

Policy Teaching via Environment Poisoning: Training-time Adversarial Attacks against Reinforcement Learning

Amin Rakhsha¹ Goran Radanovic¹ Rati Devidze¹ Xiaojin Zhu² Adish Singla¹

Abstract

We study a security threat to reinforcement learning where an attacker poisons the learning environment to force the agent into executing a target policy chosen by the attacker. As a victim, we consider RL agents whose objective is to find a policy that maximizes average reward in undiscounted infinite-horizon problem settings. The attacker can manipulate the rewards or the transition dynamics in the learning environment at training-time and is interested in doing so in a stealthy manner. We propose an optimization framework for finding an optimal stealthy attack for different measures of attack cost. We provide sufficient technical conditions under which the attack is feasible and provide lower/upper bounds on the attack cost. We instantiate our attacks in two settings: (i) an offline setting where the agent is doing planning in the poisoned environment, and (ii) an online setting where the agent is learning a policy using a regret-minimization framework with poisoned feedback. Our results show that the attacker can easily succeed in teaching any target policy to the victim under mild conditions and highlight a significant security threat to reinforcement learning agents in practice.

1. Introduction

Understanding adversarial attacks on learning algorithms is critical to finding security threats against the deployed machine learning systems and in designing novel algorithms robust to those threats. We focus on *training-time* adversarial attacks on learning algorithms, also known as data poisoning attacks (Huang et al., 2011; Biggio & Roli, 2018; Zhu, 2018). Different from *test-time* attacks where the

adversary perturbs test data to change the algorithm’s decisions, poisoning attacks manipulate the training data to change the algorithm’s decision-making policy itself.

Most of the existing work on data poisoning attacks has focused on supervised learning algorithms (Biggio et al., 2012; Mei & Zhu, 2015; Xiao et al., 2015; Alfeld et al., 2016; Li et al., 2016; Koh et al., 2018). In contemporary works, researchers have explored data poisoning attacks against stochastic multi-armed bandits (Jun et al., 2018; Liu & Shroff, 2019) and contextual bandits (Ma et al., 2018), which belong to family of online learning algorithms with limited feedback—such algorithms are widely used in real-world applications such as news article recommendation (Li et al., 2010) and web advertisements ranking (Chapelle et al., 2014). The feedback loop in online learning makes these applications easily susceptible to data poisoning, e.g., attacks in the form of click baits (Miller et al., 2011). In this paper, we focus on data poisoning attacks against reinforcement learning (RL) algorithms, an online learning paradigm for sequential decision-making under uncertainty (Sutton & Barto, 2018).¹ Given that RL algorithms are increasingly used in critical applications, including cyber-physical systems (Li & Qiu, 2019) and personal assistive devices (Rybski et al., 2007), it is of utmost importance to investigate the security threat to RL algorithms against different forms of poisoning attacks.

1.1. Overview of our Results and Contributions

In the following, we discuss a few of the types/dimensions of poisoning attacks in RL in order to highlight the novelty of our work in comparison to existing work.

Type of adversarial manipulation. Existing works on poisoning attacks against RL have studied the manipulation of rewards only (Zhang & Parkes, 2008; Zhang et al., 2009; Ma et al., 2019; Huang & Zhu, 2019). A key novelty of our work is that we study environment poisoning, i.e., manipulating rewards or manipulating transition dynamics.

¹Poisoning attacks is also mathematically equivalent to the formulation of machine teaching with teacher being the adversary (Zhu et al., 2018). However, the problem of designing optimized environments for teaching a target policy to an RL agent is not well-understood in machine teaching.

¹Max Planck Institute for Software Systems (MPI-SWS).
²University of Wisconsin-Madison. Correspondence to: Adish Singla <adishs@mpi-sws.org>.

For certain applications, it is more natural to manipulate the transition dynamics instead of the rewards, such as (i) the inventory management problem setting where state transitions are controlled by demand and supply of products in a market and (ii) conversational agents where the state is represented by the history of conversations. Through our proposed optimization framework, we show that the problem of optimally poisoning transition dynamics is a lot more challenging than that of poisoning rewards, and the attack might not be always feasible. We provide sufficient technical conditions which ensures attacker’s success and provide lower/upper bounds on the attack cost.

Objective of the learning agent. Existing works have focused on studying RL agents that maximizes cumulative reward in *discounted* problem settings. In our work, we consider RL agents that maximizes average reward in *undiscounted* infinite-horizon settings (Puterman, 1994; Mahadevan, 1996)—this is a more suitable objective for many real-world applications that have cyclic tasks or tasks without absorbing states, e.g., inventory management and scheduling problems (Tadepalli et al., 1994; Puterman, 1994), and a robot learning to avoid obstacles (Mahadevan & Connell, 1992). This subtle difference in RL objective leads to technical challenges because of the absence of the contraction property that comes from the usual discount factor (Mahadevan, 1996). The results and bounds we provide are based on several measures of the problem setting including *Hajnal measure* of the transition kernel and *diameter of the Markov chain* induced by target policy in the original unpoisoned environment. Another reason we are considering average reward objective is because there are well-studied online RL algorithms for this objective using regret-minimization framework as discussed next.

Offline planning and online learning. Existing works have focused on attacks in an *offline* setting where the adversary first poisons the reward function in the environment and then the RL agent finds a policy via *planning* (Zhang & Parkes, 2008; Zhang et al., 2009; Ma et al., 2019). In contrast, we call a setting as *online* where the adversary interacts with a *learning* agent to manipulate the feedback signals. One of the key differences in these two settings is in measuring attacker’s cost: The ℓ_∞ -norm of manipulation is commonly studied for the offline setting; for the online setting, the cumulative cost of attack over time (e.g., measured by ℓ_1 -norm of manipulations) is more relevant and has not been studied in literature. We instantiate our attacks in both the settings with appropriate notions of attack cost. For the online learning setting, we consider an agent who is learning a policy using regret-minimization framework (Auer & Ortner, 2007; Jaksch et al., 2010).

We note that our attacks are constructive, and we provide numerical simulations to support our theoretical state-

ments. Our results demonstrate that the attacker can easily succeed in teaching (forcing) the victim to execute the desired target policy at a minimal cost.

2. Environment and RL Agent

We consider a standard RL setting, based on Markov decision processes and RL agents that optimize their expected utility. In contrast to prior work on reward poisoning attacks that assume RL agents are optimizing their total discounted rewards, we focus on the average reward optimality criterion, as specified in more detail in the following subsections.

2.1. Environment, Policy, and Optimality Criteria

The environment is a Markov Decision Process (MDP) defined as $M = (S, A, R, P)$, where S is the state space, A is the action space, $R: S \times A \rightarrow \mathbb{R}$ is the reward function, and $P: S \times A \times S \rightarrow [0, 1]$ is the state transition dynamics, i.e., $P(s, a, s')$ denotes the probability of reaching state s' when taking action a in state s . We denote a deterministic policy by π , which is a map from states to actions, i.e., $\pi: S \rightarrow \mathcal{A}$.

In our work, we consider *average reward* optimality criteria in *undiscounted* infinite-horizon settings (Puterman, 1994; Mahadevan, 1996)—this optimality criteria is particularly suitable for applications that have cyclic tasks or tasks without absorbing states (see discussion in Section 1.1). Given an initial state s , the expected average reward of a policy π in MDP M is given by $\rho(\pi, M, s) := \lim_{N \rightarrow \infty} \frac{1}{N} \mathbb{E} \left[\sum_{t=0}^{N-1} R(s_t, a_t) \mid s_0 = s, \pi \right]$, where the expectation is over the rewards received by the agent when starting from initial state $s_0 = s$ and following the policy π .

Under mild conditions, the above limit exists and the value $\rho(\pi, M, s)$ is independent of the initial state s . In this paper, we will assume the condition that the MDP is *ergodic*, which implies that every policy π has a stationary distribution μ^π , and $\mu^\pi(s) > 0$ for every state s (Puterman, 1994). When ergodicity holds, the average reward or *gain* of a policy can be computed as $\rho(\pi, M) = \sum_{s \in S} \mu^\pi(s) R(s, \pi(s))$; we also denote this as ρ^π when the MDP M is clear from context. For an overview of average-reward RL, we refer the reader to (Puterman, 1994; Mahadevan, 1996; Even-Dar et al., 2005).

A policy π^* is optimal if for every other deterministic policy π we have $\rho^{\pi^*} \geq \rho^\pi$, and ϵ -robust optimal if $\rho^{\pi^*} \geq \rho^\pi + \epsilon$ also holds. Finally, we define a coefficient $\alpha = \min_{s,a,s',a'} \sum_{x \in S} \min(P(s, a, x), P(s', a', x))$. Note that $(1 - \alpha)$ is equivalent to Hajnal measure of P (Puterman, 1994).

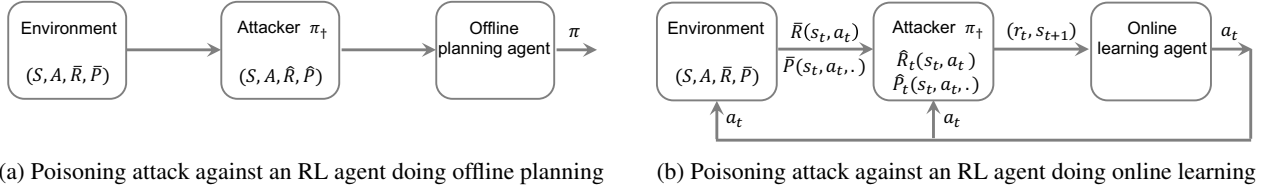


Figure 1: (a) Adversary first poisons the environment by manipulating reward function and transition dynamics, then, the RL agent finds an optimal policy via *planning* algorithms based on Dynamic Programming (Puterman, 1994). (b) Adversary interacts with an RL agent to manipulate the feedback signals; here, we consider an agent who is learning a policy using regret-minimization framework (Auer & Ortner, 2007; Jaksch et al., 2010).

2.2. RL Agent

We consider RL agents with average reward optimality criteria in the following two settings (also, see Figure 1).

Offline planning agent. In the *offline* setting, an RL agent is given an MDP M , and chooses a deterministic optimal policy $\pi^* \in \arg \max_{\pi} \rho(\pi, M)$. The optimal policy can be found via *planning* algorithms based on Dynamic Programming such as value iteration (Puterman, 1994).

Online learning agent. In the *online* setting, an RL agent does not know the MDP M (i.e., R and P are unknown). At each step t , the agent stochastically chooses an action a_t based on the previous observations, and then as feedback it obtains reward r_t and transitions to the next state s_{t+1} . In this paper, we consider an agent who is using a learning algorithm to take actions based on regret-minimization framework. Performance of such a learner in MDP M is measured by its *regret* which after T steps is given by $\text{REGRET}(T, M) = \rho^* \cdot T - \sum_{t=0}^{T-1} r_t$, where $\rho^* := \rho(\pi^*, M)$ is the optimal average reward. We only consider agents with an algorithm promising a sublinear bound for total expected regret, i.e., $\mathbb{E}[\text{REGRET}(T, M)]$ is $o(T)$. We note that well-studied algorithms with sublinear regret exist for average reward criteria, e.g., UCRL algorithm (Auer & Ortner, 2007; Jaksch et al., 2010) and algorithms based on posterior sampling method (Agrawal & Jia, 2017).

3. Attack Models and Problem Formulation

In this section, we formulate the problem of adversarial attacks on the RL agent in both the offline and online settings. In what follows, the original MDP (before poisoning) is denoted by $\bar{M} = (S, A, \bar{R}, \bar{P})$, and an *overline* is added to the corresponding quantities before poisoning, such as $\bar{\rho}^{\pi}$, $\bar{\mu}^{\pi}$, and $\bar{\alpha}$.

The attacker has a target policy $\pi_{\dagger}$ and poisons the environment with the *goal* of teaching/forcing the RL agent to executing this policy.² The attacker is interested in doing a

²Our results can be translated to attacks against the Batch RL

stealthy attack with minimal *cost* to avoid being detected.³ We assume that the attacker knows the original MDP $\bar{M}$, i.e., the original reward function $\bar{R}$ and state transition dynamics $\bar{P}$. This assumption is standard in the existing literature on poisoning attacks against RL. The attacker requires that the RL agent behaves as specified in Section 2, however the attacker doesn't know the agent's algorithm or internal parameters.

3.1. Attack Against an Offline Planning Agent

In attacks against an offline planning agent, the attacker manipulates the original MDP $\bar{M} = (S, A, \bar{R}, \bar{P})$ to a poisoned MDP $\hat{M} = (S, A, \hat{R}, \hat{P})$ which is then used by the RL agent for finding the optimal policy, see Figure 1a.

Goal of the attack. Given a margin parameter ϵ , the attacker poisons the reward function or the transition dynamics so that the target policy $\pi_{\dagger}$ is ϵ -robust optimal in the poisoned MDP $\hat{M}$, i.e., the following condition holds:

$$\rho(\pi_{\dagger}, \hat{M}) \geq \rho(\pi, \hat{M}) + \epsilon, \quad \forall \pi \neq \pi_{\dagger}. \quad (1)$$

Cost of the attack. We consider two notions of costs defined in terms of ℓ_p -norm of differences in reward functions (i.e., $\bar{R}$ and $\hat{R}$) and in transition dynamics (i.e., $\bar{P}$ and $\hat{P}$). For attacks via poisoning rewards, we quantify the cost as

$$\|\hat{R} - \bar{R}\|_p = \left(\sum_{s,a} (|\hat{R}(s,a) - \bar{R}(s,a)|)^p \right)^{1/p},$$

where we treated the reward functions as vectors of length $|S| \cdot |A|$ with values $R(s,a)$. For attacks via poisoning transition dynamics, we quantify the cost as follows:

$$\|\hat{P} - \bar{P}\|_p = \left(\sum_{s,a} \left(\sum_{s'} |\hat{P}(s,a,s') - \bar{P}(s,a,s')| \right)^p \right)^{1/p}.$$

agent studied by (Ma et al., 2019) where the attacker poisons the training data used by the agent to learn MDP parameters.

³Note that the attacks that we will study in subsequent sections will be poisoning either rewards only or transition dynamics only. In a general attack manipulating both the rewards and dynamics simultaneously, one needs to appropriately combine the costs (also, refer to discussion in Section 7).

Here, we have computed ℓ_p -norm of a vector of length $|S| \cdot |A|$ where the value for each (s, a) is given by the ℓ_1 -norm of difference of two distributions $\widehat{P}(s, a, \cdot)$ and $\overline{P}(s, a, \cdot)$.

3.2. Attack Against an Online Learning Agent

In attacks against an online learning agent, the attacker at time t manipulates the reward function $\overline{R}(s_t, a_t)$ and transition dynamics $\overline{P}(s_t, a_t, \cdot)$ for the current state s_t and agent's action a_t , see Figure 1b. Then, at time t , the (poisoned) reward r_t is obtained from $\widehat{R}_t(s_t, a_t)$ instead of $\overline{R}(s_t, a_t)$ and the (poisoned) next state s_{t+1} is sampled from $\widehat{P}_t(s_t, a_t, \cdot)$ instead of $\overline{P}(s_t, a_t, \cdot)$.

Goal of the attack. Specification of the attacker's goal in this online setting is not as straightforward as that in the offline setting, primarily because the agent might never converge to any stationary policy. In our work, at time t when the current state is s_t , we measure the mismatch of agent's action a_t w.r.t. the target policy $\pi_{\dagger}$ as $\mathbb{1}[a_t \neq \pi_{\dagger}(s_t)]$ where $\mathbb{1}[\cdot]$ denotes the indicator function. With this, we define a notion of *average mismatch* of learner's actions in time horizon T as follows:

$$\text{AVGMiss}(T) = \frac{1}{T} \cdot \left(\sum_{t=0}^{T-1} \mathbb{1}[a_t \neq \pi_{\dagger}(s_t)] \right). \quad (2)$$

The goal of the attacker is to ensure that $\text{AVGMiss}(T)$ is $o(1)$, or, alternatively, the total number of time steps where there is a mismatch is $o(T)$.

