robust policy—with as few samples as possible—that performs well even when the deployed environment deviates from the nominal one used to collect the history dataset. We consider a distributionally robust formulation of offline RL, focusing on tabular robust Markov decision processes with an uncertainty set specified by the Kullback-Leibler divergence in both finite-horizon and infinite-horizon settings. To combat with sample scarcity, a model-based algorithm that combines distributionally robust value iteration with the principle of pessimism in the face of uncertainty is proposed, by penalizing the robust value estimates with a carefully designed data-driven penalty term. Under a mild and tailored assumption of the history dataset that measures distribution shift without requiring full coverage of the state-action space, we establish the finite-sample complexity of the proposed algorithms. We further develop an information-theoretic lower bound, which suggests that learning RMDPs is at least as hard as the standard MDPs when the uncertainty level is sufficient small, and corroborates the tightness of our upper bound up to polynomial factors of the (effective) horizon length for a range of uncertainty levels. To the best of our knowledge, this provides the first provably near-optimal robust offline RL algorithm that learns under model uncertainty and partial coverage.

Keywords: offline/batch reinforcement learning, distributional robustness, pessimism, model-based reinforcement learning, finite-sample analysis

1. Introduction

Reinforcement learning (RL) concerns about finding an optimal policy that maximizes an agent’s expected total reward in an unknown environment. A fundamental challenge of deploying RL to real-world applications is the limited ability to explore or interact with the environment, due to resources, time, or safety constraints. Offline RL, or batch RL,

seeks to circumvent this challenge by resorting to history data—which are often collected by executing some possibly unknown behavior policy in the past—with the hope that the history data might already provide significant insights about the targeted optimal policy without further exploration (Levine et al., 2020).

Besides maximizing the expected total reward, perhaps an equally important goal—to say the least—for an RL agent is safety and robustness (Garcia and Fernández, 2015), especially in high-stake applications such as robotics, autonomous driving, clinical trials, financial investments, and so on (Choi et al., 2009; Schulman et al., 2013). It has been observed that a standard RL agent trained in an ideal environment might be extremely sensitive and fail catastrophically when the deployed environment is subject to small adversarial perturbations (Zhang et al., 2020). Consequently, robust RL has attracted a surge of attentions with the goal to learn an optimal policy that is robust to environment perturbations. In fact, providing robustness guarantees becomes even more relevant in the offline setting, which can be formulated as *robust offline RL*, since the history data is often inevitably collected from a timeframe where it is no longer reasonable to assume model stillness, due to the highly non-stationary and time-varying dynamics of many real-world applications. Altogether, this naturally leads to a question:

Can we learn a near-optimal policy which is robust with respect to uncertainties and variabilities of the environments using as few history samples as possible?

1.1 Challenges and premises in robust offline RL

Despite significant amount of recent activities in robust RL and offline RL, addressing model uncertainty and sample efficiency simultaneously remains challenging due to several key issues that we single out below.

- *Distribution shift.* The history data is generated by following some behavior policy in an outdated environment, which can result in a data distribution that is heavily deviated from the desired one, i.e., induced by the target policy in the deployed environment.
- *Partial and limited coverage.* The history data might only provide partial and limited coverage over the entire state-action space, where the limited sample size leads to a poor estimate of the associated model parameters, and consequently, unreliable policy learning outcomes.

Understanding the implications of—and designing algorithms that work around—these challenges play a major role in advancing the state-of-the-art of robust offline RL. In particular, two prevalent algorithmic ideas, distributional robustness and pessimism, are called out as our guiding principles.

- *Distributional robustness.* Instead of finding an optimal policy in a fixed environment, motivated by the literature in distributionally robust optimization (Delage and Ye, 2010), one might seek to find a policy that achieves the best worst-case performance for all the environments in some uncertainty set around the offline environment, as formalized in the framework of robust RL (Iyengar, 2005; Nilim and El Ghaoui, 2005).

- *Pessimism.* When the samples are scarce, it is wise to act with caution based on the principle of pessimism, where one subtracts a penalty term—representing the confidence of the corresponding estimate—from the value functions to avoid excessive risk. Encouragingly, pessimism has been recently shown as an indispensable ingredient to achieve sample efficiency in offline RL without requiring full coverage (Jin et al., 2021; Rashidinejad et al., 2021; Li et al., 2022), as long as the trajectory of the behavior policy provides sufficient overlap with that of the target policy.

While these two ideas have been proven useful for robust RL and offline RL *separately*, tackling robust offline RL needs novel ingredients that go significantly beyond a naïve combination of existing techniques. This is because, in robust offline RL, one needs to handle the distribution shift induced not only by the behavior policy, but also by model perturbations, thus the penalty term derived from the pessimism principle in standard offline RL is no longer applicable. Indeed, while the value function of standard RL depends linearly with respect to the transition kernel, the dependency between the nominal transition kernel and the robust value function unfortunately becomes highly nonlinear—even without a closed-form expression—making the control of statistical uncertainty extremely challenging in robust offline RL.

1.2 Main contributions

In this work, we provide an affirmative answer to the question raised earlier, by developing a provably efficient model-based algorithm that learns a near-optimal *distributionally-robust* policy from a minimal number of offline samples. Specifically, we consider a Robust Markov Decision Process (RMDP) with S states, A actions in both the nonstationary finite-horizon setting (with horizon length H) and the discounted infinite-horizon setting (with discount factor γ). Different from standard MDPs, RMDPs specify a family transition kernels, which lie within an uncertainty set taken as a small ball of size σ around a nominal transition kernel with respect to the Kullback-Leibler (KL) divergence. Given K episodes (resp. N transitions) of history data drawn by following some behavior policy π^b under the nominal transition kernel in the finite-horizon (resp. infinite-horizon) setting, our goal is to learn the optimal robust policy π^* in the maximin sense, which has the best worst-case value for all the transition kernels within the uncertainty set (Iyengar, 2005; Nilim and El Ghaoui, 2005). Our main results are summarized below.

- We introduce a notion called *robust single-policy clipped concentrability coefficient* $C_{\text{rob}}^* \in [1/S, \infty]$ to quantify the quality of history data, which measures the distribution shift between the behavior policy π^b and the optimal robust policy π^* in the presence of model perturbations, without requiring full coverage of the entire state-action space by the behavior policy. In contrast, prior algorithms (Yang et al., 2022; Zhou et al., 2021; Panaganti and Kalathil, 2022)—using simulator or offline data—all require full coverage of the entire state-action space.
- We propose a novel pessimistic variant of distributionally robust value iteration with a plug-in estimate of the nominal transition kernel (Iyengar, 2005; Nilim and El Ghaoui, 2005), called DRVI-LCB, by penalizing the robust value estimates with a carefully designed data-driven penalty term. We demonstrate that DRVI-LCB finds an ε -optimal

robust policy as soon as the sample size is above $\tilde{O}\left(\frac{SC_{\text{rob}}^* H^5}{P_{\min}^* \sigma^2 \varepsilon^2}\right)$ for the finite-horizon setting and $\tilde{O}\left(\frac{SC_{\text{rob}}^*}{P_{\min}^* \sigma^2 (1-\gamma)^4 \varepsilon^2}\right)$ for the infinite-horizon setting, up to some logarithmic factor after a burn-in cost independent of ε . Here, $P_{\min}^*$ is the smallest positive state transition probability of the optimal robust policy π^* under the nominal kernel.

- To complement the upper bound, we further develop information-theoretic lower bounds for a range of uncertainty levels, showing there exists some transition kernel such that at least $\Omega\left(\frac{SC_{\text{rob}}^* H^4}{\varepsilon^2}\right)$ samples (resp. $\Omega\left(\frac{SC_{\text{rob}}^*}{(1-\gamma)^3 \varepsilon^2}\right)$ samples) are needed to find an ε -optimal robust policy when the uncertainty level $\sigma \lesssim 1/H$ (resp. $\sigma \lesssim (1-\gamma)$), and at least $\Omega\left(\frac{SC_{\text{rob}}^* H^3}{P_{\min}^* \sigma^2 \varepsilon^2}\right)$ samples (resp. $\Omega\left(\frac{SC_{\text{rob}}^*}{P_{\min}^* \sigma^2 (1-\gamma)^2 \varepsilon^2}\right)$ samples) are needed to find an ε -optimal robust policy when the uncertainty level $\sigma \asymp \log(1/P_{\min}^*)$, regardless of the choice of algorithms in the finite-horizon (resp. infinite-horizon) setting. Hence, this suggests that learning RMDPs is at least as hard as the standard MDP (Li et al., 2022) when the uncertainty level is sufficiently small, and corroborates the near-optimality of DRVI-LCB with respect to all key parameters up to a polynomial factor of the horizon length H (resp. the effective horizon length $\frac{1}{1-\gamma}$) for a range of uncertainty levels ($\sigma \asymp \log(1/P_{\min}^*)$).

To the best of our knowledge, our paper is the first work to execute the principle of pessimism in a data-driven manner for robust offline RL, leading to the first provably efficient algorithm that learns under simultaneous model uncertainty and partial coverage of the history dataset. See Table 1 for a summary.

Comparison with prior art under full coverage. Prior works (Yang et al., 2022; Zhou et al., 2021; Panaganti and Kalathil, 2022) have only addressed the infinite-horizon setting under full coverage of the history data. Fortunately, our results also seamlessly cover this easier scenario, by replacing C_{rob}^* with A . Specializing our result to this setting to facilitate comparison, DRVI-LCB finds an ε -optimal robust policy with at most $\tilde{O}\left(\frac{SA}{P_{\min}^* (1-\gamma)^4 \sigma^2 \varepsilon^2}\right)$ samples, which depends linearly with respect to the size of the state space S (ignoring other parameters). In contrast, all prior works (Yang et al., 2022; Zhou et al., 2021; Panaganti and Kalathil, 2022) incur sample complexities that scale at least quadratically with respect to the size of the state space S . In addition, our bound improves the *exponential* dependency on $\frac{1}{1-\gamma}$ of Zhou et al. (2021); Panaganti and Kalathil (2022) to a *polynomial* dependency, as well as the *quadratic* dependency on $1/P_{\min}$ (which satisfies $P_{\min} \leq P_{\min}^*$) of Yang et al. (2022) to a *linear* one on $1/P_{\min}^*$. These improvements further corroborate the benefit of the proposed DRVI-LCB even under full coverage. See Table 2 for detailed comparisons.

1.3 Related works

We shall focus on the closely related works on offline RL and distributionally robust RL.

Offline RL. Focusing on the task of learning an optimal policy from offline data, a significant amount of prior arts sets to understand the sample complexity and efficacy of offline RL under different assumptions of the history dataset. A bulk of prior results requires the history data to cover all the state-action pairs, under assumptions such as uniformly

Horizon	Algorithm	Coverage	Sample complexity	Uncertainty level
infinite-horizon	DRVI-LCB (this work)	partial	$\frac{SC_{\text{rob}}^*}{P_{\min}^*(1-\gamma)^4\sigma^2\varepsilon^2}$	full range
	Lower bound (this work)	partial	$\frac{SC_{\text{rob}}^*}{(1-\gamma)^3\varepsilon^2}$	$\sigma \lesssim (1-\gamma)$
	Lower bound (this work)	partial	$\frac{SC_{\text{rob}}^*}{P_{\min}^*(1-\gamma)^2\sigma^2\varepsilon^2}$	$\sigma \asymp \log(1/P_{\min}^*)$
finite-horizon	DRVI-LCB (this work)	partial	$\frac{SC_{\text{rob}}^* H^5}{P_{\min}^* \sigma^2 \varepsilon^2}$	full range
	Lower bound (this work)	partial	$\frac{SC_{\text{rob}}^* H^4}{\varepsilon^2}$	$\sigma \lesssim 1/H$
	Lower bound (this work)	partial	$\frac{SC_{\text{rob}}^* H^3}{P_{\min}^* \sigma^2 \varepsilon^2}$	$\sigma \asymp \log(1/P_{\min}^*)$

Table 1: Our results for finding an ε -optimal robust policy in the infinite/finite-horizon robust MDPs with an uncertainty set measured with respect to the KL divergence using history data under partial coverage. The sample complexities included in the table are valid for sufficiently small ε , with all logarithmic factors omitted. Here, σ is the uncertainty level, S is the size of the state space, H is the horizon length for the finite-horizon setting, γ is the discount factor for the infinite-horizon setting, C_{rob}^* is the robust single-policy clipped concentrability coefficient, and $P_{\min}^*$ is the smallest positive state transition probability of the nominal kernel *visited by the optimal robust policy* π^* .

bounded concentrability coefficients (Chen and Jiang, 2019; Munos, 2005) and uniformly lower bounded data visitation distribution (Yin and Wang, 2021; Yin et al., 2021), where the latter assumption is also related to studies of asynchronous Q-learning (Li et al., 2021). More recently, the principle of pessimism has been investigated for offline RL in both model-based (Jin et al., 2021; Xie et al., 2021; Rashidinejad et al., 2021; Li et al., 2022) and model-free algorithms (Kumar et al., 2020; Shi et al., 2022; Yan et al., 2023; Woo et al., 2024), without the stringent requirement of full coverage. In particular, Li et al. (2022) established the near-minimax optimality of a pessimistic variant of value iteration under the single-policy clipped concentrability of history data, which inspired our algorithm design in the distributionally robust setting.

Distributionally robust RL. While distributionally robust optimization has been mainly investigated in the context of supervised learning (Rahimian and Mehrotra, 2019; Gao, 2022; Bertsimas et al., 2018; Duchi and Namkoong, 2021; Blanchet and Murthy, 2019; Sinha et al., 2018), distributionally robust dynamic programming has also attracted considerable amount of attention, e.g. Iyengar (2005); Nilim and Ghaoui (2003); Xu and Mannor (2012); Nilim and El Ghaoui (2005), where natural robust extensions to the standard Bell-

Problem type	Algorithm	Coverage	Sample complexity
infinite-horizon	DRVI (Zhou et al., 2021)	full	$\frac{S^2 A \exp\left(O\left(\frac{1}{1-\gamma}\right)\right)}{(1-\gamma)^4 \sigma^2 \varepsilon^2}$
	REVI/DRVI (Panaganti and Kalathil, 2022)	full	$\frac{S^2 A \exp\left(O\left(\frac{1}{1-\gamma}\right)\right)}{(1-\gamma)^4 \sigma^2 \varepsilon^2}$
	DRVI (Yang et al., 2022)	full	$\frac{S^2 A}{P_{\min}^2 (1-\gamma)^4 \sigma^2 \varepsilon^2}$
	DRVI-LCB (this work)	full	$\frac{SA}{P_{\min}^* (1-\gamma)^4 \sigma^2 \varepsilon^2}$
finite-horizon	DRVI-LCB (this work)	full	$\frac{SAH^5}{P_{\min}^* \sigma^2 \varepsilon^2}$

Table 2: Comparisons between our results and prior arts for finding an ε -optimal robust policy in the infinite/finite-horizon robust MDPs with an uncertainty set measured with respect to the KL divergence under full coverage of the history data. The sample complexities included in the table are valid for sufficiently small ε , with all logarithmic factors omitted. Here, σ is the uncertainty level, S is the size of the state space, A is the size of the action space, H is the horizon length for the finite-horizon setting, γ is the discount factor for the infinite-horizon setting, $P_{\min}^*$ is the smallest positive state transition probability of the nominal kernel *visited by the optimal robust policy* π^* , and $P_{\min}$ is the smallest positive state transition probability of the nominal kernel; it holds $P_{\min} \leq P_{\min}^*$.

man machineries are developed under mild assumptions. Targeting robust MDPs, empirical and theoretical works have been widely explored under different forms of uncertainty sets (Iyengar, 2005; Xu and Mannor, 2012; Wolff et al., 2012; Kaufman and Schaefer, 2013; Ho et al., 2018; Smirnova et al., 2019; Ho et al., 2021; Goyal and Grand-Clement, 2022; Derman and Mannor, 2020; Tamar et al., 2014; Badrinath and Kalathil, 2021; Abdullah et al., 2019; Hou et al., 2020; Song and Zhao, 2020; Yang, 2017; Wang et al., 2022; Ding et al., 2023). Nonetheless, the majority of prior theoretical analyses focus on planning with an exact knowledge of the uncertainty set (Iyengar, 2005; Xu and Mannor, 2012; Tamar et al., 2014), or are asymptotic in nature (Roy et al., 2017).

A number of robust RL algorithms were proposed recently with an emphasis on finite-sample performance guarantees under different data generating mechanisms. Wang and Zou (2021) proposed a robust Q-learning algorithm with an R-contamination uncertain set for the online setting, which achieves a similar bound as its non-robust counterpart. Badrinath and Kalathil (2021) proposed a model-free algorithm for the online setting with linear function approximation to cope with large state spaces. Yang et al. (2022); Panaganti and Kalathil (2022) developed sample complexities for a model-based robust RL algorithm with a variety of uncertainty sets where the data are collected using a generative model. In addition, Zhou et al. (2021) examined the uncertainty set defined by the KL divergence

for offline data with uniformly lower bounded data visitation distribution. These works all require full coverage of the state-action space, whereas ours is the first one to leverage the principle of pessimism in robust offline RL.

Since the first appearance of our paper on arXiv in August 2022, a few more papers have emerged that also tackle the sample complexity of robust RL algorithms. For example, Wang et al. (2023a,b) developed finite-sample complexity bounds for robust variants of Q-learning with the generative model when the uncertainty set is measured by KL divergence; in particular, the improved bound of variance-reduced robust Q-learning (Wang et al., 2023b) becomes independent of the size of the uncertainty set when it is sufficiently small with respect to the minimal support probability of the nominal kernel at a price of worse dependency with $1/P_{\min}^*$. Shi et al. (2023) provided near-optimal sample complexity bounds for model-based robust RL algorithms with the generative model when the uncertainty set is measured by the total variation or chi-square distances, which highlighted that different uncertainty sets can lead to drastically different sample complexities, and hence, statistical consequences. Wang et al. (2024) studied the sample complexity of offline robust RL with the total variation uncertainty set under linear function approximation, and Shi et al. (2024) extended the study of robust MDPs to the setting of robust Markov games.

1.4 Notation and paper organization

Throughout this paper, we denote by $\Delta(\mathcal{S})$ the probability simplex over a set $\mathcal{S}$, and introduce the notation $[H] := \{1, \dots, H\}$ for any positive integer $H > 0$. In addition, for any vector $x = [x(s, a)]_{(s,a) \in \mathcal{S} \times \mathcal{A}} \in \mathbb{R}^{\mathcal{S}\mathcal{A}}$ (resp. $x = [x(s)]_{s \in \mathcal{S}} \in \mathbb{R}^{\mathcal{S}}$) that constitutes certain values for each state-action pair (resp. state), we overload the notation by letting $x^2 = [x(s, a)^2]_{(s,a) \in \mathcal{S} \times \mathcal{A}}$ (resp. $x^2 = [x(s)^2]_{s \in \mathcal{S}}$). Moreover, for any two vectors $x = [x_i]_{1 \leq i \leq n}$ and $y = [y_i]_{1 \leq i \leq n}$, the notation $x \leq y$ (resp. $x \geq y$) means $x_i \leq y_i$ (resp. $x_i \geq y_i$) for all $1 \leq i \leq n$. Finally, the Kullback-Leibler (KL) divergence for any two distributions P and Q is denoted as $\text{KL}(P \parallel Q)$.

The rest of this paper is organized as follows. Section 2 provides the backgrounds and introduces the distributionally robust formulation of finite-horizon MDPs in the offline setting under partial coverage. Section 3 presents the proposed algorithm and provides sample complexity guarantees. Section 4 develops the corresponding results for the infinite-horizon setting. Section 5 demonstrate the performance of the proposed algorithm through numerical experiments. Finally, we conclude in Section 6. The detailed proofs are postponed to the appendix.

2. Problem formulation: episodic finite-horizon RMDPs

2.1 Basics of finite-horizon episodic tabular MDPs

Consider an episodic finite-horizon MDP, represented by $\mathcal{M} = (\mathcal{S}, \mathcal{A}, H, P := \{P_h\}_{h=1}^H, \{r_h\}_{h=1}^H)$, where $\mathcal{S} = \{1, \dots, S\}$ and $\mathcal{A} = \{1, \dots, A\}$ are the finite state and action spaces, respectively, H is the horizon length, $P_h : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ (resp. $r_h : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$)

denotes the probability transition kernel (resp. reward function) at step h ($1 \leq h \leq H$).¹ For any transition kernel P , we introduce the $\mathcal{S}$ -dimensional distribution vectors

$$P_{h,s,a} := P_h(\cdot | s, a) \in [0, 1]^{1 \times \mathcal{S}}, \quad \forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} \quad (1)$$

to represent the probability transition vector in state s when taking action a at step h .

Denote by $\pi = \{\pi_h\}_{h=1}^H$ as the policy or action selection rule of an agent, where $\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})$ specifies the action selection probability over the action space; when the policy is deterministic, we slightly abuse the notation and refer to $\pi_h(s)$ as the action selected by policy π in state s at step h . The value function $V^{\pi,P} = \{V_h^{\pi,P}\}_{h=1}^H$ of policy π with a transition kernel P is defined by

$$\forall (h, s) \in [H] \times \mathcal{S} : \quad V_h^{\pi,P}(s) := \mathbb{E}_{\pi,P} \left[\sum_{t=h}^H r_t(s_t, a_t) \mid s_h = s \right], \quad (2)$$

where the expectation is taken over the randomness of the trajectory $\{s_h, a_h, r_h\}_{h=1}^H$ generated by executing policy π , namely, $a_t \sim \pi_t(s_t)$, and $s_{t+1} \sim P_t(\cdot | s_t, a_t)$. Similarly, the Q-function $Q^{\pi,P} = \{Q_h^{\pi,P}\}_{h=1}^H$ of policy π is defined as

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad Q_h^{\pi,P}(s, a) := r_h(s, a) + \mathbb{E}_{\pi,P} \left[\sum_{t=h+1}^H r_t(s_t, a_t) \mid s_h = s, a_h = a \right], \quad (3)$$

where the expectation is again taken over the randomness of the trajectory.

Moreover, when the initial state s_1 is drawn from a given distribution ρ , let $d_h^{\pi,P}(s | \rho)$ and $d_h^{\pi,P}(s, a | \rho)$ denote respectively the state occupancy distribution and the state-action occupancy distribution induced by π at time step $h \in [H]$. In particular, we often dropped the dependency with respect to ρ whenever it is clear from the context, by simply writing $d_h^{\pi,P}(s) := d_h^{\pi,P}(s | \rho)$ and $d_h^{\pi,P}(s, a) := d_h^{\pi,P}(s, a | \rho)$, i.e.,

$$\forall (h, s) \in [H] \times \mathcal{S} : \quad d_h^{\pi,P}(s) := \mathbb{P}(s_h = s | s_1 \sim \rho, \pi, P), \quad (4a)$$

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad d_h^{\pi,P}(s, a) := \mathbb{P}(s_h = s | s_1 \sim \rho, \pi, P) \pi_h(a | s), \quad (4b)$$

which are conditioned on $s_1 \sim \rho$ and the event that all actions and states are drawn according to policy π and transition kernel P .

2.2 Distributionally robust MDPs

In this section, we focus on finite-horizon episodic distributionally robust MDPs (RMDPs), denoted by $\mathcal{M}_{\text{rob}} = (\mathcal{S}, \mathcal{A}, H, \mathcal{U}^\sigma(P^0), \{r_h\}_{h=1}^H)$. Different from standard MDPs, we now consider an ensemble of probability transition kernels or models within an uncertainty set centered around a nominal one $P^0 = \{P_h^0\}_{h=1}^H$, where the distance between the transition kernels is measured in terms of the Kullback-Leibler (KL) divergence. Specifically, given

1. Without loss of generality, we assume the reward function is deterministic, fixed, and normalized to be within $[0, 1]$; it is straightforward to generalize our framework to incorporate random rewards with uncertainties.

an uncertainty level $\sigma > 0$, the uncertainty set around P^0 , which satisfies the so-called (s, a) -rectangularity condition (Wiesemann et al., 2013), is specified as

$$\mathcal{U}^\sigma(P^0) := \otimes \mathcal{U}^\sigma(P_{h,s,a}^0), \quad \mathcal{U}^\sigma(P_{h,s,a}^0) := \{P_{h,s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{h,s,a} \parallel P_{h,s,a}^0) \leq \sigma\}, \quad (5)$$

where $\otimes$ denote the Cartesian product. In words, the KL divergence between the true transition probability vector and the nominal one at each state-action pair is at most σ ; moreover, the RMDP reduces to the standard MDP when $\sigma = 0$.

Instead of evaluating a policy in a fixed MDP, the performance of a policy in the RMDP is evaluated based on its worst-case—i.e., smallest—value function over all the instances in the uncertainty set. That is, we define the *robust value function* $V^{\pi,\sigma} = \{V_h^{\pi,\sigma}\}_{h=1}^H$ and the *robust Q-function* $Q^{\pi,\sigma} = \{Q_h^{\pi,\sigma}\}_{h=1}^H$ respectively as

$$\begin{aligned} \forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad V_h^{\pi,\sigma}(s) &:= \inf_{P \in \mathcal{U}^\sigma(P^0)} V_h^{\pi,P}(s), \\ Q_h^{\pi,\sigma}(s, a) &:= \inf_{P \in \mathcal{U}^\sigma(P^0)} Q_h^{\pi,P}(s, a), \end{aligned}$$

where the infimum is taken over the uncertainty set of transition kernels.

Optimal robust policy. For finite-horizon RMDPs, it has been established that there exists at least one deterministic policy that maximizes the robust value function and the robust Q-function simultaneously (Iyengar, 2005; Nilim and El Ghaoui, 2005). In view of this, we shall denote a deterministic policy $\pi^* = \{\pi_h^*\}_{h=1}^H$ as an optimal robust policy throughout this paper. The resulting *optimal robust value function* $V^{*,\sigma} = \{V_h^{*,\sigma}\}_{h=1}^H$ and *optimal robust Q-function* $Q^{*,\sigma} = \{Q_h^{*,\sigma}\}_{h=1}^H$ are denoted by

$$\forall (h, s) \in [H] \times \mathcal{S} : \quad V_h^{*,\sigma}(s) := V_h^{\pi^*,\sigma}(s) = \max_{\pi} V_h^{\pi,\sigma}(s), \quad (6a)$$

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad Q_h^{*,\sigma}(s, a) := Q_h^{\pi^*,\sigma}(s, a) = \max_{\pi} Q_h^{\pi,\sigma}(s, a). \quad (6b)$$

Similar to (4), we adopt the following short-hand notation for the occupancy distributions associated with the optimal policy:

$$\forall (h, s) \in [H] \times \mathcal{S} : \quad d_h^{*,P}(s) := d_h^{\pi^*,P}(s), \quad (7a)$$

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad d_h^{*,P}(s, a) := d_h^{\pi^*,P}(s, a) = d_h^{*,P}(s) \mathbb{1}\{a = \pi_h^*(s)\}. \quad (7b)$$

Robust Bellman equations. It turns out the Bellman’s principle of optimality can be extended naturally to its robust counterpart (Iyengar, 2005; Nilim and El Ghaoui, 2005), which plays a fundamental role in solving the RMDP. To begin with, for any policy π , the robust value function and robust Q-function satisfy the following *robust Bellman consistency equation*:

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad Q_h^{\pi,\sigma}(s, a) = r_h(s, a) + \inf_{P \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V_{h+1}^{\pi,\sigma}. \quad (8)$$

Additionally, the optimal robust Q-function obeys the *robust Bellman optimality equation*:

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A} : \quad Q_h^{*,\sigma}(s, a) = r_h(s, a) + \inf_{P \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V_{h+1}^{*,\sigma}, \quad (9)$$

which can be solved efficiently via a robust variant of value iteration when the RMDP is known (Iyengar, 2005; Nilim and El Ghaoui, 2005).

2.3 Distributionally robust offline RL

Let $\mathcal{D}$ be a history/batch dataset, which consists of a collection of K *independent* episodes generated based on executing a behavior policy $\pi^b = \{\pi_h^b\}_{h=1}^H$ in some nominal MDP $\mathcal{M}^0 = (\mathcal{S}, \mathcal{A}, H, P^0 := \{P_h^0\}_{h=1}^H, \{r_h\}_{h=1}^H)$. More specifically, for $1 \leq k \leq K$, the k -th episode $(s_1^k, a_1^k, \dots, s_H^k, a_H^k, s_{H+1}^k)$ is generated according to

$$s_1^k \sim \rho^b, \quad a_h^k \sim \pi_h^b(\cdot | s_h^k) \quad \text{and} \quad s_{h+1}^k \sim P_h^0(\cdot | s_h^k, a_h^k), \quad 1 \leq h \leq H. \quad (10)$$

Throughout the paper, ρ^b represents for some initial distribution associated with the history dataset. Then, we introduce the following short-hand notation for the occupancy distribution w.r.t. π^b :

$$\forall (h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}: \quad d_h^b(s) := d_h^{\pi^b, P^0}(s), \quad d_h^b(s, a) := d_h^{\pi^b, P^0}(s, a). \quad (11)$$

Goal. With the history dataset $\mathcal{D}$ in hand, our goal is to find a near-optimal robust policy $\hat{\pi}$, which satisfies

$$V_1^{\hat{\pi}, \sigma}(\rho) \geq V_1^{\star, \sigma}(\rho) - \varepsilon \quad (12)$$

using as few samples as possible, where ε is the target accuracy level, and

$$V_1^{\pi, \sigma}(\rho) := \mathbb{E}_{s_1 \sim \rho}[V_1^{\pi, \sigma}(s_1)] \quad \text{and} \quad V_1^{\star, \sigma}(\rho) := \mathbb{E}_{s_1 \sim \rho}[V_1^{\star, \sigma}(s_1)] \quad (13)$$

are evaluated when the initial state s_1 is drawn from a given distribution ρ .

Robust single-policy clipped concentrability. To quantify the quality of the history dataset to achieve the set goal, it is desirable to capture the distribution mismatch between the history dataset and the desired ones, inspired by the *single-policy clipped concentrability* assumption recently proposed by Li et al. (2022), we introduce a tailored assumption for robust MDPs as follows.

Assumption 1 (Robust single-policy clipped concentrability) *The behavior policy of the history dataset $\mathcal{D}$ satisfies*

$$\max_{(s, a, h, P) \in \mathcal{S} \times \mathcal{A} \times [H] \times \mathcal{U}^\sigma(P^0)} \frac{\min \{d_h^{\star, P}(s, a), \frac{1}{S}\}}{d_h^b(s, a)} \leq C_{\text{rob}}^\star \quad (14)$$

for some quantity $C_{\text{rob}}^\star \in [\frac{1}{S}, \infty]$. Here, we take $C_{\text{rob}}^\star$ to be the smallest quantity satisfying (14), and refer to it as the robust single-policy clipped concentrability coefficient. In addition, we follow the convention $0/0 = 0$.

In words, $C_{\text{rob}}^\star$ measures the worst-case discrepancy—between the optimal robust policy $\pi^\star$ (from initial state distribution ρ) in any model $P \in \mathcal{U}^\sigma(P^0)$ within the uncertainty set and the behavior policy π^b (from initial state distribution ρ^b) in the nominal model P^0 —in terms of the clipped maximum density ratio of the state-action occupancy distributions.

- *Distribution shift.* When the uncertainty level $\sigma = 0$, Assumption 1 reduces back to the single-policy clipped concentrability in Li et al. (2022) for standard offline RL, a weaker notion that can be S times smaller than the single-policy concentrability adopted in, e.g., Rashidinejad et al. (2021); Xie et al. (2021); Shi et al. (2022). On the other end, whenever $\sigma > 0$, the proposed robust single-policy clipped concentrability accounts for the distribution shift not only due to the policies in use (π^* versus π^b) with respect to the respective initial state distributions, but also the underlying environments ($P \in \mathcal{U}^\sigma(P^0)$ versus P^0), and therefore, is generally larger than that in the non-robust counterpart.
- *Partial coverage.* As long as C_{rob}^* is finite, i.e., $C_{\text{rob}}^* < \infty$, it admits the scenarios when the history dataset only provides *partial coverage* over the entire state-action space, as long as the behavior policy π^b visits the state-action pairs that are visited by the optimal robust policy π^* under at least one model in the uncertainty set.

Remark 1 To facilitate comparison with prior works assuming full coverage, we can bound C_{rob}^* when the batch dataset is generated using a simulator (Yang et al., 2022; Panaganti and Kalathil, 2022); namely, we can generate sample state transitions based on the transition kernel of the nominal MDP for all state-action pairs at all time steps. In this case, it amounts to that $d_h^b(s, a) = \frac{1}{SA}$ for all $(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$, which directly leads to the bound $C_{\text{rob}}^* = \max_{(s, a, h, P) \in \mathcal{S} \times \mathcal{A} \times [H] \times \mathcal{U}^\sigma(P^0)} \frac{\min \left\{ d_h^{*,P}(s, a), \frac{1}{S} \right\}}{d_h^b(s, a)} \leq \frac{1/S}{1/(SA)} = A$.

3. Algorithm and theory: episodic finite-horizon RMDPs

In this section, we present a model-based algorithm—namely DRVI-LCB—for robust offline RL in the finite-horizon setting, along with its performance guarantees.

3.1 Building an empirical nominal MDP

For a moment, imagine we have access to N independent sample transitions $\mathcal{D}_0 := \{(h_i, s_i, a_i, s'_i)\}_{i=1}^N$ drawn from the transition kernel P^0 of the nominal MDP $\mathcal{M}^0$, where each sample (h_i, s_i, a_i, s'_i) indicates the transition from state s_i to state s'_i when action a_i is taken at step h_i , drawn according to $s'_i \sim P_{h_i}^0(\cdot | s_i, a_i)$. It is then natural to build an empirical estimate $\hat{P}^0 = \{\hat{P}_h^0\}_{h=1}^H$ of P^0 based on the empirical frequencies of state transitions, where

$$\hat{P}_h^0(s' | s, a) := \begin{cases} \frac{1}{N_h(s, a)} \sum_{i=1}^N \mathbb{1}\{(h_i, s_i, a_i, s'_i) = (h, s, a, s')\}, & \text{if } N_h(s, a) > 0 \\ 0, & \text{else} \end{cases} \quad (15)$$

for any $(h, s, a, s') \in [H] \times \mathcal{S} \times \mathcal{A} \times \mathcal{S}$. Here, $N_h(s, a)$ denotes the total number of sample transitions from (s, a) at step h as

$$N_h(s, a) := \sum_{i=1}^N \mathbb{1}\{(h_i, s_i, a_i) = (h, s, a)\}. \quad (16)$$

Algorithm 1: Two-fold subsampling trick for the finite-horizon setting.

- 1 **input:** a dataset $\mathcal{D}$, probability δ .
- 2 **data splitting:** split $\mathcal{D}$ into two $\mathcal{D}^{\text{main}}$ and $\mathcal{D}^{\text{aux}}$, where each contain $K/2$ trajectories.
- 3 **lower bounding the number of transitions in $\mathcal{D}^{\text{main}}$:** denote the number of transitions from state s at step h in $\mathcal{D}^{\text{main}}$ (resp. $\mathcal{D}^{\text{aux}}$) as $N_h^{\text{main}}(s)$ (resp. $N_h^{\text{aux}}(s)$), construct

$$N_h^{\text{trim}}(s) := \max \left\{ N_h^{\text{aux}}(s) - 10\sqrt{N_h^{\text{aux}}(s) \log \frac{HS}{\delta}}, 0 \right\}; \quad (17)$$

- 4 **generate the subsampled dataset $\mathcal{D}^{\text{trim}}$:** randomly sample the transitions (i.e., the quadruples taking the form (s, a, h, s')) from $\mathcal{D}^{\text{main}}$ uniformly at random, such that for each $(s, h) \in \mathcal{S} \times [H]$, $\mathcal{D}^{\text{trim}}$ contains $\min\{N_h^{\text{trim}}(s), N_h^{\text{main}}(s)\}$ sample transitions.
 - 5 **output:** set $\mathcal{D}_0 = \mathcal{D}^{\text{trim}}$.
-

While it is possible to directly break down the history dataset $\mathcal{D}$ into sample transitions, unfortunately, the sample transitions from the same episode are not independent, significantly hurdling the analysis. To alleviate this, Li et al. (2022, Algorithm 2) proposed a simple two-fold subsampling scheme to preprocess the history dataset $\mathcal{D}$ and decouple the statistical dependency, resulting into a distributionally equivalent dataset $\mathcal{D}_0$ with independent samples; for completeness, we provide the procedure in Algorithm 1. We have the following lemma paraphrased from Li et al. (2022) for the obtained dataset $\mathcal{D}_0$.

Lemma 2 ((Li et al., 2022)) *With probability at least $1 - 8\delta$, the output dataset from the two-fold subsampling scheme in Li et al. (2022) is distributionally equivalent to $\mathcal{D}_0$, where $\{N_h(s, a)\}$ are independent of the sample transitions in $\mathcal{D}^0$ and obey*

$$N_h(s, a) \geq \frac{K d_h^b(s, a)}{8} - 5\sqrt{K d_h^b(s, a) \log \frac{KH}{\delta}}. \quad (18)$$

for all $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$.

Therefore, by invoking the two-fold sampling trick from Li et al. (2022), it is sufficient to treat the dataset $\mathcal{D}_0$ with independent samples onwards with Lemma 2 in place, which greatly simplifies the analysis.

3.2 DRVI-LCB: a pessimistic variant of robust value iteration

Armed with the estimate $\hat{P}^0$ of the nominal transition kernel P^0 , we are positioned to introduce our algorithm DRVI-LCB, summarized in Algorithm 2.