Cost of the attack. We consider a notion of *average cost* of attack in time horizon T denoted as $\text{AVGCOST}(T)$. For reward poisoning attacks, $\text{AVGCOST}(T)$ is

$$\frac{1}{T} \cdot \left(\sum_{t=0}^{T-1} \left(|\widehat{R}_t(s_t, a_t) - \overline{R}(s_t, a_t)| \right)^p \right)^{1/p},$$

and for dynamics poisoning attacks, $\text{AVGCOST}(T)$ is

$$\frac{1}{T} \cdot \left(\sum_{t=0}^{T-1} \left(\sum_{s'} |\widehat{P}_t(s_t, a_t, s') - \overline{P}(s_t, a_t, s')| \right)^p \right)^{1/p}.$$

Note that here the ℓ_p -norm is defined over a vector of length T with values quantifying the attack cost at each time step t . One of the key differences in measuring attacker's cost for offline and online settings is the use of appropriate norm. While the ℓ_∞ -norm of manipulation is most suitable and commonly studied for the offline setting; for the online setting, the cumulative cost of attack over time measured by ℓ_1 -norm is most relevant.

4. Attacks in Offline Setting

In this section, we introduce and analyze attacks against an offline planning agent that derives its policy using a poisoned MDP $\widehat{M}$. The attacker tries to minimally change the

original MDP $\overline{M}$, while at the same time ensuring that the target policy is optimal in the modified MDP $\widehat{M}$.

4.1. Overview of the Approach

To formulate optimization problems explicitly and provide bounds on the quality of obtainable solutions, we first define concepts that enable us to do the following: i) reduce the complexity of achieving the attacker's goal explained in Section 3.1, ii) quantify how much change is required in rewards or transitions. To find MDP $\widehat{M}$ for which the target policy is ϵ -robust optimal, one could directly utilize constraints expressed by (1). However, the number of constraints in (1) equals $(|A|^{|S|} - 1)$, i.e., it is exponential in $|S|$, making optimization problems that directly utilize them intractable. We show that it is enough to satisfy these constraints for $(|S| \cdot |A| - |S|)$ policies that we call *neighbors* of the target policy and which we define as follows:

Definition 1. For a policy π , its neighbor policy $\pi\{s; a\}$ is defined as

$$\pi\{s; a\}(x) = \begin{cases} \pi(x) & x \neq s \\ a & x = s \end{cases}.$$

The following lemma provides a simple verification criteria for examining whether a policy of interest is ϵ -robust optimal in a given MDP. Its proof can be found in the supplementary material.

Lemma 1. Policy π is ϵ -robust optimal iff we have $\rho^\pi \geq \rho^{\pi\{s; a\}} + \epsilon$ for every state s and action $a \neq \pi(s)$.

In other words, Lemma 1 implies that (sub)optimality of the target policy can be deduced by examining its neighbor policies. By using the definition of average rewards, we conclude that the attacker's strategy, which consists of modifying MDP $\overline{M} = (S, A, \overline{R}, \overline{P})$ to MDP $\widehat{M} = (S, A, \widehat{R}, \widehat{P})$, is successful if and only if for each state s and action $a \neq \pi_{\dagger}(s)$ we have:

$$\begin{aligned} \sum_{s'} \widehat{\mu}^{\pi_{\dagger}}(s') \cdot \widehat{R}(s', \pi_{\dagger}(s')) &\geq \\ \sum_{s'} \widehat{\mu}^{\pi_{\dagger}\{s; a\}}(s') \cdot \widehat{R}(s', \pi_{\dagger}\{s; a\}(s')) &+ \epsilon. \end{aligned} \quad (3)$$

To simplify the exposition of our results, we introduce the following quantity w.r.t. MDP $\overline{M}$ for all $s \in S, a \in A$:

$$\overline{\chi}_\epsilon^\pi(s, a) = \begin{cases} \left[\frac{\overline{\rho}^{\pi\{s; a\}} - \overline{\rho}^\pi + \epsilon}{\overline{\mu}^{\pi\{s; a\}}(s)} \right]^+ & \text{for } a \neq \pi(s), \\ 0 & \text{for } a = \pi(s). \end{cases} \quad (4)$$

Here, $[x]^+$ is equal to $\max\{0, x\}$. As we show in our formal results, the amount of change needed in modifying the original MDP $\overline{M}$ to MDP $\widehat{M}$ so that (3) holds is captured through $\overline{\chi}_\epsilon^\pi$.

4.2. Attacks in Offline Setting via Poisoning Rewards

The first attack we consider is an attack on an offline planning agent via poisoning rewards, where the attacker's strategy consists of choosing a new reward function which makes the target policy ϵ -robust optimal. In this attack we have $\hat{P} = \bar{P}$, and hence $\hat{\mu}^\pi$ is equal to $\bar{\mu}^\pi$ for any policy π , so $\hat{\mu}^\pi$ can be precomputed based on MDP $\bar{M}$. This allows us to formulate an optimization problem for finding an optimal reward poisoning attack using condition (3):

$$\begin{aligned} \min_R \quad & \|R - \bar{R}\|_p \\ \text{s.t.} \quad & \text{condition (3) holds with } \hat{R} = R \text{ and } \hat{\mu} = \bar{\mu}. \end{aligned} \quad (\text{P1})$$

In this optimization problem, the constraints are obtained by using condition (3) for specific instantiation of $\hat{R}$ and $\hat{\mu}$. In particular, we set $\hat{R}$ to be the optimization variable R , and $\hat{\mu}$ to be the stationary distributions of the original transition kernel $\bar{P}$ (i.e., $\bar{\mu}$).

The optimization problem (P1) is a tractable convex program with linear constraints, and we show it is always feasible. Additionally, the following results provide lower and upper bounds on the attack cost, that is, on the cost of solution $\hat{R}$.

Theorem 1. *The optimization problem (P1) is always feasible, and the cost of the optimum solution is bounded by*

$$\frac{\bar{\alpha}}{2} \cdot \|\bar{\chi}_\epsilon^{\pi^\dagger}\|_\infty \leq \|\hat{R} - \bar{R}\|_p \leq \|\bar{\chi}_\epsilon^{\pi^\dagger}\|_p.$$

The proof of the theorem can be found in the supplementary material. The lower bound on the cost of the attack is obtained by extending the proof technique of (Ma et al., 2019) to the average reward criterion, and it depends on two factors. The first factor, $\frac{\bar{\alpha}}{2}$ (see the definition in Section 2.1), is related to the Hajnal measure of the original MDP $\bar{M}$. The second factor, $\|\bar{\chi}_\epsilon^{\pi^\dagger}\|_\infty$, captures the initial disadvantage of the target policy relative to its neighbor policies, as can be seen from the inequality $\bar{\chi}_\epsilon^{\pi^\dagger}(s, a) \geq (\bar{\rho}^{\pi^\dagger\{s;a\}} - \bar{\rho}^{\pi^\dagger})$.

The upper bound is obtained by analyzing the optimum solution to a modified version of the optimization problem (P1) which enforces reward function $\hat{R}$ to be equal to $\bar{R}$ for state-action pairs that define the target policy. We further discuss the modified optimization problem in Section 5.2.

4.3. Attacks in Offline Setting via Poisoning Dynamics

Let us now consider attacks on an offline planning agent based on poisoning transition dynamics. The attacker's strategy now consists of choosing a new transition dynamics $\hat{P}$ that makes the target policy ϵ -robust optimal. Again, we would like to utilize condition (3), however, now we

cannot precompute $\hat{\mu}$ since we are modifying the transition dynamics. We have to explicitly account for that in our optimization problem, and we can do so by noting that stationary distribution $\hat{\mu}^\pi$ satisfies:

$$\hat{\mu}^\pi(s) = \sum_{s'} \hat{P}(s', \pi(s'), s) \cdot \hat{\mu}^\pi(s'). \quad (5)$$

Therefore, by using condition (3) and $\hat{R} = \bar{R}$, we can formulate an optimization problem for finding an optimal dynamics poisoning attack:

$$\begin{aligned} \min_{P, \mu^{\pi^\dagger}, \mu^{\pi^\dagger\{s;a\}}} \quad & \|P - \bar{P}\|_p \\ \text{s.t.} \quad & \mu^{\pi^\dagger} \text{ and } P \text{ satisfy (5),} \\ & \forall s, a \neq \pi^\dagger(s) : \mu^{\pi^\dagger\{s;a\}} \text{ and } P \text{ satisfy (5),} \\ & \text{condition (3) holds with } \hat{R} = \bar{R} \text{ and } \hat{\mu} = \mu, \\ & \forall s, a, s' : P(s, a, s') \geq \delta \cdot \bar{P}(s, a, s'). \end{aligned} \quad (\text{P2})$$

Here, $\delta \in (0, 1]$ in the last set of constraints is a given parameter, specifying how much one is allowed to decrease the original values of transition probabilities. $\delta > 0$ is a regularity condition which ensures that the new MDP is ergodic.⁴ In the supplementary material, we provide a more detailed discussion on δ and how to choose it. Furthermore, the constraints related to condition (3) are obtained by instantiating $\hat{R}$ and $\hat{\mu}$. In particular, $\hat{R}$ is set to be the original reward function $\bar{R}$ and $\hat{\mu}$ is set to be the optimization variables μ (i.e., $\mu^{\pi^\dagger}$ and $\mu^{\pi^\dagger\{s;a\}}$).

Since the first and the second set of constraints are quadratic equality constraints, the optimization problem is in general non-convex. Furthermore, the third set of constraints defined by condition (3) might not even be feasible, for example, when reward function is constant or almost constant (i.e., $\bar{R}(s, a) = \bar{R}(s', a')$ or more generally $|\bar{R}(s, a) - \bar{R}(s', a')| \leq \epsilon$). In Theorem 2 we provide a sufficient condition that renders the optimization problem feasible and we provide bounds on the cost of the attack. Given the nature of the constraints in (P2), we describe an approach for finding an approximate solution in Section 6.

To introduce the formal statement, let us first define relevant quantities. Let $\bar{D}^\pi$ denote the *diameter* of Markov chain induced by policy π in MDP $\bar{M}$, i.e., $\bar{D}^\pi = \max_{s, s'} \bar{T}^\pi(s, s')$, where $\bar{T}^\pi(s, s')$ is the expected time to reach s' starting from s and following policy π . Next we define *value function* of a policy and a few related quantities. Value function of policy π in MDP $\bar{M}$ is defined as

$$\bar{V}^\pi(s) = \lim_{N \rightarrow \infty} \mathbb{E} \left[\sum_{t=0}^{N-1} (\bar{R}(s_t, a_t) - \bar{\rho}_\pi) \mid s_0 = s, \pi \right],$$

⁴This follows because strictly positive trajectory probabilities in MDP $\bar{M}$ remain strictly positive in the new MDP $\hat{M}$, which further implies that all states remain recurrent and aperiodic.

where the expectation is over the rewards received by agent (relative to the offset $\bar{\rho}_\pi$) when starting from initial state $s_0 = s$ and following policy π . Furthermore, let us denote by $\bar{B}^\pi(s) = \sum_{s'} \bar{P}(s, \pi(s), s') \cdot \bar{V}^\pi(s')$ the expected value of the next state given the current state s and policy π . We also define $\bar{V}_{\min}^\pi = \min_s \bar{V}^\pi(s)$ to be the minimum value of $\bar{V}^\pi$, and $sp(\bar{V}^\pi) = \max_s \bar{V}^\pi(s) - \min_s \bar{V}^\pi(s)$ to be the span of $\bar{V}^\pi$. Finally, we introduce the following two quantities important for the theorem statement:

- $\bar{\beta}_\delta^\pi(s, a)$ defined as

$$\bar{\beta}_\delta^\pi(s, a) = \bar{R}(s, \pi(s)) - \bar{R}(s, a) + \bar{B}^\pi(s) - \bar{V}_{\min}^\pi - \delta \cdot sp(\bar{V}^\pi). \quad (6)$$

- $\bar{\Lambda}(s, a)$ defined as

$$\bar{\Lambda}(s, a) = \begin{cases} \frac{\bar{\chi}_0^{\pi^\dagger}(s, a) + \epsilon \cdot (1 + \bar{D}^{\pi^\dagger})}{\bar{\chi}_0^{\pi^\dagger}(s, a) + \bar{\beta}_\delta^\pi(s, a)} & \text{if } \bar{\chi}_\epsilon^{\pi^\dagger}(s, a) > 0, \\ 0 & \text{otherwise,} \end{cases} \quad (7)$$

where $\bar{\chi}_0^{\pi^\dagger}$ is obtained from $\bar{\chi}_\epsilon^{\pi^\dagger}$ by setting $\epsilon = 0$.

Theorem 2. *If there exists a solution $\hat{P}$ to the optimization problem (P2), its cost satisfies*

$$\|\hat{P} - \bar{P}\|_p \cdot \|\bar{V}^{\pi^\dagger}\|_\infty \geq \frac{\delta \cdot \bar{\alpha}}{2} \cdot \|\bar{\chi}_0^{\pi^\dagger}\|_\infty.$$

If for every state s and action a , it holds that either $\bar{\beta}_\delta^{\pi^\dagger}(s, a) \geq \epsilon \cdot (1 + \bar{D}^{\pi^\dagger})$ OR $\bar{\chi}_\epsilon^{\pi^\dagger}(s, a) = 0$, then the optimization problem (P2) has a solution $\hat{P}$ whose cost is upper bounded by $\|\hat{P} - \bar{P}\|_p \leq 2 \cdot \|\bar{\Lambda}\|_p$.

The proof can be found in the supplementary material. The proof technique for the first claim is similar to the one used for proving the lower bound in Theorem 1.

The sufficient condition of Theorem 2 suggests that the initial disadvantage of the target policy $\pi_\dagger$ relative to its neighbor policy $\pi_\dagger\{s; a\}$ (captured by $\bar{\chi}_\epsilon^{\pi^\dagger}(s, a)$) can be overcome if $\bar{\beta}_\delta^{\pi^\dagger}(s, a)$ is large enough. This in turn means that either $\bar{R}(s, \pi_\dagger(s))$ is sufficiently larger than $\bar{R}(s, a)$, or the future prospect of the target policy (i.e., $\bar{B}^{\pi^\dagger}(s)$) is greater than the myopic disadvantage (i.e., $\bar{R}(s, a) - \bar{R}(s, \pi_\dagger(s))$), by a margin which depends on the value function $\bar{V}^{\pi^\dagger}$. The proof for the upper bound is constructive, and it is based on an attack that treats state $s_{\text{sink}} \in \arg \min_{s'} \bar{V}^{\pi^\dagger}(s')$ as a sink state. Then, we shift transitions of state-action pairs that have $\bar{\chi}_\epsilon^{\pi^\dagger}(s, a) > 0$ towards this sink state. Notice that this is a type of an attack that does not alter the transition dynamics for the target policy. We further discuss this type of attacks in Section 5.3.

5. Attacks in Online Setting

We now turn to attacks on an agent that learns over time using the environment feedback. Unlike the planning agent from the previous section, an online learning agent derives its policy from the interaction history, i.e., tuples of the form (s_t, a_t, r_t, s_{t+1}) . To attack an online learning agent, an attacker changes the environment feedback, i.e., reward r_t or state s_{t+1} .

5.1. Overview of the Approach

The underlying idea behind our approach is to utilize the fact that a learning agent has a bounded regret, and thus chooses a suboptimal action a bounded number of times. Hence, to steer a learning agent toward selecting the target policy, it suffices for the attacker to provide the feedback (i.e., reward r_t and the next state s_{t+1}) sampled from an MDP in which the target policy is ϵ -robust optimal. Such an MDP can be derived using the results from the previous sections. We first show that this approach is sound: assuming that an ergodic MDP M has $\pi_\dagger$ as its ϵ -robust optimal policy, the expected number of steps in which a learner whose experience is drawn from MDP M deviates from $\pi_\dagger$ is of order $\mathbb{E}[\text{REGRET}(T, M)]$.

Lemma 2. *Consider an ergodic MDP M that has $\pi_\dagger$ as its ϵ -robust optimal policy, and an online learning agent whose expected regret on an MDP M is $\mathbb{E}[\text{REGRET}(T, M)]$. The average mismatch of the agent is bounded by $\mathbb{E}[\text{AVGMISS}(T)] \leq \frac{1}{T} \cdot K(T, M)$, where*

$$K(T, M) = \frac{\mu_{\max}}{\epsilon} \cdot \left(\mathbb{E}[\text{REGRET}(T, M)] + 2 \|\bar{V}^{\pi^\dagger}\|_\infty \right), \quad (8)$$

with $\mu_{\max} := \max_{s,a} \mu^{\pi_\dagger\{s;a\}}(s)$. Here, μ^π and V^π are respectively the stationary distribution and the value function of a policy π on MDP M .

The proof of the lemma can be found in the supplementary material. The lemma implies that $o(1)$ average mismatch can be achieved in expectation using the sampling based attack described above, assuming that $\hat{M}$ is ergodic and that $\pi_\dagger$ is its ϵ -robust optimal policy. Since the solutions to the optimization problems (P1) and (P2) from Section 4 satisfy this, the lemma applies to them as well.

However, the expected average cost of such an attack could be $\Omega(1)$ (non-diminishing over time) for a learner with sublinear expected regret, unless the original and the sampling MDP have equal rewards and transition probabilities for the state-action pairs of the target policy. Intuitively, if a learner follows the target policy and there exists s for which $\hat{R}(s, \pi_\dagger(s)) \neq \bar{R}(s, \pi_\dagger(s))$ or $\hat{P}(s, \pi_\dagger(s), \cdot) \neq \bar{P}(s, \pi_\dagger(s), \cdot)$, then the attacker would incur a non-zero cost whenever the learner visits s . To avoid this issue, we need

to enforce constraints on the sampling MDP specifying that the attack does not alter rewards and transitions that correspond to the state-action pairs of the target policy. This brings us to the following template that we utilize for attacks on an online learner:

- Modify the optimization problems (P1) and (P2) by respectively adding constraints $\hat{R}(s, \pi_{\dagger}(s)) = \bar{R}(s, \pi_{\dagger}(s))$ and $\hat{P}(s, \pi_{\dagger}(s), s') = \bar{P}(s, \pi_{\dagger}(s), s')$.
- Obtain the sampling MDP $\widehat{M}$ by solving the more constrained version of (P1) or (P2).
- Use the sampling MDP $\widehat{M}$ instead of the environment $\bar{M}$ during the learning process, i.e., obtain r_t from $\hat{R}(s_t, a_t)$ in the case of poisoning rewards and $s_{t+1} \sim \hat{P}(s_t, a_t, \cdot)$ in the case of poisoning dynamics (see Figure 1b).

5.2. Attacks in Online Setting via Poisoning Rewards

For the case of attacks via poisoning rewards, the sampling MDP $\widehat{M}$ differs from the original MDP $\bar{M}$ only in its reward function $\hat{R}$. Following the attack template from the above, we first find $\hat{R}$ using a modified version of (P1) which ensures that $\hat{R}(s, \pi_{\dagger}(s)) = \bar{R}(s, \pi_{\dagger}(s))$, i.e.:

$$\begin{aligned} \min_R \quad & \|R - \bar{R}\|_p \\ \text{s.t.} \quad & \forall s : R(s, \pi_{\dagger}(s)) = \bar{R}(s, \pi_{\dagger}(s)), \\ & \text{all constraints from problem (P1)}. \end{aligned} \quad (\text{P3})$$

As we show in the supplementary material, the optimal solution to the optimization problem (P3) is $\hat{R}(s, a) = \bar{R}(s, a) - \bar{\chi}_{\epsilon}^{\pi_{\dagger}}(s, a)$. Since $\bar{\chi}_{\epsilon}^{\pi_{\dagger}}(s, \pi_{\dagger}(s)) = 0$, using the definition of $\text{AVGCOST}(T)$ we conclude that an upper bound on $\mathbb{E}[\text{AVGCOST}(T)]$ could be computed using (i) the number of times that a non-target action is selected and (ii) the maximum value of $\bar{\chi}_{\epsilon}^{\pi_{\dagger}}(s, a)$. The quantity in (i) is specified in the result of Lemma 2, which brings us to the main result of this subsection.

Theorem 3. *Let $\hat{R}$ be the optimal solution to (P3). Consider the attack defined by r_t obtained from $\hat{R}(s_t, a_t)$ and $s_{t+1} \sim \hat{P}(s_t, a_t, \cdot)$, and an online learning agent whose expected regret on an MDP M is $\mathbb{E}[\text{REGRET}(T, M)]$. The average mismatch of the learner is in expectation upper bounded by*

$$\mathbb{E}[\text{AVGMIS}(T)] \leq \frac{K(T, \widehat{M})}{T},$$

where K is defined in (8). Furthermore, the average attack cost is in expectation upper bounded by

$$\mathbb{E}[\text{AVGCOST}(T)] \leq \|\bar{\chi}_{\epsilon}^{\pi_{\dagger}}\|_{\infty} \cdot \frac{(K(T, \widehat{M}))^{1/p}}{T}.$$

5.3. Attacks in Online Setting via Poisoning Dynamics

For the case of attacks via dynamics poisoning, the sampling MDP $\widehat{M}$ differs from the original MDP $\bar{M}$ in its transition dynamics $\hat{P}$. Following the attack template, we can derive the optimization problem for finding $\hat{P}$ by simply adding the constraints $P(s, \pi_{\dagger}(s), s') = \bar{P}(s, \pi_{\dagger}(s), s')$ in the optimization problem (P2), i.e.:

$$\begin{aligned} \min_{P, \mu^{\pi_{\dagger}}, \mu^{\pi_{\dagger}\{s;a\}}} \quad & \|P - \bar{P}\|_p \\ \text{s.t.} \quad & \forall s, s' : P(s, \pi_{\dagger}(s), s') = \bar{P}(s, \pi_{\dagger}(s), s'), \\ & \text{all constraints from problem (P2)}. \end{aligned} \quad (\text{P4})$$

In the supplementary material, we show that the additional constraint allows us to transform (P4) into a tractable convex program with linear constraints. In fact, the convex program has a specific structure so that its optimal solution can be obtained by solving $|S| \cdot (|A| - 1)$ simple convex problems—each of these simpler problems only involves $|S|$ variables and $|S| + 1$ linear constraints. Moreover, the convex program is feasible if the conditions of Theorem 2 are met, and its solutions satisfy $\|\hat{P}(s, a, \cdot) - \bar{P}(s, a, \cdot)\|_1 \leq 2 \cdot \bar{\Lambda}(s, a)$, where $\bar{\Lambda}$ is defined in (7). Since $\bar{\Lambda}(s, \pi_{\dagger}(s)) = 0$, this means that one can bound $\mathbb{E}[\text{AVGCOST}(T)]$ based on the number of times that a non-target action is selected and the maximum value of $\bar{\Lambda}(s, a)$. Combining this insight with Lemma 2, we obtain the following theorem.

Theorem 4. *Assume that the sufficient condition of Theorem 2 holds, and let $\hat{P}$ be the optimal solution to (P4). Consider the attack defined by r_t obtained from $\bar{R}(s_t, a_t)$ and $s_{t+1} \sim \hat{P}(s_t, a_t, \cdot)$, and an online learning agent whose expected regret on an MDP M is $\mathbb{E}[\text{REGRET}(T, M)]$. The average mismatch of the learner is in expectation upper bounded by*

$$\mathbb{E}[\text{AVGMIS}(T)] \leq \frac{K(T, \widehat{M})}{T},$$

where K is defined in (8). Furthermore, the expected attack cost is upper bounded by

$$\mathbb{E}[\text{AVGCOST}(T)] \leq 2 \cdot \|\bar{\Lambda}\|_{\infty} \cdot \frac{(K(T, \widehat{M}))^{1/p}}{T},$$

where $\bar{\Lambda}$ is defined in (7).

A direct consequence of Theorem 3 and Theorem 4 is that for a learner with a sublinear expected regret, the average mismatch and the average cost are expected to decrease over time as $\mathcal{O}\left(\frac{\mathbb{E}[\text{REGRET}(T, \widehat{M})]}{T}\right)$ and $\mathcal{O}\left(\frac{\mathbb{E}[\text{REGRET}(T, \widehat{M})]^{1/p}}{T}\right)$. Note that while we considered ℓ_p

norms with $p \geq 1$ to define the attack cost, the above results can be generalized to include the case of $p = 0$. For $p = 0$, the expected value of the average cost is simply upper bounded by $\frac{K(T, \bar{M})}{T}$, and this follows from Lemma 2.

6. Numerical Simulations

We perform numerical simulations on an environment represented as an MDP with four states and two actions, see Figure 2 for details. Even though simple, this environment provides a very rich and an intuitive problem setting to validate the theoretical statements and understand the effectiveness of the attacks by varying different parameters. We also vary the number of states in the MDP to check the efficiency of solving different optimization problems, and report run times. We further extend our experimental evaluation in the supplementary material and report results on an additional environment.

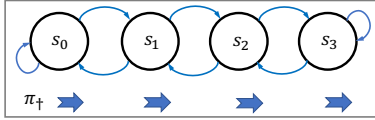


Figure 2: The environment has $|S| = 4$ states and $|A| = 2$ actions given by $\{\text{left}, \text{right}\}$. The original reward function $\bar{R}$ is action independent and has the following values: s_1 and s_2 are rewarding states with $\bar{R}(s_1, \cdot) = \bar{R}(s_2, \cdot) = 0.5$; state s_3 has negative reward of $\bar{R}(s_3, \cdot) = -0.5$, and the reward of the state s_0 given by $\bar{R}(s_0, \cdot)$ will be varied in experiments. With probability 0.9, the actions succeed in navigating the agent to left or right as shown on arrows; with probability 0.1 the agent’s next state is sampled randomly from the set S . The target policy $\pi_{\dagger}$ is to take `right` action in all states as shown in the illustration.

6.1. Attacks in the Offline Setting: Setup and Results

For the offline setting, we considered the following attack strategies: (i) RATTACK: reward attacks using $\hat{R}$ as solution to problem (P1) (ℓ_p -norm with $p = 1, 2, \infty$), (ii) NT-RATTACK: reward attacks using $\hat{R}$ as solution to problem (P3) (solution is independent of ℓ_p -norm), (iii) DATTACK: dynamics attacks using $\hat{P}$ as a solution to problem (P2) (ℓ_p -norm for $p = 1, 2, \infty$), and (iv) NT-DATTACK: dynamics attacks using $\hat{P}$ as a solution to problem (P4) (solution is independent of ℓ_p -norm). The regularity parameter δ in the problems for solving dynamic poisoning attacks is set to 0.0001.⁵

We note that optimal solutions to the problems (P1), (P3), and (P4) can be computed efficiently using standard optimization techniques (also, refer to discussions in Section 4

⁵Here, NT- prefix is used to highlight that *non-target only* manipulations are allowed in problems (P3) and (P4).

and Section 5). Problem (P2) is computationally more challenging, and we provide a simple yet effective approach towards finding an approximate solution by iteratively solving the problem (P4) as follows: First, we use a simple heuristic to obtain a pool of transition kernels $\tilde{P}$ by perturbations of $\bar{P}$ that increase the average reward of $\pi_{\dagger}$, and as second step, we use $\tilde{P}$ ’s from this pool as as input to problem (P4) instead of $\bar{P}$. So, the runtime of solving problem (P2) depends on the number of iterations we invoke problem (P4) internally. The implementation details and code are provided in supplementary materials.

We vary $\bar{R}(s_0, \cdot) \in [-5, 5]$ and vary ϵ margin $\in [0, 1]$. We use ℓ_{∞} -norm in the measure of attack cost (see Section 3.1). The results are reported as an average of 10 runs (here, DATTACK is the only stochastic algorithm because of the random initialization as discussed above).

There are two key points we want to highlight in Figure 3. First, as we increase ϵ margin, the attack problem becomes more difficult: While the reward poisoning attacks are always feasible (though with increasing attack cost), it is infeasible to do dynamics poisoning attacks for $\epsilon > 0.85$. Second, the plots also show that the solution to the generic problems (RATTACK and DATTACK for any ℓ_p -norm with $p = 1, 2, \infty$) can have much lower cost compared to solutions obtained by problems allowing non-target only manipulations (NT-RATTACK and NT-DATTACK).

For the MDP in Figure 2 with $|S| = 4$, the run times for solving problems (P1), (P2), (P3), and (P4) are roughly 0.0110s, 3.3010s, 0.0003s, and 0.0339s, respectively. We further ran experiments on a variant of the MDP with $|S| = 100$ (keeping the same linear chain structure as in Figure 2), and this led to average run times of 0.8326s, 157.2484s, 0.6153s, and 1.1705s, respectively.

6.2. Attacks in the Online Setting: Setup and Results

For the online setting, we considered the following attack strategies: (i) reward attacks using RATTACK (with ℓ_{∞} , ℓ_1 -norm) and NT-RATTACK; (ii) dynamics attacks using DATTACK (with ℓ_{∞} , ℓ_1 -norm) and NT-DATTACK; and (iii) a default setting without adversary denoted as NONE where environment feedback is sampled from the original MDP $\bar{M}$.

In the experiments, we fix $\bar{R}(s_0, \cdot) = -2.5$ and $\epsilon = 0.1$; we plot the measure of the attacker’s achieved goal in terms of AVGMIS and attacker’s cost in terms of AVGCOST for ℓ_1 -norm measured over time t (see Section 3.2). We consider an RL agent implementing the UCRL learning algorithm (Auer & Ortner, 2007), however the attacker does not use any knowledge of the agent’s learning algorithm. The results are reported as an average of 20 runs.

The results in Figure 4 show that our proposed online at-

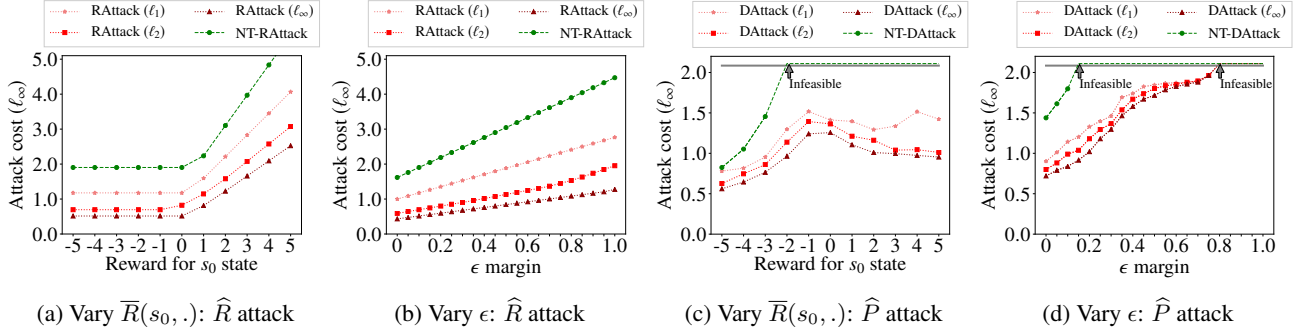


Figure 3: Results for poisoning attacks in the offline setting from Section 4. (a, b) plots show results for attack on rewards and (c, d) plots show results for attack on dynamics. Details are in Section 6.1.

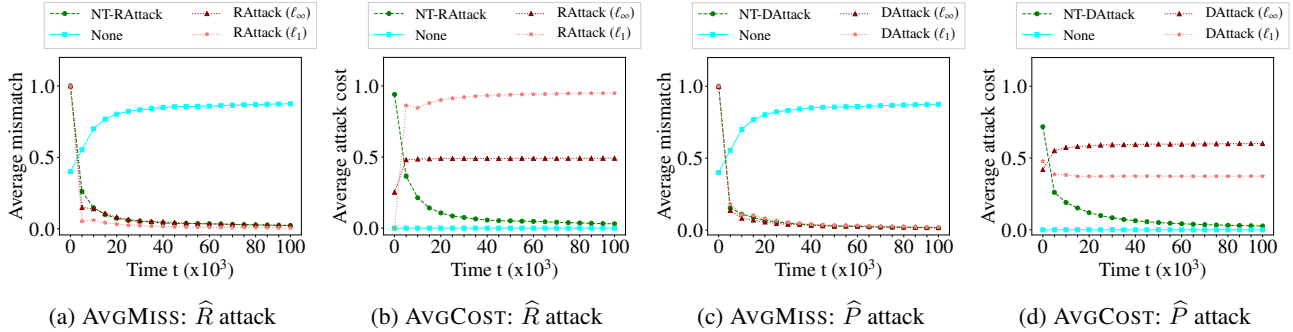


Figure 4: Results for poisoning attacks in the online setting from Section 5. (a, b) plots show results for attack on rewards and (c, d) plots show results for attack on dynamics. Details are in Section 6.2.

attacks with NT-RATTACK and NT-DATTACK are highly effective: learner is forced to follow the target policy while the attacker’s average cost is $o(1)$ (see Theorems 3, 4). In contrast, we can see that the online attacks with RATTACK and DATTACK lead to high cost for the attacker, i.e., the cumulative cost is linear w.r.t. time as anticipated in Section 5.1 (see discussions following Lemma 2).