Distributionally robust value iteration. Before proceeding, let us recall the update rule of the classical distributionally robust value iteration (DRVI), which serves as the basis

Algorithm 2: Robust value iteration with LCB (DRVI-LCB) for robust offline RL.

```

1 input: a dataset  $\mathcal{D}_0$ ; reward function  $r$ ; uncertainty level  $\sigma$ .
2 initialization:  $\hat{Q}_{H+1} = 0, \hat{V}_{H+1} = 0$ .
3 for  $h = H, \dots, 1$  do
4     Compute the empirical nominal transition kernel  $\hat{P}_h^0$  according to (15);
5     for  $s \in \mathcal{S}, a \in \mathcal{A}$  do
6         Compute the penalty term  $b_h(s, a)$  according to (22);
7         Set  $\hat{Q}_h(s, a)$  according to (21);
8     for  $s \in \mathcal{S}$  do
9         Set  $\hat{V}_h(s) = \max_a \hat{Q}_h(s, a)$  and  $\hat{\pi}_h(s) = \arg \max_a \hat{Q}_h(s, a)$ ;
10 output:  $\hat{\pi} = \{\hat{\pi}_h\}_{1 \leq h \leq H}$ .
    
```

of our algorithmic development. Given an estimate of the nominal MDP $\hat{P}^0$ and the radius σ of the uncertainty set, DRVI updates the robust value functions according to

$$\hat{Q}_h(s, a) = r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}\hat{V}_{h+1}, \quad \text{and} \quad \hat{V}_h(s) = \max_a \hat{Q}_h(s, a), \quad (19)$$

which works backwards from $h = H$ to $h = 1$, with the terminal condition $\hat{Q}_{H+1} = 0$. Due to strong duality (Hu and Hong, 2013), the update rule of the robust Q-functions in (19) can be equivalently reformulated in its dual form as

$$\hat{Q}_h(s, a) = r_h(s, a) + \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-\hat{V}_{h+1}}{\lambda} \right) \right) - \lambda \sigma \right\}, \quad (20)$$

which can be solved efficiently (Iyengar, 2005; Yang et al., 2022; Panaganti and Kalathil, 2022).

Our algorithm DRVI-LCB. Motivated by the principle of pessimism in standard offline RL (Jin et al., 2021; Xie et al., 2021; Rashidinejad et al., 2021; Li et al., 2022), we propose to perform a pessimistic variant of DRVI, where the update rule of DRVI-LCB at step h is modified as

$$\hat{Q}_h(s, a) = \max \left\{ r_h(s, a) + \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-\hat{V}_{h+1}}{\lambda} \right) \right) - \lambda \sigma \right\} - b_h(s, a), 0 \right\}. \quad (21)$$

Here, the robust Q-function estimate is adjusted by subtracting a carefully designed data-driven penalty term $b_h(s, a)$ that measures the uncertainty of the value estimates. Specifically, for some $\delta \in (0, 1)$ and any $(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$, the penalty term $b_h(s, a)$ is defined as

$$b_h(s, a) = \begin{cases} \min \left\{ c_b \frac{H}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{\hat{P}_{\min,h}(s,a)N_h(s,a)}}, H \right\} & \text{if } N_h(s, a) > 0, \\ H & \text{otherwise,} \end{cases} \quad (22)$$

where c_b is some universal constant, and

$$\hat{P}_{\min,h}(s, a) := \min_{s'} \left\{ \hat{P}_h^0(s' | s, a) : \hat{P}_h^0(s' | s, a) > 0 \right\}. \quad (23)$$

The penalty term is novel and different from the one used in standard (no-robust) offline RL (Jin et al., 2021; Xie et al., 2021; Rashidinejad et al., 2021; Li et al., 2022; Shi et al., 2022), by taking into consideration the unique problem structure pertaining to robust MDPs. In particular, it tightly upper bounds the statistical uncertainty which carries a non-linear and implicit dependency w.r.t. the estimated nominal transition kernel induced by the uncertainty set $\mathcal{U}(P^0)$, addressing unique challenges not present for the standard MDP case.

3.3 Performance guarantees

Before stating the main theorems, let us first introduce several important metrics.

- $P_{\min}^*$, which only depends on the state-action pairs covered by the optimal robust policy π^* under the nominal model P^0 :

$$P_{\min}^* := \min_{h,s,s'} \left\{ P_h^0(s' | s, \pi_h^*(s)) : P_h^0(s' | s, \pi_h^*(s)) > 0 \right\}. \quad (24)$$

In words, $P_{\min}^*$ is the smallest positive state transition probability of the optimal robust policy π^* under the nominal kernel P^0 .

- Similarly, we introduce $P_{\min}^b$ which only depends on the state-action pairs covered by the behavior policy π^b under the nominal model P^0 :

$$P_{\min}^b := \min_{h,s,a,s'} \left\{ P_h^0(s' | s, a) : d_h^b(s, a) > 0, P_h^0(s' | s, a) > 0 \right\}. \quad (25)$$

In words, $P_{\min}^b$ is the smallest positive state transition probability of the behavior policy π^b under the nominal kernel P^0 .

- Finally, let $d_{\min}^b$ denote the smallest positive state-action occupancy distribution of the behavior policy π^b under the nominal model P^0 :

$$d_{\min}^b := \min_{h,s,a} \left\{ d_h^b(s, a) : d_h^b(s, a) > 0 \right\}. \quad (26)$$

We are now positioned to present the performance guarantees of DRVI-LCB for robust offline RL.

Theorem 3 *Given an uncertainty level $\sigma > 0$, suppose that the penalty terms in Algorithm 2 are chosen as (22) for sufficiently large c_b . With probability at least $1 - \delta$, the output $\hat{\pi}$ of Algorithm 2 obeys*

$$V_1^{\star,\sigma}(\rho) - V_1^{\hat{\pi},\sigma}(\rho) \leq c_0 \frac{H^2}{\sigma} \sqrt{\frac{SC_{\text{rob}}^* \log^2(KHS/\delta)}{P_{\min}^* K}}, \quad (27)$$

as long as the number of episodes K satisfies

$$K \geq \frac{c_1 \log(KHS/\delta)}{d_{\min}^b P_{\min}^b}, \quad (28)$$

where c_0 and c_1 are some sufficiently large universal constants.

Our theorem is the first to characterize the sample complexities of robust offline RL under *partial coverage*, to the best of our knowledge (cf. Table 2). Theorem 3 shows that DRVI-LCB finds an ε -optimal robust policy as soon as the sample size $T = KH$ is above the order of

$$\underbrace{\frac{SC_{\text{rob}}^* H^5}{P_{\min}^* \sigma^2 \varepsilon^2}}_{\varepsilon\text{-dependent}} + \underbrace{\frac{H}{d_{\min}^b P_{\min}^b}}_{\text{burn-in cost}}, \quad (29)$$

up to some logarithmic factor, where the burn-in cost is independent of the accuracy level ε . For sufficiently small accuracy level ε , this results in a sample complexity of

$$\tilde{O}\left(\frac{SC_{\text{rob}}^* H^5}{P_{\min}^* \sigma^2 \varepsilon^2}\right). \quad (30)$$

Our theorem suggests that the sample efficiency of robust offline RL critically depends on the problem structure of the given RMDP (i.e. coverage of the optimal robust policy π^* as measured by $P_{\min}^*$) as well as the quality of the history dataset (as measured by C_{rob}^*). Given that C_{rob}^* can be as small as on the order of $1/S$, the sample complexity requirement can exhibit a much weaker dependency with the size of the state space S .

On the flip side, to assess the optimality of Theorem 3, we develop an information-theoretic lower bound for robust offline RL as provided in the following theorem.

Theorem 4 *For any $(H, S, C, P_{\min}^*, \sigma, \varepsilon)$ obeying $H \geq 2e^8$, $C \geq 4/S$, $P_{\min}^* \in (0, \frac{1}{H}]$, and $\varepsilon \leq \frac{H}{384e^6 \log(1/P_{\min}^*)}$, we can construct a collection of finite-horizon RMDPs $\{\mathcal{M}_\theta | \theta \in \Theta\}$, an initial state distribution ρ , and a batch dataset with K independent sample trajectories each with length H satisfying $2C \leq C_{\text{rob}}^* \leq 4C$, such that*

$$\inf_{\hat{\pi}} \max_{\theta \in \Theta} \mathbb{P}_\theta \left\{ V_1^{\star, \sigma}(\rho) - V_1^{\hat{\pi}, \sigma}(\rho) > \varepsilon \right\} \geq \frac{1}{8},$$

provided that

$$T = KH \leq \begin{cases} \frac{c_1 SC_{\text{rob}}^* H^4}{\varepsilon^2} & \text{if } 0 < \sigma \leq \frac{1}{20H}, \\ \frac{c_1 SC_{\text{rob}}^* H^3}{P_{\min}^* \sigma^2 \varepsilon^2} & \text{if } \log(1/P_{\min}^*) - 6 \leq \sigma \leq \log(1/P_{\min}^*) - 5 \end{cases}. \quad (31)$$

Here, $c_1 > 0$ is some universal constant, the infimum is taken over all estimators $\hat{\pi}$, and $\mathbb{P}_\theta$ denotes the probability when the RMDP is $\mathcal{M}_\theta$.

The messages of Theorem 4 are two-fold.

- When the uncertainty level $\sigma \lesssim 1/H$ is relatively small, Theorem 4 shows that no algorithm can succeed in finding an ε -optimal robust policy when the sample complexity falls below the order of

$$\Omega\left(\frac{SC_{\text{rob}}^* H^4}{\varepsilon^2}\right),$$

which is at least as large as the sample complexity requirement of non-robust offline RL (Li et al., 2022). Consequently, this leads to new insights regarding the statistical hardness of learning robust RMDPs with the KL uncertainty set: it can be at least as hard as the standard MDPs (which corresponds to $\sigma = 0$), for sufficiently small uncertainty levels.

- When the uncertainty level $\sigma \asymp \log(1/P_{\min}^*)$, Theorem 4 shows that no algorithm can succeed in finding an ε -optimal robust policy when the sample complexity falls below the order of

$$\Omega\left(\frac{SC_{\text{rob}}^* H^3}{P_{\min}^* \sigma^2 \varepsilon^2}\right),$$

which confirms the near-optimality of DRVI-LCB up to a factor of H^2 ignoring logarithmic factors. Therefore, DRVI-LCB is the first provable algorithm for robust offline RL with a near-optimal sample complexity without requiring the stringent full coverage assumption.

4. Robust offline RL for discounted infinite-horizon RMDPs

In this section, we turn to the studies of robust offline RL for discounted infinite-horizon MDPs.

4.1 Backgrounds on discounted infinite-horizon RMDPs

Similar to the finite-horizon setting, we consider the discounted infinite-horizon robust MDPs (RMDPs) represented by $\mathcal{M}_{\text{rob}} = \{\mathcal{S}, \mathcal{A}, \gamma, \mathcal{U}^\sigma(P^0), r\}$. Here, $\mathcal{S} = \{1, 2, \dots, S\}$ is the state space, $\mathcal{A} = \{1, 2, \dots, A\}$ is the action space, $\gamma \in [0, 1)$ is the discounted factor, and $r : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$ is the intermediate reward function. Different from the standard MDPs, $\mathcal{U}^\sigma(P^0)$ denote the set of possible transition kernels within an uncertainty set centered around a nominal kernel $P^0 : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ using the distance measured in terms of the KL divergence. In particular, given an uncertainty level $\sigma > 0$, the uncertainty set around P^0 is specified as

$$\mathcal{U}^\sigma(P^0) := \otimes \mathcal{U}^\sigma(P_{s,a}^0), \quad \mathcal{U}^\sigma(P_{s,a}^0) := \{P_{s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{s,a} \parallel P_{s,a}^0) \leq \sigma\}, \quad (32)$$

where we denote a vector of the transition kernel P or P^0 at (s, a) respectively as

$$P_{s,a} := P(\cdot \mid s, a) \in \mathbb{R}^{1 \times S}, \quad P_{s,a}^0 := P^0(\cdot \mid s, a) \in \mathbb{R}^{1 \times S}. \quad (33)$$

Note that at any time step, the adversary of the nature chooses a history-independent component within the fixed uncertainty set $\mathcal{U}^\sigma(P_{s,a}^0)$ defined in (30), conditioned only on the current state-action pair (s, a) . This is to ensure the computation tractability of finding such adversary.

Policy and robust value/Q functions. A (possibly random) stationary policy $\pi : \mathcal{S} \rightarrow \Delta(\mathcal{A})$ represents the selection rule of the agent, namely, $\pi(a | s)$ denote the probability of choosing a in state s . With some abuse of notation, let $\pi(s)$ represent the action chosen by π when π is a deterministic policy. We define the *robust value function* $V^{\pi, \sigma}$ and *robust Q-function* $Q^{\pi, \sigma}$ respectively as

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad V^{\pi, \sigma}(s) := \inf_{P \in \mathcal{U}^\sigma(P^0)} V^{\pi, P}(s), \quad Q^{\pi, \sigma}(s, a) := \inf_{P \in \mathcal{U}^\sigma(P^0)} Q^{\pi, P}(s, a),$$

where the value function $V^{\pi, P}$ and Q-function $Q^{\pi, P}$ w.r.t. policy π and transition kernel P are defined respectively by

$$\forall s \in \mathcal{S} : \quad V^{\pi, P}(s) := \mathbb{E}_{\pi, P} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s \right], \quad (34)$$

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad Q^{\pi, P}(s, a) := \mathbb{E}_{\pi, P} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s, a_0 = a \right], \quad (35)$$

where the expectation is taken over the randomness of the trajectory. In words, the robust value/Q functions characterize the worst case over all the instances in the uncertainty set.

Optimal policy and robust Bellman equation. Similar to the finite-horizon RMDPs, it is well-known that there exists at least one deterministic policy that maximizes the robust value function and Q-function simultaneously in the infinite-horizon setting as well (Iyengar, 2005; Nilim and El Ghaoui, 2005). With this in mind, we denote the optimal policy as π^* and the corresponding *optimal robust value function* (resp. *optimal robust Q-function*) as $V^{*, \sigma}$ (resp. $Q^{*, \sigma}$), namely

$$\forall s \in \mathcal{S} : \quad V^{*, \sigma}(s) := V^{\pi^*, \sigma}(s) = \max_{\pi} V^{\pi, \sigma}(s), \quad (36a)$$

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad Q^{*, \sigma}(s, a) := Q^{\pi^*, \sigma}(s, a) = \max_{\pi} Q^{\pi, \sigma}(s, a). \quad (36b)$$

In addition, we continue to admit the Bellman's optimality principle, resulting in the following *robust Bellman consistency equation* (resp. *robust Bellman optimality equation*):

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad Q^{\pi, \sigma}(s, a) = r(s, a) + \gamma \inf_{P \in \mathcal{U}^\sigma(P_{s,a}^0)} P V^{\pi, \sigma}, \quad (37a)$$

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad Q^{*, \sigma}(s, a) = r(s, a) + \gamma \inf_{P \in \mathcal{U}^\sigma(P_{s,a}^0)} P V^{*, \sigma}. \quad (37b)$$

Occupancy distributions. To begin, let ρ be some initial state distribution. We denote $d^{\pi, P}(s | \rho)$ and $d^{\pi, P}(s, a | \rho)$ respectively as the state occupancy distribution and the state-action occupancy distribution induced by policy π , namely

$$\forall s \in \mathcal{S} : \quad d^{\pi, P}(s) := (1 - \gamma) \sum_{t=0}^{\infty} \gamma^t \mathbb{P}(s_t = s \mid s_0 \sim \rho, \pi, P), \quad (38a)$$

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad d^{\pi, P}(s, a) := (1 - \gamma) \sum_{t=0}^{\infty} \gamma^t \mathbb{P}(s_t = s \mid s_0 \sim \rho, \pi, P) \pi(a | s). \quad (38b)$$

Here, the occupancy distributions are conditioned on $s_0 \sim \rho$ and the sequence of actions and states are generated based on policy π and transition kernel P . Next, applying (38) with $\pi = \pi^*$, we adopt the the following short-hand notation for the occupancy distributions associated with the optimal policy:

$$\forall s \in \mathcal{S} : \quad d^{\star, P}(s) := d^{\pi^*, P}(s), \quad (39a)$$

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad d^{\star, P}(s, a) := d^{\pi^*, P}(s, a) = d^{\star, P}(s) \mathbb{1}\{a = \pi^*(s)\}. \quad (39b)$$

4.2 Data collection and constructing the empirical MDP

Suppose that we observe a batch/history dataset $\mathcal{D} = \{(s_i, a_i, s'_i)\}_{1 \leq i \leq N}$ consisting of N sample transitions. These transitions are independently generated, where the state-action pair is drawn from some behavior distribution $d^b \in \Delta(\mathcal{S} \times \mathcal{A})$, followed by a next state drawn over the nominal transition kernel P^0 , i.e.,

$$(s_i, a_i) \stackrel{\text{i.i.d.}}{\sim} d^b \quad \text{and} \quad s'_i \stackrel{\text{i.i.d.}}{\sim} P^0(\cdot | s_i, a_i), \quad 1 \leq i \leq N. \quad (40)$$

Armed with these, we are ready to introduce the goal in the infinite-horizon setting. Given the history dataset $\mathcal{D}$ obeying Assumption 2, for some target accuracy $\varepsilon > 0$, we aim to find a near-optimal robust policy $\hat{\pi}$, which satisfies

$$V^{\hat{\pi}, \sigma}(\rho) \geq V^{\star, \sigma}(\rho) - \varepsilon \quad (41)$$

in a sample-efficient manner for some initial state distribution ρ .

Remark 5 *For simplicity, we limit ourselves to the case when the history dataset consists of independent sample transitions. It is not difficult to generalize to the Markovian data case, when we only have access to a single trajectory of data generated by following some behavior policy, by combining the two-fold subsampling trick in Li et al. (2022, Appendix D) with our analysis. We leave this extension to interested readers.*

Similar to Assumption 1, we design the following *robust single-policy clipped concentrability* assumption tailored for infinite-horizon RMDPs to characterize the quality of the history dataset.

Assumption 2 (Robust single-policy clipped concentrability for infinite-horizon MDPs) *The behavior policy of the history dataset $\mathcal{D}$ satisfies*

$$\max_{(s, a, P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^0)} \frac{\min \{d^{\star, P}(s, a), \frac{1}{S}\}}{d^b(s, a)} \leq C_{\text{rob}}^* \quad (42)$$

for some finite quantity $C_{\text{rob}}^* \in [\frac{1}{S}, \infty)$. Following the convention $0/0 = 0$, we denote C_{rob}^* to be the smallest quantity satisfying (42), and refer to it as the robust single-policy clipped concentrability coefficient.

Remark 6 *Similar to Remark 1, we can bound $C_{\text{rob}}^* \leq A$ when the batch dataset is generated using a simulator (Yang et al., 2022; Panaganti and Kalathil, 2022). By combining this bound of C_{rob}^* with the theoretical guarantees developed momentarily in Theorem 9, we obtain the comparison in Table 2.*

Algorithm 3: Robust value iteration with LCB (DRVI-LCB) for infinite-horizon RMDPs.

```

1 input: a dataset  $\mathcal{D}$ ; reward function  $r$ ; uncertainty level  $\sigma$ ; number of iterations  $M$ .
2 initialization:  $\widehat{Q}_0(s, a) = 0$ ,  $\widehat{V}_0(s) = 0$  for all  $(s, a) \in \mathcal{S} \times \mathcal{A}$ .
3 Compute the empirical nominal transition kernel  $\widehat{P}^0$  according to (44);
4 Compute the penalty term  $b(s, a)$  according to (48);
5 for  $m = 1, 2, \dots, M$  do
6   for  $s \in \mathcal{S}, a \in \mathcal{A}$  do
7     Set  $\widehat{Q}_m(s, a)$  according to (51);
8   for  $s \in \mathcal{S}$  do
9     Set  $\widehat{V}_m(s) = \max_a \widehat{Q}_m(s, a)$ ;
10 output:  $\widehat{\pi}$  s.t.  $\widehat{\pi}(s) = \arg \max_a \widehat{Q}_M(s, a)$  for all  $s \in \mathcal{S}$ .
    
```

Building an empirical nominal MDP Recalling that we have N independent samples in the dataset $\mathcal{D} = \{(s_i, a_i, s'_i)\}_{1 \leq i \leq N}$. First, we denote $N(s, a)$ as the total number of sample transitions from any state-action pair (s, a) as

$$N(s, a) := \sum_{i=1}^N \mathbf{1}\{(s_i, a_i) = (s, a)\}. \quad (43)$$

Armed with $N(s, a)$, we construct the empirical estimate $\widehat{P}^0$ of the nominal kernel P^0 by the visiting frequencies of state-action pairs as follows:

$$\widehat{P}^0(s' | s, a) := \begin{cases} \frac{1}{N(s, a)} \sum_{i=1}^N \mathbf{1}\{(s_i, a_i, s'_i) = (s, a, s')\}, & \text{if } N(s, a) > 0 \\ 0, & \text{else} \end{cases} \quad (44)$$

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

4.3 DRVI-LCB for discounted infinite-horizon RMDPs

With the estimate $\widehat{P}^0$ of the nominal transition kernel P^0 in hand, we are positioned to introduce our algorithm DRVI-LCB for infinite-horizon RMDPs, which bears some similarity with the finite-horizon version (cf. Algorithm 2), by taking the uncertainties of the value estimates into consideration throughout the value iterations. The procedure is summarized in Algorithm 3.

The pessimistic robust Bellman operator. At the core of DRVI-LCB is a pessimistic variant of the classical robust Bellman operator in the infinite-horizon setting (Zhou et al., 2021; Iyengar, 2005; Nilim and El Ghaoui, 2005), denoted as $\mathcal{T}^\sigma(\cdot) : \mathbb{R}^{\mathcal{S}\mathcal{A}} \rightarrow \mathbb{R}^{\mathcal{S}\mathcal{A}}$, which we recall as follows:

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad \mathcal{T}^\sigma(Q)(s, a) := r(s, a) + \gamma \inf_{P \in \mathcal{U}^\sigma(P_{s,a}^0)} PV, \quad \text{with} \quad V(s) := \max_a Q(s, a). \quad (45)$$

Encouragingly, the robust Bellman operator shares the nice γ -contraction property of the standard Bellman operator, ensuring fast convergence of robust value iteration by applying the robust Bellman operator (45) recursively. In the robust offline setting, instead of recursing using the population robust Bellman operator, we need to construct a pessimistic variant of the robust Bellman operator $\widehat{T}_{\text{pe}}^\sigma(\cdot)$ w.r.t. the empirical nominal kernel $\widehat{P}^0$ as follows:

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A}: \quad \widehat{T}_{\text{pe}}^\sigma(Q)(s, a) = \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P}V - b(s, a), 0 \right\}, \quad (46)$$

where $b(s, a)$ denotes the penalty term that measures the data-dependent uncertainty of the value estimates.

To specify the tailored penalty term $b(s, a)$ in (46), we first introduce an additional term

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A}: \quad \widehat{P}_{\min}(s, a) := \min_{s'} \left\{ \widehat{P}^0(s' | s, a) : \widehat{P}^0(s' | s, a) > 0 \right\}, \quad (47)$$

which in words represents the smallest positive transition probability of the estimated nominal kernel $\widehat{P}^0(s' | s, a)$. Then for some $\delta \in (0, 1)$, some universal constant $c_b > 0$, $b(s, a)$ is defined as

$$b(s, a) = \begin{cases} \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{\widehat{P}_{\min}(s, a)N(s, a)}} + \frac{4}{\sigma N(1-\gamma)}, \frac{1}{1-\gamma} \right\} + \frac{2}{\sigma N} & \text{if } N(s, a) > 0, \\ \frac{1}{1-\gamma} + \frac{2}{\sigma N} & \text{otherwise.} \end{cases} \quad (48)$$

As shall be illuminated, our proposed pessimistic robust Bellman operator $\widehat{T}_{\text{pe}}^\sigma(\cdot)$ (cf. (46)) plays an important role in DRVI-LCB. Encouragingly, despite the additional data-driven penalty term $b(s, a)$, it still enjoys the celebrated γ -contractive property, which greatly facilitates the analysis. Before continuing, we summarize the γ -contraction property below, whose proof is postponed to Appendix C.1.

Lemma 7 (γ -Contraction) *For any $\gamma \in [0, 1)$, the operator $\widehat{T}_{\text{pe}}^\sigma(\cdot)$ (cf. (46)) is a γ -contraction w.r.t. $\|\cdot\|_\infty$. Namely, for any $Q_1, Q_2 \in \mathbb{R}^{SA}$ s.t. $Q_1(s, a), Q_2(s, a) \in [0, \frac{1}{1-\gamma}]$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$, one has*

$$\left\| \widehat{T}_{\text{pe}}^\sigma(Q_1) - \widehat{T}_{\text{pe}}^\sigma(Q_2) \right\|_\infty \leq \gamma \|Q_1 - Q_2\|_\infty. \quad (49)$$

Additionally, there exists a unique fixed point $\widehat{Q}_{\text{pe}}^{*,\sigma}$ of the operator $\widehat{T}_{\text{pe}}^\sigma(\cdot)$ obeying $0 \leq \widehat{Q}_{\text{pe}}^{*,\sigma}(s, a) \leq \frac{1}{1-\gamma}$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$.

Our algorithm DRVI-LCB for infinite-horizon robust offline RL. Armed with the γ -contraction property of the pessimistic robust Bellman operator $\widehat{T}_{\text{pe}}^\sigma(\cdot)$, we are positioned to introduce DRVI-LCB for infinite-horizon RMDPs, summarized in Algorithm 3. Specifically, DRVI-LCB can be seen as a value iteration algorithm w.r.t. $\widehat{T}_{\text{pe}}^\sigma(\cdot)$ (cf. (46)), whose update rule at the m -th iteration can be formulated as

$$\widehat{Q}_m(s, a) = \widehat{T}_{\text{pe}}^\sigma(\widehat{Q}_{m-1})(s, a) = \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P}\widehat{V}_{m-1} - b(s, a), 0 \right\}, \quad (50)$$

and $\widehat{V}_m(s) = \max_a \widehat{Q}_m(s, a)$ for all $m = 1, 2, \dots, M$. In view of strong duality (Hu and Hong, 2013), the above convex problem can be translated into a dual formulation, leading to the following equivalent update rule:

$$\widehat{Q}_m(s, a) = \max \left\{ r(s, a) + \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\widehat{P}_{s,a}^0 \cdot \exp \left(\frac{-\widehat{V}_{m-1}}{\lambda} \right) \right) - \lambda \sigma \right\} - b(s, a), 0 \right\}, \quad (51)$$

which can be solved efficiently (Iyengar, 2005; Yang et al., 2022; Panaganti and Kalathil, 2022) as a one-dimensional optimization problem.

To finish the description, we initialize the estimates of Q-function ($\widehat{Q}_0$) and value function ($\widehat{V}_0$) to be zero and output the greedy policy of the final Q-estimates ($\widehat{Q}_M$) as the final policy $\widehat{\pi}$, namely,

$$\widehat{\pi}(s) = \arg \max_a \widehat{Q}_M(s, a) \quad \text{for all } s \in \mathcal{S}. \quad (52)$$

It turns out that the iterates $\{\widehat{Q}_m\}_{m \geq 0}$ of DRVI-LCB converge linearly to the fixed point $\widehat{Q}_{\text{pe}}^{*, \sigma}$ owing to the nice γ -contraction property outlined in Lemma 7. This fact is summarized in the following lemma, whose proof is postponed to Appendix C.2.

Lemma 8 *Let $\widehat{Q}_0 = 0$. The iterates of Algorithm 3 obey*

$$\forall m \geq 0: \quad \widehat{Q}_m \leq \widehat{Q}_{\text{pe}}^{*, \sigma} \quad \text{and} \quad \|\widehat{Q}_m - \widehat{Q}_{\text{pe}}^{*, \sigma}\|_{\infty} \leq \frac{\gamma^m}{1 - \gamma}. \quad (53)$$

4.4 Performance guarantees

Before introducing the main theorems, we first define several essential metrics.

- $d_{\min}^b$: the smallest positive entry of the distribution d^b , i.e.,

$$d_{\min}^b := \min_{s,a} \left\{ d^b(s, a) : d^b(s, a) > 0 \right\}. \quad (54)$$

- $P_{\min}^b$: the smallest positive state transition probability under the nominal kernel P^0 in the region covered by dataset $\mathcal{D}$, i.e.,

$$P_{\min}^b := \min_{s,a,s'} \left\{ P^0(s' | s, a) : d^b(s, a) > 0, P^0(s' | s, a) > 0 \right\}. \quad (55)$$

Note that $P_{\min}^b$ is determined only by the state-action pairs covered by the batch dataset $\mathcal{D}$.

- $P_{\min}^*$: the smallest positive state transition probability of the optimal robust policy π^* under the nominal kernel P^0 , namely

$$P_{\min}^* := \min_{s,s'} \left\{ P^0(s' | s, \pi^*(s)) : P^0(s' | s, \pi^*(s)) > 0 \right\}. \quad (56)$$

We also note that $P_{\min}^*$ is determined only by the state-action pairs covered by the optimal robust policy π^* under the nominal model P^0 .

We are now positioned to introduce the sample complexity upper bound of DRVI-LCB, together with the minimax lower bound, for solving infinite-horizon RMDPs. First, we present the performance guarantees of DRVI-LCB for robust offline RL in the infinite-horizon case, with the proof deferred to Appendix C.3.

Theorem 9 *Let c_0 and c_1 be some sufficiently large universal constants. Given an uncertainty level $\sigma > 0$, suppose that the penalty terms in Algorithm 3 are chosen as (48) for sufficiently large c_b . With probability at least $1 - \delta$, the output $\hat{\pi}$ of Algorithm 3 obeys*

$$V^{\star, \sigma}(\rho) - V^{\hat{\pi}, \sigma}(\rho) \leq \frac{c_0}{\sigma(1-\gamma)^2} \sqrt{\frac{SC_{\text{rob}}^* \log^2\left(\frac{(1+\sigma)N^3 S}{(1-\gamma)\delta}\right)}{P_{\min}^* N}}, \quad (57)$$

as long as the number of samples N satisfies

$$N \geq \frac{c_1 \log(NS/\delta)}{d_{\min}^b P_{\min}^b}. \quad (58)$$

The result directly indicates that DRVI-LCB can find an ε -optimal policy as long as the sample size in dataset $\mathcal{D}$ exceeds the order of (ignoring logarithmic factors)

$$\underbrace{\frac{SC_{\text{rob}}^*}{P_{\min}^* (1-\gamma)^4 \sigma^2 \varepsilon^2}}_{\varepsilon\text{-dependent}} + \underbrace{\frac{1}{d_{\min}^b P_{\min}^b}}_{\text{burn-in cost}}. \quad (59)$$

Note that the burn-in cost is independent with the accuracy level ε , which tells us that the sample complexity is no more than

$$\tilde{O}\left(\frac{SC_{\text{rob}}^*}{P_{\min}^* (1-\gamma)^4 \sigma^2 \varepsilon^2}\right) \quad (60)$$

as long as ε is small enough. The sample complexity of DRVI-LCB still dramatically outperforms prior works under full coverage, which has been compared in detail in Section 1.2 (cf. Table 2). In particular, our sample complexity produces an exponential improvement over Zhou et al. (2021); Panaganti and Kalathil (2022) in terms of the dependency with the effective horizon $\frac{1}{1-\gamma}$, which is especially significant for long-horizon problems. Compared with Yang et al. (2022), our sample complexity is better by at least a factor of $S/P_{\min}^*$. To achieve the claimed bound, we resort to a delicate technique called the leave-one-out analysis (Agarwal et al., 2020; Li et al., 2020, 2022), by carefully designing an auxiliary set of RMDPs to decouple the statistical dependency introduced across the iterates of pessimistic robust value iteration. This is the first time that the leave-one-out analysis is applied to understanding the sample efficiency of model-based robust RL algorithms, which is of potential independent interest to tighten the sample complexity of other robust RL problems.

To complement the upper bound, we develop an information-theoretic lower bound for robust offline RL as provided in the following theorem whose proof can be found in Appendix C.4.

Theorem 10 *For any $(S, P_{\min}^*, C_{\text{rob}}^*, \gamma, \sigma, \varepsilon)$ obeying $\frac{1}{1-\gamma} \geq 2e^8$, $P_{\min}^* \in (0, 1 - \gamma]$, $S \geq \log(1/P_{\min}^*)$, $C_{\text{rob}}^* \geq 8/S$, and $\varepsilon \leq \frac{1}{384e^6(1-\gamma)\log(1/P_{\min}^*)}$, we can construct a collection of infinite-horizon RMDPs $\{\mathcal{M}_\theta \mid \theta \in \Theta\}$, an initial state distribution ρ , and a batch dataset with N independent samples, such that*

$$\inf_{\hat{\pi}} \max_{\theta \in \Theta} P_\theta \left\{ V^{\star, \sigma}(\rho) - V^{\hat{\pi}, \sigma}(\rho) > \varepsilon \right\} \geq \frac{1}{8},$$

provided that

$$N \leq \begin{cases} \frac{c_1 SC_{\text{rob}}^*}{(1-\gamma)^3 \varepsilon^2} & \text{if } 0 < \sigma \leq \frac{1-\gamma}{20}, \\ \frac{c_1 SC_{\text{rob}}^*}{P_{\min}^* (1-\gamma)^2 \sigma^2 \varepsilon^2} & \text{if } \log(1/P_{\min}^*) - 6 \leq \sigma \leq \log(1/P_{\min}^*) - 5 \end{cases} \quad (61)$$

Here, $c_1 > 0$ is some universal constant, the infimum is taken over all estimators $\hat{\pi}$, and $\mathbb{P}_\theta$ denotes the probability when the RMDP is $\mathcal{M}_\theta$.

Similar to the finite-horizon setting, the messages of Theorem 10 are two-fold.

- When the uncertainty level $\sigma \lesssim 1 - \gamma$ is relatively small, Theorem 10 shows that no algorithm can succeed in finding an ε -optimal robust policy when the sample complexity falls below the order of

$$\Omega \left(\frac{SC_{\text{rob}}^*}{(1-\gamma)^3 \varepsilon^2} \right),$$

which is at least as large as the sample complexity requirement of non-robust offline RL (Li et al., 2022). Consequently, this again suggests that learning robust RMDPs with the KL uncertainty set can be at least as hard as the standard MDPs (which corresponds to $\sigma = 0$), for sufficiently small uncertainty levels.

- When the uncertainty level $\sigma \asymp \log(1/P_{\min}^*)$, Theorem 10 shows that no algorithm can succeed in finding an ε -optimal robust policy when the sample complexity falls below the order of

$$\Omega \left(\frac{SC_{\text{rob}}^*}{P_{\min}^* (1-\gamma)^2 \sigma^2 \varepsilon^2} \right),$$

which directly confirms that DRVI-LCB is near-optimal up to a polynomial factor of the effective horizon length $\frac{1}{1-\gamma}$ (cf. (59)). To the best of our knowledge, DRVI-LCB is the first provable algorithm with near-optimal sample complexity for infinite-horizon robust offline RL. Moreover, the requirement imposed on the history dataset is also much weaker than prior literature on robust offline RL (Yang et al., 2022; Zhou et al., 2021), without the need of full coverage of the state-action space.

5. Numerical experiments

We conduct experiments on the gambler’s problem (Sutton and Barto, 2018; Zhou et al., 2021) to evaluate the performance of the proposed algorithm DRVI-LCB, with comparisons to the robust value iteration algorithm DRVI without pessimism (Panaganti and Kalathil, 2022). Our code can be accessed at:

<https://github.com/Laixishi/Robust-RL-with-KL-divergence>.

Gambler’s problem. In the gambler’s game (Sutton and Barto, 2018; Zhou et al., 2021), a gambler bets on a sequence of coin flips, winning the stake with heads and losing with tails. Starting from some initial balance, the game ends when the gambler’s balance either reaches 50 or 0, or the total number of bets H is hit. This problem can be formulated as an episodic finite-horizon MDP, with a state space $\mathcal{S} = \{0, 1, \dots, 50\}$ and the associated possible actions $a \in \{0, 1, \dots, \min\{s, 50 - s\}\}$ at state s . Here, we set the horizon length $H = 100$. Moreover, the parameter of the transition kernel, which is the probability of heads for the coin flip, is fixed as p_{head} and remains the same in all time steps $h \in [H]$. The reward is set as 1 when the state reaches $s = 50$ and 0 for all other cases. In addition, suppose the initial state (i.e., the gambler’s initial balance) distribution ρ is taken uniformly at random within $\mathcal{S}$.

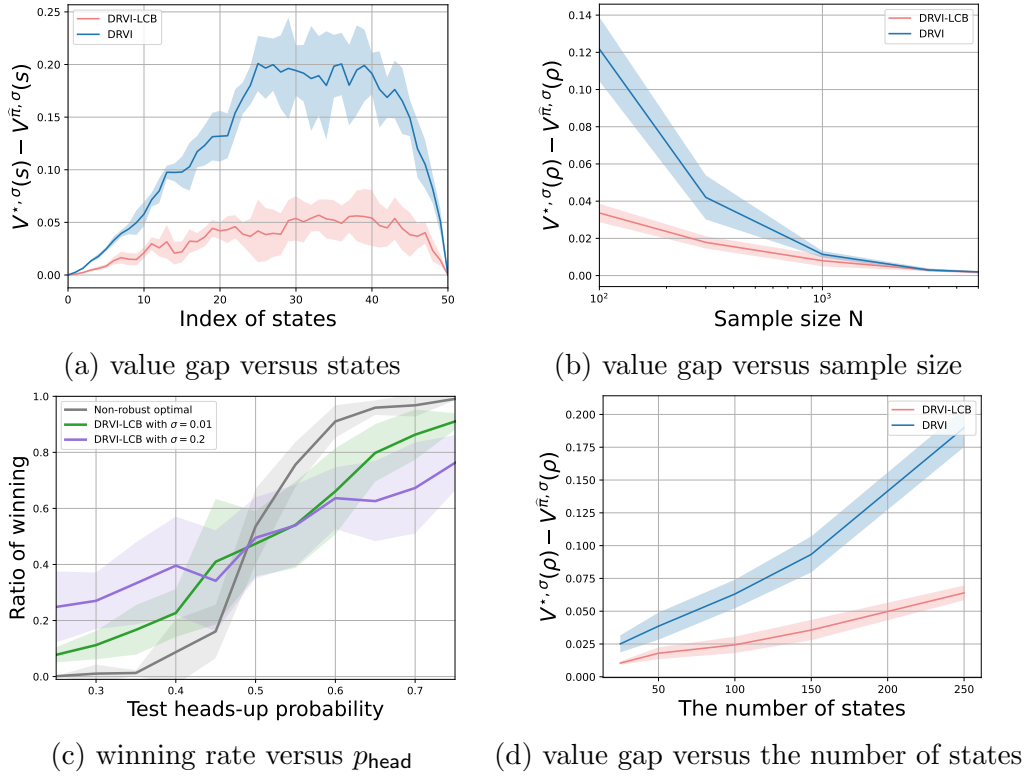


Figure 1: Performance of the proposed algorithm DRVI-LCB using independent samples per state-action pair and time step, where it shows better sample efficiency than the baseline algorithm DRVI without pessimism, as well as better robustness in the learned policy compare to its non-robust counterpart.

The benefit of pessimism. We first utilize a history dataset with N independent samples per state-action pair and time step, generated from the nominal MDP with $p_{\text{head}}^0 = 0.6$. We evaluate the performance of the learned policy $\hat{\pi}$ using our proposed method DRVI-LCB with comparison to DRVI without pessimism, where we fix the uncertainty level $\sigma = 0.1$ for learning the robust optimal policy. The experiments are repeated 10 times with the average

and standard deviations reported. To begin with, Figure 1(a) plots the sub-optimality value gap $V_1^{*,\sigma}(s) - V_1^{\hat{\pi},\sigma}(s)$ for every $s \in \mathcal{S}$, when a sample size $N = 100$ is used to learn the robust policies. It is shown that DRVI-LCB outperform the baseline DRVI uniformly over the state space when the sample size is small, corroborating the benefit of pessimism in the sample-starved regime. Furthermore, Figure 1(b) shows the sub-optimality gap $V_1^{*,\sigma}(\rho) - V_1^{\hat{\pi},\sigma}(\rho)$ with varying sample sizes $N = 100, 300, 1000, 3000, 5000$, where the initial test distribution ρ is generated randomly.² While the performance of DRVI-LCB and DRVI both improves with the increase of the sample size, the proposed algorithm DRVI-LCB achieves much better performance with fewer samples.

The benefit of distributional robustness. To corroborate the benefit of distributional robustness, we evaluate the performance of the policy learned from $N = 1000$ samples using DRVI-LCB on perturbed environments with varying model parameters $p_{\text{head}} \in [0.25, 0.75]$. We measure the practical performance based on the ratio of winning (i.e., reaching the state $s = 50$) calculated from 3000 episodes. Figure 1(c) illustrates the ratio of winning against the test probability of heads for the policies learned from DRVI-LCB with $\sigma = 0.01$ and $\sigma = 0.2$, which are benchmarked against the non-robust optimal policy of the nominal MDP using the exact model. It can be seen that the policies learned from DRVI-LCB deviate from the non-robust optimal policy as σ increases, which achieves better worst-case rates of winning across a wide range of perturbed environments. On the other end, while the non-robust policy maximizes the performance when the test environment is close to the history one used for training, its performance degenerates to be much worse than the robust policies when the probability of heads is mismatched significantly, especially when p_{head} drops below, say around, 0.5.

Impact of the number of states. We evaluate the performance of DRVI-LCB and DRVI when the number of states varies within $[25, 50, 100, 150, 200, 250]$, using a fixed sample size $N = 300$. Figure 1(d) show that DRVI-LCB performs consistently better than DRVI as the number of states S increases, where the value gap exhibits a linear scaling with respect to the number of states.

Performance using trajectory data. Instead of using independent samples, we now evaluate the proposed algorithm using a dataset consisting of K sample trajectories generated from a uniform random policy, for the same setting of Figure 1(b). Figure 2 shows the sub-optimality value gap with respect to the number of trajectories K , where the performance of both DRVI-LCB and DRVI improves as K increases, and DRVI-LCB achieves better performance especially when K is small, consistent with the observation under independent data.

6. Conclusion

To accommodate both model robustness and sample efficiency, this paper proposes a distributionally robust model-based algorithm for offline RL with the principle of pessimism. We study the finite-sample complexity of the proposed algorithm DRVI-LCB, and estab-

2. The probability distribution vector $\rho \in \Delta(\mathcal{S})$ is generated as $\rho(s) = u_s / \sum_{s \in \mathcal{S}} u_s$, where u_s is drawn independently from a uniform distribution.

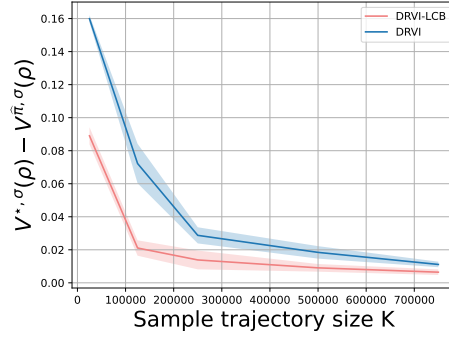


Figure 2: Performance of the proposed algorithm DRVI-LCB compared against DRVI using trajectory data.

lish an information-theoretic lower bound to benchmark its near-optimality for a range of uncertainty levels. Numerical experiments are provided to demonstrate the efficacy of the proposed algorithm. To the best of our knowledge, this provides the first provably near-optimal robust offline RL algorithm that learns under model perturbation and partial coverage. This work opens up several interesting directions.

- *Tightening the gap between upper and lower bounds.* Our upper and lower bounds still leave room for future improvements. For example, it is yet to establish the information-theoretic lower bound over the full range of the uncertainty set, and close the gap between the upper and lower bounds with respect to the horizon length.
- *Model-free algorithms for robust offline RL.* Can we design provably efficient model-free algorithms for robust offline RL with partial coverage? Recent works (Wang et al., 2023a,b) in understanding robust variants of Q-learning in the generative model might shed light on how to approach this question.
- *Choice of uncertainty sets.* Moreover, it is possible to extend our framework to handle uncertainty sets defined using other distances such as the chi-square distance and the total variation distance in a similar fashion. Shi et al. (2023) recently established near minimax-optimal sample complexities for the total variation distance and the chi-square distance in the generative model setting, paving ways to study these uncertainty sets in the offline setting.
- *Adaptive tuning of the uncertainty set.* In this work, we treat the radius of the uncertainty set as a fixed, a priori specified parameter, and study the sample complexity of learning a robust optimal policy with respect to the given uncertainty set modeling the sim-to-real gap. It is of great interest to incorporate the tuning of the uncertainty set (both its size and metric) to complete the pipeline of the algorithm design, which will require a different framework than the one adopted in the current paper.

We leave these questions to future investigations.

Acknowledgments

This work is supported in part by the grants ONR N00014-19-1-2404, NSF CCF-2106778, DMS-2134080, and CNS-2148212. L. Shi is also gratefully supported by the Leo Finzi Memorial Fellowship, Wei Shen and Xuehong Zhang Presidential Fellowship, and Liang Ji-Dian Graduate Fellowship at Carnegie Mellon University. The authors thank Gen Li for helpful discussions.

Appendix A. Preliminaries

Before starting, let us introduce some additional notation useful throughout the theoretical analysis. Let $\text{ess inf } X$ denote the essential infimum of a function/variable X .

A.1 Properties of the robust Bellman operator

To begin with, we introduce the following strong duality lemma which is widely used in distributionally robust optimization when the uncertainty set is defined with respect to the KL divergence.

Lemma 11 ((Hu and Hong, 2013), Theorem 1) *Suppose $f(x)$ has a finite moment generating function in some neighborhood around $x = 0$, then for any $\sigma > 0$ and a nominal distribution P^0 , we have*

$$\sup_{\mathcal{P} \in \mathcal{U}^\sigma(P^0)} \mathbb{E}_{X \sim \mathcal{P}}[f(X)] = \inf_{\lambda \geq 0} \left\{ \lambda \log \mathbb{E}_{X \sim P^0} \left[\exp \left(\frac{f(X)}{\lambda} \right) \right] + \lambda \sigma \right\}. \quad (62)$$

Armed with the above lemma, it is easily verified that for any positive constant M and a nominal distribution vector $P^0 \in \mathbb{R}^{1 \times S}$ supported over the state space $\mathcal{S}$, if $X(s) \in [0, M]$ for all $s \in \mathcal{S}$, then

$$\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P^0)} \mathcal{P}X = \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P^0 \exp \left(-\frac{X}{\lambda} \right) \right) - \lambda \sigma \right\}. \quad (63)$$

For convenience, we introduce the following lemma, paraphrased from Zhou et al. (2021, Lemma 4) and its proof, to further characterize several essential properties of the optimal dual value.

Lemma 12 ((Zhou et al., 2021)) *Let $X \sim P$ be a bounded random variable with $X \in [0, M]$. Let $\sigma > 0$ be any uncertainty level and the corresponding optimal dual variable be*

$$\lambda^* \in \arg \max_{\lambda \geq 0} f(\lambda, P), \quad \text{where } f(\lambda, P) := \left\{ -\lambda \log \mathbb{E}_{X \sim P} \left[\exp \left(\frac{-X}{\lambda} \right) \right] - \lambda \sigma \right\}. \quad (64)$$

Then the optimal value λ^ obeys*

$$\lambda^* \in \left[0, \frac{M}{\sigma} \right], \quad (65)$$

where $\lambda^ = 0$ if and only if*

$$\log (\mathbb{P}(X = \text{ess inf } X)) + \sigma \geq 0. \quad (66)$$

Moreover, when $\lambda^ = 0$, we have*

$$\lim_{\lambda \rightarrow 0} f(\lambda, P) = \lim_{\lambda \rightarrow 0} \left\{ -\lambda \log \mathbb{E}_{X \sim P} \left[\exp \left(\frac{-X}{\lambda} \right) \right] - \lambda \sigma \right\} = \text{ess inf } X. \quad (67)$$

A.2 Concentration inequalities

In light of Lemma 12 (cf. 67), we are interested in comparing the values of $\text{essinf} X$ when X is drawn from the population nominal distribution or its empirical estimate. This is supplied by the following lemma from Zhou et al. (2021).

Lemma 13 ((Zhou et al., 2021)) *Let $X \sim P$ be a discrete bounded random variable with $X \in [0, M]$. Let P_n denote the empirical distribution constructed from n independent samples $X_1, X_2, \dots, X_n$, and let $\hat{X} \sim P_n$. Denote $P_{\min, X}$ as the smallest positive probability $P_{\min, X} := \min\{\mathbb{P}(X = x) : x \in \text{supp}(X)\}$, where $\text{supp}(X)$ is the support of X . Then for any $\delta \in (0, 1)$, with probability at least $1 - \delta$, we have*

$$\min_{i \in [n]} X_i = \text{essinf} \hat{X} = \text{essinf} X, \quad (68)$$

as long as

$$n \geq -\frac{\log(2/\delta)}{\log(1 - P_{\min, X})}. \quad (69)$$

We next gather a few elementary facts about the Binomial distribution, which will be useful throughout the proof.

Lemma 14 (Chernoff's inequality) *Suppose $N \sim \text{Binomial}(n, p)$, where $n \geq 1$ and $p \in [0, 1]$. For some universal constant $c_f > 0$, we have*

$$\mathbb{P}(|N/n - p| \geq pt) \leq \exp(-c_f n p t^2), \quad \forall t \in [0, 1]. \quad (70)$$

Lemma 15 ((Shi et al., 2022, Lemma 8)) *Suppose $N \sim \text{Binomial}(n, p)$, where $n \geq 1$ and $p \in [0, 1]$. For any $\delta \in (0, 1)$, we have*

$$N \geq \frac{np}{8 \log(\frac{1}{\delta})} \quad \text{if } np \geq 8 \log\left(\frac{1}{\delta}\right), \quad (71a)$$

$$N \leq \begin{cases} e^2 np & \text{if } np \geq \log\left(\frac{1}{\delta}\right), \\ 2e^2 \log\left(\frac{1}{\delta}\right) & \text{if } np \leq 2 \log\left(\frac{1}{\delta}\right) \end{cases} \quad (71b)$$

hold with probability at least $1 - 4\delta$.

A.3 Kullback-Leibler (KL) divergence

We next introduce some useful facts about the Kullback-Leibler (KL) divergence for two distributions P and Q , denoted as $\text{KL}(P \parallel Q)$. Denoting $\text{Ber}(p)$ (resp. $\text{Ber}(q)$) as the Bernoulli distribution with mean p (resp. q), we introduce

$$\text{KL}(\text{Ber}(p) \parallel \text{Ber}(q)) := p \log \frac{p}{q} + (1 - p) \log \frac{1 - p}{1 - q}, \quad (72)$$

which represents the KL divergence from $\text{Ber}(p)$ to $\text{Ber}(q)$. We now introduce the following lemma.

Lemma 16 *For any $p, q \in [\frac{1}{2}, 1)$ and $p > q$, it holds that*

$$\text{KL}(\text{Ber}(p) \parallel \text{Ber}(q)) \leq \text{KL}(\text{Ber}(q) \parallel \text{Ber}(p)) \leq \frac{(p-q)^2}{p(1-p)}. \quad (73)$$

Moreover, for any $0 \leq x < y < q$, it holds

$$\text{KL}(\text{Ber}(x) \parallel \text{Ber}(q)) > \text{KL}(\text{Ber}(y) \parallel \text{Ber}(q)). \quad (74)$$

Proof The first half of this lemma is proven in Li et al. (2022, Lemma 10). For the latter half, it follows from that the function

$$f(x, q) := \text{KL}(\text{Ber}(x) \parallel \text{Ber}(q))$$

is monotonically decreasing for all $x \in (0, q]$, since its derivative with respect to x satisfies $\frac{\partial f(x, q)}{\partial x} = \log \frac{x}{q} + \log \frac{1-q}{1-x} < 0$. $\blacksquare$

Appendix B. Analysis: episodic finite-horizon RMDPs

B.1 Proof of Theorem 3

Before starting, we introduce several additional notation that will be useful in the analysis. First, we denote the state-action space covered by the behavior policy π^b in the nominal model P^0 as

$$\mathcal{C}^b = \left\{ (h, s, a) : d_h^b(s, a) > 0 \right\}. \quad (75)$$

Moreover, we recall the definition in (23) and define a similar one based on the exact nominal model P^0 as

$$P_{\min, h}(s, a) := \min_{s'} \left\{ P_h^0(s' | s, a) : P_h^0(s' | s, a) > 0 \right\}. \quad (76)$$

Clearly, by comparing with the definitions (24) and (25), it holds that

$$P_{\min}^* = \min_{h, s} P_{\min, h}(s, \pi_h^*(s)), \quad P_{\min}^b = \min_{(h, s, a) \in \mathcal{C}^b} P_{\min, h}(s, a). \quad (77)$$

For any time step $h \in [H]$, we denote the set of possible state occupancy distributions associated with the optimal policy π^* in a model within the uncertainty set $P \in \mathcal{U}^\sigma(P^0)$ as

$$\mathcal{D}_h^* := \left\{ \left[d_h^{\star, P}(s) \right]_{s \in \mathcal{S}} : P \in \mathcal{U}^\sigma(P^0) \right\} = \left\{ \left[d_h^{\star, P}(s, \pi_h^*(s)) \right]_{s \in \mathcal{S}} : P \in \mathcal{U}^\sigma(P^0) \right\}, \quad (78)$$

where the second equality is due to the fact that π^* is chosen to be deterministic.

With these in place, the proof of Theorem 3 is separated into several key steps, as outlined below.

Step 1: establishing the pessimism property. To achieve this claim, we heavily count on the following lemma whose proof can be found in Appendix B.2.

Lemma 17 *Instate the assumptions in Theorem 3. Then for all $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$, consider any vector $V \in \mathbb{R}^S$ independent of $\hat{P}_{h,s,a}^0$ obeying $\|V\|_\infty \leq H$. With probability at least $1 - \delta$, one has*

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \leq b_h(s, a) \quad (79)$$

with $b_h(s, a)$ given in (22). Moreover, for all $(h, s, a) \in \mathcal{C}^b$, with probability at least $1 - \delta$, one has

$$\frac{P_{\min,h}(s, a)}{8 \log(KHS/\delta)} \leq \hat{P}_{\min,h}(s, a) \leq e^2 P_{\min,h}(s, a). \quad (80)$$

Armed with the above lemma, with probability at least $1 - \delta$, we shall show the following relation holds

$$\forall (s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H + 1] : \quad \hat{Q}_h(s, a) \leq Q_h^{\hat{\pi}, \sigma}(s, a), \quad \hat{V}_h(s) \leq V_h^{\hat{\pi}, \sigma}(s), \quad (81)$$

which means that $\hat{Q}_h$ (resp. $\hat{V}_h$) is a pessimistic estimate of $Q_h^{\hat{\pi}, \sigma}$ (resp. $V_h^{\hat{\pi}, \sigma}$). Towards this, it is easily verified that the latter assertion concerning $V_h^{\hat{\pi}, \sigma}$ is implied by the former, since

$$\hat{V}_h(s) = \max_a \hat{Q}_h(s, a) \leq \max_a Q_h^{\hat{\pi}, \sigma}(s, a) = V_h^{\hat{\pi}, \sigma}(s). \quad (82)$$

Therefore, the remainder of this step focuses on verifying the former assertion in (81) by induction.

- To begin, the claim (81) holds at the base case when $h = H + 1$, by invoking the trivial fact $\hat{Q}_{H+1}(s, a) = Q_{H+1}^{\hat{\pi}, \sigma}(s, a) = 0$.
- Then, suppose that $\hat{Q}_{h+1}(s, a) \leq Q_{h+1}^{\hat{\pi}, \sigma}(s, a)$ holds for all $(s, a) \in \mathcal{S} \times \mathcal{A}$ at some time step $h \in [H]$, it boils down to show $\hat{Q}_h(s, a) \leq Q_h^{\hat{\pi}, \sigma}(s, a)$.

By the update rule of $\hat{Q}_h(s, a)$ in Algorithm 2 (cf. line 7), the above relation holds immediately if $\hat{Q}_h(s, a) = 0$ since $\hat{Q}_h(s, a) = 0 \leq Q_h^{\hat{\pi}, \sigma}(s, a)$. Otherwise, $\hat{Q}_h(s, a)$ is updated via

$$\begin{aligned} & \hat{Q}_h(s, a) \\ &= r_h(s, a) + \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-\hat{V}_{h+1}}{\lambda} \right) \right) - \lambda \sigma \right\} - b_h(s, a) \\ &\stackrel{(i)}{=} r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P} \hat{V}_{h+1} - b_h(s, a) \\ &\leq r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P} \hat{V}_{h+1} + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P} \hat{V}_{h+1} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P} \hat{V}_{h+1} \right| - b_h(s, a) \end{aligned}$$

$$\stackrel{(ii)}{\leq} r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V_{h+1}^{\hat{\pi},\sigma} + 0 \stackrel{(iii)}{=} Q_h^{\hat{\pi},\sigma}(s, a), \quad (83)$$

where (i) rewrites the update rule back to its primal form (cf. (19)), (ii) holds by applying (79) with the condition (28) satisfied and the induction hypothesis $\hat{V}_{h+1} \leq V_{h+1}^{\hat{\pi},\sigma}$, and lastly, (iii) follows by the robust Bellman consistency equation (8).

Putting them together, we have verified the claim (81) by induction.

Step 2: bounding $V_h^{*,\sigma}(s) - V_h^{\hat{\pi},\sigma}(s)$. With the pessimism property (81) in place, we observe that the following relation holds

$$0 \leq V_h^{*,\sigma}(s) - V_h^{\hat{\pi},\sigma}(s) \leq V_h^{*,\sigma}(s) - \hat{V}_h(s) \leq Q_h^{*,\sigma}(s, \pi_h^*(s)) - \hat{Q}_h(s, \pi_h^*(s)), \quad (84)$$

where the last inequality follows from $\hat{Q}_h(s, \pi_h^*(s)) \leq \max_a \hat{Q}_h(s, a) = \hat{V}_h(s)$. Then, by the robust Bellman optimality equation in (9) and the primal version of the update rule (cf. (19))

$$\begin{aligned} Q_h^{*,\sigma}(s, \pi_h^*(s)) &= r_h(s, \pi_h^*(s)) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}V_{h+1}^{*,\sigma}, \\ \hat{Q}_h(s, \pi_h^*(s)) &= r_h(s, \pi_h^*(s)) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} - b_h(s, \pi_h^*(s)), \end{aligned}$$

we arrive at

$$\begin{aligned} V_h^{*,\sigma}(s) - \hat{V}_h(s) &\leq Q_h^{*,\sigma}(s, \pi_h^*(s)) - \hat{Q}_h(s, \pi_h^*(s)) \\ &= \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}V_{h+1}^{*,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} + b_h(s, \pi_h^*(s)) \\ &\leq \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}V_{h+1}^{*,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} \\ &\quad + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} \right| + b_h(s, \pi_h^*(s)) \\ &\stackrel{(i)}{\leq} \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}V_{h+1}^{*,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} + 2b_h(s, \pi_h^*(s)) \\ &\stackrel{(ii)}{\leq} \hat{P}_{h,s,\pi_h^*(s)}^{\inf}(V_{h+1}^{*,\sigma} - \hat{V}_{h+1}) + 2b_h(s, \pi_h^*(s)), \end{aligned} \quad (85)$$

where (i) holds by applying Lemma 2 (cf. (79)) since $\hat{V}_{h+1}$ is independent of $P_{h,s,\pi_h^*(s)}^0$ by construction, and (ii) arises from introducing the notation

$$\hat{P}_{h,s,\pi_h^*(s)}^{\inf} := \operatorname{argmin}_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*(s)}^0)} \mathcal{P}\hat{V}_{h+1} \quad (86)$$

and consequently,

$$\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*}^0)} \mathcal{P}V_{h+1}^{\star,\sigma} \leq \hat{P}_{h,s,\pi_h^*}^{\text{inf}} V_{h+1}^{\star,\sigma}, \quad \text{and} \quad \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,\pi_h^*}^0)} \mathcal{P}\hat{V}_{h+1} = \hat{P}_{h,s,\pi_h^*}^{\text{inf}} \hat{V}_{h+1}.$$

To continue, let us introduce some additional notation for convenience. Define a sequence of matrices $\hat{P}_h^{\text{inf}} \in \mathbb{R}^{S \times S}$ and vectors $b_h^* \in \mathbb{R}^S$ for $h \in [H]$, where their s -th rows (resp. entries) are given by

$$\left[\hat{P}_h^{\text{inf}} \right]_{s,\cdot} = \hat{P}_{h,s,\pi_h^*}^{\text{inf}}, \quad \text{and} \quad b_h^*(s) = b_h(s, \pi_h^*(s)). \quad (87)$$

Applying (85) recursively over the time steps $h, h+1, \dots, H$ using the above notation gives

$$\begin{aligned} 0 \leq V_h^{\star,\sigma} - \hat{V}_h &\leq \hat{P}_h^{\text{inf}} (V_{h+1}^{\star,\sigma} - \hat{V}_{h+1}) + 2b_h^* \\ &\leq \hat{P}_h^{\text{inf}} \hat{P}_{h+1}^{\text{inf}} (V_{h+2}^{\star,\sigma} - \hat{V}_{h+2}) + 2\hat{P}_h^{\text{inf}} b_{h+1}^* + 2b_h^* \leq \dots \leq 2 \sum_{i=h}^H \left(\prod_{j=h}^{i-1} \hat{P}_j^{\text{inf}} \right) b_i^*, \end{aligned} \quad (88)$$

where we let $\left(\prod_{j=i}^{i-1} \hat{P}_j^{\text{inf}} \right) = I$ for convenience.

For any $d_h^* \in \mathcal{D}_h^*$ (cf. (78)), taking inner product with (88) leads to

$$\left\langle d_h^*, V_h^{\star,\sigma} - \hat{V}_h \right\rangle \leq \left\langle d_h^*, 2 \sum_{i=h}^H \left(\prod_{j=h}^{i-1} \hat{P}_j^{\text{inf}} \right) b_i^* \right\rangle = 2 \sum_{i=h}^H \langle d_i^*, b_i^* \rangle, \quad (89)$$

where

$$d_i^* := \left[(d_h^*)^\top \left(\prod_{j=h}^{i-1} \hat{P}_j^{\text{inf}} \right) \right]^\top \in \mathcal{D}_i^* \quad (90)$$

by the definition of $\mathcal{D}_i^*$ (cf. (78)) for all $i = h+1, \dots, H$.

Step 3: controlling $\langle d_i^*, b_i^* \rangle$ using concentrability. Since $\langle d_i^*, b_i^* \rangle = \sum_{s \in S} d_i^*(s) b_i^*(s)$, we shall divide the discussion in two different cases.

- For $s \in S$ where $\max_{P \in \mathcal{U}^\sigma(P^0)} d_i^{\star,P}(s, \pi_i^*(s)) = \max_{P \in \mathcal{U}^\sigma(P^0)} d_i^{\star,P}(s) = 0$, it follows from the definition (cf. (78)) that for any $d_i^* \in \mathcal{D}_i^*$, it satisfies that

$$d_i^*(s) = 0. \quad (91)$$

- For $s \in S$ where $\max_{P \in \mathcal{U}^\sigma(P^0)} d_i^{\star,P}(s, \pi_i^*(s)) = \max_{P \in \mathcal{U}^\sigma(P^0)} d_i^{\star,P}(s) > 0$, by the assumption in (14)

$$\max_{P \in \mathcal{U}^\sigma(P^0)} \frac{\min \{ d_i^{\star,P}(s, \pi_i^*(s)), \frac{1}{S} \}}{d_i^{\text{b}}(s, \pi_i^*(s))} = \max_{P \in \mathcal{U}^\sigma(P^0)} \frac{\min \{ d_i^{\star,P}(s), \frac{1}{S} \}}{d_i^{\text{b}}(s, \pi_i^*(s))} \leq C_{\text{rob}}^* < \infty,$$

it implies that

$$d_i^b(s, \pi_i^*(s)) > 0 \quad \text{and} \quad (i, s, \pi_i^*(s)) \in \mathcal{C}^b. \quad (92)$$

Lemma 2 tells that with probability at least $1 - 8\delta$,

$$\begin{aligned} N_i(s, \pi_i^*(s)) &\geq \frac{K d_i^b(s, \pi_i^*(s))}{8} - 5 \sqrt{K d_i^b(s, \pi_i^*(s)) \log \frac{KH}{\delta}} \stackrel{(i)}{\geq} \frac{K d_i^b(s, \pi_i^*(s))}{16} \\ &\stackrel{(ii)}{\geq} \frac{K \max_{P \in \mathcal{U}^\sigma(P^0)} \min \left\{ d_i^{*,P}(s, \pi_i^*(s)), \frac{1}{S} \right\}}{16 C_{\text{rob}}^*} \geq \frac{K \min \left\{ d_i^*(s), \frac{1}{S} \right\}}{16 C_{\text{rob}}^*}, \end{aligned} \quad (93)$$

where (i) holds due to

$$K d_i^b(s, \pi_i^*(s)) \geq c_1 \frac{d_i^b(s, \pi_i^*(s)) \log(KHS/\delta)}{d_{\min}^b P_{\min}^b} \geq \frac{c_1 \log \frac{KH}{\delta}}{P_{\min}^b} \geq c_1 \log \frac{KH}{\delta} \quad (94)$$

for some sufficiently large c_1 , where the first inequality follows from Condition (28), the second inequality follows from

$$d_{\min}^b = \min_{h,s,a} \left\{ d_h^b(s, a) : d_h^b(s, a) > 0 \right\} \leq d_i^b(s, \pi_i^*(s)) \quad (95)$$

and the last inequality follows from $P_{\min}^b \leq 1$. In addition, (ii) follows from Assumption 1.

With this in place, we observe that the pessimistic penalty (see (22)) obeys

$$\begin{aligned} b_i^*(s) &\leq c_b \frac{H}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{\widehat{P}_{\min,i}(s, \pi_i^*(s)) N_i(s, \pi_i^*(s))}} \stackrel{(i)}{\leq} 4c_b \frac{H}{\sigma} \sqrt{\frac{\log^2(\frac{KHS}{\delta})}{P_{\min,i}(s, \pi_i^*(s)) N_i(s, \pi_i^*(s))}} \\ &\leq 16c_b \frac{H}{\sigma} \sqrt{\frac{C_{\text{rob}}^* \log^2 \frac{KHS}{\delta}}{P_{\min,i}(s, \pi_i^*(s)) K \min \left\{ d_i^*(s), \frac{1}{S} \right\}}}, \end{aligned} \quad (96)$$

where (i) holds by applying (80) in view of the fact that $(i, s, \pi_i^*(s)) \in \mathcal{C}^b$ by (92), and the last inequality holds by (93).

Combining the results in the above two cases leads to

$$\begin{aligned} \sum_{s \in \mathcal{S}} d_i^*(s) b_i^*(s) &\leq \sum_{s \in \mathcal{S}} 16 d_i^*(s) c_b \frac{H}{\sigma} \sqrt{\frac{C_{\text{rob}}^* \log^2 \frac{KHS}{\delta}}{P_{\min,i}(s, \pi_i^*(s)) K \min \left\{ d_i^*(s), \frac{1}{S} \right\}}} \\ &\stackrel{(i)}{\leq} 16c_b \frac{H}{\sigma} \sqrt{\sum_{s \in \mathcal{S}} d_i^*(s) \frac{C_{\text{rob}}^* \log^2 \frac{KHS}{\delta}}{P_{\min,i}(s, \pi_i^*(s)) K \min \left\{ d_i^*(s), \frac{1}{S} \right\}}} \sqrt{\sum_{s \in \mathcal{S}} d_i^*(s)} \\ &\leq 32c_b \frac{H}{\sigma} \sqrt{\frac{S C_{\text{rob}}^* \log^2 \frac{KHS}{\delta}}{P_{\min,i}(s, \pi_i^*(s)) K}}, \end{aligned} \quad (97)$$

where (i) follows from the Cauchy-Schwarz inequality and the last inequality hold by the trivial fact

$$\sum_{s \in \mathcal{S}} \frac{d_i^*(s)}{\min \left\{ d_i^*(s), \frac{1}{S} \right\}} \leq \sum_{s \in \mathcal{S}} d_i^*(s) \left(\frac{1}{d_i^*(s)} + \frac{1}{1/S} \right) = \sum_{s \in \mathcal{S}} 1 + \frac{1}{S} \sum_{s \in \mathcal{S}} d_i^*(s) \leq 2S. \quad (98)$$

Step 4: finishing up the proof. Then, inserting (97) back into (89) with $h = 1$ shows

$$\left\langle d_1^*, V_1^{*,\sigma} - \widehat{V}_1 \right\rangle \leq 2 \sum_{i=1}^H \langle d_i^*, b_i^* \rangle \leq \sum_{i=1}^H 64c_b \frac{H}{\sigma} \sqrt{\frac{SC_{\text{rob}}^* \log^2 \frac{KH}{\delta}}{P_{\min,i}(s, \pi_i^*(s))K}} \leq c_2 \frac{H^2}{\sigma} \sqrt{\frac{SC_{\text{rob}}^* \log^2 \frac{KH}{\delta}}{P_{\min}^* K}}, \quad (99)$$

where the last inequality holds by plugging in the relation $P_{\min}^* \leq P_{\min,i}(s, \pi_i^*(s))$ for $i = 1, \dots, H$ by the definition in (24) (see also (77)), and choosing c_2 to be large enough. The proof is completed.

B.2 Proof of Lemma 17

To begin, we shall introduce the following fact that

$$\forall (h, s, a) \in \mathcal{C}^b : \quad N_h(s, a) \geq \frac{c_1 \log \frac{KHS}{\delta}}{16P_{\min,h}(s, a)} \geq -\frac{\log \frac{2KHS}{\delta}}{\log(1 - P_{\min,h}(s, a))}, \quad (100)$$

as long as Condition (28) holds. The proof is postponed to Appendix B.2.3. With this in mind, we shall first establish the simpler bound (80) and then move on to show (79).

B.2.1 PROOF OF (80)

To begin, recall that (100) is satisfied for all $(h, s, a) \in \mathcal{C}^b$. By Lemma 15 and the union bound, it holds that with probability at least $1 - \delta$ that for all $(h, s, a) \in \mathcal{C}^b$:

$$\forall s' \in \mathcal{S} : \quad P_h^0(s' | s, a) \geq \frac{\widehat{P}_h^0(s' | s, a)}{e^2} \geq \frac{P_h^0(s' | s, a)}{8e^2 \log(\frac{KHS}{\delta})}. \quad (101)$$

To characterize the relation between $P_{\min,h}(s, a)$ and $\widehat{P}_{\min,h}(s, a)$ for any $(h, s, a) \in \mathcal{C}^b$, we suppose—without loss of generality—that $P_{\min,h}(s, a) = P_h^0(s_1 | s, a)$ and $\widehat{P}_{\min,h}(s, a) = \widehat{P}_h^0(s_2 | s, a)$ for some $s_1, s_2 \in \mathcal{S}$. Then, it follows that

$$\begin{aligned} P_{\min,h}(s, a) = P_h^0(s_1 | s, a) &\stackrel{(i)}{\geq} \frac{\widehat{P}_h^0(s_1 | s, a)}{e^2} \geq \frac{\widehat{P}_{\min,h}(s, a)}{e^2} = \frac{\widehat{P}_h^0(s_2 | s, a)}{e^2} \\ &\stackrel{(ii)}{\geq} \frac{P_h^0(s_2 | s, a)}{8e^2 \log(\frac{KHS}{\delta})} \geq \frac{P_{\min,h}(s, a)}{8e^2 \log(\frac{KHS}{\delta})}, \end{aligned}$$

where (i) and (ii) follow from (101).

B.2.2 PROOF OF (79)

The main goal of (79) is to control the gap between robust Bellman operations based on the nominal transition kernel $P_{h,s,a}^0$ and the estimated kernel $\widehat{P}_{h,s,a}^0$ by the constructed penalty term. Towards this, first consider $(h, s, a) \notin \mathcal{C}^b$, which corresponds to the state-action pairs (s, a) that haven't been visited at step h by the behavior policy. In other words, $N_h(s, a) = 0$. In this case, (79) can be easily verified that

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \stackrel{(i)}{=} \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \leq \|V\|_\infty \stackrel{(ii)}{\leq} H \stackrel{(iii)}{=} b_h(s, a), \quad (102)$$

where (i) follows from the fact $\widehat{P}_{h,s,a}^0 = 0$ when $N_h(s, a) = 0$ (see (15)), (ii) arises from the assumption $\|V\|_\infty \leq H$, and (iii) holds by the definition of $b_h(s, a)$ in (22). Therefore, the remainder of the proof will focus on verifying (79) for $(h, s, a) \in \mathcal{C}^b$. Rewriting the term of interest via duality (cf. Lemma 11) yields

$$\begin{aligned} & \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \\ &= \left| \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\widehat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} - \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \right|. \end{aligned} \quad (103)$$

Denoting

$$\widehat{\lambda}_{h,s,a}^* := \arg \max_{\lambda \geq 0} \left\{ -\lambda \log \left(\widehat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\}, \quad (104a)$$

$$\lambda_{h,s,a}^* := \arg \max_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\}, \quad (104b)$$

Lemma 12 (cf. (65)) then gives that

$$\lambda_{h,s,a}^* \in \left[0, \frac{H}{\sigma} \right], \quad \widehat{\lambda}_{h,s,a}^* \in \left[0, \frac{H}{\sigma} \right], \quad (105)$$

due to $\|V\|_\infty \leq H$. We shall control (103) in three different cases separately: (a) $\lambda_{h,s,a}^* = 0$ and $\widehat{\lambda}_{h,s,a}^* = 0$; (b) $\lambda_{h,s,a}^* > 0$ and $\widehat{\lambda}_{h,s,a}^* = 0$ or $\lambda_{h,s,a}^* = 0$ and $\widehat{\lambda}_{h,s,a}^* > 0$; and (c) $\lambda_{h,s,a}^* \neq 0$ or $\widehat{\lambda}_{h,s,a}^* \neq 0$.

Case (a): $\lambda_{h,s,a}^* = 0$ and $\widehat{\lambda}_{h,s,a}^* = 0$. Applying Lemma 12 and Lemma 13 to (103) gives that, with probability at least $1 - \frac{\delta}{KH}$,

$$\begin{aligned} & \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \stackrel{(i)}{=} \left| \text{essinf}_{s \sim \widehat{P}_{h,s,a}^0} V(s) - \text{essinf}_{s \sim P_{h,s,a}^0} V(s) \right| \\ & \stackrel{(ii)}{=} \left| \text{essinf}_{s \sim P_{h,s,a}^0} V(s) - \text{essinf}_{s \sim P_{h,s,a}^0} V(s) \right| \\ & = 0 \leq b_h(s, a). \end{aligned} \quad (106)$$

where (i) holds by Lemma 12 (cf. (67)) and (ii) arises from Lemma 13 (cf. (68)) given (100).