7. Conclusion

We studied a security threat to reinforcement learning (RL) where an attacker poisons the environment, thereby forcing the agent into executing a target policy. Our work provides theoretical underpinnings of environment poisoning against RL along several new attack dimensions, including (i) adversarial manipulation of the transition dynamics, (ii) attack against RL agents maximizing average reward in undiscounted infinite horizon, and (iii) analyzing different attack costs for offline planning and online learning settings.

There are several promising directions for future work. These include expanding the attack models (e.g., attacking rewards and transitions simultaneously) and broadening the set of attack goals (e.g., under partial specification of target policy). At the same time, relaxing the assumptions

on the attacker knowledge of the underlying MDP could lead to more robust attack strategies. Another interesting future direction would be to make the studied attack models more scalable, e.g., applicable to continuous and large environments. Another interesting topic would be to devise attack strategies against RL agents that use transfer learning approaches, especially in multi-agent RL systems, see (Da Silva & Costa, 2019).

While the paper provides a separate treatment for reward and transitions attack models, simultaneously attacking rewards and transitions might lead to more cost-effective solutions. One possible way to formulate the optimization problem for the simultaneous attack model is to define the cost function as a weighted sum of the cost functions used in (P1) and (P2), and combine the corresponding constraints.

While the experimental results demonstrate the effectiveness of the studied attack models, they do not reveal which types of learning algorithms are most vulnerable to the attack strategies studied in the paper. Further experimentation using a diverse set of the state of the art learning algorithms could reveal this, and provide some guidance in designing defensive strategies and novel RL algorithms robust to manipulations.

Acknowledgements

Xiaojin Zhu is supported in part by NSF 1545481, 1623605, 1704117, 1836978 and the MADLab AF Center of Excellence FA9550-18-1-0166.

References

- Agrawal, S. and Jia, R. Optimistic posterior sampling for reinforcement learning: worst-case regret bounds. In *Advances in Neural Information Processing Systems*, 2017.
- Alfeld, S., Zhu, X., and Barford, P. Data poisoning attacks against autoregressive models. In *AAAI*, 2016.
- Asmuth, J., Littman, M. L., and Zinkov, R. Potential-based shaping in model-based reinforcement learning. In *AAAI*, pp. 604–609, 2008.
- Auer, P. and Ortner, R. Logarithmic online regret bounds for undiscounted reinforcement learning. In *Advances in Neural Information Processing Systems*, pp. 49–56, 2007.
- Biggio, B. and Roli, F. Wild patterns: Ten years after the rise of adversarial machine learning. *Pattern Recognition*, 84:317–331, 2018.
- Biggio, B., Nelson, B., and Laskov, P. Poisoning attacks against support vector machines. In *ICML*, 2012.
- Brown, D. S. and Niekum, S. Machine teaching for inverse reinforcement learning: Algorithms and applications. In *AAAI*, pp. 7749–7758, 2019.
- Cakmak, M., Lopes, M., et al. Algorithmic and human teaching of sequential decision tasks. In *AAAI*, 2012.
- Chapelle, O., Manavoglu, E., and Rosales, R. Simple and scalable response prediction for display advertising. *ACM TIST*, 5(4):61:1–61:34, 2014.
- Chen, Y., Singla, A., Mac Aodha, O., Perona, P., and Yue, Y. Understanding the role of adaptivity in machine teaching: The case of version space learners. In *NeurIPS*, pp. 1483–1493, 2018.
- Da Silva, F. L. and Costa, A. H. R. A survey on transfer learning for multiagent reinforcement learning systems. *Journal of Artificial Intelligence Research*, 64:645–703, 2019.
- Devidze, R., Mansouri, F., Haug, L., Chen, Y., and Singla, A. Understanding the power and limitations of teaching with imperfect knowledge. In *IJCAI*, pp. 2647–2654, 2020.
- Durrett, R. and Durrett, R. *Essentials of stochastic processes*, volume 1. Springer, 1999.
- Even-Dar, E., Kakade, S. M., and Mansour, Y. Experts in a markov decision process. In *NIPS*, pp. 401–408, 2005.
- Goldman, S. A. and Kearns, M. J. On the complexity of teaching. *Journal of Computer and System Sciences*, 50(1):20–31, 1995.
- Hadfield-Menell, D., Russell, S. J., Abbeel, P., and Dragan, A. Cooperative inverse reinforcement learning. In *Advances in Neural Information Processing Systems*, pp. 3909–3917, 2016.
- Haug, L., Tschitschek, S., and Singla, A. Teaching inverse reinforcement learners via features and demonstrations. In *NeurIPS*, 2018.
- Huang, L., Joseph, A. D., Nelson, B., Rubinstein, B. I., and Tygar, J. D. Adversarial machine learning. In *Proceedings of the 4th ACM workshop on Security and artificial intelligence*, pp. 43–58, 2011.
- Huang, S., Papernot, N., Goodfellow, I., Duan, Y., and Abbeel, P. Adversarial attacks on neural network policies. *arXiv preprint arXiv:1702.02284*, 2017.
- Huang, Y. and Zhu, Q. Deceptive reinforcement learning under adversarial manipulations on cost signals. In *GameSec*, pp. 217–237, 2019.
- Jaksch, T., Ortner, R., and Auer, P. Near-optimal regret bounds for reinforcement learning. *Journal of Machine Learning Research*, 11(Apr):1563–1600, 2010.
- Jun, K., Li, L., Ma, Y., and Zhu, X. J. Adversarial attacks on stochastic bandits. In *NeurIPS*, pp. 3644–3653, 2018.
- Kamalaruban, P., Devidze, R., Cevher, V., and Singla, A. Interactive teaching algorithms for inverse reinforcement learning. In *IJCAI*, pp. 2692–2700, 2019.
- Koh, P. W., Steinhardt, J., and Liang, P. Stronger data poisoning attacks break data sanitization defenses. *CoRR*, abs/1811.00741, 2018.
- Li, B., Wang, Y., Singh, A., and Vorobeychik, Y. Data poisoning attacks on factorization-based collaborative filtering. In *Advances in neural information processing systems*, pp. 1885–1893, 2016.
- Li, C. and Qiu, M. *Reinforcement Learning for Cyber-Physical Systems: with Cybersecurity Case Studies*. Chapman and Hall/CRC, 2019.
- Li, L., Chu, W., Langford, J., and Schapire, R. E. A contextual-bandit approach to personalized news article recommendation. In *WWW*, pp. 661–670, 2010.
- Lin, Y., Hong, Z., Liao, Y., Shih, M., Liu, M., and Sun, M. Tactics of adversarial attack on deep reinforcement learning agents. In *IJCAI*, pp. 3756–3762, 2017.

- Liu, F. and Shroff, N. B. Data poisoning attacks on stochastic bandits. In *ICML*, pp. 4042–4050, 2019.
- Ma, Y., Jun, K., Li, L., and Zhu, X. Data poisoning attacks in contextual bandits. In *GameSec*, pp. 186–204, 2018.
- Ma, Y., Zhang, X., Sun, W., and Zhu, J. Policy poisoning in batch reinforcement learning and control. In *NeurIPS*, pp. 14543–14553, 2019.
- Mahadevan, S. Average reward reinforcement learning: Foundations, algorithms, and empirical results. *Machine Learning*, 22(1-3):159–195, 1996.
- Mahadevan, S. and Connell, J. Automatic programming of behavior-based robots using reinforcement learning. *Artif. Intell.*, 55(2):311–365, 1992.
- Mansouri, F., Chen, Y., Vartanian, A., Zhu, J., and Singla, A. Preference-based batch and sequential teaching: Towards a unified view of models. In *NeurIPS*, pp. 9195–9205, 2019.
- Mei, S. and Zhu, X. Using machine teaching to identify optimal training-set attacks on machine learners. In *AAAI*, pp. 2871–2877, 2015.
- Miller, B., Pearce, P., Grier, C., Kreibich, C., and Paxson, V. Whats clicking what? techniques and innovations of todays clickbots. In *International Conference on Detection of Intrusions and Malware, and Vulnerability Assessment*, pp. 164–183. Springer, 2011.
- Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., Graves, A., Riedmiller, M., Fidjeland, A. K., Ostrovski, G., et al. Human-level control through deep reinforcement learning. *Nature*, 518(7540):529–533, 2015.
- Ng, A. Y., Harada, D., and Russell, S. Policy invariance under reward transformations: Theory and application to reward shaping. In *ICML*, 1999.
- Osa, T., Pajarinen, J., Neumann, G., Bagnell, J. A., Abbeel, P., Peters, J., et al. An algorithmic perspective on imitation learning. *Foundations and Trends® in Robotics*, 7(1-2):1–179, 2018.
- Peltola, T., Çelikok, M. M., Dae, P., and Kaski, S. Machine teaching of active sequential learners. In *NeurIPS*, 2019.
- Puterman, M. L. *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. John Wiley & Sons, Inc., 1st edition, 1994. ISBN 0471619779.
- Rybski, P. E., Yoon, K., Stolarz, J., and Veloso, M. M. Interactive robot task training through dialog and demonstration. In *Proceedings of the International Conference on Human-robot interaction*, pp. 49–56, 2007.
- Schulman, J., Levine, S., Abbeel, P., Jordan, M., and Moritz, P. Trust region policy optimization. In *ICML*, pp. 1889–1897, 2015.
- Singla, A., Bogunovic, I., Bartók, G., Karbasi, A., and Krause, A. On actively teaching the crowd to classify. In *NIPS Workshop on Data Driven Education*, 2013.
- Singla, A., Bogunovic, I., Bartók, G., Karbasi, A., and Krause, A. Near-optimally teaching the crowd to classify. In *ICML*, 2014.
- Sutton, R. S. and Barto, A. G. *Reinforcement learning: An introduction*. MIT press, 2018.
- Tadepalli, P., Ok, D., et al. H-learning: A reinforcement learning method to optimize undiscounted average reward. 1994.
- Tretschk, E., Oh, S. J., and Fritz, M. Sequential attacks on agents for long-term adversarial goals. *arXiv preprint arXiv:1805.12487*, 2018.
- Tschiatschek, S., Ghosh, A., Haug, L., Devidze, R., and Singla, A. Learner-aware teaching: Inverse reinforcement learning with preferences and constraints. In *NeurIPS*, 2019.
- Walsh, T. J. and Goschin, S. Dynamic teaching in sequential decision making environments. In *UAI*, pp. 863–872, 2012.
- Xiao, H., Biggio, B., Brown, G., Fumera, G., Eckert, C., and Roli, F. Is feature selection secure against training data poisoning? In *ICML*, pp. 1689–1698, 2015.
- Zhang, H. and Parkes, D. C. Value-based policy teaching with active indirect elicitation. In *AAAI*, 2008.
- Zhang, H., Parkes, D. C., and Chen, Y. Policy teaching through reward function learning. In *EC*, 2009.
- Zhu, X. Machine teaching: An inverse problem to machine learning and an approach toward optimal education. In *AAAI*, pp. 4083–4087, 2015.
- Zhu, X. An optimal control view of adversarial machine learning. *CoRR*, abs/1811.04422, 2018.
- Zhu, X., Singla, A., Zilles, S., and Rafferty, A. N. An overview of machine teaching. *CoRR*, abs/1801.05927, 2018.

A. List of Appendices

In this section we provide a brief description of the content provided in the appendices of the paper.

- Appendix B contains additional related work.
- Appendix C provides implementation details and additional experimental evaluation on a different environment.
- Appendix D introduces some useful quantities and a lemma which are useful for proofs.
- Appendix E contains proof of Lemma 1 and some general results used in future proofs.
- Appendix F contains proof of Theorem 1 for offline attacks via poisoning rewards.
- Appendix G contains proof of Theorem 2 for offline attacks via poisoning dynamics and related discussions.
- Appendix H contains proof of Lemma 2.
- Appendix I contains proof of Theorem 3 for online attacks via poisoning rewards.
- Appendix J contains proof of Theorem 4 for online attacks via poisoning dynamics.

B. Additional Related Work

Test-time attacks against RL. A growing body of contemporary works have studied test-time attacks against RL, in particular, on RL algorithms with neural network policies (Mnih et al., 2015; Schulman et al., 2015). These attacks are typically done by adding noise in the observed state (e.g., a camera image) to fool the neural network policy into taking malicious actions (Huang et al., 2017; Lin et al., 2017; Tretschk et al., 2018). Different attack goals have been considered in these works, e.g., guiding the agent to some adversarial states or forcing agent to take actions that maximizes adversary’s own rewards. Our work is technically quite different and is focused on training-time attacks where the goal is to force the agent to learn a target policy.

Teaching an RL agent. Poisoning attacks is mathematically equivalent to the formulation of machine teaching with teacher being the adversary (Goldman & Kearns, 1995; Singla et al., 2013; 2014; Zhu, 2015; Zhu et al., 2018; Chen et al., 2018; Mansouri et al., 2019; Peltola et al., 2019; Devidze et al., 2020). In particular, there have been a number of recent works on teaching an RL agent via providing an optimized curriculum of demonstrations (Cakmak et al., 2012; Walsh & Goschin, 2012; Hadfield-Menell et al., 2016; Haug et al., 2018; Kamalaruban et al., 2019; Tschitschek et al., 2019; Brown & Niekum, 2019). However, these works have focused on imitation-learning based RL agents who learn from provided demonstrations without any reward feedback (Osa et al., 2018). Given that we consider RL agents who find policies based on rewards, our work is technically very different from theirs. There is also a related literature on changing the behavior of an RL agent via *reward shaping* (Ng et al., 1999; Asmuth et al., 2008); here the reward function is changed to only speed up the convergence of the learning algorithm while ensuring that the optimal policy in the modified environment is unchanged.

C. Implementation Details and Additional Experiments (Section 6)

In this section, we provide implementation details and report experimental results on a different environment.⁶

C.1. Implementation Details

The optimal solutions to the problems (P1), (P3), and (P4) can be computed efficiently using standard optimization techniques. In the source code, the solvers for these optimization problems are implemented as the following functions:

- problem (P1) in the function `general_attack_on_reward()`, see `teacher.py`
- problem (P3) in the function `non_target_attack_on_reward()`, see `teacher.py`
- problem (P4) in the function `non_target_attack_on_dynamics()`, see `teacher.py`

Our approach for solving problem (P2) is implemented in the function `general_attack_on_dynamics()`, see `teacher.py`. Details are provided below:

- As a first step, we use a simple heuristic to obtain a pool of transition kernels $\tilde{P}$ by perturbations of $\bar{P}$ that increase the average reward of $\pi_{\dagger}$. Here, the transition kernel $\tilde{P}$ differs from $\bar{P}$ only for the actions taken by the target policy, i.e., $(s, \pi_{\dagger}(s)) \forall s \in S$. This pool is created in the function `generate_pool()`.

⁶The source code is available at https://github.com/adishs/icml2020_rl-policy-teaching_code.

- As the second step, we take each of $\tilde{P}$'s from this pool as input to problem (P4) instead of $\bar{P}$ which in turn gives us a corresponding pool of solutions. Note here the problem (P4) modifies only the actions which are not taken by the target policy, i.e., $(s, a) \forall s \in S, a \neq \pi_{\dagger}(s)$. This is done in the function `solve_pool()`.
- As the final step, we pick a solution from this pool of solutions with the minimal cost. This is done in the function `get_P_with_smallest_norm()`.

C.2. Additional Experiments

We perform additional numerical simulations on an environment represented as an MDP with nine states and two actions per state, see Figure 5 for details. This environment is inspired from a navigation task and is slightly more complex than the environment in Figure 2.