Case (b): $\lambda_{h,s,a}^* > 0$ and $\widehat{\lambda}_{h,s,a}^* = 0$ or $\lambda_{h,s,a}^* = 0$ and $\widehat{\lambda}_{h,s,a}^* > 0$. Towards this, note that two trivial facts are implied by the definition (104):

$$\sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \geq -\widehat{\lambda}_{h,s,a}^* \log \left(P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\widehat{\lambda}_{h,s,a}^*} \right) \right) - \widehat{\lambda}_{h,s,a}^* \sigma, \quad (107a)$$

$$\sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \geq -\lambda_{h,s,a}^* \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda_{h,s,a}^*} \right) \right) - \lambda_{h,s,a}^* \sigma. \quad (107b)$$

To continue, first, we consider a subcase when $\lambda_{h,s,a}^* = 0$ and $\hat{\lambda}_{h,s,a}^* > 0$. With probability at least $1 - \frac{\delta}{KH}$, it follows from Lemma 12 (cf. (67)) and Lemma 13 (cf. (68)) that

$$\begin{aligned} \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} &\geq \lim_{\lambda \rightarrow 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \\ &= \text{essinf}_{s \sim \hat{P}_{h,s,a}^0} V(s) = \text{essinf}_{s \sim P_{h,s,a}^0} V(s) \\ &= \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\}, \end{aligned} \quad (108)$$

leading to

$$\begin{aligned} &\left| \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} - \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \right| \\ &\stackrel{(i)}{\leq} \left(-\hat{\lambda}_{h,s,a}^* \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\hat{\lambda}_{h,s,a}^*} \right) \right) - \hat{\lambda}_{h,s,a}^* \sigma \right) \\ &\quad - \left(-\hat{\lambda}_{h,s,a}^* \log \left(P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\hat{\lambda}_{h,s,a}^*} \right) \right) - \hat{\lambda}_{h,s,a}^* \sigma \right) \\ &\leq \hat{\lambda}_{h,s,a}^* \left| \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\hat{\lambda}_{h,s,a}^*} \right) \right) - \log \left(P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\hat{\lambda}_{h,s,a}^*} \right) \right) \right|, \end{aligned} \quad (109)$$

where (i) follows from the definition of $\hat{\lambda}_{h,s,a}^*$ in (104) and the fact in (107a).

We pause to claim that with probability at least $1 - \delta$, the following bound holds

$$\forall (h, s, a) \in \mathcal{C}^b, V \in \mathbb{R}^S : \quad \frac{\left| \left(\hat{P}_{h,s,a}^0 - P_{h,s,a}^0 \right) \cdot \exp \left(\frac{-V}{\lambda} \right) \right|}{P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda} \right)} \leq \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_{\min, h}(s, a)}} \leq \frac{1}{2}. \quad (110)$$

The proof is postponed to Appendix B.2.4. With (110) in place, we can further bound (109) (which is plugged into (103)) as

$$\begin{aligned} &\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \\ &\leq \hat{\lambda}_{h,s,a}^* \left| \log \left(1 + \frac{\left(\hat{P}_{h,s,a}^0 - P_{h,s,a}^0 \right) \cdot \exp \left(\frac{-V}{\lambda} \right)}{P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda} \right)} \right) \right| \\ &\stackrel{(i)}{\leq} 2\hat{\lambda}_{h,s,a}^* \frac{\left| \left(\hat{P}_{h,s,a}^0 - P_{h,s,a}^0 \right) \cdot \exp \left(\frac{-V}{\lambda} \right) \right|}{P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda} \right)} \end{aligned}$$

$$\begin{aligned}
&\stackrel{(ii)}{\leq} \frac{2H}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_{\min, h}(s, a)}} \\
&\leq \frac{2eH}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) \hat{P}_{\min, h}(s, a)}} \leq c_b \frac{H}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{\hat{P}_{\min, h}(s, a) N_h(s, a)}}, \tag{111}
\end{aligned}$$

where (i) follows from $\log(1+x) \leq 2|x|$ for any $|x| \leq \frac{1}{2}$ in view of (110), (ii) follows from (105) as well as (110), and the last line follows from (80) and choosing c_b to be sufficiently large.

Moreover, note that it can be easily verified that

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \leq H$$

due to the assumption $\|V\|_\infty \leq H$. Plugging in the definition of $b_h(s, a)$ in (22), combined with the above bounds, we have that with probability at least $1 - \delta$,

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \leq \min \left\{ c_b \frac{H}{\sigma} \sqrt{\frac{\log(\frac{KHS}{\delta})}{N_h(s, a) \hat{P}_{\min, h}(s, a)}}, H \right\} =: b_h(s, a). \tag{112}$$

The other subcase when $\lambda_{h,s,a}^* > 0$ and $\hat{\lambda}_{h,s,a}^* = 0$ follows similarly from the bound

$$\begin{aligned}
&\left| \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} - \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \right| \\
&\leq \lambda_{h,s,a}^* \left| \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda_{h,s,a}^*} \right) \right) - \log \left(P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda_{h,s,a}^*} \right) \right) \right|, \tag{113}
\end{aligned}$$

and therefore, will be omitted for simplicity.

Case (c): $\lambda_{h,s,a}^* > 0$ and $\hat{\lambda}_{h,s,a}^* > 0$. It follows that

$$\begin{aligned}
&\left| \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(\hat{P}_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} - \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{h,s,a}^0 \exp \left(\frac{-V}{\lambda} \right) \right) - \lambda \sigma \right\} \right| \\
&\stackrel{(i)}{\leq} \max \left\{ \left(-\hat{\lambda}_{h,s,a}^* \log \left(\hat{P}_{h,s,a}^0 \cdot e^{\frac{-V}{\hat{\lambda}_{h,s,a}^*}} \right) - \hat{\lambda}_{h,s,a}^* \sigma \right) - \left(-\hat{\lambda}_{h,s,a}^* \log \left(P_{h,s,a}^0 \cdot e^{\frac{-V}{\hat{\lambda}_{h,s,a}^*}} \right) - \hat{\lambda}_{h,s,a}^* \sigma \right), \right. \\
&\quad \left. \left(-\lambda_{h,s,a}^* \log \left(P_{h,s,a}^0 \cdot e^{\frac{-V}{\lambda_{h,s,a}^*}} \right) - \lambda_{h,s,a}^* \sigma \right) - \left(-\lambda_{h,s,a}^* \log \left(\hat{P}_{h,s,a}^0 \cdot e^{\frac{-V}{\lambda_{h,s,a}^*}} \right) - \lambda_{h,s,a}^* \sigma \right) \right\} \\
&\leq \max_{\lambda \in \{\lambda_{h,s,a}^*, \hat{\lambda}_{h,s,a}^*\}} \lambda \left| \log \left(\hat{P}_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda} \right) \right) - \log \left(P_{h,s,a}^0 \cdot \exp \left(\frac{-V}{\lambda} \right) \right) \right|, \tag{114}
\end{aligned}$$

where (i) can be verified by applying the facts in (107). Hence, the above term (114) can be controlled again in a similar manner as (109); we omit the details for simplicity.

Summing up. Combining the previous results in different cases by the union bound, with probability at least $1 - 10\delta$, it is satisfied that for all $(h, s, a) \in \mathcal{C}^b$:

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{h,s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^0)} \mathcal{P}V \right| \leq b_h(s, a),$$

which concludes the proof.

B.2.3 PROOF OF (100)

Observe that for all $(h, s, a) \in \mathcal{C}^b$:

$$Kd_h^b(s, a) \stackrel{(i)}{\geq} \frac{c_1 d_h^b(s, a) \log(KHS/\delta)}{d_{\min}^b P_{\min}^b} \stackrel{(ii)}{\geq} \frac{c_1 \log(KHS/\delta)}{P_{\min}^b} \stackrel{(iii)}{\geq} \frac{c_1 \log(KHS/\delta)}{P_{\min,h}(s, a)}, \quad (115)$$

where (i) follows from Condition (28), (ii) follows from the definition that $d_{\min}^b \leq d_h^b(s, a)$ for $(h, s, a) \in \mathcal{C}^b$, and (iii) comes from (77).

Lemma 2 then tells that with probability at least $1 - 8\delta$,

$$\begin{aligned} N_h(s, a) &\geq \frac{Kd_h^b(s, a)}{8} - 5\sqrt{Kd_h^b(s, a) \log \frac{KH}{\delta}} \\ &\geq \frac{Kd_h^b(s, a)}{16} \geq \frac{c_1 \log \frac{KH}{\delta}}{16P_{\min,h}(s, a)}, \end{aligned} \quad (116)$$

where the second line follows from the above relation as long as c_1 is sufficiently large. The last inequality of (100) then follows from

$$\frac{c_1 \log \frac{KHS}{\delta}}{16P_{\min,h}(s, a)} \geq -\frac{\log \frac{2KHS}{\delta}}{\log(1 - P_{\min,h}(s, a))}, \quad (117)$$

since $x \leq -\log(1 - x)$ for all $x \in [0, 1]$.

B.2.4 PROOF OF (110)

Denoting

$$\text{supp}(P_{h,s,a}^0) := \{s' \in \mathcal{S} : P_h^0(s' | s, a) > 0\}$$

as the support of $P_{h,s,a}^0$, we observe that

$$\begin{aligned} \left| \frac{(\hat{P}_{h,s,a}^0 - P_{h,s,a}^0) \cdot \exp\left(\frac{-V}{\lambda}\right)}{P_{h,s,a}^0 \cdot \exp\left(\frac{-V}{\lambda}\right)} \right| &\leq \frac{\sum_{s' \in \text{supp}(P_{h,s,a}^0)} |\hat{P}_h^0(s' | s, a) - P_h^0(s' | s, a)| \exp\left(\frac{-V(s')}{\lambda}\right)}{\sum_{s' \in \text{supp}(P_{h,s,a}^0)} P_h^0(s' | s, a) \exp\left(\frac{-V(s')}{\lambda}\right)} \\ &\leq \max_{s' \in \text{supp}(P_{h,s,a}^0)} \frac{|\hat{P}_h^0(s' | s, a) - P_h^0(s' | s, a)|}{P_h^0(s' | s, a)}, \end{aligned} \quad (118)$$

where the second line follows from $\sum_i a_i = \sum_i b_i \frac{a_i}{b_i} \leq (\max_i \frac{a_i}{b_i}) \sum_i b_i$ for any positive sequences $\{a_i, b_i\}_i$ obeying $a_i, b_i > 0$.

To continue, note that for any $(h, s, a) \in \mathcal{C}^b$ and $s' \in \text{supp}(P_{h,s,a}^0)$, $N_h(s, a)\hat{P}_h^0(s' | s, a)$ follows the binomial distribution $\text{Binomial}(N_h(s, a), P_h^0(s' | s, a))$. Thus, applying Lemma 14 with $t = \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_h^0(s' | s, a)}}$ yields

$$\mathbb{P}\left(\left|\hat{P}_h^0(s' | s, a) - P_h^0(s' | s, a)\right| \geq P_h^0(s' | s, a)t\right) \leq \exp(-c_f N_h(s, a) P_h^0(s' | s, a)t^2) \leq \frac{\delta}{KHS}, \quad (119)$$

as soon as $t \leq \frac{1}{2}$, which can be verified by the fact (100) and $P_{\min, h}(s, a) \leq P_h^0(s' | s, a)$ (cf. (76)), namely,

$$N_h(s, a) \geq \frac{c_1 \log \frac{KHS}{\delta}}{16 P_{\min, h}(s, a)} \geq \frac{\log(\frac{KHS}{\delta})}{4c_f P_{\min, h}(s, a)} \geq \frac{\log(\frac{KHS}{\delta})}{4c_f P_h^0(s' | s, a)} \quad (120)$$

as long as c_1 is sufficiently large.

Applying (119) and taking the union bound over $s \in \text{supp}(P_{h,s,a}^0)$ lead to that with probability at least $1 - \frac{\delta}{KH}$,

$$\begin{aligned} \max_{s' \in \text{supp}(P_{h,s,a}^0)} \frac{\left|\hat{P}_h^0(s' | s, a) - P_h^0(s' | s, a)\right|}{P_h^0(s' | s, a)} &\leq \max_{s' \in \text{supp}(P_{h,s,a}^0)} \frac{P_h^0(s' | s, a) \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_h^0(s' | s, a)}}}{P_h^0(s' | s, a)} \\ &= \max_{s' \in \text{supp}(P_{h,s,a}^0)} \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_h^0(s' | s, a)}} \\ &\leq \sqrt{\frac{\log(\frac{KHS}{\delta})}{c_f N_h(s, a) P_{\min, h}(s, a)}} \leq \frac{1}{2}, \end{aligned}$$

where the last line uses again (120). Plugging this back into (118) and applying the union bound over $(h, s, a) \in \mathcal{C}^b$ then completes the proof.

B.3 Proof of Theorem 4

The proof of Theorem 4 is inspired by the construction in Li et al. (2022) for standard MDPs, but is considerably more involved to handle the uncertainty set unique in robust MDPs. In particular, we construct two different classes of hard instances for different range of the uncertainty level σ to achieve a tighter σ -dependent lower bound. In what follows, we start with the lower bound for the case when the uncertainty level is relatively small, by first constructing some hard instances and then characterizing the sample complexity requirements over these instances. We then move onto the case when the uncertainty level is relatively large, and carry out a similar argument.

B.3.1 CONSTRUCTION OF HARD PROBLEM INSTANCES: SMALL UNCERTAINTY LEVEL

Construction of a collection of hard MDPs To begin, let's consider a collection $\Theta \subseteq \{0, 1\}^H$, consisting of vectors with H dimensions. The Gilbert-Varshamov lemma

(Gilbert, 1952) tells that there exists a set $\Theta \subseteq \{0, 1\}^H$ such that:

$$|\Theta| \geq e^{H/8} \quad \text{and} \quad \|\theta - \tilde{\theta}\|_1 \geq \frac{H}{8} \quad \text{for any } \theta, \tilde{\theta} \in \Theta \text{ obeying } \theta \neq \tilde{\theta}. \quad (121)$$

Armed with Θ , we then generate a collection of MDPs

$$\text{MDP}(\Theta) = \left\{ \mathcal{M}^\theta = (\mathcal{S}, \mathcal{A}, P^\theta = \{P_h^{\theta_h}\}_{h=1}^H, \{r_h\}_{h=1}^H, H) \mid \theta = [\theta_h]_{1 \leq h \leq H} \in \Theta \right\}, \quad (122)$$

where

$$\mathcal{S} = \{0, 1, \dots, S-1\}, \quad \text{and} \quad \mathcal{A} = \{0, 1\}.$$

The transition kernel $P^\theta = \{P_h^{\theta_h}\}_{h=1}^H$ of the MDP $\mathcal{M}^\theta$ is specified as follows:

$$P_h^{\theta_h}(s' \mid s, a) = \begin{cases} p\mathbb{1}(s' = 0) + (1-p)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, \theta_h) \\ q\mathbb{1}(s' = 0) + (1-q)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, 1 - \theta_h) \\ \mathbb{1}(s' = 1) & \text{if } s = 1 \\ q\mathbb{1}(s' = s) + (1-q)\mathbb{1}(s' = 1) & \text{if } s > 1 \end{cases} \quad (123)$$

for any $(s, a, s', h) \in \mathcal{S} \times \mathcal{A} \times \mathcal{S} \times [H]$. Here, p and q are set according to

$$p = 1 - \frac{c_1}{H} \quad \text{and} \quad q = p - \frac{c_2 \varepsilon}{H^2} \quad (124)$$

for $c_1 = 1/8$ and some c_2 that satisfies

$$\frac{c_2 \varepsilon}{H^2} \leq \frac{c_1}{2H} \leq \frac{1}{8}. \quad (125)$$

Clearly, it follows that

$$p > q \geq \frac{1}{2} \quad (126)$$

by construction. Furthermore, the MDP will stay in the state subset $\{0, 1\}$ if its initial state falls in $\{0, 1\}$. The reward function of these MDPs is set as

$$r_h(s, a) = \begin{cases} 1 & \text{if } s = 0 \\ 0 & \text{if } s > 0 \end{cases} \quad (127)$$

for any $(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$.

Uncertainty set of the transition kernels. Denote the transition kernel vector as

$$P_{h,s,a}^\theta := P_h^\theta(\cdot \mid s, a) \in [0, 1]^{1 \times S}. \quad (128)$$

For any $(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$, the perturbation of the transition kernels in $\mathcal{M}^\theta$ is restricted to the following uncertainty set

$$\mathcal{U}^\sigma(P^\theta) := \otimes \mathcal{U}^\sigma(P_{h,s,a}^\theta), \quad \mathcal{U}^\sigma(P_{h,s,a}^\theta) := \left\{ P_{h,s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{h,s,a} \parallel P_{h,s,a}^\theta) \leq \sigma \right\} \quad (129)$$

with the uncertainty level σ satisfying

$$0 < \sigma \leq \frac{1}{20H}. \quad (130)$$

Before continuing, we shall introduce some notation for convenience. For any $P_h^{\theta_h}(\cdot | s, a)$ in (123), we define the limit of the perturbed kernel transiting to the next state s' from the current state-action pair (s, a) by

$$\underline{P}_h^{\theta_h}(s' | s, a) := \inf_{P_{h,s,a} \in \mathcal{U}^\sigma(P_{h,s,a}^{\theta_h})} P_h(s' | s, a), \quad (131)$$

and in particular, denote

$$\underline{p}_h := \underline{P}_h^{\theta_h}(0 | 0, \theta_h), \quad \underline{q}_h := \underline{P}_h^{\theta_h}(0 | 0, 1 - \theta_h). \quad (132)$$

Armed with the above definitions, we introduce the following lemma which implies some useful properties of the uncertainty set.

Lemma 18 *The perturbed transition kernels obey*

$$\underline{p}_1 = \underline{p}_2 = \cdots \underline{p}_H, \quad \underline{q}_1 = \underline{q}_2 = \cdots \underline{q}_H. \quad (133)$$

Denoting $\underline{p}_1 := p^*$ and $\underline{q}_1 := q^*$, we have that when the uncertainty level σ satisfies (130),

$$\underline{p}_\star \geq \underline{q}_\star \geq 1 - \frac{c_3}{H} \quad \text{and} \quad \underline{p}_\star - \underline{q}_\star \geq p - q \geq 0 \quad (134)$$

for constant $c_3 = 2c_1 = \frac{1}{4}$.

Proof The proof is postponed to Appendix B.3.5. ■

Value functions and optimal policies. We take a moment to derive the corresponding value functions and identify the optimal policies. With some abuse of notation, for any MDP $\mathcal{M}_\theta$, we denote $\pi^{\star, \theta} = \{\pi_h^{\star, \theta}\}_{h=1}^H$ as the optimal policy, and let $V_h^{\pi, \sigma, \theta}$ (resp. $V_h^{\star, \sigma, \theta}$) represent the robust value function of policy π (resp. $\pi^{\star, \theta}$) at step h with uncertainty radius σ . Armed with these notation, we introduce the following lemma which collects the properties concerning the value functions and optimal policies.

Lemma 19 *Consider any $\theta \in \Theta$ and any policy π . Then it holds that*

$$V_h^{\pi, \sigma, \theta}(0) = 1 + x_h^{\pi, \theta} V_{h+1}^{\pi, \sigma, \theta}(0) \quad (135)$$

for any $h \in [H]$, where

$$x_h^{\pi, \theta} = \underline{p}_\star \pi_h(\theta_h | 0) + \underline{q}_\star \pi_h(1 - \theta_h | 0). \quad (136)$$

In addition, for any $h \in [H]$ and $s \in \mathcal{S} \setminus \{0\}$, the optimal policies and the optimal value functions obey

$$\pi_h^{\star, \theta}(\theta_h | 0) = 1, \quad V_h^{\star, \sigma, \theta}(0) \geq \frac{2}{3}(H + 1 - h), \quad (137a)$$

$$\pi_h^{\star, \theta}(\theta_h | s) = 1, \quad V_h^{\star, \sigma, \theta}(s) = 0, \quad (137b)$$

provided that $0 < c_1 \leq 1/2$.

Proof See Appendix B.3.6. ■

Construction of the history/batch dataset. In the nominal environment $\mathcal{M}_\theta$, a batch dataset is generated consisting of K *independent* sample trajectories each of length H , where each trajectory is generated according to (10), based on the following initial state distribution ρ^b and behavior policy $\pi^b = \{\pi_h^b\}_{h=1}^H$:

$$\rho^b(s) = \mu(s) \quad \text{and} \quad \pi_h^b(a|s) = \frac{1}{2}, \quad \forall (s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]. \quad (138)$$

Here, $\mu(s)$ is defined as the following state distribution supported on the state subset $\{0, 1\}$:

$$\mu(s) = \frac{1}{CS} \mathbb{1}(s=0) + \left(1 - \frac{1}{CS}\right) \mathbb{1}(s=1), \quad (139)$$

where $\mathbb{1}(\cdot)$ is the indicator function, and $C > 0$ is some constant that will determine the concentrability coefficient C_{rob}^* (as we shall detail momentarily) and obeys

$$\frac{1}{CS} \leq \frac{1}{4}. \quad (140)$$

As it turns out, for any MDP $\mathcal{M}_\theta$, the occupancy distributions of the above batch dataset are the same (due to symmetry) and admit the following simple characterization:

$$d_1^{b, P^\theta}(0, a) = \frac{1}{2} \mu(0), \quad \forall a \in \mathcal{A}, \quad (141a)$$

$$\frac{\mu(s)}{2} \leq d_h^{b, P^\theta}(s) \leq 2\mu(s), \quad \frac{\mu(s)}{4} \leq d_h^{b, P^\theta}(s, a) \leq \mu(s), \quad \forall (s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]. \quad (141b)$$

In addition, we choose the following initial state distribution

$$\rho(s) = \begin{cases} 1, & \text{if } s = 0 \\ 0, & \text{if } s > 0 \end{cases}. \quad (142)$$

With this choice of ρ , the single-policy clipped concentrability coefficient C_{rob}^* and the quantity C are intimately connected as follows:

$$2C \leq C_{\text{rob}}^* \leq 4C. \quad (143)$$

The proof of the claim (141) and (143) are postponed to Appendix B.3.7.

B.3.2 ESTABLISHING THE MINIMAX LOWER BOUND: SMALL UNCERTAINTY LEVEL

Towards this, we first make the following claim: for an arbitrary policy π obeying

$$\sum_{h=1}^H \|\pi_h(\cdot|0) - \pi_h^{\star, \theta}(\cdot|0)\|_1 \geq \frac{H}{8}, \quad (144)$$

one has

$$\langle \rho, V_1^{\star, \sigma, \theta} - V_1^{\pi, \sigma, \theta} \rangle > \varepsilon. \quad (145)$$

We shall postpone the proof of this claim to Appendix B.3.8.

Armed with the above claim and following the same arguments in (Li et al., 2022, Section C.3.2), we complete the proof by observing: for some small enough constant c_4 , as long as the sample size is beneath

$$N = KH \leq \frac{c_4 C_{\text{rob}}^* S H^4}{\varepsilon^2}, \quad (146)$$

then we necessarily have

$$\inf_{\hat{\pi}} \max_{\theta \in \Theta} \mathbb{P}_{\theta} \left\{ V_{\theta}^{\star, \sigma}(\rho) - V_{\theta}^{\hat{\pi}, \sigma}(\rho) \geq \varepsilon \right\} \geq \frac{1}{4}, \quad (147)$$

where $\mathbb{P}_{\theta}$ denote the probability conditioned on that the MDP is $\mathcal{M}_{\theta}$. We omit the details for brevity and complete the proof; interested readers can referred to (Li et al., 2022, Section C.3.2).

B.3.3 CONSTRUCTION OF HARD PROBLEM INSTANCES: LARGE UNCERTAINTY LEVEL

We now move onto the case when the uncertainty level is relatively large, we construct another class of hard instances, which is almost the same as the previous one except for the transition kernel.

Construction of a collection of hard MDPs. Let us introduce two MDPs

$$\left\{ \mathcal{M}_{\phi} = (\mathcal{S}, \mathcal{A}, P^{\phi} = \{P_h^{\phi}\}_{h=1}^H, \{r_h\}_{h=1}^H, H) \mid \phi = \{0, 1\} \right\}, \quad (148)$$

where the state space is $\mathcal{S} = \{0, 1, \dots, S-1\}$, and the action space is $\mathcal{A} = \{0, 1\}$. The transition kernel P^{ϕ} of the constructed MDP $\mathcal{M}_{\phi}$ is defined as

$$P_1^{\phi}(s' \mid s, a) = \begin{cases} p\mathbb{1}(s' = 0) + (1-p)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, \phi) \\ q\mathbb{1}(s' = 0) + (1-q)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, 1-\phi) \\ \mathbb{1}(s' = 1) & \text{if } s = 1 \\ q\mathbb{1}(s' = s) + (1-q)\mathbb{1}(s' = 1) & \text{if } s > 1 \end{cases} \quad (149a)$$

and

$$P_h^{\phi}(s' \mid s, a) = \mathbb{1}(s' = s), \quad \forall (h, s, a) \in \{2, \dots, H\} \times \mathcal{S} \times \mathcal{A}. \quad (149b)$$

In words, except at step $h = 1$, the MDP always stays in the same state. Additionally, the MDP will always stay in the state subset $\{0, 1\}$ if the initial distribution is supported only on $\{0, 1\}$, in view of (149). Here, p and q are set to be

$$p = 1 - \alpha \quad \text{and} \quad q = 1 - \alpha - \Delta \quad (150)$$

for some $H \geq 2e^8$, α and Δ obeying

$$0 < \alpha \leq \frac{1}{H} \leq \frac{1}{2e^8} \quad \text{and} \quad \Delta \leq \frac{\alpha}{2} \leq \frac{1}{2H} \leq \frac{1}{4e^8}, \quad (151)$$

where β is set as

$$\beta := \frac{\log \frac{1}{\alpha + \Delta}}{2} \geq \frac{\log(2H/3)}{2} \geq 4. \quad (152)$$

The assumption (151) immediately indicates the facts

$$1 > p > q \geq \frac{1}{2}. \quad (153)$$

Moreover, for any $(h, s, a) \in [H] \times \mathcal{S} \times \mathcal{A}$, the reward function is defined as

$$r_h(s, a) = \begin{cases} 1 & \text{if } s = 0 \\ 0 & \text{otherwise} \end{cases}. \quad (154)$$

Construction of the history/batch dataset. We utilize the same batch dataset described in Appendix B.3.1 and choose the same initial state distribution ρ in (142). As a result, for any MDP $\mathcal{M}_\phi$, the occupancy distributions of the above batch dataset are the same (due to symmetry) and admit the following simple characterization:

$$d_1^{\mathbf{b}, P^\phi}(0, a) = \frac{1}{2}\mu(0), \quad \forall a \in \mathcal{A}, \quad (155a)$$

$$\frac{\mu(s)}{2} \leq d_h^{\mathbf{b}, P^\phi}(s) \leq 2\mu(s), \quad \frac{\mu(s)}{4} \leq d_h^{\mathbf{b}, P^\phi}(s, a) \leq \mu(s), \quad \forall (s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]. \quad (155b)$$

The proof of the claim (155) is postponed to Appendix B.3.9.

Uncertainty set of the transition kernels. Denote the transition kernel vector as

$$P_{h,s,a}^\phi := P_h^\phi(\cdot | s, a) \in [0, 1]^{1 \times \mathcal{S}}. \quad (156)$$

For any $(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H]$, the perturbation of the transition kernels in $\mathcal{M}_\phi$ is restricted to the following uncertainty set

$$\mathcal{U}^\sigma(P^\phi) := \otimes \mathcal{U}^\sigma(P_{h,s,a}^\phi), \quad \mathcal{U}^\sigma(P_{h,s,a}^\phi) := \left\{ P_{h,s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{h,s,a} \parallel P_{h,s,a}^\phi) \leq \sigma \right\}, \quad (157)$$

where the radius of the uncertainty set σ obeys:

$$\left(1 - \frac{3}{\beta}\right) \log\left(\frac{1}{\alpha + \Delta}\right) \leq \sigma \leq \left(1 - \frac{2}{\beta}\right) \log\left(\frac{1}{\alpha + \Delta}\right). \quad (158)$$

Before continuing, we shall introduce some notation for convenience. For any $P_h^\phi(\cdot | s, a)$ in (149), we define the limit of the perturbed kernel transiting to the next state s' from the current state-action pair (s, a) by

$$\underline{P}_h^\phi(s' | s, a) := \inf_{P_{h,s,a} \in \mathcal{U}^\sigma(P_{h,s,a}^\phi)} P_h(s' | s, a), \quad (159)$$

and in particular, denote

$$\underline{p} := \underline{P}_1^\phi(0 | 0, \phi), \quad \underline{q} := \underline{P}_1^\phi(0 | 0, 1 - \phi). \quad (160)$$

Armed with the above definitions, we introduce the following lemma which implies some useful properties of the uncertainty set.

Lemma 20 *When β satisfies (152) and the uncertainty level σ satisfies (158), the perturbed transition kernels obey*

$$\underline{p} \geq \underline{q} \geq \frac{1}{\beta}. \quad (161)$$

Proof See Appendix B.3.10. ■

Value functions and optimal policies. Similar to Appendix B.3.1, for any MDP $\mathcal{M}_\phi$, we denote $\pi^{*,\phi} = \{\pi_h^{*,\phi}\}_{h=1}^H$ as the optimal policy, and let $V_h^{\pi,\sigma,\phi}$ (resp. $V_h^{*,\sigma,\phi}$) represent the robust value function of policy π (resp. $\pi^{*,\phi}$) at step h with uncertainty radius σ . Then we introduce the following lemma which collects the properties concerning the value functions and optimal policies.

Lemma 21 *For any $\phi \in \{0, 1\}$ and any policy π , defining*

$$z_\phi^\pi := \underline{p}\pi_1(\phi|0) + \underline{q}\pi_1(1-\phi|0), \quad (162)$$

it holds that

$$V_1^{\pi,\sigma,\phi}(0) = 1 + z_\phi^\pi(H-1). \quad (163)$$

In addition, the optimal policies and the optimal value functions obey

$$V_1^{*,\sigma,\phi}(0) = 1 + \underline{p}(H-1), \quad (164a)$$

$$\forall h \in [H] \setminus \{1\}: \quad V_h^{*,\sigma,\phi}(0) = H - h + 1, \quad (164b)$$

$$\forall h \in [H]: \quad \pi_h^{*,\phi}(\phi|0) = 1, \quad \pi_h^{*,\phi}(\phi|1) = 1, \quad V_h^{*,\sigma,\phi}(1) = 0. \quad (164c)$$

The robust single-policy clipped concentrability coefficient C_{rob}^ obeys*

$$2C \leq C_{\text{rob}}^* \leq 4C. \quad (165)$$

Proof See Appendix B.3.11. ■

In view of Lemma 21, we note that the smallest positive state transition probability of the optimal policy π^* under any MDP $\mathcal{M}_\phi$ with $\phi \in \{0, 1\}$ thus can be given by

$$P_{\min}^* := \min_{h,s,s'} \left\{ P_h^\phi(s'|s, \pi_h^{*,\phi}(s)) : P_h^\phi(s'|s, \pi_h^{*,\phi}(s)) > 0 \right\} = P_1^\phi(1|0, 1-\phi) = 1-p, \quad (166)$$

which obeys

$$\alpha = P_{\min}^* \in (0, 1/H]$$

according to (150) and (151).

B.3.4 ESTABLISHING THE MINIMAX LOWER BOUND: LARGE UNCERTAINTY LEVEL

We are now ready to establish the sample complexity lower bound. With the choice of the initial distribution ρ in (142), for any policy estimator $\hat{\pi}$ computed based on the batch dataset, we plan to control the quantity

$$\langle \rho, V_1^{*,\sigma,\phi} - V_1^{\hat{\pi},\sigma,\phi} \rangle = V_1^{*,\sigma,\phi}(0) - V_1^{\hat{\pi},\sigma,\phi}(0).$$

Step 1: converting the goal to estimate ϕ . We make the following claim which shall be verified in Appendix B.3.12: given $\varepsilon \leq \frac{H}{384e^6 \log(\frac{1}{\alpha})} \leq \frac{H}{384e^6 \log(\frac{1}{\alpha+\Delta})}$, choosing

$$\Delta = \frac{128e^6 \sigma(1-q)\varepsilon}{H} = \frac{128e^6 \sigma(\alpha+\Delta)\varepsilon}{H} \leq \frac{128e^6(\alpha+\Delta)\varepsilon \log\left(\frac{1}{\alpha+\Delta}\right)}{H} \leq \frac{\alpha}{2}, \quad (167)$$

which satisfies (151) with the aid of (158) and (150), it holds that for any policy $\hat{\pi}$,

$$\langle \rho, V_1^{\star, \sigma, \phi} - V_1^{\hat{\pi}, \sigma, \phi} \rangle \geq 2\varepsilon(1 - \hat{\pi}_1(\phi | 0)). \quad (168)$$

Armed with this relation between the policy $\hat{\pi}$ and its sub-optimality gap, we are positioned to construct an estimate of ϕ . We denote $\mathbb{P}_\phi$ as the probability distribution when the MDP is $\mathcal{M}_\phi$, for any $\phi \in \{0, 1\}$.

Suppose for the moment that a policy estimate $\hat{\pi}$ achieves

$$\mathbb{P}_\phi \left\{ \langle \rho, V_1^{\star, \sigma, \phi} - V_1^{\hat{\pi}, \sigma, \phi} \rangle \leq \varepsilon \right\} \geq \frac{7}{8}, \quad (169)$$

then in view of (168), we necessarily have $\hat{\pi}_1(\phi | 0) \geq \frac{1}{2}$ with probability at least $\frac{7}{8}$. With this in mind, we are motivated to construct the following estimate $\hat{\phi}$ for $\phi \in \{0, 1\}$:

$$\hat{\phi} = \arg \max_{a \in \{0, 1\}} \hat{\pi}_1(a | 0), \quad (170)$$

which obeys

$$\mathbb{P}_\phi \{\hat{\phi} = \phi\} \geq \mathbb{P}_\phi \{\hat{\pi}_1(\phi | 0) > 1/2\} \geq \frac{7}{8}. \quad (171)$$

In what follows, we would like to show (171) cannot happen without enough samples, which would in turn contradict (168).

Step 2: probability of error in testing two hypotheses. Armed with the above preparation, we shall focus on differentiating the two hypotheses $\phi \in \{0, 1\}$. Towards this, consider the minimax probability of error defined as follows:

$$p_e := \inf_{\psi} \max \{ \mathbb{P}_0(\psi \neq 0), \mathbb{P}_1(\psi \neq 1) \}, \quad (172)$$

where the infimum is taken over all possible tests ψ constructed from the batch dataset.

Let $\mu^{\mathbf{b}, \phi}$ (resp. $\mu_h^{\mathbf{b}, \phi}(s_h)$) be the distribution of a sample trajectory $\{s_h, a_h\}_{h=1}^H$ (resp. a sample (a_h, s_{h+1}) conditional on s_h) for the MDP $\mathcal{M}_\phi$. Following standard results from Tsybakov and Zaiats (2009, Theorem 2.2) and the additivity of the KL divergence (cf. Tsybakov and Zaiats (2009, Page 85)), we obtain

$$\begin{aligned} p_e &\geq \frac{1}{4} \exp \left(-K \text{KL}(\mu^{\mathbf{b}, 0} \parallel \mu^{\mathbf{b}, 1}) \right) \\ &\geq \frac{1}{4} \exp \left\{ -\frac{1}{2} K \mu(0) \left(\text{KL}(P_1^0(\cdot | 0, 0) \parallel P_1^1(\cdot | 0, 0)) + \text{KL}(P_1^0(\cdot | 0, 1) \parallel P_1^1(\cdot | 0, 1)) \right) \right\}, \end{aligned} \quad (173)$$

where we also use the independence of the K trajectories in the batch dataset in the first line. Here, the second line arises from the chain rule of the KL divergence (Duchi, 2018, Lemma 5.2.8) and the Markov property of the sample trajectories (recall that $d_h^{\mathbf{b}, P^0} = d_h^{\mathbf{b}, P^1}$) according to

$$\begin{aligned} \text{KL}(\mu^{\mathbf{b}, 0} \parallel \mu^{\mathbf{b}, 1}) &= \sum_{h=1}^H \mathbb{E}_{s_h \sim d_h^{\mathbf{b}, P^0}} \left[\text{KL}(\mu_h^{\mathbf{b}, 0}(s_h) \parallel \mu_h^{\mathbf{b}, 1}(s_h)) \right] \\ &= \sum_{a \in \{0, 1\}} d_1^{\mathbf{b}, P^0}(0, a) \text{KL}(P_1^0(\cdot \mid 0, a) \parallel P_1^1(\cdot \mid 0, a)) \\ &= \frac{1}{2} \mu(0) \sum_{a \in \{0, 1\}} \text{KL}(P_1^0(\cdot \mid 0, a) \parallel P_1^1(\cdot \mid 0, a)), \end{aligned}$$

where the penultimate equality holds by the fact that $P_h^0(\cdot \mid s, a)$ and $P_h^1(\cdot \mid s, a)$ only differ when $h = 1$ and $s = 0$, and the last equality follows from (155).