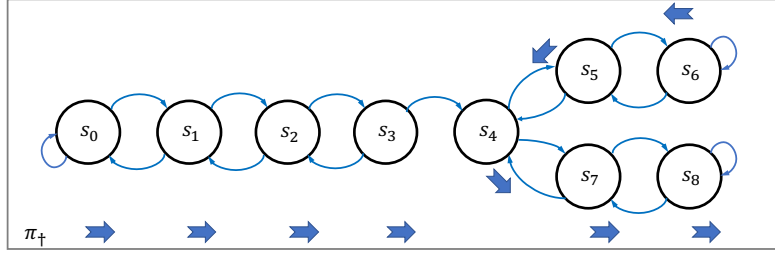


Figure 5: The environment has $|S| = 9$ states and $|A| = 2$ actions per state as illustrated. The original reward function $\bar{R}$ is action independent and has the following values: $\bar{R}(s_1, \cdot) = \bar{R}(s_2, \cdot) = \bar{R}(s_3, \cdot) = -2.5$, $\bar{R}(s_4, \cdot) = \bar{R}(s_5, \cdot) = 1.0$, $\bar{R}(s_6, \cdot) = \bar{R}(s_7, \cdot) = \bar{R}(s_8, \cdot) = 0$, and the reward of the state s_0 given by $\bar{R}(s_0, \cdot)$ will be varied in experiments. With probability 0.9, the actions succeed in navigating the agent as shown on arrows; with probability 0.1 the agent's next state is sampled randomly from the set S . The target policy $\pi_{\dagger}$ is to take actions as shown with bold arrows in the illustration.

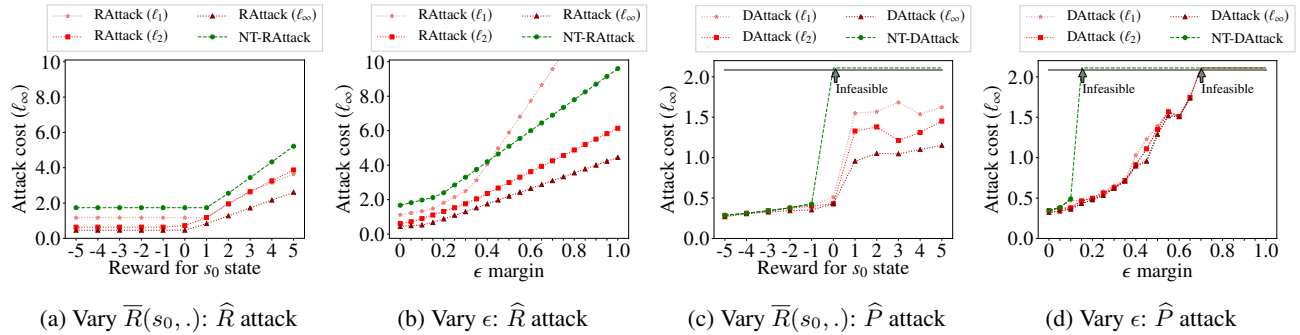


Figure 6: Results for poisoning attacks in the offline setting from Section 4. (a, b) plots show results for attack on rewards and (c, d) plots show results for attack on dynamics.

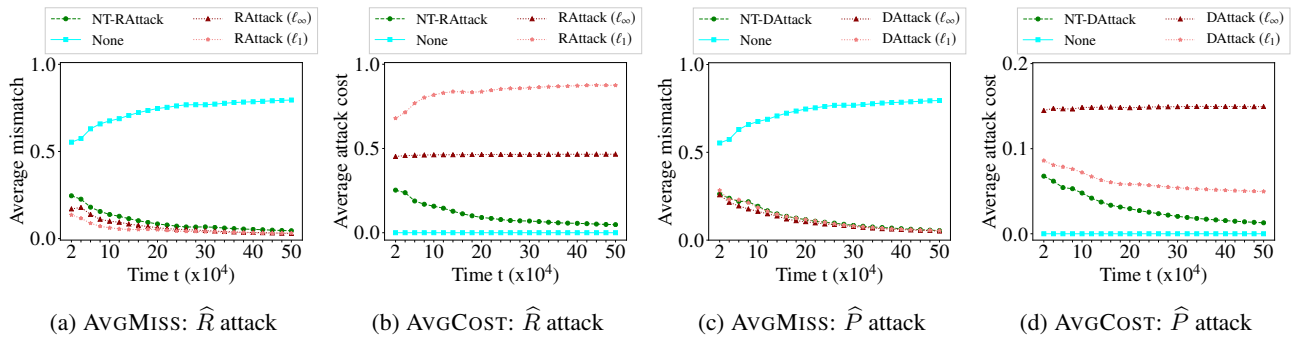


Figure 7: Results for poisoning attacks in the online setting from Section 5. (a, b) plots show results for attack on rewards and (c, d) plots show results for attack on dynamics.

The experimental setup and attack strategies used for these additional experiments are exactly same as that used in Section 6 for both the offline and the online settings. To recall, for the offline setting, we vary $\bar{R}(s_0, \cdot) \in [-5, 5]$ and vary ϵ margin $\in [0, 1]$. We use ℓ_∞ -norm in the measure of attack cost. The results are reported as an average of 10 runs. For the online setting, we fix $\bar{R}(s_0, \cdot) = -2.5$ and $\epsilon = 0.1$. We plot the measure of the attacker’s achieved goal in terms of AVGMiss and attacker’s cost in terms of AVGCOST for ℓ_1 -norm measured over time t . The results are reported as an average of 20 runs.

While the scales of cost and mismatch are different, the key takeaway results from Figure 6 (offline setting) and Figure 7 (online setting) are same as discussed in Section 6.1 and Section 6.2 respectively.

D. Background

Throughout many proofs, we will use relative values defined in average reward reinforcement learning. Relative values or bias values of policy π are defined as following⁷:

$$Q^\pi(s, a) = \lim_{N \rightarrow \infty} \mathbb{E} \left[\sum_{t=0}^{N-1} (R(s_t, a_t) - \rho^\pi) \mid s_0 = s, a_0 = a, \pi \right],$$

where s_0 and a_0 are the initial state and we follow policy π after taking action a_0 in state s_0 . We will often refer to relative values as Q values. These values satisfy the following recurrence equation

$$Q^\pi(s, a) = R(s, a) - \rho^\pi + \sum_{s' \in S} P(s, a, s') Q^\pi(s', \pi(s')).$$

By definition we have $V^\pi(s) = Q^\pi(s, \pi(s))$. Based on Corollary 8.2.7 in (Puterman, 1994), these values can be calculated by solving set of equations

$$V^\pi(s) = R(s, \pi(s)) - \rho^\pi + \sum_{s' \in S} P(s, \pi(s), s') V^\pi(s') \quad (9)$$

$$\sum_{s \in S} \mu^\pi(s) V^\pi(s) = 0 \quad (10)$$

for $V^\pi(s)$ values and setting

$$Q^\pi(s, a) = R(s, a) - \rho^\pi + \sum_{s' \in S} P(s, a, s') V^\pi(s'). \quad (11)$$

We rely on the following result of Even-Dar et al. (2005).

Lemma 3. (Lemma 7 in (Even-Dar et al., 2005)) For two policies π and π' we have:

$$\rho^\pi - \rho^{\pi'} = \sum_{s \in S} \mu^{\pi'}(s) (Q^\pi(s, \pi(s)) - Q^\pi(s, \pi'(s))).$$

E. Proofs for Offline Attacks: Lemma 1 (Section 4.1)

We prove Lemma 1 through several intermediate results. The first one is a direct consequence of Lemma 3.

Corollary 1. For any policy π and its neighbor policy $\pi\{s; a\}$ we have:

$$\rho^\pi - \rho^{\pi\{s; a\}} = \mu^{\pi\{s; a\}}(s) (Q^\pi(s, \pi(s)) - Q^\pi(s, a)).$$

Next, we provide a sufficient condition for a policy π to be uniquely optimal.

Lemma 4. If we have $\rho^\pi \geq \rho^{\pi\{s; a\}} + \epsilon$ for every state s and action $a \neq \pi(s)$, and $\epsilon > 0$, then π is the only optimal policy.

⁷ This limit assumes that all policies are aperiodic. For periodic policies, we need to use the Cesaro limit (Mahadevan, 1996).

Proof. Let arbitrary $s \in S$ and $a \in \mathcal{A}$ be such that $a \neq \pi(s)$. Based on corollary 1, we have

$$Q^\pi(s, \pi(s)) - Q^\pi(s, a) = \frac{\rho^\pi - \rho^{\pi\{s;a\}}}{\mu^{\pi\{s;a\}}(s)} > 0. \quad (12)$$

Note that in this paper we are focusing on recurrent MDPs, thus, we know $\mu^{\pi\{s;a\}}(s) > 0$. Now let $\pi' \neq \pi$ be a policy with $\pi'(s') = a' \neq \pi(s')$. Using (12) we have

$$\rho^\pi - \rho^{\pi'} = \sum_{s \in S} \mu^{\pi'}(s) (Q^\pi(s, \pi(s)) - Q^\pi(s, \pi'(s))) \geq \mu^{\pi'}(s') (Q^\pi(s', a') - Q^\pi(s', \pi(s'))) > 0$$

We again used the fact that MDP is recurrent to say $\mu^{\pi'}(s') > 0$. $\square$

Proof of Lemma 1

Proof. The necessity of the condition follows directly from the definition of ϵ -robust policies. Let us focus on its sufficiency.

Consider deterministic policies π , and denote the Hamming distance between two policies π_1 and π_2 by $D_H(\pi_1, \pi_2)$, i.e., $D_H(\pi_1, \pi_2) = \sum_{s \in S} \mathbb{1}[\pi_1(s) \neq \pi_2(s)]$ where $\mathbb{1}[\cdot]$ denotes the indicator function. Assume that the condition of the proposition holds for policy π^* , i.e., that $\rho^{\pi^*} \geq \rho^{\pi_1} + \epsilon$ for all π_1 s.t. $D_H(\pi_1, \pi^*) = 1$.

Lemma 4 implies that π^* is uniquely optimal. Now, consider policy π_k s.t. $D_H(\pi_k, \pi^*) = k > 1$. Since π^* is (uniquely) optimal and the MDP is recurrent ($\mu^{\pi^*}(s) > 0$), we have that

$$\rho^{\pi^*} - \rho^{\pi_k} = \sum_{s \in S} \mu^{\pi^*}(s) \cdot [Q^{\pi_k}(s, \pi^*(s)) - Q^{\pi_k}(s, \pi_k(s))] > 0,$$

which implies that there exists $s_k \in S$ s.t. $[Q^{\pi_k}(s_k, \pi^*(s_k)) - Q^{\pi_k}(s_k, \pi_k(s_k))] > 0$. Define policy π_{k-1} as

$$\pi_{k-1}(s) = \begin{cases} \pi^*(s) & \text{if } s = s_k \\ \pi_k(s) & \text{otw.} \end{cases}$$

We have that

$$\begin{aligned} \rho^{\pi_{k-1}} &= \rho^{\pi_k} + \sum_{s \in S} \mu^{\pi_{k-1}}(s) \cdot [Q^{\pi_k}(s, \pi_{k-1}(s)) - Q^{\pi_k}(s, \pi_k(s))] \\ &= \rho^{\pi_k} + \mu^{\pi_{k-1}}(s_k) \cdot [Q^{\pi_k}(s_k, \pi^*(s_k)) - Q^{\pi_k}(s_k, \pi_k(s_k))] \geq \rho^{\pi^*}. \end{aligned}$$

Therefore, by induction, we know that there exists a policy π_1 such that $\rho^{\pi_k} \leq \rho^{\pi_1}$ and $D_H(\pi_1, \pi^*) = 1$. Utilizing our initial assumption, we obtain that $\rho^{\pi^*} \geq \rho^{\pi_1} + \epsilon$ for all $\pi \neq \pi^*$, which proves that π^* is ϵ -robust optimal if the condition of the proposition holds. $\square$

F. Proofs for Offline Attacks: Poisoning Rewards (Section 4.2)

We break the proof of Theorem 1 into two parts: In Appendix F.1, we prove the lower bound in the theorem, and in Appendix F.2 we show the upper bound and feasibility claim for the attack.

F.1. Proofs for the Lower Bound in Theorem 1

To prove the lower bound in Theorem 1, we will use a proof technique that is similar to the one presented in (Ma et al., 2019), but adapted to our setting (the average reward criterion).

First, let us define operator F as

$$F(Q, R, \rho, P, \pi)(s, a) = R(s, a) - \rho + \sum_{s' \in S} P(s, a, s') V^\pi(s'), \quad (13)$$

or in vector notation

$$F(Q, R, \rho, P, \pi) = R - \rho \cdot \mathbf{1} + P \cdot V^\pi,$$

where $V^\pi(s') = Q(s', \pi(s'))$ (and π is a deterministic policy). Furthermore, we defined the span of X as $sp(X) = \max_i X(i) - \min_i X(i)$ - as argued in (Puterman, 1994), sp is a seminorm.

Lemma 5. *The following holds:*

$$sp(F(Q_1, R, \rho, P, \pi) - F(Q_2, R, \rho, P, \pi)) \leq (1 - \alpha) \cdot sp(Q_1 - Q_2),$$

where

$$\alpha = \min_{s, a, s', a'} \sum_{x \in S} \min\{P(s, a, x), P(s', a', x)\}.$$

Proof. We have that

$$sp(F(Q_1, R, \rho, P, \pi) - F(Q_2, R, \rho, P, \pi)) = sp(P \cdot (V_1^\pi - V_2^\pi)).$$

Following the proof of Proposition 6.6.1 in (Puterman, 1994), we obtain that for $b(x, s, a, s', a') = \min\{P(s, a, x), P(s', a', x)\}$

$$\begin{aligned} & \sum_{x \in S} P(s, a, x) \cdot (V_1^\pi(x) - V_2^\pi(x)) - \sum_{x \in S} P(s', a', x) \cdot (V_1^\pi(x) - V_2^\pi(x)) \\ &= \sum_{x \in S} (P(s, a, x) - b(x, s, a, s', a')) \cdot (V_1^\pi(x) - V_2^\pi(x)) \\ & \quad - \sum_{x \in S} (P(s', a', x) - b(x, s, a, s', a')) \cdot (V_1^\pi(x) - V_2^\pi(x)) \\ &\leq \sum_{x \in S} (P(s, a, x) - b(x, s, a, s', a')) \cdot \max_{x'} (V_1^\pi(x') - V_2^\pi(x')) \\ & \quad - \sum_{x \in S} (P(s', a', x) - b(x, s, a, s', a')) \cdot \min_{x'} (V_1^\pi(x') - V_2^\pi(x')) \\ &= (1 - \sum_{x \in S} b(x, s, a, s', a')) \cdot sp(V_1^\pi - V_2^\pi) \leq (1 - \alpha) \cdot sp(V_1^\pi - V_2^\pi) \end{aligned}$$

Therefore

$$sp(F(Q_1, R, \rho, P, \pi) - F(Q_2, R, \rho, P, \pi)) = sp(P \cdot (V_1^\pi - V_2^\pi)) \leq (1 - \alpha) \cdot sp(V_1^\pi - V_2^\pi).$$

Now, notice that for $s_{\max} = \arg \max_s [V_1^\pi(s) - V_2^\pi(s)]$ and $s_{\min} = \arg \min_s [V_1^\pi(s) - V_2^\pi(s)]$ we have that

$$\begin{aligned} V_1^\pi(s_{\max}) - V_2^\pi(s_{\max}) &= Q_1(s_{\max}, \pi(s_{\max})) - Q_2(s_{\max}, \pi(s_{\max})) \leq \max_{s, a} [Q_1(s, a) - Q_2(s, a)] \\ V_1^\pi(s_{\min}) - V_2^\pi(s_{\min}) &= Q_1(s_{\min}, \pi(s_{\min})) - Q_2(s_{\min}, \pi(s_{\min})) \geq \min_{s, a} [Q_1(s, a) - Q_2(s, a)] \end{aligned}$$

Therefore $sp(V_1^\pi - V_2^\pi) \leq sp(Q_1 - Q_2)$, which implies that

$$sp(F(Q_1, R, \rho, P, \pi) - F(Q_2, R, \rho, P, \pi)) \leq (1 - \alpha) \cdot sp(Q_1 - Q_2)$$

□

To obtain the statement of the theorem, we will need to relate $sp(Q_1 - Q_2)$ to $\|R_1 - R_2\|_\infty$. The following lemma provides a more general result, relating Q-values, reward functions and transition matrices.

Lemma 6. Let Q_1^π and V_1^π denote Q and V values of policy π in MDP $M_1 = (S, \mathcal{A}, R_1, P_1)$ and Q_2^π denote Q values of policy π in MDP $M_2 = (S, \mathcal{A}, R_2, P_2)$. The following holds:

$$\|R_1 - R_2\|_\infty + \|P_1 - P_2\|_\infty \cdot \|V_1^\pi\|_\infty \geq \frac{\alpha_2}{2} \cdot sp(Q_1^\pi - Q_2^\pi).$$

where

$$\alpha_2 = \min_{s,a,s',a'} \sum_{x \in S} \min\{P_2(s, a, x), P_2(s', a', x)\}.$$

Proof. Notice that Q_1^π and Q_2^π satisfy

$$\begin{aligned} Q_1^\pi(s, a) &= F(Q_1^\pi, R_1, \rho_1^\pi, P_1, \pi) \\ Q_2^\pi(s, a) &= F(Q_2^\pi, R_2, \rho_2^\pi, P_2, \pi), \end{aligned}$$

where ρ_1^π and ρ_2^π denote the average rewards of policy π in M_1 and M_2 , respectively. We obtain

$$\begin{aligned} sp(Q_1^\pi - Q_2^\pi) &= sp(F(Q_1^\pi, R_1, \rho_1^\pi, P_1, \pi) - F(Q_2^\pi, R_2, \rho_2^\pi, P_2, \pi)) \\ &= sp(F(Q_1^\pi, R_1, \rho_1^\pi, P_1, \pi) - F(Q_1^\pi, R_2, \rho_2^\pi, P_1, \pi) \\ &\quad + F(Q_1^\pi, R_2, \rho_2^\pi, P_1, \pi) - F(Q_1^\pi, R_2, \rho_2^\pi, P_2, \pi) \\ &\quad + F(Q_1^\pi, R_2, \rho_2^\pi, P_2, \pi) - F(Q_2^\pi, R_2, \rho_2^\pi, P_2, \pi)) \\ &\leq sp(F(Q_1^\pi, R_1, \rho_1^\pi, P_1, \pi) - F(Q_1^\pi, R_2, \rho_2^\pi, P_1, \pi)) \\ &\quad + sp(F(Q_1^\pi, R_2, \rho_2^\pi, P_1, \pi) - F(Q_1^\pi, R_2, \rho_2^\pi, P_2, \pi)) \\ &\quad + sp(F(Q_1^\pi, R_2, \rho_2^\pi, P_2, \pi) - F(Q_2^\pi, R_2, \rho_2^\pi, P_2, \pi)) \\ &\leq sp(R_1 - \rho_1^\pi \cdot \mathbf{1} - R_2 + \rho_2^\pi \cdot \mathbf{1}) + sp((P_1 - P_2) \cdot V_1^\pi) + (1 - \alpha_2) \cdot sp(Q_1^\pi - Q_2^\pi) \end{aligned}$$

where the last inequality is due to Lemma 5 (i.e., $sp(F(Q_1^\pi, R_2, \rho_2^\pi, P_2, \pi) - F(Q_2^\pi, R_2, \rho_2^\pi, P_2, \pi)) \leq (1 - \alpha_2) \cdot sp(Q_1^\pi - Q_2^\pi)$). Due to the properties of sp , we have

$$sp(R_1 - \rho_1^\pi \cdot \mathbf{1} - R_2 + \rho_2^\pi \cdot \mathbf{1}) = sp(R_1 - R_2) \leq 2 \cdot \|R_1 - R_2\|_\infty,$$

and

$$\begin{aligned} sp((P_1 - P_2) \cdot V_1^\pi) &\leq 2 \cdot \|(P_1 - P_2) \cdot V_1^\pi\|_\infty \\ &= 2 \cdot \max_{s,a} \left| \sum_{s'} (P_1(s, a, s') - P_2(s, a, s')) \cdot V_1^\pi(s') \right| \\ &\leq 2 \cdot \max_{s,a} \sum_{s'} |(P_1(s, a, s') - P_2(s, a, s')) \cdot V_1^\pi(s')| \\ &\leq 2 \cdot \max_{s,a} \sum_{s'} |P_1(s, a, s') - P_2(s, a, s')| \cdot \max_{s''} |V_1^\pi(s'')| \\ &\leq 2 \cdot \|P_1 - P_2\|_\infty \cdot \|V_1^\pi\|_\infty, \end{aligned}$$

where $\|P_1 - P_2\|_\infty = \max_{s,a} \sum_{s'} |P_1(s, a, s') - P_2(s, a, s')|$. Putting this together with the upper bound on $sp(Q_1^\pi - Q_2^\pi)$, we obtain

$$2 \cdot \|R_1 - R_2\|_\infty + 2 \cdot \|P_1 - P_2\|_\infty \cdot \|V_1^\pi\|_\infty \geq \alpha_2 \cdot sp(Q_1^\pi - Q_2^\pi),$$

which proves the claim. $\square$

We are now ready to prove the lower bound in Theorem 1.

Proof of the Lower Bound in Theorem 1:

Proof. From Lemma 6, it follows that a lower bound can be obtained by bounding $sp(\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger})$ from below, where $\hat{Q}^{\pi^\dagger}$ are Q -values of the target policy in the modified MDP, and $\bar{Q}^{\pi^\dagger}$ are Q values of the target policy in the initial MDP.

Let s' and a' be a state action pair that satisfy: $s', a' = \arg \max_{s,a} \bar{\chi}_\epsilon^{\pi^\dagger}(s, a)$. W.l.o.g. we can assume that $\bar{\chi}_\epsilon^{\pi^\dagger}(s', a') > 0$, since otherwise no modifications are needed to the original MDP and the lower bound trivially holds. We have

$$\begin{aligned}
 sp(\bar{Q}^{\pi^\dagger} - \hat{Q}^{\pi^\dagger}) &= sp(\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger}) = \max_{s,a} [\hat{Q}^{\pi^\dagger}(s, a) - \bar{Q}^{\pi^\dagger}(s, a)] - \min_{s,a} [\hat{Q}^{\pi^\dagger}(s, a) - \bar{Q}^{\pi^\dagger}(s, a)] \\
 &\geq \hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - (\hat{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', a')) \\
 &= (\hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - \hat{Q}^{\pi^\dagger}(s', a')) + (\bar{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s'))) \\
 &\geq \frac{\epsilon}{\bar{\mu}^{\pi_\dagger\{s';a'\}}(s')} + (\bar{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s'))) \\
 &= \frac{\bar{\rho}^{\pi_\dagger\{s';a'\}} - \bar{\rho}^{\pi_\dagger} + \epsilon}{\bar{\mu}^{\pi_\dagger\{s';a'\}}(s')} = \bar{\chi}_\epsilon^{\pi^\dagger}(s', a'),
 \end{aligned}$$

where we used the fact that $\hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) \geq \hat{Q}^{\pi^\dagger}(s', a') + \frac{\epsilon}{\bar{\mu}^{\pi_\dagger\{s';a'\}}(s')}$ (because $\pi_\dagger$ is robustly optimal in the modified MDP and the transition kernel did not change) and Lemma 3 (Corollary 1) to relate Q values to average rewards ρ . Finally, using Lemma 6 (and noting that we did not change transition matrix $\bar{P}$), we obtain

$$\|\hat{R} - \bar{R}\|_\infty \geq \frac{\bar{\alpha}}{2} \cdot \|\bar{\chi}_\epsilon^{\pi^\dagger}\|_\infty$$

The inequality $\|\hat{R} - \bar{R}\|_p \geq \|\hat{R} - \bar{R}\|_\infty$ implies the lower bound. $\square$

F.2. Proofs for the Upper Bound and Feasibility in Theorem 1

Lemma 7. $\hat{R} = \bar{R} - \bar{\chi}_\epsilon^{\pi^\dagger}$ is the solution of the following problem:

$$\begin{aligned}
 \min_{\hat{R}} \quad & \|\hat{R} - \bar{R}\|_p \\
 \text{s.t.} \quad & \sum_{s'} \bar{\mu}^{\pi^\dagger}(s') \cdot \hat{R}(s', \pi_\dagger(s')) \geq \sum_{s'} \bar{\mu}^{\pi_\dagger\{s;a\}}(s') \cdot \hat{R}(s', \pi_\dagger\{s;a\}(s')) + \epsilon \quad \forall s, a \neq \pi_\dagger(s), \\
 & \hat{R}(s, \pi_\dagger(s)) = \bar{R}(s, \pi_\dagger(s)) \quad \forall s.
 \end{aligned} \tag{P6}$$

Proof. Since the transitions are not changed, we have

$$\begin{aligned}
 \sum_{s'} \hat{R}(s', \pi_\dagger\{s;a\}(s')) \cdot \bar{\mu}^{\pi_\dagger\{s;a\}}(s') &= \sum_{s'} \hat{R}(s', \pi_\dagger\{s;a\}(s')) \cdot \bar{\mu}^{\pi_\dagger\{s;a\}}(s') \\
 &= \sum_{s'} \bar{R}(s', \pi_\dagger\{s;a\}(s')) \cdot \bar{\mu}^{\pi_\dagger\{s;a\}}(s') - \bar{\chi}_\epsilon^{\pi^\dagger}(s, a) \cdot \bar{\mu}^{\pi_\dagger\{s;a\}}(s) \\
 &= \bar{\rho}^{\pi_\dagger\{s;a\}} - \bar{\chi}_\epsilon^{\pi^\dagger}(s, a) \cdot \bar{\mu}^{\pi_\dagger\{s;a\}}(s) \\
 &\leq \bar{\rho}^{\pi_\dagger\{s;a\}} - (\bar{\rho}^{\pi_\dagger\{s;a\}} - \bar{\rho}^{\pi_\dagger} + \epsilon) \\
 &= \bar{\rho}^{\pi_\dagger} - \epsilon \\
 &= \hat{\rho}^{\pi^\dagger} - \epsilon \\
 &= \sum_{s'} \bar{\mu}^{\pi^\dagger}(s') \cdot \hat{R}(s', \pi_\dagger(s')) - \epsilon
 \end{aligned}$$

which shows that the attack satisfies the first constraint of the optimization problem. The second constraint is also satisfied as we have $\bar{\chi}_\epsilon^{\pi^\dagger}(s, \pi_\dagger(s)) = 0$ $\square$

Proof of Feasibility and the Upper Bound in Theorem 1:

Proof. Consider the attack with $\hat{R} = \bar{R} - \bar{\chi}_\epsilon^{\pi^\dagger}$. As we showed in Lemma 7, it is the solution for problem (P6) which has all the constraints in (P1). Thus, this attack is also a solution for (P1) which shows the feasibility of this optimization problem.

Furthermore, this attack also provides us with an upper bound on the quality of the obtained solution, i.e., for this attack:

$$\left\| \hat{R} - \bar{R} \right\|_p = \left(\sum_{s,a} \left(\hat{R}(s,a) - \bar{R}(s,a) \right)^p \right)^{1/p} = \left(\sum_{s,a} \bar{\chi}_\epsilon^{\pi^\dagger}(s,a)^p \right)^{1/p} = \left\| \bar{\chi}_\epsilon^{\pi^\dagger} \right\|_p,$$

and we know that an optimal attack achieves a lower cost (i.e., a lower value of $\left\| \hat{R} - \bar{R} \right\|_p$). $\square$

G. Proofs for Offline Attacks: Poisoning Dynamics (Section 4.3)

To prove Theorem 2, we first introduce a more constrained problem which has a tractable reformulation in Appendix G.1. This problem enables us to prove our feasibility condition in Appendix G.3 and will also be the basis of our online dynamic attacks. Then, in Appendix G.3, we prove the theorem.

G.1. More Constrained Problem and Reformulation

Consider the following problem

$$\begin{aligned} \min_{P, \mu^{\pi^\dagger}, \mu^{\pi^\dagger\{s;a\}}} \quad & \left\| P - \bar{P} \right\|_p \\ \text{s.t.} \quad & \mu^\pi(s) = \sum_{s'} P(s', \pi(s'), s) \cdot \mu^\pi(s') \quad \forall s, \\ & \mu^{\pi^\dagger\{s;a\}}(s') = \sum_{s''} P(s'', \pi^\dagger\{s;a\}(s''), s') \cdot \mu^{\pi^\dagger\{s;a\}}(s'') \quad \forall s', s, a \neq \pi^\dagger(s), \\ & \sum_{s'} \mu^{\pi^\dagger}(s') \cdot \bar{R}(s', \pi^\dagger(s')) \geq \sum_{s'} \mu^{\pi^\dagger\{s;a\}}(s') \cdot \bar{R}(s', \pi^\dagger\{s;a\}(s')) + \epsilon \quad \forall s, a \neq \pi^\dagger(s), \\ & P(s, a, s') \geq \delta \cdot \bar{P}(s, a, s') \quad \forall s, a, s', \\ & P(s, \pi^\dagger(s), s') = \bar{P}(s, \pi^\dagger(s), s') \quad \forall s, s'. \end{aligned} \tag{P7}$$

This is the exact same problem (P2) with an additional condition $P(s, \pi^\dagger(s), s') = \bar{P}(s, \pi^\dagger(s), s')$ for all s, s' . This problem is also used in the online attacks (see problem (P4)). We will show that this problem has a simple reformulation as the following:

$$\begin{aligned} \min_P \quad & \left\| P - \bar{P} \right\|_p \\ \text{s.t.} \quad & \bar{R}(s, \pi^\dagger(s)) + \bar{B}^{\pi^\dagger}(s) - \bar{R}(s, a) \geq \epsilon + \sum_{s'} P(s, a, s') \cdot \bar{U}^{\pi^\dagger}(s, s') \quad \forall s, a \neq \pi^\dagger(s), s', \\ & P(s, a, s') \geq \delta \cdot \bar{P}(s, a, s') \quad \forall s, a, s', \\ & P(s, \pi^\dagger(s), s') = \bar{P}(s, \pi^\dagger(s), s') \quad \forall s, s', \end{aligned} \tag{P8}$$

where

$$\bar{U}^{\pi^\dagger}(s, s') = \bar{V}^{\pi^\dagger}(s') + \epsilon \cdot \bar{T}^{\pi^\dagger}(s', s) \cdot \mathbb{1}[s' \neq s].$$

Note that $\bar{U}^{\pi^\dagger}$ is obtained from the initial MDP $\bar{M}$.

Before we show these problems are equivalent, let us first characterize the minimum value of $\mu^{\pi^\dagger\{s;a\}}(s)$. To do so, we utilize the diameter of Markov chain defined by MDP $\bar{M}$ and policy $\pi^\dagger$, i.e., $\bar{D}^{\pi^\dagger} = \max_{s,s'} \bar{T}^{\pi^\dagger}(s, s')$. Here, $\bar{T}^{\pi^\dagger}(s, s')$ is the expected time to reach s' starting from s in MDP $\bar{M}$ under policy $\pi^\dagger$. We have:

Lemma 8. Assume that for transition kernel P , we have $P(s, \pi^\dagger(s), s') = \bar{P}(s, \pi^\dagger(s), s')$ for all s, s' , and μ^π denotes the stationary distribution of π under P . Then, for $\mu^{\pi^\dagger\{s;a\}}(s)$ we have

$$\mu^{\pi^\dagger\{s;a\}}(s) = \frac{1}{1 + \sum_{s' \neq s} P(s, a, s') \bar{T}^{\pi^\dagger}(s', s)} \geq \frac{1}{1 + \bar{D}^{\pi^\dagger}},$$

where $\bar{D}^{\pi^\dagger} = \max_{s', s''} \bar{T}^{\pi^\dagger}(s', s'')$ is the diameter of MDP $\bar{M}$ with policy $\pi^\dagger$.