It remains to control the KL divergence terms in (173). Given $p \geq q \geq 1/2$ (cf. (153)), applying Lemma 16 (cf. (73)) yields

$$\begin{aligned} \text{KL}(P_1^0(\cdot \mid 0, 0) \parallel P_1^1(\cdot \mid 0, 0)) &= \text{KL}(p \parallel q) \leq \frac{(p - q)^2}{(1 - p)p} \stackrel{(i)}{=} \frac{\Delta^2}{p(1 - p)} \\ &\stackrel{(ii)}{=} \frac{128^2 e^{12} \sigma^2 (1 - q)^2 \varepsilon^2}{H^2 p(1 - p)} \\ &\stackrel{(iii)}{\leq} \frac{c_1 \sigma^2 P_{\min}^* \varepsilon^2}{H^2}, \end{aligned} \tag{174}$$

where (i) follows from the definition (150), (ii) holds by plugging in the expression of Δ in (167), (iii) arises from $1 - q \leq 2(1 - p) = 2P_{\min}^*$ (see (151) and (166)), $p > \frac{1}{2}$, as long as c_1 is a large enough constant. It can be shown that $\text{KL}(P_1^0(\cdot \mid 0, 1) \parallel P_1^1(\cdot \mid 0, 1))$ can be upper bounded in the same way. Substituting (174) back into (173) demonstrates that: if the sample size is chosen as

$$KH \leq \frac{H^3 S C_{\text{rob}}^* \log 2}{4c_1 P_{\min}^* \sigma^2 \varepsilon^2}, \tag{175}$$

then one necessarily has

$$\begin{aligned} p_e &\geq \frac{1}{4} \exp \left\{ -\frac{1}{2} K \mu(0) \cdot 2 \frac{c_1 \sigma^2 P_{\min}^* \varepsilon^2}{H^2} \right\} \stackrel{(i)}{=} \frac{1}{4} \exp \left\{ -K \frac{c_1 \sigma^2 P_{\min}^* \varepsilon^2}{S C H^2} \right\} \\ &\stackrel{(ii)}{\geq} \frac{1}{4} \exp \left\{ -K \frac{4c_1 \sigma^2 P_{\min}^* \varepsilon^2}{S C_{\text{rob}}^* H^2} \right\} \geq \frac{1}{8}, \end{aligned} \tag{176}$$

where (i) follows from (139) and (ii) holds by (165).

Step 3: putting things together. Finally, suppose that there exists an estimator $\hat{\pi}$ such that

$$\mathbb{P}_0 \{ \langle \rho, V_1^{\star, \sigma, 0} - V_1^{\hat{\pi}, \sigma, 0} \rangle > \varepsilon \} < \frac{1}{8} \quad \text{and} \quad \mathbb{P}_1 \{ \langle \rho, V_1^{\star, \sigma, 1} - V_1^{\hat{\pi}, \sigma, 1} \rangle > \varepsilon \} < \frac{1}{8}.$$

Then Step 1 tells us that the estimator $\hat{\phi}$ defined in (170) must satisfy

$$\mathbb{P}_0(\hat{\phi} \neq 0) < \frac{1}{8} \quad \text{and} \quad \mathbb{P}_1(\hat{\phi} \neq 1) < \frac{1}{8},$$

which cannot happen under the sample size condition in (175) to avoid contradiction with (176). The proof is thus finished.

B.3.5 PROOF OF LEMMA 18

First, (133) can be easily verified by the definition of $\underline{p}_h$ and $\underline{q}_h$ in (132) and the transition $P_h^{\theta_h}$ in (123).

Proof of the first inequality in (134). It is observed that

$$\begin{aligned} & \text{KL} \left(\text{Ber} \left(1 - \frac{2c_1}{H} \right) \parallel \text{Ber}(q) \right) \\ &= \left(1 - \frac{2c_1}{H} \right) \log \left(\frac{1 - \frac{2c_1}{H}}{q} \right) + \frac{2c_1}{H} \log \left(\frac{\frac{2c_1}{H}}{1 - q} \right) \\ &= \left(1 - \frac{2c_1}{H} \right) \log \left(1 + \frac{1 - q - \frac{2c_1}{H}}{q} \right) + \frac{2c_1}{H} \log \left(1 + \frac{q - 1 + \frac{2c_1}{H}}{1 - q} \right) \\ &\stackrel{(i)}{\geq} \left(1 - \frac{2c_1}{H} \right) \frac{\frac{1 - q - \frac{2c_1}{H}}{q}}{2q} \frac{1 + q - \frac{2c_1}{H}}{1 - \frac{2c_1}{H}} + \frac{2c_1}{H} \cdot 2 \frac{q - 1 + \frac{2c_1}{H}}{1 - q + \frac{2c_1}{H}} \\ &= \left(q - 1 + \frac{2c_1}{H} \right) \left[-\frac{1 + q - \frac{2c_1}{H}}{2q^2} + \frac{4c_1}{H} \frac{1}{1 - q + \frac{2c_1}{H}} \right] \\ &\stackrel{(ii)}{\geq} \frac{c_1}{2H} \left[\frac{-1}{(1 - \frac{3c_1}{2H})^2} + \frac{8}{7} \right] \geq \frac{1}{20H} \geq \sigma, \end{aligned} \tag{177}$$

where (i) holds by $\log(1 + x) \geq \frac{x}{2(1+x)}$ when $0 \leq x < \infty$ and $\log(1 + x) \geq \frac{x}{2} \frac{2+x}{1+x}$ when $-1 < x \leq 0$ (Topsøe, 2007), and the penultimate inequality holds by $1 - \frac{3c_1}{2H} \geq 1 - \frac{3}{256} \geq \frac{63}{64}$. Here, (ii) can be verified by

$$\begin{aligned} & -\frac{1 + q - \frac{2c_1}{H}}{2q^2} \geq -\frac{1 + q}{2q^2} \geq \frac{-2}{2q^2} \stackrel{(iii)}{\geq} \frac{-1}{(1 - \frac{3c_1}{2H})^2}, \\ & \frac{4c_1}{H} \frac{1}{1 - q + \frac{2c_1}{H}} \stackrel{(iv)}{\geq} \frac{4c_1}{H} \frac{1}{\frac{3c_1}{2H} + \frac{2c_1}{H}} \geq \frac{8}{7}, \end{aligned}$$

where (iii) holds by $q \geq 1 - \frac{3c_1}{2H}$, and (iv) arises from $0 \leq 1 - q \leq \frac{3c_1}{2H}$.

With above fact and $\text{KL} \left(\text{Ber} \left(\underline{q}_\star \right) \parallel \text{Ber}(q) \right) = \sigma$ in mind, applying Lemma 16 leads to $\underline{q}_\star \geq 1 - \frac{2c_1}{H}$.

Proof of the second inequality in (134). First, observing the first claim in (134), combined with Lemma 16, we know that for any uncertainty level $\sigma \leq \frac{1}{20H}$, there exists a unique $\underline{q}_\star$ obeying $\underline{q}_\star \geq 1 - \frac{2c_1}{H} > \frac{1}{2}$ such that

$$\sigma = \text{KL} \left(\text{Ber}(\underline{q}_\star) \parallel \text{Ber}(q) \right). \quad (178)$$

Then let us define the following function for $0 < x \leq \frac{1-q}{3}$ (i.e., $p = q + x$)

$$g(x, q) = \text{KL} \left(\text{Ber}(\underline{q}_\star + x) \parallel \text{Ber}(q + x) \right) = (\underline{q}_\star + x) \log \frac{\underline{q}_\star + x}{q + x} + (1 - \underline{q}_\star - x) \log \frac{1 - \underline{q}_\star - x}{1 - q - x} \quad (179)$$

The first derivative $\nabla_x g(x, q)$ is

$$\begin{aligned} & \nabla_x g(x, q) \\ &= \log \left(\frac{\underline{q}_\star + x}{q + x} \right) + (q + x) \frac{q + x - (\underline{q}_\star + x)}{(q + x)^2} - \log \left(\frac{1 - \underline{q}_\star - x}{1 - q - x} \right) \\ & \quad + (1 - q - x) \frac{-(1 - q - x) + (1 - \underline{q}_\star - x)}{(1 - q - x)^2} \\ &= \log \left(\frac{\underline{q}_\star + x}{q + x} \right) - \log \left(\frac{1 - \underline{q}_\star - x}{1 - q - x} \right) + \frac{q - \underline{q}_\star}{q + x} + \frac{q - \underline{q}_\star}{1 - q - x} \\ &= \log \left(1 + \frac{\underline{q}_\star - q}{q + x} \right) - \log \left(1 + \frac{q - \underline{q}_\star}{1 - q - x} \right) + \frac{q - \underline{q}_\star}{q + x} + \frac{q - \underline{q}_\star}{1 - q - x} \\ &\stackrel{(i)}{\geq} \frac{\underline{q}_\star - q}{2(q + x)} \frac{2 + \frac{\underline{q}_\star - q}{q + x}}{1 + \frac{\underline{q}_\star - q}{q + x}} - \frac{q - \underline{q}_\star}{2(1 - q - x)} \frac{2 + \frac{q - \underline{q}_\star}{1 - q - x}}{1 + \frac{q - \underline{q}_\star}{1 - q - x}} + \frac{q - \underline{q}_\star}{q + x} + \frac{q - \underline{q}_\star}{1 - q - x} \\ &= \frac{q - \underline{q}_\star}{2(q + x)} \left(2 - \frac{2q + 2x + \underline{q}_\star - q}{q + x + \underline{q}_\star - q} \right) + \frac{q - \underline{q}_\star}{2(1 - q - x)} \left(2 - \frac{2 - 2q - 2x + q - \underline{q}_\star}{1 - q - x + q - \underline{q}_\star} \right) \\ &= \frac{q - \underline{q}_\star}{2(q + x)} \frac{\underline{q}_\star - q}{x + \underline{q}_\star} + \frac{q - \underline{q}_\star}{2(1 - q - x)} \frac{q - \underline{q}_\star}{1 - \underline{q}_\star - x} \\ &= \frac{(q - \underline{q}_\star)^2 \left[(q + x)(\underline{q}_\star + x) - (1 - q - x)(1 - \underline{q}_\star - x) \right]}{2(q + x)(\underline{q}_\star + x)(1 - q - x)(1 - \underline{q}_\star - x)} \geq 0 \end{aligned} \quad (180)$$

where (i) holds by $\log(1 + x) \leq \frac{x}{2} \frac{2+x}{1+x}$ when $0 \leq x < \infty$ and $\log(1 + x) \geq \frac{x}{2} \frac{2+x}{1+x}$ when $-1 < x \leq 0$ (Topsøe, 2007), and the last inequality always holds for any $0 < x \leq \frac{1-q}{3}$ and $\underline{q}_\star \geq \frac{1}{2}$.

The above fact shows that $g(p - q, q) \geq g(0, q) = \sigma$, and thus

$$\text{KL} \left(\text{Ber}(p - q + \underline{q}_\star) \parallel \text{Ber}(p) \right) \geq \sigma, \quad (181)$$

which complete the proof by observing that $\underline{p}_\star \geq p - q + \underline{q}_\star$ via applying Lemma 16 with (178).

B.3.6 PROOF OF LEMMA 19

Ordering the value function for different states. First, note that for any policy π at the final step $H + 1$, we have

$$\forall s \in \mathcal{S} : \quad V_{H+1}^{\pi, \sigma, \theta}(s) = 0. \quad (182)$$

Then for any $\theta \in \Theta$ and any policy π , it is easily verified that

$$\begin{aligned} V_H^{\pi, \sigma, \theta}(0) &= \mathbb{E}_{a \sim \pi_h(\cdot | 0)} \left[r_h(0, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,0,a}^{\theta_h})} \mathcal{P} V_{H+1}^{\pi, \sigma, \theta} \right] = 1, \\ V_H^{\pi, \sigma, \theta}(s) &= \mathbb{E}_{a \sim \pi_h(\cdot | s)} \left[r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^{\theta_h})} \mathcal{P} V_{H+1}^{\pi, \sigma, \theta} \right] = 0, \quad \forall s \in \mathcal{S} \setminus \{0\}, \end{aligned} \quad (183)$$

which directly indicates that

$$\forall s \in \mathcal{S} \setminus \{0, 1\} : \quad V_H^{\pi, \sigma, \theta}(0) > V_H^{\pi, \sigma, \theta}(s) = 0. \quad (184)$$

Then we provide the following claim which will be proved momentarily using induction: For any $\theta \in \Theta$ and any policy π , the following equation holds

$$\forall (h, s) \in [H] \times \mathcal{S} \setminus \{0\} : \quad V_h^{\pi, \sigma, \theta}(0) > V_h^{\pi, \sigma, \theta}(s) = 0. \quad (185)$$

The above result leads to the following immediate fact for state $s \in \mathcal{S} \setminus \{0\}$:

$$\forall (h, s) \in [H] \times \mathcal{S} \setminus \{0\} : \quad V_h^{\star, \sigma, \theta}(s) = \max_{\pi} V_h^{\pi, \sigma, \theta}(s) = 0, \quad (186)$$

since (185) holds for any π and $h \in [H]$. Therefore, for any state $s \in \mathcal{S} \setminus \{0\}$, without loss of generality, we choose the optimal policy obeying

$$\forall (h, s) \in [H] \times \mathcal{S} \setminus \{0\} : \quad \pi_h^{\star, \theta}(\theta_h | s) = 1. \quad (187)$$

Then the rest of the proof will focus on deriving the value function and optimal policy over state $s = 0$. To begin with, recalling the value function in (192)

$$V_h^{\pi, \sigma, \theta}(0) = 1 + x_h^{\pi, \theta} V_{h+1}^{\pi, \sigma, \theta}(0) \quad (188)$$

and observing that the function $V_h^{\pi, \sigma, \theta}(0)$ is increasing in $x_h^{\pi, \theta}$ and that $x_h^{\pi, \theta}$ is increasing in $\pi_h(\theta_h | 0)$ (due to the fact $\underline{p}_\star \geq \underline{q}_\star$ in (134)). As a result, the optimal policy obeys

$$\pi_h^{\star, \theta}(\theta_h | 0) = 1 \quad (189)$$

at state 0. Plugging it back to (188) gives

$$\begin{aligned} V_h^{\star, \sigma, \theta}(0) &= 1 + x_h^{\pi^{\star, \theta}, \sigma, \theta} V_{h+1}^{\star, \sigma, \theta}(0) \\ &\stackrel{(i)}{=} 1 + \underline{p}_\star V_{h+1}^{\star, \sigma, \theta}(0) \geq \sum_{j=0}^{H-h} \underline{p}_\star^j \geq \sum_{j=0}^{H-h} \left(1 - \frac{c_3}{H}\right)^j = \frac{1 - \left(1 - \frac{c_3}{H}\right)^{H-h+1}}{c_3/H} \end{aligned}$$

$$\geq \frac{2}{3}(H+1-h). \quad (190)$$

where (i) holds by $x_h^{\pi^*,\theta} = \underline{p}_* \pi_h^{\pi^*,\theta}(\theta_h | 0) + \underline{q}_* \pi_h^{\pi^*,\theta}(1 - \theta_h | 0) = \underline{p}_*$. Here, the last inequality holds since we observe that

$$\left(1 - \frac{c_3}{H}\right)^{H-h+1} \leq \exp\left(-\frac{c_3}{H}(H-h+1)\right) \leq 1 - \frac{2c_3(H-h+1)}{3H},$$

as long as $c_3 \leq 0.5$, which follows due to the elementary inequalities $1 - x \leq \exp(-x)$ for any $x \geq 0$ and $\exp(-x) \leq 1 - 2x/3$ for any $0 \leq x \leq 1/2$.

Proof of claim in (185). We shall (185) through induction. Towards this, assuming that at time step $h+1$, the following holds

$$\forall s \in \mathcal{S} \setminus \{0, 1\} : \quad V_{h+1}^{\pi,\sigma,\theta}(0) > V_{h+1}^{\pi,\sigma,\theta}(s) = 0. \quad (191)$$

Observing that the base case when $h = H$ has already been confirmed in (184), now we move on to prove the same property for time step h .

To start with, the robust value function of state 0 at step h satisfies

$$\begin{aligned} V_h^{\pi,\sigma,\theta}(0) &= \mathbb{E}_{a \sim \pi_h(\cdot | 0)} \left[r_h(0, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,0,a}^{\theta_h})} \mathcal{P} V_{h+1}^{\pi,\sigma,\theta} \right] \\ &\stackrel{(i)}{=} 1 + \pi_h(\theta_h | 0) \left(\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,0,\theta}^{\theta_h})} \mathcal{P} V_{h+1}^{\pi,\sigma,\theta} \right) + \pi_h(1 - \theta_h | 0) \left(\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,0,1-\theta_h}^{\theta_h})} \mathcal{P} V_{h+1}^{\pi,\sigma,\theta} \right) \\ &\stackrel{(ii)}{=} 1 + \pi_h(\theta_h | 0) \left[\underline{p} V_{h+1}^{\pi,\sigma,\theta}(0) + (1 - \underline{p}) V_{h+1}^{\pi,\sigma,\theta}(1) \right] \\ &\quad + \pi_h(1 - \theta_h | 0) \left[\underline{q} V_{h+1}^{\pi,\sigma,\theta}(0) + (1 - \underline{q}) V_{h+1}^{\pi,\sigma,\theta}(1) \right] \\ &\stackrel{(iii)}{=} 1 + V_{h+1}^{\pi,\sigma,\theta}(1) + x_h^{\pi,\theta} \left[V_{h+1}^{\pi,\sigma,\theta}(0) - V_{h+1}^{\pi,\sigma,\theta}(1) \right] \\ &= 1 + x_h^{\pi,\theta} V_{h+1}^{\pi,\sigma,\theta}(0) \end{aligned} \quad (192)$$

where (i) uses the definition of the reward function in (127), (ii) uses the induction assumption in (191) so that the infimum is attained by picking the choice specified in (132) with a smallest probability mass imposed on the transition to state 0. Finally, we plug in the definition (136) of $x_h^{\pi,\theta}$ in (iii), and the last line follows from (191).

$$V_h^{\pi,\sigma,\theta}(s) = \mathbb{E}_{a \sim \pi_h(\cdot | s)} \left[r_h(s, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,s,a}^{\theta_h})} \mathcal{P} V_{h+1}^{\pi,\sigma,\theta} \right] = 0,$$

where the last inequality holds by the reward and transition function in (127) and (123) with the induction assumption (191).

Combining (192) and (193), we complete the proof:

$$\forall s \in \mathcal{S} \setminus \{0\} : \quad V_h^{\pi,\sigma,\theta_h}(0) \geq 1 > V_h^{\pi,\sigma,\theta}(s) = 0. \quad (193)$$

B.3.7 PROOF OF CLAIM (141) AND (143)

Proof of the claim (141). With the initial state distribution and behavior policy defined in (138), we have for any MDP $\mathcal{M}_\theta$,

$$d_1^{\mathbf{b}, P^\theta}(s) = \rho^{\mathbf{b}}(s) = \mu(s),$$

which leads to

$$\forall a \in \mathcal{A}: \quad d_1^{\mathbf{b}, P^\theta}(0, a) = \frac{1}{2}\mu(0). \quad (194)$$

In view of (149a), the state occupancy distribution at any step $h = 2, 3, \dots, H$ obeys

$$\begin{aligned} d_h^{\mathbf{b}, P^\theta}(0) &\geq \mathbb{P}\left\{s_h = 0 \mid s_{h-1} = 0; \pi^{\mathbf{b}}\right\} \geq d_{h-1}^{\mathbf{b}, P^\theta}(0) \left[\pi_{h-1}^{\mathbf{b}}(\theta_{h-1} \mid 0) \underline{p}_\star + \pi_{h-1}^{\mathbf{b}}(1 - \theta_{h-1} \mid 0) \underline{q}_\star\right] \\ &\geq d_{h-1}^{\mathbf{b}, P^\theta}(0) \underline{q}_\star \geq \dots \geq d_1^{\mathbf{b}, P^\theta}(0) \prod_{j=0}^{h-1} \underline{q}_\star \geq d_1^{\mathbf{b}, P^\theta}(0) \left(1 - \frac{c_3}{H}\right)^H > \frac{\mu(0)}{2}, \end{aligned} \quad (195)$$

where the last line makes use of the properties $\underline{q}_\star \geq 1 - c_3/H$ in Lemma 18 and

$$\left(1 - \frac{c_3}{H}\right)^H \geq \left(1 - \frac{1}{2H}\right)^H > \frac{1}{2}$$

provided that $0 < c_3 = 1/4 < 1/2$. In addition, as state 1 is an absorbing state and state 0 will only transfer to itself or state 1 at each time step, we directly achieve that

$$d_h^{\mathbf{b}, P^\theta}(0) \leq d_{h-1}^{\mathbf{b}, P^\theta}(0) \leq \dots \leq d_1^{\mathbf{b}, P^\theta}(0) \leq \mu(0). \quad (196)$$

For state 1, as it is absorbing, we directly have

$$d_h^{\mathbf{b}, P^\theta}(1) = \mathbb{P}\left\{s_h = 1 \mid s_{h-1} = 1; \pi^{\mathbf{b}}\right\} \geq d_{h-1}^{\mathbf{b}, P^\theta}(1) \geq \dots \geq d_1^{\mathbf{b}, P^\theta}(1) = \mu(1). \quad (197)$$

According to the assumption in (140), it is easily verified that

$$d_h^{\mathbf{b}, P^\theta}(1) \leq 1 \leq 2\mu(1). \quad (198)$$

Finally, combining (195), (196), (197), (198), the definitions of $P_h^{\theta_h}(\cdot \mid s, a)$ in (123) and the Markov property, we arrive at for any $(h, s) \in [H] \times \mathcal{S}$,

$$\frac{\mu(s)}{2} \leq d_h^{\mathbf{b}, P^\theta}(s) \leq 2\mu(s), \quad (199)$$

which directly leads to

$$\frac{\mu(s)}{4} \leq d_h^{\mathbf{b}, P^\theta}(s, a) = \pi_1^{\mathbf{b}}(a \mid s) d_h^{\mathbf{b}, P^\theta}(s) \leq \mu(s). \quad (200)$$

Proof of the claim (143). Examining the definition of C_{rob}^* in (14), we make the following observations.

- For $h = 1$, we have

$$\begin{aligned} \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\theta)} \frac{\min \{d_1^{*,P}(s,a), \frac{1}{S}\}}{d_1^{\text{b},P^\theta}(s,a)} &\stackrel{(i)}{=} \max_{P \in \mathcal{U}^\sigma(P^\theta)} \frac{\min \{d_1^{*,P}(0,\theta_h), \frac{1}{S}\}}{d_1^{\text{b},P^\theta}(0,\theta_h)} \\ &\stackrel{(ii)}{=} \max_{P \in \mathcal{U}^\sigma(P^\theta)} \frac{1}{S d_1^{\text{b},P^\theta}(0,\theta_h)} \stackrel{(iii)}{=} \frac{2}{S\mu(0)} = 2C, \end{aligned} \quad (201)$$

where (i) holds by $d_1^{*,P}(s) = \rho(s) = 0$ for all $s \in \mathcal{S} \setminus \{0\}$ (see (142)) and $\pi_h^{*,\theta}(\theta_h | 0) = 1$ for all $h \in [H]$, (ii) follows from the fact $d_1^{*,P}(0,\theta) = 1$, (iii) is verified in (141), and the last equality arises from the definition in (139).

- Similarly, for $h = 2, 3, \dots, H$, we arrive at

$$\begin{aligned} \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\theta)} \frac{\min \{d_h^{*,P}(s,a), \frac{1}{S}\}}{d_h^{\text{b},P^\theta}(s,a)} &\stackrel{(i)}{=} \max_{s \in \{0,1\}, P \in \mathcal{U}^\sigma(P^\theta)} \frac{\min \{d_h^{*,P}(s,\theta_h), \frac{1}{S}\}}{d_h^{\text{b},P^\theta}(s,\theta_h)} \\ &\leq \max_{s \in \{0,1\}, P \in \mathcal{U}^\sigma(P^\theta)} \frac{1}{S d_h^{\text{b},P^\theta}(s,\theta_h)} \stackrel{(ii)}{\leq} \frac{4}{S\mu(0)} = 4C, \end{aligned} \quad (202)$$

where (i) holds by the optimal policy in (137) and the trivial fact that $d_h^{*,P}(s) = 0$ for all $s \in \mathcal{S} \setminus \{0,1\}$ (see (142) and (123)), (ii) arises from (141), and the last equality comes from (139).

Combining the above cases, we complete the proof by

$$2C \leq C_{\text{rob}}^* = \max_{(h,s,a,P) \in [H] \times \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\theta)} \frac{\min \{d_h^{*,P}(s,a), \frac{1}{S}\}}{d_h^{\text{b},P^\theta}(s,a)} \leq 4C.$$

B.3.8 PROOF OF THE CLAIM (145)

By virtue of (136) and (137), we see that $x_h^{\pi^{*,\theta},\theta} = \underline{p}_*$ for all $h \in [H]$, which combined with (135) gives

$$\begin{aligned} \langle \rho, V_h^{*,\sigma,\theta} - V_h^{\pi,\sigma,\theta} \rangle &= V_h^{*,\theta}(0) - V_h^{\pi,\theta}(0) \\ &= \underline{p}_* V_{h+1}^{*,\sigma,\theta}(0) - x_h^{\pi,\theta} V_{h+1}^{\pi,\theta}(0) \\ &= x_h^{\pi,\theta} \left(V_{h+1}^{*,\sigma,\theta}(0) - V_{h+1}^{\pi,\sigma,\theta}(0) \right) + (\underline{p}_* - x_h^{\pi,\theta}) V_{h+1}^{*,\sigma,\theta}(0) \\ &\stackrel{(i)}{\geq} \underline{q}_* \left(V_{h+1}^{*,\sigma,\theta}(0) - V_{h+1}^{\pi,\sigma,\theta}(0) \right) + (\underline{p}_* - x_h^{\pi,\theta}) V_{h+1}^{*,\sigma,\theta}(0) \\ &\stackrel{(ii)}{\geq} \underline{q}_* \left(V_{h+1}^{*,\sigma,\theta}(0) - V_{h+1}^{\pi,\sigma,\theta}(0) \right) + \frac{1}{2}(p-q) \|\pi_h^{*,\theta}(\cdot | 0) - \pi_h(\cdot | 0)\|_1 V_{h+1}^{*,\sigma,\theta}(0) \end{aligned}$$

$$\stackrel{(iii)}{\geq} \underline{q}_\star \left(V_{h+1}^{\star, \sigma, \theta}(0) - V_{h+1}^{\pi, \sigma, \theta}(0) \right) + \frac{c_2 \varepsilon}{3H^2} (H+1-h) \left\| \pi_h^{\star, \theta}(\cdot | 0) - \pi_h(\cdot | 0) \right\|_1 \quad (203)$$

where (i) follows from the fact that $x_h^\pi \geq \underline{q}_\star$ for any π and $h \in [H]$, and (iii) holds by the facts (137) and the choice (124) of (p, q) . Here, (ii) arises from

$$\begin{aligned} \underline{p}_\star - x_h^{\pi, \theta} &= (\underline{p}_\star - \underline{q}_\star)(1 - \pi_h(\theta_h | 0)) \\ &\geq (p - q)(1 - \pi_h(\theta_h | 0)) \\ &= \frac{1}{2}(p - q)(1 - \pi_h(\theta_h | 0) + \pi_h(1 - \theta_h | 0)) = \frac{1}{2}(p - q) \left\| \pi_h^{\star, \theta}(\cdot | 0) - \pi_h(\cdot | 0) \right\|_1, \end{aligned} \quad (204)$$

where the first inequality holds by applying Lemma 18. With the fact of (203) in mind, combined with the fact $\underline{q}_\star \geq 1 - \frac{c_3}{H}$, following the same proof pipeline of (Li et al., 2022, (276) to (278)) leads to

$$\langle \rho, V_h^{\star, \sigma, \theta} - V_h^{\pi, \sigma, \theta} \rangle > \varepsilon. \quad (205)$$

We omit the proof here for conciseness.

B.3.9 PROOF OF (155)

With the initial state distribution and behavior policy defined in (138), we have for any MDP $\mathcal{M}_\phi$ with $\phi \in \{0, 1\}$,

$$d_1^{\mathbf{b}, P^\phi}(s) = \rho^{\mathbf{b}}(s) = \mu(s),$$

which leads to

$$\forall a \in \mathcal{A}: \quad d_1^{\mathbf{b}, P^\phi}(0, a) = \frac{1}{2}\mu(0). \quad (206)$$

In view of (149a), the state occupancy distribution at step $h = 2$ obeys

$$d_2^{\mathbf{b}, P^\phi}(0) = \mathbb{P} \left\{ s_2 = 0 \mid s_1 \sim d_1^{\mathbf{b}, P^\phi}; \pi^{\mathbf{b}} \right\} = \mu(0) \left[\pi_1^{\mathbf{b}}(\phi | 0)p + \pi_1^{\mathbf{b}}(1 - \phi | 0)q \right] = \frac{(p + q)\mu(0)}{2},$$

and

$$\begin{aligned} d_2^{\mathbf{b}, P^\phi}(1) &= \mathbb{P} \left\{ s_2 = 1 \mid s_1 \sim d_1^{\mathbf{b}, P^\phi}; \pi^{\mathbf{b}} \right\} \\ &= \mu(0) \left[\pi_1^{\mathbf{b}}(\phi | 0)(1 - p) + \pi_1^{\mathbf{b}}(1 - \phi | 0)(1 - q) \right] + \mu(1) = \mu(1) + \frac{(2 - p - q)\mu(0)}{2}. \end{aligned}$$

With the above result in mind and recalling the assumption in (153), we arrive at

$$\frac{\mu(0)}{2} \leq d_2^{\mathbf{b}, P^\phi}(0) \leq \mu(0), \quad \mu(1) \leq d_2^{\mathbf{b}, P^\phi}(1) \stackrel{(i)}{\leq} 2\mu(1), \quad (207)$$

where (i) holds by applying (153) and (140) (which implies $\mu(0) \leq \mu(1)$ by the assumption in (140))

$$d_2^{\mathbf{b}, P^\phi}(1) = \mu(1) + \frac{(2-p-q)\mu(0)}{2} \leq \mu(1) + \mu(0) \leq 2\mu(1).$$

Finally, from the definitions of $P_h^\phi(\cdot | s, a)$ in (149b) and the Markov property, we arrive at for any $(h, s) \in [H] \times \mathcal{S}$,

$$\frac{\mu(s)}{2} \leq d_h^{\mathbf{b}, P^\phi}(s) \leq 2\mu(s), \quad (208)$$

which directly leads to

$$\frac{\mu(s)}{4} \leq d_h^{\mathbf{b}, P^\phi}(s, a) = \pi_1^{\mathbf{b}}(a | s) d_h^{\mathbf{b}, P^\phi}(s) \leq \mu(s). \quad (209)$$

B.3.10 PROOF OF LEMMA 20

Note that $\underline{p} \geq \underline{q}$ can be easily verified since $p > q$, which indicates that the first assertion is true. So we will focus on the second assertion in (161). Towards this, invoking the definition in (72), let σ' be the KL divergence from $\text{Ber}(\frac{1}{\beta})$ to $\text{Ber}(q)$, defined as follows

$$\begin{aligned} \sigma' &:= \text{KL} \left(\text{Ber} \left(\frac{1}{\beta} \right) \parallel \text{Ber}(q) \right) = \frac{1}{\beta} \log \frac{\frac{1}{\beta}}{q} + \left(1 - \frac{1}{\beta} \right) \log \frac{\left(1 - \frac{1}{\beta} \right)}{1-q} \\ &= \left(\frac{1}{\beta} \right) \log \left(\frac{1}{\beta} \right) - \left(\frac{1}{\beta} \right) \log(q) + \left(1 - \frac{1}{\beta} \right) \log \left(\frac{1}{\alpha + \Delta} \right) + \left(1 - \frac{1}{\beta} \right) \log \left(1 - \frac{1}{\beta} \right), \end{aligned} \quad (210)$$

where the second line uses the definition of q in (150). We claim that σ' satisfies the following relation with σ , which will be proven at the end of this proof:

$$0 < \sigma \leq \left(1 - \frac{2}{\beta} \right) \log \left(\frac{1}{\alpha + \Delta} \right) \leq \sigma' \leq \left(1 - \frac{1}{\beta} \right) \log \left(\frac{1}{\alpha + \Delta} \right). \quad (211)$$

Recalling the definition of the transition kernel in (149a)

$$P_1^\phi(0 | 0, 1 - \phi) = q, \quad P_1^\phi(1 | 0, 1 - \phi) = 1 - q, \quad P_1^\phi(s | 0, 1 - \phi) = 0, \quad \forall s \in \mathcal{S} \setminus \{0, 1\},$$

the uncertainty set of the transition kernel with radius σ is thus given as

$$\begin{aligned} \mathcal{U}^\sigma(P_{1,0,1-\phi}^\phi) &= \left\{ P_{1,0,1-\phi} \in \Delta(\mathcal{S}) : P(0 | 0, 1 - \phi) = q', P(1 | 0, 1 - \phi) = 1 - q', \right. \\ &\quad \left. \text{KL}(\text{Ber}(q') \parallel \text{Ber}(q)) \leq \sigma \right\}. \end{aligned} \quad (212)$$

Recalling the definition of $\underline{q}$ in (160), we can bound

$$\underline{q} = \inf_{P_{1,0,1-\phi} \in \mathcal{U}^\sigma(P_{1,0,1-\phi}^\phi)} P(0 | 0, 1 - \phi) = \inf_{q' : \text{KL}(\text{Ber}(q') \parallel \text{Ber}(q)) \leq \sigma} q'$$

$$\stackrel{(i)}{\geq} \inf_{q': \text{KL}(\text{Ber}(q') \| \text{Ber}(q)) \leq \sigma'} q' = \frac{1}{\beta},$$

where (i) holds by $\sigma \leq \sigma'$ (cf. (211)) and the last equality follows from applying Lemma 16 (cf. (74)) and (210) to arrive at

$$\forall 0 \leq q' < \frac{1}{\beta} : \quad \text{KL}(\text{Ber}(q') \| \text{Ber}(q)) > \text{KL}\left(\text{Ber}\left(\frac{1}{\beta}\right) \| \text{Ber}(q)\right) = \sigma'.$$

Proof of (211). To control σ' , we plug in the assumptions in (153) and $\beta \geq 4$ and arrive at the trivial facts

$$\left(\frac{1}{\beta}\right) \log\left(\frac{1}{\beta}\right) - \left(\frac{1}{\beta}\right) \log(q) < 0, \quad \left(1 - \frac{1}{\beta}\right) \log\left(1 - \frac{1}{\beta}\right) < 0.$$

The above facts directly lead to

$$\sigma' \leq \left(1 - \frac{1}{\beta}\right) \log\left(\frac{1}{\alpha + \Delta}\right). \quad (213)$$

Similarly, observing

$$-1 \leq \left(\frac{1}{\beta}\right) \log\left(\frac{1}{\beta}\right) + \left(1 - \frac{1}{\beta}\right) \log\left(1 - \frac{1}{\beta}\right) \leq 0, \quad -\left(\frac{1}{\beta}\right) \log(q) \geq 0,$$

we arrive at

$$\sigma' \geq -1 + \left(1 - \frac{1}{\beta}\right) \log\left(\frac{1}{\alpha + \Delta}\right) \geq \left(1 - \frac{2}{\beta}\right) \log\left(\frac{1}{\alpha + \Delta}\right) \quad (214)$$

as long as $\log\left(\frac{1}{\alpha + \Delta}\right) \geq \beta$ (cf. (152)). With (213) and (214) in hand, it is straightforward to see that the choice of the uncertainty radius σ in (158) obeys the advertised bound (211).

B.3.11 PROOF OF LEMMA 21

For notational conciseness, we shall drop the superscript ϕ and use the shorthand $V_h^{\pi, \sigma} = V_h^{\pi, \sigma, \phi}$ and $V_h^{\star, \sigma} = V_h^{\star, \sigma, \phi}$ whenever it is clear from the context. We begin by deriving the robust value function for any policy π . Starting with state 1, at any step $h \in [H]$, it obeys

$$V_h^{\pi, \sigma}(1) = \mathbb{E}_{a \sim \pi_h(\cdot | 1)} \left[r_h(1, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,1,a}^\phi)} \mathcal{P} V_{h+1}^{\pi, \sigma} \right] = 0 + V_{h+1}^{\pi, \sigma}(1),$$

where the first equality follows from the robust Bellman consistency equation (cf. (8)), and the second equality follows from the observation that the distribution $P_{h,1,a}^\phi$ is supported solely on state 1 in view of (149a), therefore $\mathcal{U}^\sigma(P_{h,1,a}^\phi) = P_{h,1,a}^\phi$. Leveraging the terminal condition $V_{H+1}^{\pi, \sigma}(1) = 0$, and recursively applying the previous relation, we have

$$V_h^{\star, \sigma}(1) = V_h^{\pi, \sigma}(1) = 0, \quad \forall h \in [H]. \quad (215)$$

Similarly, turning to state 0, at any step $h > 1$, the robust value function satisfies

$$V_h^{\pi,\sigma}(0) = \mathbb{E}_{a \sim \pi_h(\cdot|0)} \left[r_h(0, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{h,0,a}^\phi)} \mathcal{P} V_{h+1}^{\pi,\sigma} \right] = 1 + V_{h+1}^{\pi,\sigma}(0),$$

which again uses the fact that the distribution $P_{h,0,a}^\phi$ is supported solely on state 0 in view of (149b), therefore $\mathcal{U}^\sigma(P_{h,0,a}^\phi) = P_{h,0,a}^\phi$. Leveraging the terminal condition $V_{H+1}^{\pi,\sigma}(0) = 0$, and recursively applying the previous relation, we have

$$V_h^{\star,\sigma}(0) = V_h^{\pi,\sigma}(0) = H - h + 1, \quad 2 \leq h \leq H. \quad (216)$$

Taking (215) and (216) together, it follows that

$$\forall 2 \leq h \leq H : \quad V_h^{\pi,\sigma}(0) > V_h^{\pi,\sigma}(1). \quad (217)$$

Consequently, the robust value function of state 0 at step $h = 1$ satisfies

$$\begin{aligned} V_1^{\pi,\sigma}(0) &= \mathbb{E}_{a \sim \pi_1(\cdot|0)} \left[r_1(0, a) + \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,a}^\phi)} \mathcal{P} V_2^{\pi,\sigma} \right] \\ &\stackrel{(i)}{=} 1 + \pi_1(\phi|0) \left(\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,\phi}^\phi)} \mathcal{P} V_2^{\pi,\sigma} \right) + \pi_1(1 - \phi|0) \left(\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,1-\phi}^\phi)} \mathcal{P} V_2^{\pi,\sigma} \right) \\ &\stackrel{(ii)}{=} 1 + \pi_1(\phi|0) \left[\underline{p} V_2^{\pi,\sigma}(0) + (1 - \underline{p}) V_2^{\pi,\sigma}(1) \right] + \pi_1(1 - \phi|0) \left[\underline{q} V_2^{\pi,\sigma}(0) \right. \\ &\quad \left. + (1 - \underline{q}) V_2^{\pi,\sigma}(1) \right] \\ &\stackrel{(iii)}{=} 1 + V_2^{\pi,\sigma}(1) + z_\phi^\pi [V_2^{\pi,\sigma}(0) - V_2^{\pi,\sigma}(1)] \\ &= 1 + z_\phi^\pi V_2^{\pi,\sigma}(0) \end{aligned} \quad (218)$$

where (i) uses the definition of the reward function in (154), (ii) uses (217) so that the infimum is attained by picking the choice specified in (160) with a smallest probability mass imposed on the transition to state 0. Finally, we plug in the definition (162) of z_ϕ^π in (iii), and the last line follows from (215).

Therefore, taking $\pi = \pi^{\star,\phi}$ in the previous relation directly leads to

$$V_1^{\star,\sigma}(0) = 1 + z_\phi^{\pi^{\star,\phi}} V_2^{\star,\sigma}(0) = 1 + z_\phi^{\pi^{\star,\phi}} (H - 1), \quad (219)$$

where the second equality follows from (216). Observing that the function $(H - 1)z$ is increasing in z and that z_ϕ^π is increasing in $\pi_1(\phi|0)$ (due to the fact $\underline{p} \geq \underline{q}$ in (161)). As a result, the optimal policy obeys

$$\pi_1^{\star,\phi}(\phi|0) = 1 \quad (220)$$

at state 0, and plugging back to (219) gives

$$V_1^{\star,\sigma}(0) = 1 + z_\phi^{\pi^{\star,\phi}} (H - 1) = 1 + \underline{p} (H - 1),$$

where $z_\phi^{\pi^{\star,\phi}} = p \pi_1^{\star,\phi}(\phi|0) + q \pi_1^{\star,\phi}(1 - \phi|0) = \underline{p}$. For the rest of the states, without loss of generality, we choose the optimal policy obeying

$$\forall h \in [H] : \quad \pi_h^{\star,\phi}(\phi|0) = 1, \quad \pi_h^{\star,\phi}(\phi|1) = 1. \quad (221)$$

Proof of claim (165). Given that $\pi_h^{\star,\phi}(\phi|0) = 1$ for all $h \in [H]$ and $\rho(0) = 1$, for any $P \in \mathcal{U}^\sigma(P^\phi)$, we have

$$\begin{aligned} d_2^{\star,P}(0, \phi) &= d_2^{\star,P}(0) \pi_2^{\star,\phi}(\phi|0) = d_2^{\star,P}(0) = \mathbb{P}_{s_2 \sim P(\cdot|s_1, \pi_1^{\star,\phi}(s_1))} \{s_2 = 0 | s_1 \sim \rho; \pi^{\star,\phi}\} \\ &= P_1(0|0, \phi) \stackrel{(i)}{\geq} \underline{P}_1^\phi(0|0, \phi) \stackrel{(ii)}{=} \underline{p} \geq \frac{1}{\beta}, \end{aligned} \quad (222)$$

which (i) holds by plugging in the definition (159), (ii) follows from the definition (160), and the final inequality arises from Lemma 20. Hence, for all $2 \leq h \leq H$, by the Markov property and $P_h^\phi(0|0, \phi) = 1$, we have

$$d_h^{\star,P}(0, \phi) = d_2^{\star,P}(0, \phi) \geq \frac{1}{\beta}. \quad (223)$$

Examining the definition of $C_{\text{rob}}^\star$ in (14), we make the following observations.

- For $h = 1$, we have

$$\begin{aligned} \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_1^{\star,P}(s, a), \frac{1}{S}\}}{d_1^{\text{b},P^\phi}(s, a)} &\stackrel{(i)}{=} \max_{P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_1^{\star,P}(0, \phi), \frac{1}{S}\}}{d_1^{\text{b},P^\phi}(0, \phi)} \\ &\stackrel{(ii)}{=} \max_{P \in \mathcal{U}^\sigma(P^\phi)} \frac{1}{S d_1^{\text{b},P^\phi}(0, \phi)} \stackrel{(iii)}{=} \frac{2}{S \mu(0)} = 2C, \end{aligned} \quad (224)$$

where (i) holds by $d_1^{\star,P}(s) = \rho(s) = 0$ for all $s \in \mathcal{S} \setminus \{0\}$ (see (142)) and $\pi_h^{\star,\phi}(\phi|0) = 1$ for all $h \in [H]$, (ii) follows from the fact $d_1^{\star,P}(0, \phi) = 1$, (iii) is verified in (155), and the last equality arises from the definition in (139).

- Similarly, for $h = 2$, we arrive at

$$\begin{aligned} \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_2^{\star,P}(s, a), \frac{1}{S}\}}{d_2^{\text{b},P^\phi}(s, a)} &\stackrel{(i)}{=} \max_{s \in \{0,1\}, P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_2^{\star,P}(s, \phi), \frac{1}{S}\}}{d_2^{\text{b},P^\phi}(s, \phi)} \\ &\leq \max_{s \in \{0,1\}, P \in \mathcal{U}^\sigma(P^\phi)} \frac{1}{S d_2^{\text{b},P^\phi}(s, \phi)} \stackrel{(ii)}{\leq} \frac{4}{S \mu(0)} = 4C, \end{aligned} \quad (225)$$

where (i) holds by the optimal policy in (164) and the trivial fact that $d_2^{\star,P}(s) = 0$ for all $s \in \mathcal{S} \setminus \{0,1\}$ (see (142) and (149a)), (ii) arises from (155), and the last equality comes from (139).

- For all other steps $h = 3, \dots, H$, observing from the deterministic transition kernels in (149b), it can be easily verified that

$$\max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_h^{\star,P}(s, a), \frac{1}{S}\}}{d_h^{\text{b},P^\phi}(s, a)} = \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_h^{\star,P}(s, a), \frac{1}{S}\}}{d_h^{\text{b},P^\phi}(s, a)} \leq 4C. \quad (226)$$

Combining the above cases, we complete the proof by

$$2C \leq C_{\text{rob}}^* = \max_{(h,s,a,P) \in [H] \times \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d_h^{\star,P}(s,a), \frac{1}{S}\}}{d_h^{\text{b},P^\phi}(s,a)} \leq 4C.$$

B.3.12 PROOF OF THE CLAIM (168)

Recall that by virtue of (162) and (164), we arrive at

$$z_\phi^* := z_\phi^{\pi^*,\phi} = \underline{p}\pi_1^{\star,\phi}(\phi|0) + \underline{q}\pi_1^{\star,\phi}(1-\phi|0) = \underline{p}.$$

Applying (163) yields

$$\langle \rho, V_1^{\star,\sigma,\phi} - V_1^{\pi,\sigma,\phi} \rangle = V_h^{\star,\sigma,\phi}(0) - V_h^{\pi,\sigma,\phi}(0) = (\underline{p} - z_\phi^\pi)(H-1) = (\underline{p} - \underline{q})(H-1)(1 - \pi_1(\phi|0)), \quad (227)$$

where the last equality uses the definition (162). Therefore, it boils down to control $\underline{p} - \underline{q}$.

To continue, we define an auxiliary value function vector $\bar{V} \in \mathbb{R}^{S \times 1}$ obeying

$$\bar{V}(0) = H - 1 \quad \text{and} \quad \bar{V}(s) = 0, \quad \forall s \in \mathcal{S} \setminus \{0\}. \quad (228)$$

With this in hand, applying Lemma 11 gives

$$\begin{aligned} & (H-1)(\underline{p} - \underline{q}) \\ & \stackrel{(i)}{=} \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,\phi}^\phi)} \mathcal{P}\bar{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,1-\phi}^\phi)} \mathcal{P}\bar{V} \\ & = \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{1,0,\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda} \right) \right) - \lambda \sigma \right\} - \sup_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{1,0,1-\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda} \right) \right) - \lambda \sigma \right\} \\ & \stackrel{(ii)}{\geq} \left\{ -\lambda^* \log \left(P_{1,0,\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda^*} \right) \right) - \lambda^* \sigma \right\} - \left\{ -\lambda^* \log \left(P_{1,0,1-\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda^*} \right) \right) - \lambda^* \sigma \right\} \\ & = -\lambda^* \left[\log \left(P_{1,0,\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda^*} \right) \right) - \log \left(P_{1,0,1-\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda^*} \right) \right) \right], \quad (229) \end{aligned}$$

where (i) follows from (see the definition of $\underline{p}$ in (160))

$$\begin{aligned} \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,\phi}^\phi)} \mathcal{P}\bar{V} &= \underline{P}_1^\phi(0|0, \phi) \bar{V}(0) = (H-1)\underline{p}, \\ \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{1,0,1-\phi}^\phi)} \mathcal{P}\bar{V} &= \underline{P}_1^\phi(0|0, 1-\phi) \bar{V}(0) = (H-1)\underline{q}. \end{aligned}$$

Here, (ii) holds by letting

$$\lambda^* := \arg \max_{\lambda \geq 0} f(\lambda) := \arg \max_{\lambda \geq 0} \left\{ -\lambda \log \left(P_{1,0,1-\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda} \right) \right) - \lambda \sigma \right\}. \quad (230)$$

The rest of the proof is then to control (229). We start with the observation that $\lambda^* > 0$; this is because in view of Lemma 12 (cf. (66)), it suffices to verify that

$$\log(1-q) + \sigma \stackrel{(i)}{\leq} \log(\alpha + \Delta) + \left(1 - \frac{2}{\beta}\right) \log \left(\frac{1}{\alpha + \Delta} \right) = -\frac{2}{\beta} \log \left(\frac{1}{\alpha + \Delta} \right) < 0, \quad (231)$$

where (i) holds by (158). We now claim the following bound for λ^* holds, whose proof is postponed to the end:

$$\frac{H}{16\sigma} \leq \frac{H-1}{\log\left(\frac{\beta}{\alpha+\Delta}\right)} \leq \lambda^* \leq \frac{H-1}{\left(1-\frac{3}{\beta}\right)\log\left(\frac{1}{\alpha+\Delta}\right)}, \quad (232)$$

which immediately implies the following by taking exponential maps given $\lambda^* > 0$:

$$\frac{\alpha+\Delta}{\beta} \leq e^{-(H-1)/\lambda^*} \leq (\alpha+\Delta)^{1-3/\beta}. \quad (233)$$

Moving to the second term of (229), it follows that

$$\begin{aligned} & \log\left(P_{1,0,\phi}^\phi \cdot \exp\left(\frac{-\bar{V}}{\lambda^*}\right)\right) - \log\left(P_{1,0,1-\phi}^\phi \cdot \exp\left(\frac{-\bar{V}}{\lambda^*}\right)\right) \\ & \stackrel{(i)}{=} \log \frac{pe^{-(H-1)/\lambda^*} + (1-p)}{qe^{-(H-1)/\lambda^*} + (1-q)} \\ & = \log\left(1 + \frac{(p-q)(e^{-(H-1)/\lambda^*} - 1)}{qe^{-(H-1)/\lambda^*} + (1-q)}\right) \\ & \stackrel{(ii)}{<} -\frac{\Delta(1 - e^{-(H-1)/\lambda^*})}{qe^{-(H-1)/\lambda^*} + (1-q)} \\ & \stackrel{(iii)}{\leq} -\frac{1}{2} \frac{\Delta}{\left(\frac{1}{\alpha+\Delta}\right)^{\frac{3}{\beta}}(1-q) + (1-q)} \\ & \leq -\frac{\Delta}{4e^6(1-q)}, \end{aligned} \quad (234)$$

where (i) follows from the definitions in (149) and (228), (ii) holds by $\log(1+x) < x$ for $x \in (-1, \infty)$, (iii) can be verified by (233), $\beta \geq 4$, and (151):

$$1 - e^{-(H-1)/\lambda^*} \geq 1 - (\alpha+\Delta)^{1-3/\beta} \geq 1 - (\alpha+\Delta)^{1/4} \geq 1 - \left(\frac{3}{2H}\right)^{1/4} \geq \frac{1}{2},$$

and the last line uses $\left(\frac{1}{\alpha+\Delta}\right)^{\frac{3}{\beta}} = \left(\frac{1}{\alpha+\Delta}\right)^{6/\log(\frac{1}{\alpha+\Delta})} = e^6$ by the definition of β in (152). Plugging (232) and (234) back into (229) and (227), we arrive at

$$\begin{aligned} \langle \rho, V_1^{\star,\sigma,\phi} - V_1^{\pi,\sigma,\phi} \rangle &= (H-1)(\underline{p} - \underline{q})(1 - \pi_1(\phi|0)) \\ &\stackrel{(i)}{\geq} \frac{H\Delta}{64e^6\sigma(1-q)}(1 - \pi_1(\phi|0)) = 2\varepsilon(1 - \pi_1(\phi|0)), \end{aligned}$$

where (i) holds by (232) and the last equality follows directly from the choice of Δ in (167).

Proof of inequality (232). Applying (65) in Lemma 12 to λ^* in (230) leads to the upper bound in (232):

$$\lambda^* \leq \frac{H-1}{\sigma} \leq \frac{H-1}{\left(1 - \frac{3}{\beta}\right) \log\left(\frac{1}{\alpha+\Delta}\right)}, \quad (235)$$

where the last inequality holds by (158). As a result, we shall focus on showing the lower bounds in (232) in the remainder of the proof.

Recalling the definition of q in (150), we can reparameterize $1-q$ using two positive variables c_q and λ_q (whose choices will be made clearer soon) as follows:

$$1-q = \alpha = c_q e^{-(H-1)/\lambda_q}. \quad (236)$$

Deriving the first derivative of the function of interest $f(\lambda)$ in (230) as follows:

$$\begin{aligned} \nabla_\lambda f(\lambda) &= \nabla_\lambda \left(-\lambda \log \left(P_{1,0,1-\phi}^\phi \cdot \exp \left(\frac{-\bar{V}}{\lambda} \right) \right) - \lambda \sigma \right) \\ &\stackrel{(i)}{=} \nabla_\lambda \left(-\lambda \log \left(q e^{-(H-1)/\lambda} + 1 - q \right) - \lambda \sigma \right) \\ &= -\sigma - \log \left(q e^{-(H-1)/\lambda} + 1 - q \right) - \frac{1}{\lambda} \cdot \frac{q(H-1)e^{-(H-1)/\lambda}}{q e^{-(H-1)/\lambda} + 1 - q}, \end{aligned} \quad (237)$$

where (i) holds by the chosen transition kernels in (149) and the last line arises from basic calculus. To continue, when $\lambda = \lambda_q$, the derivative of the function $f(\lambda)$ can be expressed as

$$\begin{aligned} \nabla_\lambda f(\lambda) \big|_{\lambda=\lambda_q} &= -\sigma - \log \left((1-q) \frac{q}{c_q} + 1 - q \right) + \frac{(1-q) \frac{q}{c_q} \log \frac{1-q}{c_q}}{(1-q) \frac{q}{c_q} + 1 - q} \\ &= -\sigma - \log(1-q) - \log \left(1 + \frac{q}{c_q} \right) + \frac{\frac{q}{c_q} \log \frac{1-q}{c_q}}{\frac{q}{c_q} + 1} \\ &= -\sigma - \log(1-q) \left(1 - \frac{q/c_q}{q/c_q + 1} \right) - \log \left(1 + \frac{q}{c_q} \right) - \frac{\frac{q}{c_q} \log(c_q)}{1 + q/c_q} \\ &\stackrel{(i)}{=} -\sigma + \log \left(\frac{1}{\alpha + \Delta} \right) \left(1 - \frac{q/c_q}{q/c_q + 1} \right) - \log \left(1 + \frac{q}{c_q} \right) - \frac{\frac{q}{c_q} \log(c_q)}{1 + q/c_q} \quad (238) \\ &\stackrel{(ii)}{\geq} \log \left(\frac{1}{\alpha + \Delta} \right) \left(\frac{2}{\beta} - \frac{q/c_q}{q/c_q + 1} \right) - \log \left(1 + \frac{q}{c_q} \right) - \frac{\frac{q}{c_q} \log(c_q)}{1 + q/c_q} \\ &\stackrel{(iii)}{\geq} \frac{1}{\beta} \log \left(\frac{1}{\alpha + \Delta} \right) - \log \left(1 + \frac{1}{\beta} \right) - 1 \\ &\geq \frac{1}{\beta} \log \left(\frac{1}{\alpha + \Delta} \right) - 2 = 0, \end{aligned} \quad (239)$$

where (i) holds by (236), (ii) follows from the bound of σ in (158), (iii) arises from letting $c_q = \beta \geq 4$ and noting the fact $1/2 \leq q < 1$ (see (153)), leading to

$$\frac{1}{2\beta} \leq \frac{q}{c_q} < \frac{1}{\beta}, \quad \frac{q/c_q}{q/c_q + 1} \leq \frac{1}{\beta}, \quad \frac{\frac{q}{c_q} \log(c_q)}{1 + q/c_q} < 1. \quad (240)$$

Finally, the last line holds by $1/\beta \leq \frac{1}{4}$ and $\log\left(\frac{1}{\alpha+\Delta}\right) = 2\beta$ (see (152)).

To proceed, note that the function $f(\lambda)$ is concave with respect to λ . Therefore, observing $\nabla_\lambda f(\lambda)|_{\lambda=\lambda_q} \geq 0$ with $c_q = \beta$, we have $\lambda_q \leq \lambda^*$, which implies (see (236))

$$1 - q = \alpha + \Delta = \beta e^{-(H-1)/\lambda_q} \leq \beta e^{-(H-1)/\lambda^*}. \quad (241)$$

The above assertion directly gives

$$\lambda^* \geq \frac{H-1}{\log\left(\frac{\beta}{\alpha+\Delta}\right)}.$$

The proof is completed by noticing

$$\frac{H-1}{\log\left(\frac{\beta}{\alpha+\Delta}\right)} = \frac{H-1}{\log\left(\frac{1}{\alpha+\Delta}\right) + \log \beta} \stackrel{(i)}{\geq} \frac{H-1}{2 \log\left(\frac{1}{\alpha+\Delta}\right)} \geq \frac{H}{16\sigma},$$

where (i) follows from (152), and the last inequality follows from (158) and the fact $\beta \in [4, \infty)$.

Appendix C. Analysis: discounted infinite-horizon RMDPs

C.1 Proof of Lemma 7

We shall first show that the operator $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$ (cf. (46)) is a γ -contraction, which will in turn imply the existence of the unique fixed point of $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$. Before starting, suppose that the entries of $Q_1, Q_2 \in \mathbb{R}^{S \times A}$ are all bounded in $[0, \frac{1}{1-\gamma}]$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$. Denote that

$$\forall s \in \mathcal{S} : \quad V_1(s) := \max_a Q_1(s, a), \quad V_2(s) := \max_a Q_2(s, a). \quad (242)$$

Proof of γ -contraction. We first show that $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$ is a γ -contraction. Towards this, instead of $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$, we begin with a simpler operator $\widetilde{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$, defined as follows:

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q)(s, a) = r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P}V - b(s, a), \quad (243)$$

which consequently leads to

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A} : \quad \widehat{\mathcal{T}}_{\text{pe}}^\sigma(Q)(s, a) = \max \left\{ \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q)(s, a), 0 \right\}. \quad (244)$$

It follows straightforwardly that

$$\left\| \widehat{\mathcal{T}}_{\text{pe}}^\sigma(Q_1) - \widehat{\mathcal{T}}_{\text{pe}}^\sigma(Q_2) \right\|_\infty \leq \left\| \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q_1) - \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q_2) \right\|_\infty, \quad (245)$$

and hence it suffices to establish the γ -contraction of $\widetilde{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$. With this in mind, we observe that

$$\left\| \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q_1) - \widetilde{\mathcal{T}}_{\text{pe}}^\sigma(Q_2) \right\|_\infty = \gamma \left\| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P}V_1 - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P}V_2 \right\|_\infty \stackrel{(i)}{\leq} \gamma \|V_1 - V_2\|_\infty$$

$$\begin{aligned}
& \stackrel{(ii)}{=} \gamma \max_s \left| \max_a Q_1(s, a) - \max_a Q_2(s, a) \right| \\
& \leq \gamma \max_{(s,a)} |Q_1(s, a) - Q_2(s, a)| = \gamma \|Q_1 - Q_2\|_\infty,
\end{aligned} \tag{246}$$

where the first equality holds by the definition of $\tilde{T}_{pe}^\sigma(\cdot)$ (cf. (243)), (i) follows from that the infimum operator is a 1-contraction w.r.t. $\|\cdot\|_\infty$ and $\|\mathcal{P}V_1 - \mathcal{P}V_2\|_\infty \leq \|V_1 - V_2\|_\infty$ for all $\mathcal{P} \in \Delta(\mathcal{S})$, (ii) arises from the definitions in (242), and the last inequality is due to the maximum operator is also a 1-contraction w.r.t. $\|\cdot\|_\infty$. Combining the above two inequalities establish the desired statement.

Existence of the unique fixed point. To continue, we shall first claim that there exists at least one fixed point of $\hat{T}_{pe}^\sigma(\cdot)$. This is a standard argument, which we omit for brevity; interested readers are encouraged to refer to, e.g. Li et al. (2022), for details.

To prove the uniqueness of the fixed points of $\hat{T}_{pe}^\sigma(\cdot)$, suppose that there exist two fixed points Q' and Q'' obeying $Q' = \hat{T}_{pe}^\sigma(Q')$ and $Q'' = \hat{T}_{pe}^\sigma(Q'')$. Moreover, the definition of $\hat{T}_{pe}^\sigma(\cdot)$ directly implies $0 \leq Q', Q'' \leq \frac{1}{1-\gamma}$, since for any $0 \leq Q \leq \frac{1}{1-\gamma}$, it follows that $0 \leq \hat{T}_{pe}^\sigma(Q) \leq \frac{1}{1-\gamma}$. By the γ -contraction property, it follows that

$$\|Q' - Q''\|_\infty = \|\hat{T}_{pe}^\sigma(Q') - \hat{T}_{pe}^\sigma(Q'')\|_\infty \leq \gamma \|Q' - Q''\|_\infty. \tag{247}$$

However, (247) can't happen given $\gamma \in [\frac{1}{2}, 1)$, indicating the uniqueness of the fixed points of $\hat{T}_{pe}^\sigma(\cdot)$.

C.2 Proof of Lemma 8

To begin with, considering any Q, Q' obeying $Q \leq Q'$, and $0 \leq Q, Q' \leq \frac{1}{1-\gamma}$. We observe that the operator $\hat{T}_{pe}^\sigma(\cdot)$ (cf. (46)) has the monotone non-decreasing property, namely,

$$\begin{aligned}
\hat{T}_{pe}^\sigma(Q)(s, a) &= \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}V - b(s, a), 0 \right\} \\
&= \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P} \max_{a'} Q(\cdot, a') - b(s, a), 0 \right\} \\
&\leq \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P} \max_{a'} Q'(\cdot, a') - b(s, a), 0 \right\} = \hat{T}_{pe}^\sigma(Q')(s, a),
\end{aligned} \tag{248}$$

where the last line uses $Q \leq Q'$. Recalling the fixed point $\hat{Q}_{pe}^{*,\sigma}$ of $\hat{T}_{pe}^\sigma(\cdot)$, armed with (248) and the initialization $\hat{Q}_0 = 0$, we arrive at

$$\hat{Q}_1 = \hat{T}_{pe}^\sigma(\hat{Q}_0) \leq \hat{T}_{pe}^\sigma(\hat{Q}_{pe}^{*,\sigma}) = \hat{Q}_{pe}^{*,\sigma},$$

where the inequality follows from $\hat{Q}_0 = 0 \leq \hat{Q}_{pe}^{*,\sigma}$. Implementing the above result recursively gives

$$\forall m \geq 0: \quad \hat{Q}_m \leq \hat{Q}_{pe}^{*,\sigma}.$$

Applying the γ -contraction property in Lemma 7 thus yields that for any $m \geq 0$,

$$\begin{aligned} \|\widehat{Q}_m - \widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty &= \left\| \widehat{\mathcal{T}}_{\text{pe}}^\sigma(\widehat{Q}_{m-1}) - \widehat{\mathcal{T}}_{\text{pe}}^\sigma(\widehat{Q}_{\text{pe}}^{*,\sigma}) \right\|_\infty \\ &\leq \gamma \|\widehat{Q}_{m-1} - \widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty \leq \dots \leq \gamma^m \|\widehat{Q}_0 - \widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty = \gamma^m \|\widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty \leq \frac{\gamma^m}{1-\gamma}, \end{aligned}$$

where the last inequality holds by the fact $\|\widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty \leq \frac{1}{1-\gamma}$ (see Lemma 7).

C.3 Proof of Theorem 9

To begin, we introduce some additional notation that will be useful throughout the analysis. We denote the state-action space covered by the batch dataset $\mathcal{D}$ as

$$\mathcal{C}^{\text{b}} = \left\{ (s, a) : d^{\text{b}}(s, a) > 0 \right\}. \quad (249)$$

In addition, recalling the definition in (47), we define a similar one based on the true nominal model P^0 as

$$P_{\min}(s, a) := \min_{s'} \left\{ P^0(s' | s, a) : P^0(s' | s, a) > 0 \right\}, \quad (250)$$

which directly indicates that

$$P_{\min}^* = \min_s P_{\min}(s, \pi^*(s)), \quad P_{\min}^{\text{b}} = \min_{(s,a) \in \mathcal{C}^{\text{b}}} P_{\min}(s, a). \quad (251)$$

Next, we denote the set of possible state occupancy distributions associated with the optimal policy π^* in a model within the uncertainty set $P \in \mathcal{U}^\sigma(P^0)$ as

$$\mathcal{D}^* := \left\{ [d^{*,P}(s)]_{s \in \mathcal{S}} : P \in \mathcal{U}^\sigma(P^0) \right\} = \left\{ [d^{*,P}(s, \pi^*(s))]_{s \in \mathcal{S}} : P \in \mathcal{U}^\sigma(P^0) \right\}, \quad (252)$$

where the second equality is due to the fact that π^* is chosen to be deterministic.

We are now ready to embark on the proof of Theorem 9. We first introduce a fact that is used throughout the proof; the proof is postponed to Appendix C.3.2:

$$\forall (s, a) \in \mathcal{C}^{\text{b}} : \quad N(s, a) \geq \frac{N d^{\text{b}}(s, a)}{12} \geq \frac{c_1 \log(NS/\delta)}{12 P_{\min}(s, a)} \geq -\frac{\log \frac{2NS}{\delta}}{\log(1 - P_{\min}(s, a))} \quad (253)$$

as long as (58) holds.

For notation simplicity, denote the output Q-function and value function from Algorithm 3 as $\widehat{Q} = \widehat{Q}_M$ and $\widehat{V} = \widehat{V}_M$. Invoking Lemma 8 with $M \geq \frac{\log \frac{\sigma N}{1-\gamma}}{\log \frac{1}{\gamma}}$ directly leads to

$$\|\widehat{Q} - \widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty \leq \frac{1}{\sigma N} \quad (254)$$

and therefore

$$\|\widehat{V} - \widehat{V}_{\text{pe}}^{*,\sigma}\|_\infty \leq \max_s \left| \max_a \widehat{Q}(s, a) - \max_a \widehat{Q}_{\text{pe}}^{*,\sigma}(s, a) \right| \leq \|\widehat{Q} - \widehat{Q}_{\text{pe}}^{*,\sigma}\|_\infty \leq \frac{1}{\sigma N}. \quad (255)$$

The proof of Theorem 9 is separated into several key steps as follows.

Step 1: controlling the uncertainty via leave-one-out analysis. Given access to only a finite number of samples for estimating the nominal transition kernel P^0 , we need to efficiently control

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} \right|$$

across the robust value iterations, where $\hat{V}$ is statistically dependent on $\hat{P}_{s,a}^0$ (since $\hat{P}_{s,a}^0$ will be reused in the update rule (cf. (51)) for all the iterations). A naive treatment via the standard covering arguments will unfortunately lead to rather loose bounds (Zhou et al., 2021; Panaganti and Kalathil, 2022; Yang et al., 2022). To overcome this challenge, we resort to the leave-one-out analysis—pioneered by Agarwal et al. (2020); Li et al. (2020, 2022) in the context of model-based RL—to decouple the statistical dependency. The results are summarized in the following lemma, with the proof provided in Appendix C.3.1.

Lemma 22 *Instate the assumptions in Theorem 9. Then for all vector $\tilde{V}$ obeying $\|\tilde{V} - \hat{V}_{\text{pe}}^{*,\sigma}\|_\infty \leq \frac{1}{\sigma N}$ and $\|\tilde{V}\|_\infty \leq \frac{1}{1-\gamma}$, with probability at least $1 - \delta$, one has*

$$\begin{aligned} & \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\tilde{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\tilde{V} \right| \\ & \leq \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta})}{\hat{P}_{\min}(s,a)N(s,a)}} + \frac{4}{N\sigma(1-\gamma)}, \frac{1}{1-\gamma} \right\} \end{aligned} \quad (256)$$

for all $(s,a) \in \mathcal{S} \times \mathcal{A}$. In addition, for all $(s,a) \in \mathcal{C}^b$, with probability at least $1 - \delta$, one has

$$\frac{P_{\min}(s,a)}{8 \log(NS/\delta)} \leq \hat{P}_{\min}(s,a) \leq e^2 P_{\min}(s,a). \quad (257)$$

Step 2: establishing the pessimism property. Armed with Lemma 22, we aim to show the key property that

$$\forall (s,a) \in \mathcal{S} \times \mathcal{A}: \quad \hat{Q}(s,a) \leq Q^{\hat{\pi},\sigma}(s,a), \quad \hat{V}(s) \leq V^{\hat{\pi},\sigma}(s). \quad (258)$$

Similar to the finite-horizon setting, it suffices to focus on verifying the former assertion in (258). Towards this, we first recall that the fixed point $\hat{Q}_{\text{pe}}^{*,\sigma}$ of the pessimistic robust Bellman operator $\hat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$ (cf. (46)) obeys

$$\hat{Q}_{\text{pe}}^{*,\sigma} = \hat{\mathcal{T}}_{\text{pe}}^\sigma(\hat{Q}_{\text{pe}}^{*,\sigma}) = \max \left\{ r(s,a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V}_{\text{pe}}^{*,\sigma} - b(s,a), 0 \right\}. \quad (259)$$

If $\hat{Q}_{\text{pe}}^{*,\sigma}(s,a) = 0$. Given the initialization $\hat{Q}_0 = 0$, invoking Lemma 8 gives

$$\hat{Q}(s,a) = \hat{Q}_M(s,a) \leq \hat{Q}_{\text{pe}}^{*,\sigma}(s,a) = 0.$$

As a result, $Q^{\hat{\pi},\sigma}(s,a) \geq 0 = \hat{Q}(s,a)$ as desired. Therefore, it boils down to examine the case when $\hat{Q}_{\text{pe}}^{*,\sigma}(s,a) > 0$. One has

$$\hat{Q}(s,a) \stackrel{(i)}{\leq} \hat{Q}_{\text{pe}}^{*,\sigma}(s,a) + \frac{1}{\sigma N} = r(s,a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V}_{\text{pe}}^{*,\sigma} - b(s,a) + \frac{1}{\sigma N}$$

$$\begin{aligned}
 &\leq r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V} - b(s, a) + \frac{1}{\sigma N} + \gamma \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V}_{\text{pe}}^{*,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V} \right| \\
 &\stackrel{\text{(ii)}}{\leq} r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V} - b(s, a) + \frac{2}{\sigma N} \\
 &\leq r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} - b(s, a) + \frac{2}{\sigma N} + \gamma \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,a}^0)} \mathcal{P}\hat{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} \right| \\
 &\leq r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V}, \tag{260}
 \end{aligned}$$

where (i) follows from (254), (ii) arises from (255) and the basic fact that infimum operator is 1-contraction w.r.t $\|\cdot\|_\infty$, and the last inequality holds by the definition of $b(s, a)$ (cf. (48)) and Lemma 22. Putting the above inequality together with the robust Bellman equation (cf. (37a)) pertaining to $Q^{\hat{\pi},\sigma}(s, a)$, we arrive at

$$\begin{aligned}
 Q^{\hat{\pi},\sigma}(s, a) - \hat{Q}(s, a) &\geq r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}V^{\hat{\pi},\sigma} - \left(r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} \right) \\
 &= \gamma \left(\inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}V^{\hat{\pi},\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} \right) \\
 &\stackrel{\text{(i)}}{=} \gamma \left(\tilde{P}_{s,a} V^{\hat{\pi},\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}\hat{V} \right) \geq \gamma \tilde{P}_{s,a} (V^{\hat{\pi},\sigma} - \hat{V}),
 \end{aligned}$$

where (i) holds by setting $\tilde{P}_{s,a} = \operatorname{argmin}_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P}V^{\hat{\pi},\sigma}$. Consequently, one has

$$\begin{aligned}
 \min_{s,a} [Q^{\hat{\pi},\sigma}(s, a) - \hat{Q}(s, a)] &\geq \min_{s,a} [\gamma \tilde{P}_{s,a} (V^{\hat{\pi},\sigma} - \hat{V})] \stackrel{\text{(i)}}{\geq} \gamma \min_s [V^{\hat{\pi},\sigma}(s) - \hat{V}(s)] \\
 &= \gamma \min_s [Q^{\hat{\pi},\sigma}(s, \hat{\pi}(s)) - \hat{Q}(s, \hat{\pi}(s))] \\
 &\geq \gamma \min_{s,a} [Q^{\hat{\pi},\sigma}(s, a) - \hat{Q}(s, a)], \tag{261}
 \end{aligned}$$

where (i) follows from $\tilde{P}_{s,a} \in \Delta(\mathcal{S})$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$. Noting that $0 \leq \gamma < 1$, we conclude $Q^{\hat{\pi},\sigma}(s, a) - \hat{Q}(s, a) \geq 0$ for all $(s, a) \in \mathcal{S} \times \mathcal{A}$. This establishes the claim (258).

Step 3: bounding $V^{*,\sigma}(\rho) - V^{\hat{\pi},\sigma}(\rho)$. In view of the pessimistic property (cf. (258)), it follows that

$$V^{*,\sigma}(s) - V^{\hat{\pi},\sigma}(s) \leq V^{*,\sigma}(s) - \hat{V}(s). \tag{262}$$

Towards this, note that

$$\begin{aligned}
 \hat{V}(s) &= \max_a \hat{Q}(s, a) \geq \hat{Q}(s, \pi^*(s)) \stackrel{\text{(i)}}{\geq} \hat{Q}_{\text{pe}}^{*,\sigma}(s, \pi^*(s)) - \frac{1}{\sigma N} \\
 &\stackrel{\text{(ii)}}{\geq} r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s,\pi^*(s)}^0)} \mathcal{P}\hat{V}_{\text{pe}}^{*,\sigma} - b(s, \pi^*(s)) - \frac{1}{\sigma N}
 \end{aligned}$$

$$\begin{aligned}
&\geq r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} - b(s, \pi^*(s)) - \frac{1}{\sigma N} \\
&\quad - \gamma \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V}_{\text{pe}}^{\star, \sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} \right| \\
&\stackrel{\text{(iii)}}{\geq} r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} - b(s, \pi^*(s)) - \frac{2}{\sigma N} \\
&\geq r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} - b(s, \pi^*(s)) - \frac{2}{\sigma N} \\
&\quad - \gamma \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\hat{P}_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} \right| \\
&\geq r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} - 2b(s, \pi^*(s)), \tag{263}
\end{aligned}$$

where (i) follows from (254), (ii) holds by applying (259), (iii) arises from (255), and the basic fact that the infimum operator is a 1-contraction w.r.t. $\|\cdot\|_\infty$, and the final inequality holds by the definition of $b(s, a)$ (see (48)) and Lemma 22.