Proof. By using the fact that $\frac{1}{\mu^{\pi_{\dagger}\{s;a\}}(s)} = T^{\pi_{\dagger}\{s;a\}}(s, s)$ (see Theorem 1.21 in (Durrett & Durrett, 1999)), where $T^{\pi_{\dagger}}(s', s)$ is the expected time to reach s starting from s' in the corresponding Markov chain, we obtain

$$\begin{aligned} \frac{1}{\mu^{\pi_{\dagger}\{s;a\}}(s)} &= T^{\pi_{\dagger}\{s;a\}}(s, s) \\ &= 1 + \sum_{s' \neq s} P(s, a, s') \cdot T^{\pi_{\dagger}\{s;a\}}(s', s) \\ &= 1 + \sum_{s' \neq s} P(s, a, s') \cdot T^{\pi_{\dagger}}(s', s) \\ &= 1 + \sum_{s' \neq s} P(s, a, s') \cdot \bar{T}^{\pi_{\dagger}}(s', s) \\ &\leq 1 + \max_{s'} \bar{T}^{\pi_{\dagger}}(s', s) \\ &\leq 1 + \bar{D}^{\pi_{\dagger}}. \end{aligned}$$

The last inequality obtained from the definition of diameter $\bar{D}^{\pi_{\dagger}}$, i.e., $\bar{D}^{\pi_{\dagger}} = \max_{s', s''} \bar{T}^{\pi_{\dagger}}(s', s'')$ and we have used the fact that transitions of $\pi_{\dagger}$ are not changed. $\square$

Lemma 9. *Problem (P8) is a reformulation of (P7).*

Proof. Let ρ^{π} , Q^{π} , and μ^{π} denote the average reward, Q-values, and stationary distribution of π on MDP $M = (S, A, \bar{R}, P)$. The last two constraints of (P7) are repeated in (P8). The rest of conditions in (P7) are equivalent to $\rho^{\pi_{\dagger}} \geq \rho^{\pi_{\dagger}\{s;a\}} + \epsilon$ for all $s, a \neq \pi_{\dagger}(s)$ which we should prove it is equivalent to the remaining constraint in (P8). By using the relation between average rewards ρ and Q values from Lemma 3 (Corollary 1), we can write this condition as:

$$Q^{\pi_{\dagger}}(s, \pi_{\dagger}(s)) - Q^{\pi_{\dagger}}(s, a) \geq \frac{\epsilon}{\mu^{\pi_{\dagger}\{s;a\}}(s)}, \forall a \neq \pi_{\dagger}(s).$$

Due to Lemma 8, this can be rewritten as

$$Q^{\pi_{\dagger}}(s, \pi_{\dagger}(s)) - Q^{\pi_{\dagger}}(s, a) \geq \epsilon + \epsilon \cdot \sum_{s' \neq s} P(s, a, s') \cdot \bar{T}^{\pi_{\dagger}}(s', s), \forall a \neq \pi_{\dagger}(s).$$

By using the recurrence relation for Q values, the definition of $\bar{B}^{\pi_{\dagger}}$, and the fact that $P(s, \pi_{\dagger}(s), \cdot) = \bar{P}(s, \pi_{\dagger}(s), \cdot)$, we obtain that for all s and a , it is equivalent to the following

$$\epsilon \leq \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - \sum_{s'} P(s, a, s') \cdot \left[\bar{V}^{\pi_{\dagger}}(s') + \epsilon \cdot \bar{T}^{\pi_{\dagger}}(s', s) \cdot \mathbb{1}[s' \neq s] \right] \quad (14)$$

$$= \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - \sum_{s'} P(s, a, s') \cdot \bar{U}^{\pi_{\dagger}}(s, s'). \quad (15)$$

which is the condition in (P8) $\square$

G.2. Feasibility Analysis

In this section we analyze feasibility of dynamics offline attacks as formulated in problem (P2). To give a sufficient condition for feasibility of (P2), it suffices to give a sufficient condition for feasibility of (P7) as it is a more constrained problem. In Lemma 10, we give a sufficient and necessary condition for feasibility of (P7). The condition stated in Theorem 2 is a simpler one showed in Corollary 2 using Lemma 10.

Lemma 10. *The optimization problem (P7) has a solution $\hat{P}$ if and only if for all state s and actions $a \neq \pi_{\dagger}(s)$ we have $C(s, a) \geq \epsilon$, where*

$$C(s, a) = \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - (1 - \delta) \cdot \min_{s'} \bar{U}^{\pi_{\dagger}}(s, s') - \delta \sum_{s'} \bar{P}(s, a, s') \cdot \bar{U}^{\pi_{\dagger}}(s, s').$$

Proof. We define $C'(s, a)$ as

$$C'(s, a) = \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - \sum_{s'} P(s, a, s') \cdot \bar{U}^{\pi_{\dagger}}(s, s').$$

The constraints of (P8) which is a reformulation of (P7) imply that we should have $C'(s, a) \geq \epsilon$. Due to the constraint $P(s, a, s') \geq \delta \cdot \bar{P}(s, a, s')$, we can write

$$C'(s, a) \leq \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - (1 - \delta) \cdot \min_{s'} \bar{U}^{\pi_{\dagger}}(s, s') - \delta \cdot \sum_{s'} \bar{P}(s, a, s') \cdot \bar{U}^{\pi_{\dagger}}(s, s') = C(s, a).$$

Therefore, for the attack to be feasible, it has to hold that $C(s, a) \geq \epsilon$. To see that $C(s, a) \geq \epsilon$ is also a sufficient condition, let us consider transition kernel $\hat{P}$ defined as

$$\hat{P}(s, a, s') = \begin{cases} \bar{P}(s, \pi_{\dagger}(s), s') & \text{if } a = \pi_{\dagger}(s) \\ 1 - \delta + \delta \cdot \bar{P}(s, a, s') & \text{if } a \neq \pi_{\dagger}(s) \text{ and } s' = \arg \min_{s''} U(s, s'') \\ \delta \cdot \bar{P}(s, a, s') & \text{otherwise} \end{cases}$$

For state s and action $a \neq \pi_{\dagger}(s)$, let $s_{\min} = \arg \min_{s''} U(s, s'')$. From the definition of Q values and the fact that we only change transitions for non-target state-action pairs, we obtain that

$$\begin{aligned} \hat{Q}^{\pi_{\dagger}}(s, \pi_{\dagger}(s)) - \hat{Q}^{\pi_{\dagger}}(s, a) &= \bar{R}(s, \pi_{\dagger}(s)) - \bar{\rho}^{\pi_{\dagger}} + \sum_{s'} \bar{P}(s, \pi_{\dagger}(s), s') \bar{V}^{\pi_{\dagger}}(s') \\ &\quad - \bar{R}(s, a) + \bar{\rho}^{\pi_{\dagger}} - \sum_{s'} \hat{P}(s, a, s') \bar{V}^{\pi_{\dagger}}(s') \\ &= \bar{R}(s, \pi_{\dagger}(s)) - \bar{R}(s, a) + \bar{B}^{\pi_{\dagger}}(s) - (1 - \delta) \cdot \bar{V}^{\pi_{\dagger}}(s_{\min}) - \delta \cdot \sum_{s'} \bar{P}(s, a, s') \bar{V}^{\pi_{\dagger}}(s') \\ &= \bar{R}(s, \pi_{\dagger}(s)) - \bar{R}(s, a) + \bar{B}^{\pi_{\dagger}}(s) - (1 - \delta) \cdot \bar{U}^{\pi_{\dagger}}(s, s_{\min}) - \delta \cdot \sum_{s'} \bar{P}(s, a, s') \bar{U}^{\pi_{\dagger}}(s, s') \\ &\quad + (1 - \delta) \cdot \epsilon \cdot \bar{T}^{\pi_{\dagger}}(s_{\min}, s) \cdot \mathbb{1}_{s_{\min} \neq s} + \delta \cdot \sum_{s'} \bar{P}(s, a, s') \cdot \epsilon \cdot \bar{T}^{\pi_{\dagger}}(s', s) \cdot \mathbb{1}_{s' \neq s} \\ &= C(s, a) + \sum_{s' \neq s} \hat{P}(s, a, s') \cdot \epsilon \cdot \bar{T}^{\pi_{\dagger}}(s', s) \\ &\geq \epsilon + \sum_{s' \neq s} \hat{P}(s, a, s') \cdot \epsilon \cdot \bar{T}^{\pi_{\dagger}}(s', s) = \frac{\epsilon}{\hat{\mu}^{\pi_{\dagger}}\{s; a\}}, \end{aligned}$$

where the last equality is due to Lemma 8. Using the relation between average rewards ρ and Q values from Lemma 3 (Corollary 1), we obtain the claim. $\square$

Finally, from Lemma 10 we derive a simpler to express sufficient condition that we use in for the main theorem of this section.

Corollary 2. *The optimization problem (P7) has a solution $\hat{P}$ if for all state s and actions $a \neq \pi_{\dagger}(s)$ we have $\bar{\beta}_{\delta}^{\pi_{\dagger}}(s, a) \geq \epsilon \cdot (1 + \bar{D}^{\pi_{\dagger}})$.*

Proof. Note that $C(s, a)$ from Lemma 10 is bounded by:

$$\begin{aligned} C(s, a) &= \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - (1 - \delta) \cdot \min_{s'} \bar{U}^{\pi_{\dagger}}(s, s') - \epsilon_P \sum_{s'} \bar{P}(s, a, s') \cdot \bar{U}^{\pi_{\dagger}}(s, s') \\ &\geq \bar{R}(s, \pi_{\dagger}(s)) + \bar{B}^{\pi_{\dagger}}(s) - \bar{R}(s, a) - (1 - \delta) \cdot (\min_{s'} \bar{V}^{\pi_{\dagger}}(s') + \epsilon \cdot \bar{D}^{\pi_{\dagger}}) - \delta \cdot (\max_{s'} \bar{V}^{\pi_{\dagger}}(s') + \epsilon \cdot \bar{D}^{\pi_{\dagger}}) \\ &= \bar{\beta}_{\delta}^{\pi_{\dagger}}(s, a) - \epsilon \bar{D}^{\pi_{\dagger}}, \end{aligned}$$

where in the last inequality we applied the definition of $\bar{\beta}_{\delta}^{\pi_{\dagger}}(s, a)$, i.e.

$$\bar{\beta}_{\delta}^{\pi_{\dagger}}(s, a) = \bar{R}(s, \pi_{\dagger}(s)) - \bar{R}(s, a) + \bar{B}^{\pi_{\dagger}}(s) - \bar{V}_{\min}^{\pi_{\dagger}} - \delta \cdot sp(\bar{V}^{\pi_{\dagger}}).$$

Using the condition $\bar{\beta}_\delta^{\pi^\dagger}(s, a) \geq \epsilon \cdot (1 + \bar{D}^{\pi^\dagger})$, we obtain

$$C(s, a) \geq \epsilon \cdot (1 + \bar{D}^{\pi^\dagger}) - \epsilon \bar{D}^{\pi^\dagger} = \epsilon,$$

which by Lemma 10 implies the statement. $\square$

G.3. Proof of Theorem 2

Proof. Feasibility and the upper bound: The sufficient condition in the statement of the theorem follows directly from Corollary 2 by noting that a solution to the optimization problem (P7) is also a solution to the optimization problem (P2). To obtain the upper bound, let us consider an attack of the following form:

$$\hat{P}(s, a, s') = (1 - \lambda(s, a)) \cdot \bar{P}(s, a, s') + \lambda(s, a) \cdot P_s(s, a, s') \quad (16)$$

where

$$P_s(s, a, s') = \begin{cases} \bar{P}(s, a, s') & \text{if } a = \pi_T(s) \\ 1 - \delta + \delta \cdot \bar{P}(s, a, s') & \text{if } a \neq \pi_T(s) \text{ and } s' = \arg \min_{s''} \bar{V}^{\pi^\dagger}(s'') \\ \delta \cdot \bar{P}(s, a, s') & \text{otw.} \end{cases}$$

Here, $\lambda(s, a)$ is defined for each state-action pair separately. For the above attack, note that

$$\|\hat{P} - \bar{P}\|_p = \left(\sum_{s,a} \left(\sum_{s'} |\hat{P}(s, a, s') - \bar{P}(s, a, s')| \right)^p \right)^{1/p} \leq 2 \cdot \left(\sum_{s,a} \lambda^p(s, a) \right)^{1/p}.$$

Hence, we want to find the minimum $\lambda(s, a)$ (across all states s and $a \neq \pi_\dagger(s)$) for which the optimization problem is feasible.

Due to Lemma 1, for the attack to be feasible, we need to satisfy

$$\hat{\rho}^{\pi^\dagger} \geq \hat{\rho}^{\pi^\dagger\{s;a\}} + \epsilon, \forall a \neq \pi_\dagger(s).$$

Using Lemma 3 (Corollary 1), we can rewrite this condition in terms of Q as

$$\hat{Q}^{\pi^\dagger}(s, a) - \hat{Q}^{\pi^\dagger}(s, \pi_\dagger(s)) \leq -\frac{\epsilon}{\hat{\mu}^{\pi^\dagger\{s;a\}}(s)}, \forall a \neq \pi_\dagger(s). \quad (17)$$

Moreover, $\hat{V}^{\pi^\dagger} = \bar{V}^{\pi^\dagger}$ for the considered attack because it does not change the transitions of the target policy. Therefore, we have

$$\begin{aligned} \hat{Q}^{\pi^\dagger}(s, a) - \hat{Q}^{\pi^\dagger}(s, \pi_\dagger(s)) &= \bar{R}(s, a) - \hat{\rho}^{\pi^\dagger} + \sum_{s'} \hat{P}(s, a, s') \cdot \bar{V}^{\pi^\dagger}(s') \\ &\quad - \bar{R}(s, \pi_\dagger(s)) + \hat{\rho}^{\pi^\dagger} - \sum_{s'} \hat{P}(s, \pi_\dagger(s), s') \cdot \bar{V}^{\pi^\dagger}(s') \\ &= \bar{R}(s, a) + \sum_{s'} [(1 - \lambda(s, a)) \cdot \bar{P}(s, a, s') + \lambda(s, a) \cdot P_s(s, a, s')] \cdot \bar{V}^{\pi^\dagger}(s') \\ &\quad - \bar{R}(s, \pi_\dagger(s)) - \sum_{s'} \bar{P}(s, \pi_\dagger(s), s') \cdot \bar{V}^{\pi^\dagger}(s') \\ &= (1 - \lambda(s, a)) \cdot [\bar{R}(s, a) - \bar{\rho}^{\pi^\dagger} + \sum_{s'} \bar{P}(s, a, s') \cdot \bar{V}^{\pi^\dagger}(s')] \\ &\quad - \bar{R}(s, \pi_\dagger(s)) + \bar{\rho}^{\pi^\dagger} - \sum_{s'} \bar{P}(s, \pi_\dagger(s), s') \cdot \bar{V}^{\pi^\dagger}(s') \\ &\quad + \lambda(s, a) \cdot [\bar{R}(s, a) + (1 - \delta) \cdot \min_{s'} \bar{V}^{\pi^\dagger}(s') + \delta \cdot \sum_{s'} \bar{P}(s, a, s') \cdot \bar{V}^{\pi^\dagger}(s')] \\ &\quad - \bar{R}(s, \pi_\dagger(s)) - \sum_{s'} \bar{P}(s, \pi_\dagger(s), s') \cdot \bar{V}^{\pi^\dagger}(s') \end{aligned}$$

$$\begin{aligned}
 &\leq (1 - \lambda(s, a)) \cdot [\bar{Q}^{\pi^\dagger}(s, a) - \bar{Q}^{\pi^\dagger}(s, \pi_\dagger(s))] \\
 &\quad + \lambda(s, a) \cdot [\bar{R}(s, a) - \bar{R}(s, \pi_\dagger(s)) + \bar{V}_{\min}^{\pi^\dagger} - \bar{B}^{\pi^\dagger}(s) + \delta \cdot (\max_{s'} \bar{V}^{\pi^\dagger}(s') - \bar{V}_{\min}^{\pi^\dagger})] \\
 &= (1 - \lambda(s, a)) \cdot \frac{\bar{\rho}^{\pi^\dagger\{s; a\}} - \bar{\rho}^{\pi^\dagger}}{\bar{\mu}^{\pi^\dagger\{s; a\}}(s)} - \lambda(s, a) \cdot \bar{\beta}_\delta^{\pi^\dagger}(s, a) \\
 &= (1 - \lambda(s, a)) \cdot \bar{\chi}_0^{\pi^\dagger}(s, a) - \lambda(s, a) \cdot \bar{\beta}_\delta^{\pi^\dagger}(s, a).
 \end{aligned}$$