To continue, invoking the robust Bellman optimality equation in (37b) gives

$$V^{\star, \sigma}(s) = Q^{\star, \sigma}(s, \pi^*(s)) = r(s, \pi^*(s)) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}V^{\star, \sigma}.$$

Combining the above relation with (263), we arrive at

$$\begin{aligned}
V^{\star, \sigma}(s) - \hat{V}(s) &\leq \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}V^{\star, \sigma} - \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V} + 2b(s, \pi^*(s)) \\
&\leq \gamma \hat{P}_{s, \pi^*(s)}^{\text{inf}} (V^{\star, \sigma} - \hat{V}) + 2b(s, \pi^*(s)), \tag{264}
\end{aligned}$$

where the final inequality holds evidently, by introducing

$$\hat{P}_{s, \pi^*(s)}^{\text{inf}} := \operatorname{argmin}_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s, \pi^*(s)}^0)} \mathcal{P}\hat{V}. \tag{265}$$

Before continuing, for convenience, let us introduce a matrix $\hat{P}^{\text{inf}} \in \mathbb{R}^{S \times S}$ and a vector $b^* \in \mathbb{R}^S$, where their s -th rows (resp. entries) are defined as

$$\left[\hat{P}^{\text{inf}} \right]_{s, \cdot} = \hat{P}_{s, \pi^*(s)}^{\text{inf}}, \quad \text{and} \quad b^*(s) = b(s, \pi^*(s)). \tag{266}$$

With these notation in hand, averaging (264) over the initial state distribution ρ leads to

$$\begin{aligned}
V^{\star, \sigma}(\rho) - \hat{V}(\rho) &= \sum_{s \in \mathcal{S}} \rho(s) (V^{\star, \sigma}(s) - \hat{V}(s)) \\
&\leq \gamma \sum_{s \in \mathcal{S}} \rho(s) \hat{P}_{s, \pi^*(s)}^{\text{inf}} (V^{\star, \sigma} - \hat{V}) + 2 \sum_{s \in \mathcal{S}} \rho(s) b(s, \pi^*(s))
\end{aligned}$$

$$= \gamma \rho^\top \hat{P}^{\text{inf}} (V^{*,\sigma} - \hat{V}) + 2\rho^\top b^*. \quad (267)$$

Applying the above result recursively gives

$$\begin{aligned} V^{*,\sigma}(\rho) - \hat{V}(\rho) &\leq \gamma \rho^\top \hat{P}^{\text{inf}} (V^{*,\sigma} - \hat{V}) + 2\rho^\top b^* \\ &\leq \gamma \left(\gamma \rho^\top \hat{P}^{\text{inf}} \right) \hat{P}^{\text{inf}} (V^{*,\sigma} - \hat{V}) + 2 \left(\gamma \rho^\top \hat{P}^{\text{inf}} \right) b^* + 2\rho^\top b^* \\ &\leq \dots \leq \left\{ \lim_{i \rightarrow \infty} \gamma^i \rho^\top \left(\hat{P}^{\text{inf}} \right)^i (V^{*,\sigma} - \hat{V}) \right\} + 2\rho^\top \sum_{i=0}^{\infty} \gamma^i \left(\hat{P}^{\text{inf}} \right)^i b^* \\ &\stackrel{(i)}{\leq} 2\rho^\top \sum_{i=0}^{\infty} \gamma^i \left(\hat{P}^{\text{inf}} \right)^i b^* = 2\rho^\top \left(I - \gamma \hat{P}^{\text{inf}} \right)^{-1} b^*, \end{aligned} \quad (268)$$

where (i) holds by $|\rho^\top \left(\hat{P}^{\text{inf}} \right)^i (V^{*,\sigma} - \hat{V})| \leq \frac{1}{1-\gamma}$ for all $i \geq 0$, and that

$\lim_{i \rightarrow \infty} \gamma^i \rho^\top \left(\hat{P}^{\text{inf}} \right)^i (V^{*,\sigma} - \hat{V}) = 0$ since $\lim_{i \rightarrow \infty} \gamma^i = 0$ for all $0 \leq \gamma < 1$.

To further characterize the above performance gap, invoking the definition of $d^{*,P}$ (cf. (38) and (39a)), we arrive at

$$\left(d^{*,\hat{P}^{\text{inf}}} \right)^\top = (1-\gamma) \rho^\top \sum_{t=0}^{\infty} \gamma^t \left(\hat{P}^{\text{inf}} \right)^t = (1-\gamma) \rho^\top \left(I - \gamma \hat{P}^{\text{inf}} \right)^{-1}. \quad (269)$$

Plugging the above expression back into (268), and combining with (262), yields

$$V^{*,\sigma}(\rho) - V^{\hat{\pi},\sigma}(\rho) \leq V^{*,\sigma}(\rho) - \hat{V}(\rho) \leq \frac{2}{1-\gamma} \left\langle d^{*,\hat{P}^{\text{inf}}}, b^* \right\rangle. \quad (270)$$

Step 4: controlling $\left\langle d^{*,\hat{P}^{\text{inf}}}, b^* \right\rangle$ using concentrability. Note that $\hat{P}^{\text{inf}} \in \mathcal{U}^\sigma(P^0)$ (see (265) and (266)), which in words means $\hat{P}^{\text{inf}}$ is some transition kernel inside $\mathcal{U}^\sigma(P^0)$ — the uncertainty set around the nominal kernel P^0 . Similar to the finite-horizon case, observing that we can express $\left\langle d^{*,\hat{P}^{\text{inf}}}, b^* \right\rangle = \sum_{s \in \mathcal{S}} d^{*,\hat{P}^{\text{inf}}}(s) b^*(s)$, we divide the states into two cases and control them separately.

- **Case 1:** $s \in \mathcal{S}$ where $\max_{P \in \mathcal{U}^\sigma(P^0)} d^{*,P}(s, \pi^*(s)) = 0$. Since $\hat{P}^{\text{inf}} \in \mathcal{U}^\sigma(P^0)$, one has

$$0 \leq d^{*,\hat{P}^{\text{inf}}}(s) = d^{*,\hat{P}^{\text{inf}}}(s, \pi^*(s)) \leq \max_{P \in \mathcal{U}^\sigma(P^0)} d^{*,P}(s, \pi^*(s)) = 0,$$

which consequently indicates

$$d^{*,\hat{P}^{\text{inf}}}(s) = 0. \quad (271)$$

- **Case 2:** $s \in \mathcal{S}$ where $\max_{P \in \mathcal{U}^\sigma(P^0)} d^{*,P}(s, \pi^*(s)) > 0$. For any such state s , we claim that

$$d^b(s, \pi^*(s)) > 0 \quad \text{and} \quad (s, \pi^*(s)) \in \mathcal{C}^b. \quad (272)$$

This is due to Assumption 2, which requires C_{rob}^* to be finite given the numerator is positive:

$$\max_{P \in \mathcal{U}^\sigma(P^0)} \frac{\min \{d^{\star, P}(s, \pi^\star(s)), \frac{1}{S}\}}{d^{\text{b}}(s, \pi^\star(s))} = \max_{P \in \mathcal{U}^\sigma(P^0)} \frac{\min \{d^{\star, P}(s), \frac{1}{S}\}}{d^{\text{b}}(s, a)} \leq C_{\text{rob}}^* < \infty. \quad (273)$$

To continue, invoking the fact in (253) with $(s, \pi^\star(s)) \in \mathcal{C}^{\text{b}}$ gives

$$\begin{aligned} N(s, \pi^\star(s)) &\geq \frac{Nd^{\text{b}}(s, \pi^\star(s))}{12} \\ &\stackrel{(i)}{\geq} \frac{N \max_{P \in \mathcal{U}^\sigma(P^0)} \min \{d^{\star, P}(s, \pi^\star(s)), \frac{1}{S}\}}{12C_{\text{rob}}^*} \geq \frac{N \min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}}{12C_{\text{rob}}^*}, \end{aligned} \quad (274)$$

where (i) holds by Assumption 2, and the last inequality holds by $\hat{P}^{\text{inf}} \in \mathcal{U}^\sigma(P^0)$. With this in mind, we can control the pessimistic penalty $b^\star(s)$ (cf. (48)) by

$$\begin{aligned} b^\star(s) &\leq \frac{c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{\hat{P}_{\min}(s, \pi^\star(s))N(s, \pi^\star(s))}} + \frac{4}{\sigma N(1-\gamma)} + \frac{2}{\sigma N} \\ &\stackrel{(i)}{\leq} \frac{4c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\frac{\log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}(s, \pi^\star(s))N(s, \pi^\star(s))}} + \frac{4}{\sigma N(1-\gamma)} + \frac{2}{\sigma N} \\ &\leq \frac{16c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\frac{C_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}(s, \pi^\star(s))N \min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}}} + \frac{6}{\sigma N(1-\gamma)} \\ &\leq \frac{20c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\frac{C_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}(s, \pi^\star(s))N \min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}}}, \end{aligned}$$

where (i) arises from (257), the penultimate inequality follows from (274), and the last inequality holds as long as c_{b} is large enough.

Summing up the above two cases, we arrive at

$$\begin{aligned} \langle d^{\star, \hat{P}^{\text{inf}}}, b^\star \rangle &= \sum_{s \in \mathcal{S}} d^{\star, \hat{P}^{\text{inf}}}(s) b^\star(s) \\ &\leq \sum_{s \in \mathcal{S}} d^{\star, \hat{P}^{\text{inf}}}(s) \frac{20c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\frac{C_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}(s, \pi^\star(s))N \min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}}} \\ &\stackrel{(i)}{\leq} \frac{20c_{\text{b}}}{\sigma(1-\gamma)} \sqrt{\sum_{s \in \mathcal{S}} d^{\star, \hat{P}^{\text{inf}}}(s) \frac{C_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}(s, \pi^\star(s))N \min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}}} \sqrt{\sum_{s \in \mathcal{S}} d^{\star, \hat{P}^{\text{inf}}}(s)} \end{aligned}$$

$$\leq \frac{40c_b}{\sigma(1-\gamma)} \sqrt{\frac{SC_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}^* N}}, \quad (275)$$

where (i) arises from Cauchy-Schwarz inequality, and the last inequality holds since $P_{\min}(s, \pi^*(s)) \geq P_{\min}^*$ for all $s \in \mathcal{S}$ (see (251)) and the following fact (which has been established in (98)):

$$\sum_{s \in \mathcal{S}} \frac{d^{\star, \hat{P}^{\text{inf}}}(s)}{\min \{d^{\star, \hat{P}^{\text{inf}}}(s), \frac{1}{S}\}} \leq 2S.$$

Finally, inserting (275) back into (270), with probability at least $1 - 2\delta$, one has

$$V^{\star, \sigma}(\rho) - V^{\hat{\pi}, \sigma}(\rho) \leq \frac{2}{1-\gamma} \langle d^{\star, \hat{P}^{\text{inf}}}, b^{\star} \rangle \leq \frac{80c_b}{\sigma(1-\gamma)^2} \sqrt{\frac{SC_{\text{rob}}^* \log^2 \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}^* N}},$$

which concludes the proof.

C.3.1 PROOF OF LEMMA 22

We first note that the second assertion in (257) is the counterpart of (80), which can be verified following the same argument in Appendix B.2.1. For brevity, we omit its proof, and shall focus on verifying (256).

To begin with, we consider the situation when $N(s, a) = 0$. In this case, (256) can be easily verified since

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(\hat{P}_{s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{s,a}^0)} \mathcal{P}V \right| \stackrel{(i)}{=} \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{s,a}^0)} \mathcal{P}V \leq \|V\|_{\infty} \stackrel{(ii)}{\leq} \frac{1}{1-\gamma}, \quad (276)$$

where (i) follows from the fact $\hat{P}_{s,a}^0 = 0$ when $N(s, a) = 0$ (see (44)), and (ii) arises from the assumption $\|V\|_{\infty} \leq \frac{1}{1-\gamma}$. Consequently, in the remainder of the proof, we focus on verifying (256) when $N(s, a) > 0$. Let us first introduce the counterpart of the claim (79) in Lemma 17 as follows.

Lemma 23 *For all $(s, a) \in \mathcal{S} \times \mathcal{A}$ with $N(s, a) > 0$, consider any vector $V \in \mathbb{R}^S$ independent of $\hat{P}_{s,a}^0$ obeying $\|V\|_{\infty} \leq \frac{1}{1-\gamma}$. With probability at least $1 - \delta$, one has*

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(\hat{P}_{s,a}^0)} \mathcal{P}V - \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{s,a}^0)} \mathcal{P}V \right| \leq \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log(\frac{NS}{\delta})}{\hat{P}_{\min}(s, a)N(s, a)}}. \quad (277)$$

Proof The proof follows from the same arguments in Appendix B.2.2, with small modifications to adapt to the infinite-horizon setting; we omit the details for conciseness. ■

Armed with the above point-wise concentration bound, we are now ready to derive the uniform concentration bound desired as in Lemma 22, counting on a leave-one-out argument divided into the following steps. The crux of the analysis is to construct a set of auxiliary RMDPs, each different from the empirical RMDP only at a single state but possessing crucial statistical independence that facilitates the concentration arguments, which can then be transferred back to the empirical RMDP via a simple triangle inequality.

Step 1: construction of auxiliary RMDPs with state-absorbing empirical nominal transitions. Denote the empirical infinite-horizon robust MDP with the nominal transition kernel $\widehat{P}^0$ as $\widehat{\mathcal{M}}_{\text{rob}}$. Then, for each state s and each scalar $u \geq 0$, we can construct an auxiliary robust MDP $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ so that it is the same as $\widehat{\mathcal{M}}_{\text{rob}}$ except the properties in state s . To be precise, let the nominal transition kernel and reward function of $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ be $P^{s,u}$ and $r^{s,u}$, which are given respectively as

$$\begin{cases} P^{s,u}(s' | s, a) = \mathbb{1}(s' = s) & \text{for all } (s', a) \in \mathcal{S} \times \mathcal{A}, \\ P^{s,u}(\cdot | \tilde{s}, a) = \widehat{P}^0(\cdot | \tilde{s}, a) & \text{for all } (\tilde{s}, a) \in \mathcal{S} \times \mathcal{A} \text{ and } \tilde{s} \neq s, \end{cases} \quad (278)$$

and

$$\begin{cases} r^{s,u}(s, a) = u & \text{for all } a \in \mathcal{A}, \\ r^{s,u}(\tilde{s}, a) = r(\tilde{s}, a) & \text{for all } (\tilde{s}, a) \in \mathcal{S} \times \mathcal{A} \text{ and } \tilde{s} \neq s. \end{cases} \quad (279)$$

Clearly, state s of the auxiliary $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ is absorbing, meaning that the state stays at s once entering it. This removes the randomness of $\widehat{P}_{s,a}^0$ for all $a \in \mathcal{A}$ in state s , a key property we will leverage later.

With the robust MDP $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ in hand, we still need to complete the design by defining the corresponding penalty term for all $(\tilde{s}, a) \in \mathcal{S} \times \mathcal{A}$, which is given as follows

$$b^{s,u}(\tilde{s}, a) := \begin{cases} \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}^{s,u}(s,a)N(\tilde{s},a)}} + \frac{4}{N\sigma(1-\gamma)}, \frac{1}{1-\gamma} \right\} + \frac{2}{\sigma N} & \text{if } N(\tilde{s}, a) > 0, \\ \frac{1}{1-\gamma} + \frac{2}{\sigma N} & \text{otherwise,} \end{cases} \quad (280)$$

where $P_{\min}^{s,u}(\tilde{s}, a)$ is defined as the smallest positive state transition probability over the nominal kernel $P^{s,u}(\cdot | \tilde{s}, a)$:

$$\forall (\tilde{s}, a) \in \mathcal{S} \times \mathcal{A}: \quad P_{\min}^{s,u}(\tilde{s}, a) := \min_{s'} \left\{ P^{s,u}(s' | \tilde{s}, a) : P^{s,u}(s' | \tilde{s}, a) > 0 \right\}. \quad (281)$$

In view of (278) and (47), it holds that $P_{\min}^{s,u}(\tilde{s}, a) = \widehat{P}_{\min}(\tilde{s}, a)$, and therefore $b^{s,u}(\tilde{s}, a) = b(\tilde{s}, a)$, when $\tilde{s} \neq s$ for any $u \geq 0$. Armed with the above definitions, the pessimistic robust Bellman operator $\widehat{\mathcal{T}}_{s,u}^\sigma(Q)(\cdot)$ of the RMDP $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ is defined as

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A}: \quad \widehat{\mathcal{T}}_{s,u}^\sigma(Q)(s, a) = \max \left\{ r(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^{s,u})} \mathcal{P}V - b^{s,u}(s, a), 0 \right\}. \quad (282)$$

Step 2: fixed-point equivalence between $\widehat{\mathcal{M}}_{\text{rob}}$ and the auxiliary RMDP $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$. Recall that $\widehat{Q}_{\text{pe}}^{\star,\sigma}$ is the unique fixed point of $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$ with the corresponding value $\widehat{V}_{\text{pe}}^{\star,\sigma}$. We claim that there exists some choice of u such that the fixed point of $\widehat{\mathcal{T}}_{s,u}^\sigma(Q)(\cdot)$ coincides with that of $\widehat{\mathcal{T}}_{\text{pe}}^\sigma(\cdot)$. In particular, given a state s , we show the following choice of u suffices:

$$u^\star := (1-\gamma)\widehat{V}_{\text{pe}}^{\star,\sigma}(s) + \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}^{s,u}(s,a)N(s,a)}} + \frac{4}{N\sigma(1-\gamma)}, \frac{1}{1-\gamma} \right\} + \frac{2}{\sigma N}. \quad (283)$$

Towards this, we shall break our arguments in two different cases.

- **For state $s' \neq s$.** In this case, for any $a \in \mathcal{A}$, it can be verified that

$$\begin{aligned}
 & \max \left\{ r^{s,u^*}(s', a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s',a}^{s,u^*})} \mathcal{P} \widehat{V}_{\text{pe}}^{*,\sigma} - b^{s,u^*}(s', a), 0 \right\} \\
 &= \max \left\{ r(s', a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s',a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{*,\sigma} - b(s', a), 0 \right\} \\
 &= \widehat{\mathcal{T}}_{\text{pe}}^\sigma(\widehat{Q}_{\text{pe}}^{*,\sigma})(s', a) = \widehat{Q}_{\text{pe}}^{*,\sigma}(s', a),
 \end{aligned} \tag{284}$$

where the second line follows from the definitions in (279) and (278) as well as $b^{s,u^*}(s', a) = b(s', a)$ when $s' \neq s$, the last line arises from the definition of the pessimistic Bellman operator (46), and that $\widehat{Q}_{\text{pe}}^{*,\sigma}$ is the fixed point.

- **For state s .** In this case, for any u and $a \in \mathcal{A}$, observing that $P^{s,u}(s' | s, a)$ has only one positive entry equal to 1 (cf. (278)), applying (281) yields

$$P_{\min}^{s,u}(s, a) = 1. \tag{285}$$

Plugging the above fact into (280) leads to

$$b^{s,u}(s, a) = \begin{cases} \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log\left(\frac{2(1+\sigma)N^3s}{(1-\gamma)\delta}\right)}{N(s,a)}} + \frac{4}{N\sigma(1-\gamma)}, \frac{1}{1-\gamma} \right\} + \frac{2}{\sigma N} & \text{if } N(s, a) > 0, \\ \frac{1}{1-\gamma} & \text{otherwise} \end{cases} \tag{286}$$

for all $a \in \mathcal{A}$. As a result, we have for any $a \in \mathcal{A}$:

$$\begin{aligned}
 & \max \left\{ r^{s,u^*}(s, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^{s,u^*})} \mathcal{P} \widehat{V}_{\text{pe}}^{*,\sigma} - b^{s,u^*}(s, a), 0 \right\} \\
 &= \max \left\{ u^* + \gamma \widehat{V}_{\text{pe}}^{*,\sigma}(s) - b^{s,u^*}(s, a), 0 \right\} \\
 &= \max \left\{ (1-\gamma) \widehat{V}_{\text{pe}}^{*,\sigma}(s) + \gamma \widehat{V}_{\text{pe}}^{*,\sigma}(s), 0 \right\} = \widehat{V}_{\text{pe}}^{*,\sigma}(s),
 \end{aligned} \tag{287}$$

where the second line follows from the fact that $P_{s,a}^{s,u^*}$ is a singleton distribution at state s , and hence $\mathcal{U}^\sigma(P_{s,a}^{s,u^*}) = P_{s,a}^{s,u^*}$ by the definition of the KL uncertainty set, and the second line follows from plugging in the definition of u^* in (283) and $b^{s,u^*}(s, a)$ in (286).

Summing up the above two cases, we establish that there exists a fixed point $\widehat{Q}_{s,u^*}^{*,\sigma}$ of the operator $\widehat{\mathcal{T}}_{s,u^*}^\sigma(\cdot)$ if we let

$$\begin{cases} \widehat{Q}_{s,u^*}^{*,\sigma}(s, a) = \widehat{V}_{\text{pe}}^{*,\sigma}(s) & \text{for all } a \in \mathcal{A}, \\ \widehat{Q}_{s,u^*}^{*,\sigma}(s', a) = \widehat{Q}_{\text{pe}}^{*,\sigma}(s', a) & \text{for all } s' \neq s \text{ and } a \in \mathcal{A}. \end{cases} \tag{288}$$

Consequently, we confirm the existence of a fixed point of the operator $\widehat{\mathcal{T}}_{s,u^*}^\sigma(\cdot)$. In addition, its corresponding value function $\widehat{V}_{s,u^*}^{*,\sigma}$ also coincides with $\widehat{V}_{\text{pe}}^{*,\sigma}$.

Step 3: building an ε -net for all reward values u . It is easily verified that the reward u^* obeys

$$u^* \leq 1 + \min \left\{ \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{P_{\min}^{s,u}(s,a)N(s,a)}} + \frac{4}{\sigma N(1-\gamma)}, \frac{1}{1-\gamma} \right\} + \frac{2}{\sigma N} \leq \frac{2}{\sigma} + \frac{2}{1-\gamma}. \quad (289)$$

As a result, we construct an ε -net (Vershynin, 2018) of the line segment within the range $[0, \frac{2}{\sigma} + \frac{2}{1-\gamma}]$ with $\varepsilon = \frac{1}{\sigma N}$ as follows:

$$\mathcal{U}_\varepsilon := \left\{ \frac{i}{\sigma N} \mid 1 \leq i \leq \left\lceil \sigma N \left(\frac{2}{\sigma} + \frac{2}{1-\gamma} \right) \right\rceil \right\}. \quad (290)$$

Armed with this covering net $\mathcal{U}_\varepsilon$, we can construct an auxiliary robust MDP $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ and its corresponding pessimistic robust Bellman operator for each $u \in \mathcal{U}_\varepsilon$ (see Step 1). Following the same arguments in the proof of Lemma 7 (cf. Appendix C.1), for each $u \in \mathcal{U}_\varepsilon$, it can be verified that there exists a unique fixed point $\widehat{Q}_{s,u}^{*,\sigma}$ of the operator $\widehat{\mathcal{T}}_{s,u}^\sigma(\cdot)$, which satisfies $0 \leq \widehat{Q}_{s,u}^{*,\sigma} \leq \frac{1}{1-\gamma} \cdot 1$. In turn, the corresponding value function also satisfies $\|\widehat{V}_{s,u}^{*,\sigma}\|_\infty \leq \frac{1}{1-\gamma}$.

In view of the definitions in (278) and (279), for all $u \in \mathcal{U}_\varepsilon$, $\widehat{\mathcal{M}}_{\text{rob}}^{s,u}$ is statistically independent from $\widehat{P}_{s,a}^0$, which indicates the independence between $\widehat{V}_{s,u}^{*,\sigma}$ and $\widehat{P}_{s,a}^0$. This makes it possible to invoke Lemma 23, and taking the union bound over all samples N and $u \in \mathcal{U}_\varepsilon$ give that, with probability at least $1 - \delta$,

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u}^{*,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u}^{*,\sigma} \right| \leq \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log \left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta} \right)}{\widehat{P}_{\min}(s,a)N(s,a)}} \quad (291)$$

hold simultaneously for all $(s, a, u) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}_\varepsilon$ with $N(s, a) > 0$.

Step 4: a covering argument. Recalling that $u^* \in [0, \frac{2}{\sigma} + \frac{2}{1-\gamma}]$ (see (289)), we can always find some $\tilde{u} \in \mathcal{U}_\varepsilon$ such that $|\tilde{u} - u^*| \leq \frac{1}{\sigma N}$. Consequently, plugging in the operator in (282) yields

$$\forall Q \in \mathbb{R}^{SA}: \quad \left\| \widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(Q) - \widehat{\mathcal{T}}_{s,u^*}^\sigma(Q) \right\|_\infty \stackrel{(i)}{\leq} |\tilde{u} - u^*| \leq \frac{1}{\sigma N}, \quad (292)$$

where (i) holds by $b^{s,\tilde{u}}(s, a) = b^{s,u^*}(s, a)$ for s (see (286)) and $b^{s,\tilde{u}}(s', a) = b^{s,u^*}(s', a) = b(s', a)$ for all $s' \neq s$.

With this in mind, we observe that the fixed points of $\widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(\cdot)$ and $\widehat{\mathcal{T}}_{s,u^*}^\sigma(\cdot)$ obey

$$\begin{aligned} \left\| \widehat{Q}_{s,\tilde{u}}^{*,\sigma} - \widehat{Q}_{s,u^*}^{*,\sigma} \right\|_\infty &= \left\| \widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(\widehat{Q}_{s,\tilde{u}}^{*,\sigma}) - \widehat{\mathcal{T}}_{s,u^*}^\sigma(\widehat{Q}_{s,u^*}^{*,\sigma}) \right\|_\infty \\ &\leq \left\| \widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(\widehat{Q}_{s,\tilde{u}}^{*,\sigma}) - \widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(\widehat{Q}_{s,u^*}^{*,\sigma}) \right\|_\infty + \left\| \widehat{\mathcal{T}}_{s,\tilde{u}}^\sigma(\widehat{Q}_{s,u^*}^{*,\sigma}) - \widehat{\mathcal{T}}_{s,u^*}^\sigma(\widehat{Q}_{s,u^*}^{*,\sigma}) \right\|_\infty \\ &\leq \gamma \left\| \widehat{Q}_{s,\tilde{u}}^{*,\sigma} - \widehat{Q}_{s,u^*}^{*,\sigma} \right\|_\infty + \frac{1}{\sigma N}, \end{aligned} \quad (293)$$

which directly indicates that

$$\left\| \widehat{Q}_{s,\tilde{u}}^{\star,\sigma} - \widehat{Q}_{s,u^\star}^{\star,\sigma} \right\|_\infty \leq \frac{1}{(1-\gamma)\sigma N} \quad (294)$$

and

$$\left\| \widehat{V}_{s,\tilde{u}}^{\star,\sigma} - \widehat{V}_{s,u^\star}^{\star,\sigma} \right\|_\infty \leq \left\| \widehat{Q}_{s,\tilde{u}}^{\star,\sigma} - \widehat{Q}_{s,u^\star}^{\star,\sigma} \right\|_\infty \leq \frac{1}{(1-\gamma)\sigma N}. \quad (295)$$

Armed with the above facts, invoking the identity $\widehat{V}_{\text{pe}}^{\star,\sigma} = \widehat{V}_{s,u^\star}^{\star,\sigma}$ established in Step 2 gives

$$\begin{aligned} & \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} \right| = \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u^\star}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u^\star}^{\star,\sigma} \right| \\ & \stackrel{(i)}{\leq} \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} \right| \\ & \quad + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u^\star}^{\star,\sigma} \right| + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,u^\star}^{\star,\sigma} \right| \\ & \stackrel{(ii)}{\leq} \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{s,\tilde{u}}^{\star,\sigma} \right| + \frac{2}{N\sigma(1-\gamma)} \\ & \leq \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log\left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta}\right)}{\widehat{P}_{\min}(s,a)N(s,a)}} + \frac{2}{N\sigma(1-\gamma)}, \end{aligned} \quad (296)$$

where (i) holds by applying the triangle inequality, (ii) arises from (295) and the basic fact that infimum operator is a 1-contraction w.r.t. $\|\cdot\|_\infty$, and the final inequality follows from (291).

Step 5: finishing up. Now we are positioned to finish up the proof. For all vector $\tilde{V}$ obeying $\|\tilde{V} - \widehat{V}_{\text{pe}}^{\star,\sigma}\|_\infty \leq \frac{1}{\sigma N}$ and $\|\tilde{V}\|_\infty \leq \frac{1}{1-\gamma}$, we apply the triangle inequality and invoke (296) to reach

$$\begin{aligned} & \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \tilde{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \tilde{V} \right| \leq \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} \right| \\ & \quad + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \tilde{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} \right| + \left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \tilde{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \widehat{V}_{\text{pe}}^{\star,\sigma} \right| \\ & \leq \frac{c_b}{\sigma(1-\gamma)} \sqrt{\frac{\log\left(\frac{2(1+\sigma)N^3S}{(1-\gamma)\delta}\right)}{\widehat{P}_{\min}(s,a)N(s,a)}} + \frac{4}{N\sigma(1-\gamma)}. \end{aligned} \quad (297)$$

Finally, we complete the proof by verifying that

$$\left| \inf_{\mathcal{P} \in \mathcal{U}^\sigma(\widehat{P}_{s,a}^0)} \mathcal{P} \tilde{V} - \inf_{\mathcal{P} \in \mathcal{U}^\sigma(P_{s,a}^0)} \mathcal{P} \tilde{V} \right| \leq \|\tilde{V}\|_\infty \leq \frac{1}{1-\gamma}. \quad (298)$$

C.3.2 PROOF OF (253)

For all $(s, a) \in \mathcal{C}^b$, one has

$$Nd^b(s, a) \stackrel{(i)}{\geq} \frac{c_1 d^b(s, a) \log(NS/\delta)}{d_{\min}^b P_{\min}^b} \stackrel{(ii)}{\geq} \frac{c_1 \log(NS/\delta)}{P_{\min}^b} \stackrel{(iii)}{\geq} \frac{c_1 \log(NS/\delta)}{P_{\min}(s, a)}, \quad (299)$$

where (i) follows from the condition (58), (ii) arises from the definition that $d_{\min}^b \leq d^b(s, a)$ for all $(s, a) \in \mathcal{C}^b$, and (iii) follows from the definition in (251). In particular, when c_1 is large enough, one has $\frac{2}{3} \log \frac{NS}{\delta} < \frac{Nd^b(s, a)}{12}$. To continue, we recall a key property of $N(s, a)$ (cf. (43)) in the following lemma.

Lemma 24 ((Li et al., 2022, Lemma 7)) *Fix $\delta \in (0, 1)$. With probability at least $1 - \delta$, the quantities $\{N(s, a)\}$ in (43) obey*

$$\max \left\{ N(s, a), \frac{2}{3} \log \frac{NS}{\delta} \right\} \geq \frac{Nd^b(s, a)}{12} \quad (300)$$

simultaneously for all $(s, a) \in \mathcal{S} \times \mathcal{A}$.

Consequently, Lemma 24 tells us that with probability at least $1 - \delta$,

$$N(s, a) \geq \frac{Nd^b(s, a)}{12} \geq \frac{c_1 \log(NS/\delta)}{12P_{\min}(s, a)} \quad (301)$$

as long as c_1 is large enough. Last but not least, taking the basic fact $x \leq -\log(1 - x)$ for all $x \in [0, 1]$, the last inequality of (253) can be verified by

$$\frac{c_1 \log(NS/\delta)}{12P_{\min}(s, a)} \geq -\frac{\log \frac{2NS}{\delta}}{\log(1 - P_{\min}(s, a))}. \quad (302)$$

C.4 Proof of Theorem 10

Similar to the finite-horizon case, we shall develop the lower bounds for the two cases when the uncertainty levels σ vary separately.

C.4.1 CONSTRUCTION OF HARD PROBLEM INSTANCES: SMALL UNCERTAINTY LEVEL

We first construct some hard discounted infinite-horizon RMDP instances and then characterize the sample complexity requirements over these instances.

Construction of a collection of hard MDPs. Suppose there are two MDPs

$$\left\{ \mathcal{M}_\theta = (\mathcal{S}, \mathcal{A}, P^\theta, r, \gamma) \mid \theta = \{0, 1\} \right\}.$$

Here, $\mathcal{S} = \{0, 1, \dots, S-1\}$, and $\mathcal{A} = \{0, 1\}$. The transition kernel P^θ of the MDP $\mathcal{M}^\theta$ is specified as follows:

$$P^\theta(s' \mid s, a) = \begin{cases} p\mathbb{1}(s' = 0) + (1-p)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, \theta) \\ q\mathbb{1}(s' = 0) + (1-q)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, 1-\theta) \\ q\mathbb{1}(s' = s) + (1-q)\mathbb{1}(s' = 1) & \text{if } s > 0 \end{cases}, \quad (303)$$

for any $(s, a, s') \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}$, where p and q are set to be

$$p = 1 - c_1(1 - \gamma), \quad q = 1 - c_1(1 - \gamma) - c_2(1 - \gamma)^2\varepsilon, \quad (304)$$

which satisfies

$$\frac{2}{3} \leq \gamma < 1 \quad \text{and} \quad c_2(1 - \gamma)^2\varepsilon \leq \frac{c_1(1 - \gamma)}{2} \leq \frac{1}{8}. \quad (305)$$

for some $c_1 \leq \frac{1}{4}$ and small enough c_2 . In view of the assumptions (305), one has

$$p > q \geq \frac{1}{2}. \quad (306)$$

Finally, we define the reward function as

$$r(s, a) = \begin{cases} 1 & \text{if } s = 0 \text{ or } s = 2 \\ 0 & \text{otherwise} \end{cases}. \quad (307)$$

Construction of the history/batch dataset. Define a useful state distribution (only supported on the state subset $\{0, 1, 2\}$) as

$$\mu(s) = \frac{1}{CS} \mathbb{1}(s = 0) + \left(1 - \frac{1}{CS}\right) \mathbb{1}(s = 1), \quad (308)$$

where $C > 0$ is some constant that determines the robust concentrability coefficient C_{rob}^* (which will be made clear soon) and obeys

$$\frac{1}{CS} \leq \frac{1}{4}. \quad (309)$$

A batch dataset—consists of N i.i.d samples $\{(s_i, a_i, s'_i)\}_{1 \leq i \leq N}$ —is generated over the nominal environment $\mathcal{M}_\theta$ according to (40), with the behavior distribution chosen to be:

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A}: \quad d^b(s, a) = \frac{\mu(s)}{2}. \quad (310)$$

Additionally, we choose the following initial state distribution:

$$\rho(s) = \begin{cases} 1, & \text{if } s = 0 \\ 0, & \text{otherwise} \end{cases}. \quad (311)$$

Uncertainty set of the transition kernels. We next describe the radius σ of the uncertainty set in our construction of the robust MDPs, along with some useful properties, which are similar to the finite-horizon case. The perturbed transition kernels in $\mathcal{M}_\theta$ is limited to the following uncertainty set

$$\mathcal{U}^\sigma(P^\theta) := \otimes \mathcal{U}^\sigma(P_{s,a}^\theta), \quad \mathcal{U}^\sigma(P_{s,a}^\theta) := \left\{ P_{s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{s,a} \parallel P_{s,a}^\theta) \leq \sigma \right\}, \quad (312)$$

where $P_{s,a}^\theta := P^\theta(\cdot | s, a) \in [0, 1]^{1 \times S}$. Moreover, the radius of the uncertainty set σ obeys

$$0 \leq \sigma \leq \frac{1 - \gamma}{20}. \quad (313)$$

For any $(s, a, s') \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}$, we denote the infimum entry of the perturbed transition kernel $P_{s,a} \in \mathcal{U}^\sigma(P_{s,a}^\theta)$ moving to the next state s' as

$$\underline{P}^\theta(s' | s, a) := \inf_{P_{s,a} \in \mathcal{U}^\sigma(P_{s,a}^\theta)} P(s' | s, a). \quad (314)$$

As shall be seen, the transition from state 0 to state 2 plays an important role in the analysis, for convenience, we denote

$$\underline{p}_\star := \underline{P}^\theta(0 | 0, \theta), \quad \underline{q}_\star := \underline{P}^\theta(0 | 0, 1 - \theta). \quad (315)$$

With these definitions in place, we summarize some useful properties of the uncertainty set in the following lemma, which parallels Lemma 18 in the finite-horizon case.

Lemma 25 *Suppose the uncertainty level σ satisfies (313). The perturbed transition kernels obey*

$$\underline{p}_\star \geq \underline{q}_\star \geq 1 - c_3(1 - \gamma) \quad \text{and} \quad \underline{p}_\star - \underline{q}_\star \geq p - q \geq 0 \quad (316)$$

for constant $c_3 = 2c_1 \leq \frac{1}{4}$.

Proof The proof follows from the same arguments as the proof for Lemma 18 in Appendix B.3.5 by replacing H with $\frac{1}{1-\gamma}$; we omit the details for brevity. $\blacksquare$

Value functions and optimal policies. Now we are positioned to derive the corresponding robust value functions and identify the optimal policies. For any MDP $\mathcal{M}_\theta$ with the above uncertainty set, denote $\pi_\theta^\star$ as the optimal policy. In addition, we denote the robust value function of any policy π (resp. the optimal policy $\pi_\theta^\star$) as $V_\theta^{\pi, \sigma}$ (resp. $V_\theta^{\star, \sigma}$). Then, we introduce the following lemma which describes some important properties of the robust value functions and optimal policies.

Lemma 26 *For any $\theta = \{0, 1\}$ and any policy π , one has*

$$V_\theta^{\pi, \sigma}(0) = \frac{1}{1 - \gamma x_\theta^\pi}, \quad (317)$$

where x_θ^π is defined as

$$x_\theta^\pi := \underline{p}_\star \pi(\theta | 0) + \underline{q}_\star \pi(1 - \theta | 0). \quad (318)$$

In addition, the optimal value functions and the optimal policies obey

$$V_\theta^{\star, \sigma}(0) = \frac{1}{1 - \gamma \underline{p}_\star}, \quad V_\theta^{\star, \sigma}(s) = 0 \quad \text{for } s = 1 \text{ or } s > 2, \quad (319a)$$

$$\pi_\theta^\star(\theta | s) = 1, \quad \text{for } s \in \mathcal{S}. \quad (319b)$$

Moreover, the robust single-policy clipped concentrability coefficient $C_{\text{rob}}^\star$ obeys

$$C_{\text{rob}}^\star = 2C. \quad (320)$$

Proof See Appendix C.4.5. $\blacksquare$

C.4.2 ESTABLISHING THE MINIMAX LOWER BOUND: SMALL UNCERTAINTY LEVEL

Towards this, we first introduce the following lemma, which parallels the claim in (167)-(168) in the finite-horizon case.

Lemma 27 *For any policy $\hat{\pi}$,*

$$V_{\theta}^{\star,\sigma}(0) - V_{\theta}^{\hat{\pi},\sigma}(0) \geq 2\varepsilon(1 - \hat{\pi}(\theta | 0)).$$

Proof This lemma can be directly verified by controlling $V_{\theta}^{\star,\sigma}(0) - V_{\theta}^{\hat{\pi},\sigma}(0)$ with the help of Lemma 25 and Lemma 26; we omit the details for brevity. $\blacksquare$

Armed with this lemma, following the same arguments in Appendix B.3.4, we can complete the proof by observing that: let c_1 be some sufficient large constant, as long as the sample size is beneath

$$N \leq \frac{c_4 SC_{\text{rob}}^{\star}}{(1 - \gamma)^3 \varepsilon^2}, \quad (321)$$

then we necessarily have

$$\inf_{\hat{\pi}} \max_{\theta \in \{0,1\}} \mathbb{P}_{\theta} \left\{ V_{\theta}^{\star,\sigma}(\rho) - V_{\theta}^{\hat{\pi},\sigma}(\rho) \geq \varepsilon \right\} \geq \frac{1}{8}, \quad (322)$$

where $\mathbb{P}_{\theta}$ denote the probability conditioned on that the MDP is $\mathcal{M}_{\theta}$. We omit the details for brevity and complete the proof.

C.4.3 CONSTRUCTION OF HARD PROBLEM INSTANCES: LARGE UNCERTAINTY LEVEL

Construction of a collection of hard MDPs. Suppose there are two MDPs

$$\left\{ \mathcal{M}_{\phi} = (\mathcal{S}, \mathcal{A}, P^{\phi}, r, \gamma) \mid \phi = \{0, 1\} \right\}.$$

Here, γ is the discount parameter, $\mathcal{S} = \{0, 1, \dots, S-1\}$ is the state space, and $\mathcal{A} = \{0, 1\}$ is the action space. The transition kernel P^{ϕ} of either constructed MDP $\mathcal{M}_{\phi}$ is defined as

$$P^{\phi}(s' | s, a) = \begin{cases} p\mathbb{1}(s' = 2) + (1 - p)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, \phi) \\ q\mathbb{1}(s' = 2) + (1 - q)\mathbb{1}(s' = 1) & \text{if } (s, a) = (0, 1 - \phi) \\ \mathbb{1}(s' = s) & \text{if } s = 1 \text{ or } s = 2 \\ q\mathbb{1}(s' = s) + (1 - q)\mathbb{1}(s' = 1) & \text{if } s > 2 \end{cases}, \quad (323)$$

where p and q are set as

$$p = 1 - \alpha \quad \text{and} \quad q = 1 - \alpha - \Delta \quad (324)$$

for some γ , α and Δ obeying

$$0 < \alpha \leq 1 - \gamma \leq 1/(2e^8) \leq \frac{1}{2} \quad \text{and} \quad \Delta \leq \frac{\alpha}{2}. \quad (325)$$

Here, α and Δ are some values that will be introduced later. Consequently, applying (324) directly leads to

$$1 \geq p \geq q \geq \gamma \geq \frac{1}{2}. \quad (326)$$

Note that state 1 and 2 are absorbing states. In addition, if the initial distribution is supported on states $\{0, 1, 2\}$, the MDP will always stay in the state $\{1, 2\}$ after the first transition.

Finally, we define the reward function as

$$r(s, a) = \begin{cases} 1 & \text{if } s = 0 \text{ or } s = 2 \\ 0 & \text{otherwise} \end{cases}. \quad (327)$$

Construction of the history/batch dataset. Define a useful state distribution (only supported on the state subset $\{0, 1, 2\}$) as

$$\mu(s) = \frac{1}{CS} \mathbb{1}(s = 0) + \frac{1}{CS} \mathbb{1}(s = 2) + \left(1 - \frac{2}{CS}\right) \mathbb{1}(s = 1), \quad (328)$$

where $C > 0$ is some constant that determines the robust concentrability coefficient C_{rob}^* (which will be made clear soon) and obeys

$$\frac{1}{CS} \leq \frac{1}{4}. \quad (329)$$

A batch dataset—consists of N i.i.d samples $\{(s_i, a_i, s'_i)\}_{1 \leq i \leq N}$ —is generated over the nominal environment $\mathcal{M}_\phi$ according to (40), with the behavior distribution chosen to be:

$$\forall (s, a) \in \mathcal{S} \times \mathcal{A}: \quad d^b(s, a) = \frac{\mu(s)}{2}. \quad (330)$$

Additionally, we choose the following initial state distribution:

$$\rho(s) = \begin{cases} 1, & \text{if } s = 0 \\ 0, & \text{otherwise} \end{cases}. \quad (331)$$

Uncertainty set of the transition kernels. We next describe the radius σ of the uncertainty set in our construction of the robust MDPs, along with some useful properties, which are similar to the finite-horizon case. To begin with, with slight abuse of notation, we introduce an important constant β defined as

$$\beta := \frac{1}{2} \log \frac{1}{\alpha + \Delta} \geq 4. \quad (332)$$

The perturbed transition kernels in $\mathcal{M}_\phi$ is limited to the following uncertainty set

$$\mathcal{U}^\sigma(P^\phi) := \otimes \mathcal{U}^\sigma(P_{s,a}^\phi), \quad \mathcal{U}^\sigma(P_{s,a}^\phi) := \left\{ P_{s,a} \in \Delta(\mathcal{S}) : \text{KL}(P_{s,a} \parallel P_{s,a}^\phi) \leq \sigma \right\}, \quad (333)$$

where $P_{s,a}^\phi := P^\phi(\cdot | s, a) \in [0, 1]^{1 \times \mathcal{S}}$. Moreover, the radius of the uncertainty set σ obeys

$$\left(1 - \frac{3}{\beta}\right) \log \frac{1}{\alpha + \Delta} \leq \sigma \leq \left(1 - \frac{2}{\beta}\right) \log \frac{1}{\alpha + \Delta}. \quad (334)$$

For any $(s, a, s') \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}$, we denote the infimum entry of the perturbed transition kernel $P_{s,a} \in \mathcal{U}^\sigma(P_{s,a}^\phi)$ moving to the next state s' as

$$\underline{P}^\phi(s' | s, a) := \inf_{P_{s,a} \in \mathcal{U}^\sigma(P_{s,a}^\phi)} P(s' | s, a). \quad (335)$$

As shall be seen, the transition from state 0 to state 2 plays an important role in the analysis, for convenience, we denote

$$\underline{p} := \underline{P}^\phi(2 | 0, \phi), \quad \underline{q} := \underline{P}^\phi(2 | 0, 1 - \phi). \quad (336)$$

With these definitions in place, we summarize some useful properties of the uncertainty set in the following lemma, which parallels Lemma 20 in the finite-horizon case.

Lemma 28 *Suppose β satisfies (332) and the uncertainty level σ satisfies (334). The perturbed transition kernels obey*

$$\underline{p} \geq \underline{q} \geq \frac{1}{\beta}. \quad (337)$$

Proof The proof follows from the same arguments as Appendix B.3.10 by replacing H with $\frac{1}{1-\gamma}$; we omit the details for brevity. $\blacksquare$

Value functions and optimal policies. Now we are positioned to derive the corresponding robust value functions and identify the optimal policies. For any MDP $\mathcal{M}_\phi$ with the above uncertainty set, denote π_ϕ^* as the optimal policy. In addition, we denote the robust value function of any policy π (resp. the optimal policy π_ϕ^*) as $V_\phi^{\pi,\sigma}$ (resp. $V_\phi^{\star,\sigma}$). Then, we introduce the following lemma which describes some important properties of the robust value functions and optimal policies.

Lemma 29 *For any $\phi = \{0, 1\}$ and any policy π , one has*

$$V_\phi^{\pi,\sigma}(0) = 1 + \frac{\gamma}{1-\gamma} z_\phi^\pi, \quad (338)$$

where z_ϕ^π is defined as

$$z_\phi^\pi := \underline{p}\pi(\phi | 0) + \underline{q}\pi(1 - \phi | 0). \quad (339)$$

In addition, the optimal value functions and the optimal policies obey

$$V_\phi^{\star,\sigma}(0) = 1 + \frac{\gamma}{1-\gamma} \underline{p}, \quad V_\phi^{\star,\sigma}(2) = \frac{1}{1-\gamma}, \quad V_\phi^{\star,\sigma}(s) = 0 \quad \text{for } s = 1 \text{ or } s > 2, \quad (340a)$$

$$\pi_\phi^*(\phi | s) = 1, \quad \text{for } s \in \mathcal{S}. \quad (340b)$$

Moreover, choosing $S \geq 2\beta$, the robust single-policy clipped concentrability coefficient C_{rob}^* obeys

$$C_{\text{rob}}^* = 2C. \quad (341)$$

Proof See Appendix C.4.6. $\blacksquare$

C.4.4 ESTABLISHING THE MINIMAX LOWER BOUND: LARGE UNCERTAINTY LEVEL

Now we are positioned to provide the sample complexity lower bound. In view of Lemma 29, the smallest positive state transition probability of the optimal policy π_ϕ^* under any nominal transition kernel P^ϕ with $\phi \in \{0, 1\}$ satisfies:

$$P_{\min}^* := \min_{s, s'} \left\{ P^\phi(s' | s, \pi_\phi^*(s)) : P^\phi(s' | s, \pi_\phi^*(s)) > 0 \right\} = P^\phi(1|0, \phi) = 1 - p. \quad (342)$$

Our goal is to control the quantity w.r.t. any policy estimator $\hat{\pi}$ based on the batch dataset and the chosen initial distribution ρ in (331), which gives

$$V_\phi^{\star, \sigma}(\rho) - V_\phi^{\hat{\pi}, \sigma}(\rho) = V_\phi^{\star, \sigma}(0) - V_\phi^{\hat{\pi}, \sigma}(0). \quad (343)$$

Towards this, we first introduce the following lemma, which parallels the claim in (167)-(168) in the finite-horizon case.

Lemma 30 *Given $\varepsilon \leq \frac{1}{384e^6(1-\gamma)\log(\frac{1}{\alpha})} \leq \frac{1}{384e^6(1-\gamma)\log(\frac{1}{\alpha+\Delta})}$, choosing $\Delta = 128e^6\sigma(1-q)\varepsilon(1-\gamma) \leq 128e^6(\alpha+\Delta)\varepsilon\log\left(\frac{1}{\alpha+\Delta}\right)(1-\gamma) \leq \frac{\alpha}{2}$, one has for any policy $\hat{\pi}$,*

$$V_\phi^{\star, \sigma}(0) - V_\phi^{\hat{\pi}, \sigma}(0) \geq 2\varepsilon(1 - \hat{\pi}(\phi|0)).$$

Proof This lemma follows from the same arguments as Appendix B.3.12 except replacing H with $\frac{1}{1-\gamma}$ under the additional condition $\gamma \geq \frac{1}{2}$; we omit the details for brevity. ■

Armed with this lemma, following the same arguments in Appendix B.3.4, we can complete the proof by observing that: let c_1 be some sufficient large constant, as long as the sample size is beneath

$$N \leq \frac{SC_{\text{rob}}^* \log 2}{4c_1 P_{\min}^* \sigma^2 (1-\gamma)^2 \varepsilon^2}, \quad (344)$$

then we necessarily have

$$\inf_{\hat{\pi}} \max_{\phi \in \{0,1\}} \mathbb{P}_\phi \left\{ V_\phi^{\star, \sigma}(\rho) - V_\phi^{\hat{\pi}, \sigma}(\rho) \geq \varepsilon \right\} \geq \frac{1}{8}, \quad (345)$$

where $\mathbb{P}_\phi$ denote the probability conditioned on that the MDP is $\mathcal{M}_\phi$. We omit the details for brevity and complete the proof.

C.4.5 PROOF OF LEMMA 26

First, it is easily verified that for any policy π ,

$$\forall s \in \mathcal{S} \setminus \{0\} : V_\theta^{\pi, \sigma}(s) = \sum_{t=0}^{\infty} \gamma^t \cdot 0 = 0 \quad (346)$$

since the reward function $r(s, a) = 0$ for $(s, a) \in \mathcal{S} \setminus \{0\} \times \mathcal{A}$ in (307).

To continue, we observe that the robust value function of state 0 satisfies

$$\begin{aligned} V_{\theta}^{\pi, \sigma}(0) &= \mathbb{E}_{a \sim \pi(\cdot | 0)} \left[r(0, a) + \gamma \inf_{P \in \mathcal{U}^{\sigma}(P_{0,a}^{\theta})} \mathcal{P} V_{\theta}^{\pi, \sigma} \right] \\ &\stackrel{(i)}{=} 1 + \gamma \pi(\theta | 0) \inf_{P \in \mathcal{U}^{\sigma}(P_{0,\theta}^{\theta})} \mathcal{P} V_{\theta}^{\pi, \sigma} + \gamma \pi(1 - \theta | 0) \inf_{P \in \mathcal{U}^{\sigma}(P_{0,1-\theta}^{\theta})} \mathcal{P} V_{\theta}^{\pi, \sigma} \end{aligned} \quad (347)$$

$$\begin{aligned} &\stackrel{(ii)}{=} 1 + \gamma \pi(\theta | 0) \left[\underline{p}_{\star} V_{\theta}^{\pi, \sigma}(0) + (1 - \underline{p}_{\star}) V_{\theta}^{\pi, \sigma}(1) \right] \\ &\quad + \gamma \pi(1 - \theta | 0) \left[\underline{q}_{\star} V_{\theta}^{\pi, \sigma}(0) + (1 - \underline{q}_{\star}) V_{\theta}^{\pi, \sigma}(1) \right] \\ &\stackrel{(iii)}{=} 1 + \gamma x_{\theta}^{\pi} [V_{\theta}^{\pi, \sigma}(0) - V_{\theta}^{\pi, \sigma}(1)] \\ &= \frac{1}{1 - \gamma x_{\theta}^{\pi}}, \end{aligned} \quad (348)$$

where (i) holds by the reward function defined in (307). To see (ii), note that (347) indicates $V_{\theta}^{\pi, \sigma}(0) \geq 1 \geq V_{\theta}^{\pi, \sigma}(1) = 0$, so that the infimum is obtained by picking the smallest possible mass on the transition to state 2, provided by the definition in (315), and (iii) follows by plugging in the definition of x_{θ}^{π} in (318).

Consequently, observing that the function $\frac{1}{1 - \gamma x}$ is increasing in x and x_{θ}^{π} is also increasing in $\pi(\theta | 0)$ (see the fact $\underline{p}_{\star} \geq \underline{q}_{\star}$ in (316)), the optimal policy in state 0 thus obeys

$$\pi_{\theta}^{\star}(\theta | 0) = 1. \quad (349)$$

Therefore,

$$x_{\theta}^{\star} := x_{\theta}^{\pi_{\theta}^{\star}} = \underline{p}_{\star} \pi_{\theta}^{\star}(\theta | 0) + \underline{q}_{\star} \pi_{\theta}^{\star}(1 - \theta | 0) = \underline{p}_{\star}, \quad (350)$$

which combined with (348) yields

$$V_{\theta}^{\star, \sigma}(0) = \frac{1}{1 - \gamma \underline{p}_{\star}}. \quad (351)$$

Regarding the optimal policy for the remaining states $s > 0$, since the action does not influence the state transition, without loss of generality, we choose the optimal policy to obey

$$\forall s > 0 : \quad \pi_{\theta}^{\star}(\theta | s) = 1. \quad (352)$$

Proof of (320). To begin with, for any MDP $\mathcal{M}_{\theta}$ with $\theta \in \{0, 1\}$, recall the definition of $C_{\text{rob}}^{\star}$ as

$$C_{\text{rob}}^{\star} = \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star, P}(s, a), \frac{1}{\underline{S}}\}}{d^{\text{b}}(s, a)}. \quad (353)$$

Given $\pi_{\theta}^{\star}(\theta | s) = 1$ for all $s \in \mathcal{S}$ and the initial distribution $\rho(0) = 1$, for any $P \in \mathcal{U}^{\sigma}(P^{\theta})$, we arrive at

$$d^{\star, P}(0, \theta) = (1 - \gamma) \rho(0) \sum_{t=0}^{\infty} \gamma^t \left(\underline{P}^{\theta}(0 | 0, \theta) \right)^t$$

$$\stackrel{(i)}{=} (1-\gamma) \sum_{t=0}^{\infty} \gamma^t \underline{p}_{\star}^t = \frac{1-\gamma}{1-\gamma \underline{p}_{\star}} \stackrel{(ii)}{\geq} \frac{1-\gamma}{1-\gamma(1-c_3(1-\gamma))} \geq \frac{1}{2}. \quad (354)$$

where (i) holds by (315) and (ii) follows from (316). In addition, we have

$$d^{\star,P}(0, 1-\theta) = 0 \quad \text{and} \quad \forall s > 1: \quad d^{\star,P}(s, a) = 0, \quad (355)$$

since $\rho(0) = 1$ and state 0 and 1 are absorbing states for all policy and all $P \in \mathcal{U}^{\sigma}(P^{\theta})$.

Armed with the above facts, we observe that

$$\max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star,P}(s, a), \frac{1}{S}\}}{d^b(s, a)} = \max_{s \in \{0,1\}, P \in \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star,P}(s, \theta), \frac{1}{S}\}}{d^b(s, \theta)} \quad (356)$$

which follows from the properties of the optimal policy in (319).

Consequently, we control $C_{\text{rob}}^{\star}$ in states separately:

$$\max_{P \in \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star,P}(0, \theta), \frac{1}{S}\}}{d^b(0, \theta)} \stackrel{(i)}{=} \frac{1}{S d^b(0, \theta)} \stackrel{(ii)}{=} \frac{2}{S \mu(0)} = 2C, \quad (357a)$$

$$\max_{P \in \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star,P}(1, \theta), \frac{1}{S}\}}{d^b(1, \theta)} \leq \frac{1}{S d^b(1, \theta)} \stackrel{(iii)}{=} \frac{2}{S(1 - \frac{1}{CS})} \stackrel{(iv)}{\leq} \frac{4}{S} \stackrel{(v)}{\leq} C, \quad (357b)$$

where (i) holds by (354) and $S \geq 2$, (ii) and (iii) follow from the definitions in (310) and (308), and (iv) and (v) and arise from the assumption in (309). Plugging the above results back into (356) directly completes the proof of

$$C_{\text{rob}}^{\star} = \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^{\sigma}(P^{\theta})} \frac{\min \{d^{\star,P}(s, a), \frac{1}{S}\}}{d^b(s, a)} = 2C.$$

C.4.6 PROOF OF LEMMA 29

For any $\mathcal{M}_{\phi}$ with $\phi \in \{0, 1\}$, we first characterize the robust value function for any policy π over different states. due to state absorbing, the uncertainty set becomes a singleton containing the nominal distribution at state $s = 1$ and $s = 2$. It is easily observed that for any policy π , the robust value functions at state $s = 1$ and $s = 2$ obey

$$V_{\phi}^{\pi, \sigma}(1) = \sum_{t=0}^{\infty} \gamma^t \cdot 0 = 0, \quad (358a)$$

$$V_{\phi}^{\pi, \sigma}(2) = \sum_{t=0}^{\infty} \gamma^t \cdot 1 = \frac{1}{1-\gamma}, \quad (358b)$$

since $r(1, a) = 0$ and $r(2, a) = 1$. In addition, for state $s > 2$, the perturbed transition kernel is supported on itself and state 1, both of which receive a reward of 0 by design (327), leading to

$$V_{\phi}^{\pi, \sigma}(s) = \sum_{t=0}^{\infty} \gamma^t \cdot 0 = 0, \quad \text{for } s > 2. \quad (358c)$$

Moving onto the remaining states, the robust value function of state 0 satisfies

$$\begin{aligned}
 V_{\phi}^{\pi, \sigma}(0) &= \mathbb{E}_{a \sim \pi(\cdot | 0)} \left[r(0, a) + \gamma \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{0,a}^{\phi})} \mathcal{P} V_{\phi}^{\pi, \sigma} \right] \\
 &\stackrel{(i)}{=} 1 + \gamma \pi(\phi | 0) \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{0,\phi}^{\phi})} \mathcal{P} V_{\phi}^{\pi, \sigma} + \gamma \pi(1 - \phi | 0) \inf_{\mathcal{P} \in \mathcal{U}^{\sigma}(P_{0,1-\phi}^{\phi})} \mathcal{P} V_{\phi}^{\pi, \sigma} \\
 &\stackrel{(ii)}{=} 1 + \gamma \pi(\phi | 0) \left[\underline{p} V_{\phi}^{\pi, \sigma}(2) + (1 - \underline{p}) V_{\phi}^{\pi, \sigma}(1) \right] \\
 &\quad + \gamma \pi(1 - \phi | 0) \left[\underline{q} V_{\phi}^{\pi, \sigma}(2) + (1 - \underline{q}) V_{\phi}^{\pi, \sigma}(1) \right] \\
 &\stackrel{(iii)}{=} 1 + \gamma V_{\phi}^{\pi, \sigma}(1) + \gamma z_{\phi}^{\pi} \left[V_{\phi}^{\pi, \sigma}(2) - V_{\phi}^{\pi, \sigma}(1) \right] \\
 &= 1 + \frac{\gamma}{1 - \gamma} z_{\phi}^{\pi},
 \end{aligned} \tag{359}$$

where (i) holds by the reward function defined in (327). To see (ii), note that (358) indicates $V_{\phi}^{\pi, \sigma}(2) \geq V_{\phi}^{\pi, \sigma}(1)$, so that the infimum is obtained by picking the smallest possible mass on the transition to state 2, provided by the definition in (336). Last but not least, (iii) follows by plugging in the definition of z_{ϕ}^{π} in (339), and the last identity is due to (358). Consequently, taking $\pi = \pi_{\phi}^{\star}$, we directly arrive at

$$V_{\phi}^{\star, \sigma}(0) = 1 + \frac{\gamma}{1 - \gamma} z_{\phi}^{\pi^{\star}}. \tag{360}$$

Observing that the function $z_{\frac{\gamma}{1-\gamma}}$ is increasing in z and z_{ϕ}^{π} is also increasing in $\pi(\phi | 0)$ (see the fact $\underline{p} \geq \underline{q}$ in (337)), the optimal policy in state 0 thus obeys

$$\pi_{\phi}^{\star}(\phi | 0) = 1. \tag{361}$$

Finally, plugging the above fact back into (339) leads to

$$z_{\phi}^{\star} := z_{\phi}^{\pi^{\star}} = \underline{p} \pi_{\phi}^{\star}(\phi | 0) + \underline{q} \pi_{\phi}^{\star}(1 - \phi | 0) = \underline{p}, \tag{362}$$

which combined with (360) yields

$$V_{\phi}^{\star, \sigma}(0) = 1 + \frac{\gamma}{1 - \gamma} \underline{p}. \tag{363}$$

Regarding the optimal policy for the remaining states $s > 0$, since the action does not influence the state transition, without loss of generality, we choose the optimal policy to obey

$$\forall s > 0 : \quad \pi_{\phi}^{\star}(\phi | s) = 1. \tag{364}$$

Proof of (341). To begin with, for any MDP $\mathcal{M}_{\phi}$ with $\phi \in \{0, 1\}$, recall the definition of $C_{\text{rob}}^{\star}$ as

$$C_{\text{rob}}^{\star} = \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^{\sigma}(P^{\phi})} \frac{\min \{d^{\star, P}(s, a), \frac{1}{S}\}}{d^b(s, a)}. \tag{365}$$

Given $\pi_\phi^*(\phi | s) = 1$ for all $s \in \mathcal{S}$ and the initial distribution $\rho(0) = 1$, for any $P \in \mathcal{U}^\sigma(P^\phi)$, we arrive at

$$d^{*,P}(0, \phi) = (1 - \gamma)\rho(0)\pi_\phi^*(\phi | 0) = (1 - \gamma), \quad (366)$$

which holds due to that the agent transits from state 0 to other states at the first step and then will never go back to state 0. In addition, one has for any $P \in \mathcal{U}^\sigma(P^\phi)$,

$$\begin{aligned} d^{*,P}(2, \phi) &= (1 - \gamma)P(2 | 0, \phi) \sum_{t=1}^{\infty} \gamma^t (P(2 | 2, \phi))^t \\ &= (1 - \gamma)P(2 | 0, \phi) \sum_{t=1}^{\infty} \gamma^t \stackrel{(i)}{\geq} \gamma \underline{p} \geq \frac{1}{2\beta}, \end{aligned} \quad (367)$$

where (i) holds by (366) and the final inequality follows from (337) and $\gamma \geq 1/2$. Armed with the above facts, we observe that

$$\max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(s, a), \frac{1}{S}\}}{d^b(s, a)} = \max_{s \in \{0,1,2\}, P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(s, \phi), \frac{1}{S}\}}{d^b(s, \phi)} \quad (368)$$

which follows from the properties of the optimal policy in (364) and consequently $d^{*,P}(s) = d^{*,P}(s, \phi) = 0$ for all $s > 2$ and all $P \in \mathcal{U}^\sigma(P^\phi)$.

To continue, we control the term in states $\{0, 1, 2\}$ separately:

$$\max_{P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(2, \phi), \frac{1}{S}\}}{d^b(2, \phi)} \stackrel{(i)}{=} \frac{1}{S d^b(2, \phi)} \stackrel{(ii)}{=} \frac{2}{S \mu(2)} = 2C, \quad (369a)$$

$$\max_{P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(0, \phi), \frac{1}{S}\}}{d^b(0, \phi)} \leq \frac{1}{S d^b(0, \phi)} \stackrel{(iii)}{=} \frac{2}{S \mu(0)} = 2C, \quad (369b)$$

$$\max_{P \in \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(1, \phi), \frac{1}{S}\}}{d^b(1, \phi)} \leq \frac{1}{S d^b(1, \phi)} \stackrel{(iv)}{=} \frac{2}{S(1 - \frac{2}{CS})} \stackrel{(v)}{\leq} \frac{4}{S} \stackrel{(vi)}{\leq} C, \quad (369c)$$

where (i) holds by (367) and $S \geq 2\beta$, (ii), (iii) and (iv) follow from the definitions in (330) and (328), (v) and (vi) arise from the assumption in (329). Plugging the above results back into (368) directly completes the proof of

$$C_{\text{rob}}^* = \max_{(s,a,P) \in \mathcal{S} \times \mathcal{A} \times \mathcal{U}^\sigma(P^\phi)} \frac{\min \{d^{*,P}(s, a), \frac{1}{S}\}}{d^b(s, a)} = 2C.$$

References

- Mohammed Amin Abdullah, Hang Ren, Haitham Bou Ammar, Vladimir Milenkovic, Rui Luo, Mingtian Zhang, and Jun Wang. Wasserstein robust reinforcement learning. *arXiv preprint arXiv:1907.13196*, 2019.
- Alekh Agarwal, Sham Kakade, and Lin F Yang. Model-based reinforcement learning with a generative model is minimax optimal. *Conference on Learning Theory*, pages 67–83, 2020.

- Kishan Panaganti Badrinath and Dileep Kalathil. Robust reinforcement learning using least squares policy iteration with provable performance guarantees. In *International Conference on Machine Learning*, pages 511–520. PMLR, 2021.
- Dimitris Bertsimas, Vishal Gupta, and Nathan Kallus. Data-driven robust optimization. *Mathematical Programming*, 167(2):235–292, 2018.
- Jose Blanchet and Karthyek Murthy. Quantifying distributional model risk via optimal transport. *Mathematics of Operations Research*, 44(2):565–600, 2019.
- Jinglin Chen and Nan Jiang. Information-theoretic considerations in batch reinforcement learning. In *International Conference on Machine Learning*, pages 1042–1051. PMLR, 2019.
- James J Choi, David Laibson, Brigitte C Madrian, and Andrew Metrick. Reinforcement learning and savings behavior. *The Journal of finance*, 64(6):2515–2534, 2009.
- Erick Delage and Yinyu Ye. Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Operations research*, 58(3):595–612, 2010.
- Esther Derman and Shie Mannor. Distributional robustness and regularization in reinforcement learning. *arXiv preprint arXiv:2003.02894*, 2020.
- Wenhao Ding, Laixi Shi, Yuejie Chi, and Ding Zhao. Seeing is not believing: Robust reinforcement learning against spurious correlation. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- John C Duchi. Introductory lectures on stochastic optimization. *The Mathematics of Data*, 25:99–186, 2018.
- John C Duchi and Hongseok Namkoong. Learning models with uniform performance via distributionally robust optimization. *The Annals of Statistics*, 49(3):1378–1406, 2021.
- Rui Gao. Finite-sample guarantees for Wasserstein distributionally robust optimization: Breaking the curse of dimensionality. *Operations Research*, 2022.
- Javier Garcia and Fernando Fernández. A comprehensive survey on safe reinforcement learning. *Journal of Machine Learning Research*, 16(1):1437–1480, 2015.
- Edgar N Gilbert. A comparison of signalling alphabets. *The Bell system technical journal*, 31(3):504–522, 1952.
- Vineet Goyal and Julien Grand-Clement. Robust markov decision processes: Beyond rectangularity. *Mathematics of Operations Research*, 2022.
- Chin Pang Ho, Marek Petrik, and Wolfram Wiesemann. Fast Bellman updates for robust MDPs. In *International Conference on Machine Learning*, pages 1979–1988. PMLR, 2018.
- Chin Pang Ho, Marek Petrik, and Wolfram Wiesemann. Partial policy iteration for l_1 -robust markov decision processes. *Journal of Machine Learning Research*, 22(275):1–46, 2021.

- Linfang Hou, Liang Pang, Xin Hong, Yanyan Lan, Zhiming Ma, and Dawei Yin. Robust reinforcement learning with Wasserstein constraint. *arXiv preprint arXiv:2006.00945*, 2020.
- Zhaolin Hu and L Jeff Hong. Kullback-leibler divergence constrained distributionally robust optimization. *Available at Optimization Online*, pages 1695–1724, 2013.
- Garud N Iyengar. Robust dynamic programming. *Mathematics of Operations Research*, 30(2):257–280, 2005.
- Ying Jin, Zhuoran Yang, and Zhaoran Wang. Is pessimism provably efficient for offline RL? In *International Conference on Machine Learning*, pages 5084–5096, 2021.
- David L Kaufman and Andrew J Schaefer. Robust modified policy iteration. *INFORMS Journal on Computing*, 25(3):396–410, 2013.
- Aviral Kumar, Aurick Zhou, George Tucker, and Sergey Levine. Conservative Q-learning for offline reinforcement learning. In H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems*, volume 33, pages 1179–1191. Curran Associates, Inc., 2020.
- Sergey Levine, Aviral Kumar, George Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- Gen Li, Yuting Wei, Yuejie Chi, Yuantao Gu, and Yuxin Chen. Breaking the sample size barrier in model-based reinforcement learning with a generative model. In *Advances in Neural Information Processing Systems*, volume 33, 2020.
- Gen Li, Yuting Wei, Yuejie Chi, Yuantao Gu, and Yuxin Chen. Sample complexity of asynchronous Q-learning: Sharper analysis and variance reduction. *IEEE Transactions on Information Theory*, 68(1):448–473, 2021.
- Gen Li, Laixi Shi, Yuxin Chen, Yuejie Chi, and Yuting Wei. Settling the sample complexity of model-based offline reinforcement learning. *arXiv preprint arXiv:2204.05275*, 2022.
- Rémi Munos. Error bounds for approximate value iteration. In *Proceedings of the National Conference on Artificial Intelligence*, volume 20, page 1006. Menlo Park, CA; Cambridge, MA; London; AAAI Press; MIT Press; 1999, 2005.
- Arnab Nilim and Laurent El Ghaoui. Robust control of Markov decision processes with uncertain transition matrices. *Operations Research*, 53(5):780–798, 2005.
- Arnab Nilim and Laurent Ghaoui. Robustness in markov decision problems with uncertain transition matrices. *Advances in neural information processing systems*, 16, 2003.
- 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.

- Hamed Rahimian and Sanjay Mehrotra. Distributionally robust optimization: A review. *arXiv preprint arXiv:1908.05659*, 2019.
- Paria Rashidinejad, Banghua Zhu, Cong Ma, Jiantao Jiao, and Stuart Russell. Bridging offline reinforcement learning and imitation learning: A tale of pessimism. *Neural Information Processing Systems (NeurIPS)*, 2021.
- Aurko Roy, Huan Xu, and Sebastian Pokutta. Reinforcement learning under model mismatch. *Advances in neural information processing systems*, 30, 2017.
- John Schulman, Jonathan Ho, Alex X Lee, Ibrahim Awwal, Henry Bradlow, and Pieter Abbeel. Finding locally optimal, collision-free trajectories with sequential convex optimization. In *Robotics: science and systems*, volume 9, pages 1–10. Citeseer, 2013.
- Laixi Shi, Gen Li, Yuting Wei, Yuxin Chen, and Yuejie Chi. Pessimistic Q-learning for offline reinforcement learning: Towards optimal sample complexity. In *Proceedings of the 39th International Conference on Machine Learning*, volume 162, pages 19967–20025. PMLR, 2022.
- 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.
- Laixi Shi, Eric Mazumdar, Yuejie Chi, and Adam Wierman. Sample-efficient robust multi-agent reinforcement learning in the face of environmental uncertainty. *arXiv preprint arXiv:2404.18909*, 2024.
- Aman Sinha, Hongseok Namkoong, and John Duchi. Certifying some distributional robustness with principled adversarial training. In *International Conference on Learning Representations*, 2018.
- Elena Smirnova, Elvis Dohmatob, and Jérémie Mary. Distributionally robust reinforcement learning. *arXiv preprint arXiv:1902.08708*, 2019.
- Jun Song and Chaoyue Zhao. Optimistic distributionally robust policy optimization. *arXiv preprint arXiv:2006.07815*, 2020.
- Richard S Sutton and Andrew G Barto. *Reinforcement learning: An introduction*. MIT press, 2018.
- Aviv Tamar, Shie Mannor, and Huan Xu. Scaling up robust MDPs using function approximation. In *International conference on machine learning*, pages 181–189. PMLR, 2014.
- FLEMMING Topsøe. Some bounds for the logarithmic function. *Inequality theory and applications*, 4:137, 2007.
- A. B. Tsybakov and V. Zaiats. *Introduction to nonparametric estimation*, volume 11. Springer, 2009.

- Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge university press, 2018.
- He Wang, Laixi Shi, and Yuejie Chi. Sample complexity of offline distributionally robust linear markov decision processes. *arXiv preprint arXiv:2403.12946*, 2024.
- Jie Wang, Rui Gao, and Hongyuan Zha. Reliable off-policy evaluation for reinforcement learning. *Operations Research*, 2022.
- Shengbo Wang, Nian Si, Jose Blanchet, and Zhengyuan Zhou. A finite sample complexity bound for distributionally robust Q-learning. In *International Conference on Artificial Intelligence and Statistics*, pages 3370–3398. PMLR, 2023a.
- Shengbo Wang, Nian Si, Jose Blanchet, and Zhengyuan Zhou. Sample complexity of variance-reduced distributionally robust Q-learning. *arXiv preprint arXiv:2305.18420*, 2023b.
- Yue Wang and Shaofeng Zou. Online robust reinforcement learning with model uncertainty. *Advances in Neural Information Processing Systems*, 34, 2021.
- Wolfram Wiesemann, Daniel Kuhn, and Berç Rustem. Robust markov decision processes. *Mathematics of Operations Research*, 38(1):153–183, 2013.
- Eric M Wolff, Ufuk Topcu, and Richard M Murray. Robust control of uncertain markov decision processes with temporal logic specifications. In *2012 IEEE 51st IEEE Conference on Decision and Control (CDC)*, pages 3372–3379. IEEE, 2012.
- Jiin Woo, Laixi Shi, Gauri Joshi, and Yuejie Chi. Federated offline reinforcement learning: Collaborative single-policy coverage suffices. *arXiv preprint arXiv:2402.05876*, 2024.
- Tengyang Xie, Nan Jiang, Huan Wang, Caiming Xiong, and Yu Bai. Policy finetuning: Bridging sample-efficient offline and online reinforcement learning. *Advances in neural information processing systems*, 34, 2021.
- Huan Xu and Shie Mannor. Distributionally robust Markov decision processes. *Mathematics of Operations Research*, 37(2):288–300, 2012.
- Yuling Yan, Gen Li, Yuxin Chen, and Jianqing Fan. The efficacy of pessimism in asynchronous Q-learning. *IEEE Transactions on Information Theory*, 2023.
- Insoon Yang. A convex optimization approach to distributionally robust markov decision processes with Wasserstein distance. *IEEE Control Systems Letters*, 1(1):164–169, 2017.
- 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.
- Ming Yin and Yu-Xiang Wang. Optimal uniform OPE and model-based offline reinforcement learning in time-homogeneous, reward-free and task-agnostic settings. *arXiv preprint arXiv:2105.06029*, 2021.

- Ming Yin, Yu Bai, and Yu-Xiang Wang. Near-optimal offline reinforcement learning via double variance reduction. *Advances in neural information processing systems*, 34, 2021.
- Huan Zhang, Hongge Chen, Chaowei Xiao, Bo Li, Mingyan Liu, Duane Boning, and Chong-Jui Hsieh. Robust deep reinforcement learning against adversarial perturbations on state observations. *Advances in Neural Information Processing Systems*, 33:21024–21037, 2020.
- Zhengqing Zhou, Qinxun Bai, Zhengyuan Zhou, Linhai Qiu, Jose Blanchet, and Peter Glynn. Finite-sample regret bound for distributionally robust offline tabular reinforcement learning. In *International Conference on Artificial Intelligence and Statistics*, pages 3331–3339. PMLR, 2021.