By Lemma 8, $\hat{\mu}^{\pi^\dagger\{s; a\}}(s) \geq \frac{1}{1 + \bar{D}^{\pi^\dagger}}$. Hence, we know that the attack is successful when $\hat{Q}^{\pi^\dagger}(s, a) - \hat{Q}^{\pi^\dagger}(s, \pi_\dagger(s)) \leq -\epsilon \cdot (1 + \bar{D}^{\pi^\dagger})$, which by the above inequality implies that a sufficient condition for the attack to be successful is

$$\lambda(s, a) \cdot \bar{\chi}_0^{\pi^\dagger}(s, a) + \lambda(s, a) \cdot \bar{\beta}_\delta^{\pi^\dagger}(s, a) \geq \epsilon \cdot (1 + \bar{D}^{\pi^\dagger}) + \bar{\chi}_0^{\pi^\dagger}(s, a).$$

Furthermore, we also know that if $\bar{\rho}^{\pi^\dagger} \geq \bar{\rho}^{\pi^\dagger\{s; a\}} + \epsilon$, we can set $\lambda(s, a) = 0$. Therefore, the attack is successful if $\lambda(s, a)$ satisfies

$$\lambda(s, a) \geq \frac{\bar{\chi}_0^{\pi^\dagger}(s, a) + \epsilon \cdot (1 + \bar{D}^{\pi^\dagger})}{\bar{\chi}_0^{\pi^\dagger}(s, a) + \bar{\beta}_\delta^{\pi^\dagger}(s, a)} \cdot \mathbf{1}[\bar{\chi}_\epsilon^{\pi^\dagger}(s, a) > 0].$$

By choosing the lower bound for $\lambda(s, a)$, we obtain the upper bound in the statement of the theorem, i.e.:

$$\|\hat{P} - \bar{P}\|_p \leq 2 \cdot \|\Lambda\|_p, \quad (18)$$

where $\Lambda(s, a) = \frac{\bar{\chi}_0^{\pi^\dagger}(s, a) + \epsilon \cdot (1 + \bar{D}^{\pi^\dagger})}{\bar{\chi}_0^{\pi^\dagger}(s, a) + \bar{\beta}_\delta^{\pi^\dagger}(s, a)} \cdot \mathbf{1}[\bar{\chi}_\epsilon^{\pi^\dagger}(s, a) > 0]$.

Lower bound: To prove the lower bound, we follow the same arguments as in the proof of Theorem 1 (the proof for the lower bound). Let s' and a' be a state action pair that satisfy: $s', a' = \arg \max_{s, a} \bar{\chi}^{\pi^\dagger}(s, a)$. Let's consider the case when $\bar{\chi}^{\pi^\dagger}(s', a') > 0$. We have

$$\begin{aligned}
 sp(\bar{Q}^{\pi^\dagger} - \hat{Q}^{\pi^\dagger}) &= sp(\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger}) = \max_{s, a} [\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger}] - \min_{s, a} [\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger}] \\
 &\geq \hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - (\hat{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', a')) \\
 &= (\hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) - \hat{Q}^{\pi^\dagger}(s', a')) + (\bar{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s'))) \\
 &\geq \frac{\epsilon}{\hat{\mu}^{\pi^\dagger\{s'; a'\}}(s')} + (\bar{Q}^{\pi^\dagger}(s', a') - \bar{Q}^{\pi^\dagger}(s', \pi_\dagger(s'))) \\
 &\geq \epsilon + \frac{\bar{\rho}^{\pi^\dagger\{s'; a'\}} - \bar{\rho}^{\pi^\dagger}}{\bar{\mu}^{\pi^\dagger\{s'; a'\}}(s')} > \bar{\chi}_0^{\pi^\dagger}(s', a'),
 \end{aligned}$$

where we used the fact that $\hat{Q}^{\pi^\dagger}(s', \pi_\dagger(s')) \geq \hat{Q}^{\pi^\dagger}(s', a') + \frac{\epsilon}{\hat{\mu}^{\pi^\dagger\{s'; a'\}}(s')}$ (because $\pi_\dagger$ is robustly optimal in the modified MDP) and Lemma 3 (Corollary 1) to relate Q values to average rewards ρ . When $\bar{\chi}^{\pi^\dagger}(s', a') = 0$, we know that $sp(\bar{Q}^{\pi^\dagger} - \hat{Q}^{\pi^\dagger}) \geq 0$ due to the properties of sp . Therefore, $sp(\bar{Q}^{\pi^\dagger} - \hat{Q}^{\pi^\dagger}) = sp(\hat{Q}^{\pi^\dagger} - \bar{Q}^{\pi^\dagger}) \geq \|\bar{\chi}_0^{\pi^\dagger}\|_\infty$.

By using Lemma 6 (but now noting that we did not change rewards $\bar{R}$), we obtain

$$\|\hat{P} - \bar{P}\|_\infty \cdot \|\bar{V}^{\pi^\dagger}\|_\infty \geq \frac{\hat{\alpha}}{2} \cdot \|\bar{\chi}_0^{\pi^\dagger}\|_\infty.$$

Factor $\hat{\alpha}$ can be bounded by

$$\begin{aligned}
 \hat{\alpha} &= \min_{s, a, s', a'} \sum_x \min\{\hat{P}(s, a, x), \hat{P}(s', a', x)\} \\
 &\geq \min_{s, a, s', a'} \sum_x \min\{\delta \cdot \bar{P}(s, a, x), \delta \cdot \bar{P}(s', a', x)\}
 \end{aligned}$$

$$\begin{aligned}
 &= \delta \cdot \min_{s,a,s',a'} \sum_x \min\{\bar{P}(s,a,x), \bar{P}(s',a',x)\} \\
 &= \delta \cdot \bar{\alpha}.
 \end{aligned}$$

Putting this together with the previous expression, we obtain that

$$\|\hat{P} - \bar{P}\|_\infty \cdot \|\bar{V}^{\pi_\dagger}\|_\infty \geq \frac{\delta \cdot \bar{\alpha}}{2} \cdot \|\bar{\chi}_0^{\pi_\dagger}\|_\infty.$$

The inequality $\|\hat{P} - \bar{P}\|_p \geq \|\hat{P} - \bar{P}\|_\infty$ implies the lower bound.

□

G.4. Choosing δ

While δ can be set to small values, making the corresponding constraint in (P2) a relatively weak condition, it is important to note that its value controls parameters of MDP $\widehat{M}$ that are important for practical considerations in the offline setting. Moreover, since δ is a parameter in the optimization problems (P4) and (P7) (and (P8)), its value is also important for the online setting.

In the case of attacks on a learning agent, our results have dependency on the agent's regret, which in turn depends on the properties of MDP $\widehat{M}$. For example, if the agent adopts UCRL as its learning procedure, its regret will depend on the diameter of $\widehat{M}$. Hence, δ should be adjusted based on time horizon T , so that the parameters of MDP $\widehat{M}$ relevant for the agent's regret do not outweigh time horizon T .

In the case of attacks on a planning agent, setting δ to small values could result in a solution $\widehat{M}$ that has a large mixing time, in which case average reward ρ might not approximate well the average of obtained rewards in a finite horizon (e.g., see (Even-Dar et al., 2005)). This means that the choice of δ should account for the finiteness of time horizon in practice.

We leave a more detailed analysis that includes these considerations for future work.

H. Proofs for Online Attacks: Lemma 2 (Section 5.1)

Proof of Lemma 2 Assume the learner follows an algorithm ALG which chooses the action from a distribution based on the previous observations. Theorem 5.5.1 in (Puterman, 1994) shows that a history-independent algorithm Π exists that chooses the action a_t from the distribution q_{t,s_t} where the distributions $q_{t,s}$ are fixed and we have

$$\forall s, a, t : \mathbb{P}_{\text{ALG}}[s_t = s, a_t = a] = \mathbb{P}_{\Pi}[s_t = s, a_t = a], \quad (19)$$

$$q_{t,s}(a) = \mathbb{P}_{\text{ALG}}[a_t = a | s_t = s]. \quad (20)$$

Here, $\mathbb{P}_{\text{ALG}}$ and $\mathbb{P}_{\Pi}$ denote probabilities in the cases that the learner follows ALG and Π , respectively. Equation (19) means Π has the same expected regret and mismatches as ALG. Thus it suffices to prove the theorem for Π .

First, we extend the definitions of P , R , and Q defined in Appendix D as

$$\begin{aligned}
 Q(s, d) &= \mathbb{E}_{a \sim d}[Q(s, a)] \\
 R(s, d) &= \mathbb{E}_{a \sim d}[R(s, a)] \\
 P(s, d, s') &= \mathbb{E}_{a \sim d}[P(s, a, s')]
 \end{aligned}$$

for distribution d on the actions.

Now let $M = (S, \mathcal{A}, R, P)$ be the environment. Denote the Q values, V values, and average reward of π on M by Q^π , V^π , and ρ^π , respectively. Also let ρ^* be average reward of $\pi_\dagger$ as $\pi_\dagger$ is ϵ -robust optimal on M . By Corollary 1, for any $s, a \neq \pi_\dagger(s)$ we have

$$V^{\pi_\dagger}(s) - Q^{\pi_\dagger}(s, a) = \frac{1}{\mu^{\pi_\dagger\{s;a\}}(s)} (\rho^* - \rho^{\pi_\dagger\{s;a\}})$$

$$\geq \frac{\epsilon}{\mu_{\max}}$$

Defining $\eta = \epsilon / \mu_{\max}$, we can write

$$\begin{aligned} Q^{\pi_{\dagger}}(s, q_{t,s}) &= \sum_{a \neq \pi_{\dagger}(s)} q_{t,s}(a) Q^{\pi_{\dagger}}(s, a) + q_{t,s}(\pi_{\dagger}(s)) V^{\pi_{\dagger}}(s) \\ &\leq \sum_{a \neq \pi_{\dagger}(s)} q_{t,s}(a) (V^{\pi_{\dagger}}(s) - \eta) + q_{t,s}(\pi_{\dagger}(s)) V^{\pi_{\dagger}}(s) \\ &= V^{\pi_{\dagger}}(s) - \eta \left(1 - q_{t,s}(\pi_{\dagger}(s)) \right) \\ &= V^{\pi_{\dagger}}(s) - \eta e_{t,s}, \end{aligned}$$

defining $e_{t,s} = 1 - q_{t,s}(\pi_{\dagger}(s))$. Using the definition of $Q^{\pi_{\dagger}}$ we have

$$R(s, q_{t,s}) + \sum_{s'} P(s, q_{t,s}, s') V^{\pi_{\dagger}}(s') - V^{\pi_{\dagger}}(s) + \eta e_{t,s} \leq \rho^*.$$

This can be written in vector notation as

$$R_{q_t} + (P_{q_t} - I) V^{\pi_{\dagger}} + \eta e_t \leq \rho^* \mathbf{1}.$$

Now let $d_t(s) = \mathbb{P}_{\Pi}[s_t = s]$, and d_t be the row vector of that. Thus, $d_{t+1} = d_t P_{q_t}$. Multiplying the last inequality by d_t from left gives

$$\begin{aligned} d_t R_{q_t} + (d_{t+1} - d_t) V^{\pi_{\dagger}} + \eta d_t e_t &\leq \rho^* \\ \Rightarrow \eta d_t e_t &\leq \rho^* - \mathbb{E}[r_t] + (d_t - d_{t+1}) V^{\pi_{\dagger}}. \end{aligned}$$

Summing the inequality for $t = 0$ to $T - 1$:

$$\begin{aligned} \eta \sum_{t=0}^{T-1} d_t e_t &\leq T \rho^* - \mathbb{E} \left[\sum_{t=0}^{T-1} r_t \right] + (d_0 - d_T) V^{\pi_{\dagger}} \\ &\leq \mathbb{E} [\text{REGRET}(T, M)] + 2 \|V^{\pi_{\dagger}}\|_{\infty} \end{aligned}$$

Now note that

$$\begin{aligned} \sum_{t=0}^{T-1} d_t e_t &= \sum_{t=0}^{T-1} \sum_s d_t(s) e_{t,s} \\ &= \sum_{t=0}^{T-1} \sum_s \mathbb{P}_{\Pi}[s_t = s] \left(1 - q_{t,s}(\pi_{\dagger}(s)) \right) \\ &= \sum_{t=0}^{T-1} \sum_s \mathbb{P}_{\Pi}[s_t = s] \mathbb{P}_{\Pi}[a_t \neq \pi_{\dagger}(s_t) | s_t = s] \\ &= \sum_{t=0}^{T-1} \mathbb{P}_{\Pi}[a_t \neq \pi_{\dagger}(s_t)] \\ &= T \mathbb{E} [\text{AvgMiss}(T)]. \end{aligned}$$

Therefore, we have

$$\begin{aligned} \mathbb{E} [\text{AvgMiss}(T)] &\leq \frac{\mathbb{E} [\text{REGRET}(T, M)] + 2 \|V^{\pi_{\dagger}}\|_{\infty}}{T \eta} \\ &= \frac{\mu_{\max}}{\epsilon \cdot T} \left(\mathbb{E} [\text{REGRET}(T, M)] + 2 \|V^{\pi_{\dagger}}\|_{\infty} \right). \end{aligned}$$

I. Proofs For Online Attacks: Poisoning Rewards (Section 5.2)

Proof of Theorem 3 From Lemma 7, $\widehat{R} = \overline{R} - \overline{\chi}_\epsilon^{\pi^\dagger}$ is a solution for (P3) which is the same as (P1) with an additional constraint. This means that $\widehat{R}$ is also a solution for (P1). Thus, $\pi_\dagger$ is ϵ -robust optimal in $\widehat{M} = (S, A, \widehat{R}, \overline{P})$. It is also obvious that the attack makes the learner learn in the MDP $\widehat{M}$. Lemma 2 immediately proves the bound on expected average mismatches:

$$\mathbb{E} [\text{AvgMiss}(T)] \leq \frac{K(T, \widehat{M})}{T}.$$

Since $\overline{\chi}_\epsilon^{\pi^\dagger}(s, a) = 0$ for $a = \pi_T(s)$, we have

$$\begin{aligned} \mathbb{E} \left[\sum_{t=0}^{T-1} (\widehat{R}_t(s_t, a_t) - \overline{R}(s_t, a_t))^p \right] &= \mathbb{E} \left[\sum_{t=0}^{T-1} \overline{\chi}_\epsilon^{\pi^\dagger}(s_t, a_t)^p \right] \\ &= \mathbb{E} \left[\sum_{t=0}^{T-1} \mathbb{1}[a_t \neq \pi(s_t)] \overline{\chi}_\epsilon^{\pi^\dagger}(s_t, a_t)^p \right] \\ &\leq \|\overline{\chi}_\epsilon^{\pi^\dagger}\|_\infty^p T \mathbb{E} [\text{AvgMiss}(T)] \\ &\leq \|\overline{\chi}_\epsilon^{\pi^\dagger}\|_\infty^p K(T, \widehat{M}), \end{aligned}$$

As $p \geq 1$, note that the function $f(x) = x^{1/p}$ is concave, so by Jensen inequality, for a random variable X we have $\mathbb{E}[f(X)] \leq f(\mathbb{E}[X])$. We can write

$$\begin{aligned} \mathbb{E} [\text{AvgCost}(T)] &= \frac{1}{T} \mathbb{E} \left[\left(\sum_{t=1}^T (\widehat{R}_t(s_t, a_t) - \overline{R}(s_t, a_t))^p \right)^{1/p} \right] \\ &\leq \frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T (\widehat{R}_t(s_t, a_t) - \overline{R}(s_t, a_t))^p \right]^{1/p} \\ &\leq \frac{\|\overline{\chi}_\epsilon^{\pi^\dagger}\|_\infty}{T} K(T, \widehat{M})^{1/p} \end{aligned}$$

J. Proofs For Online Attacks: Poisoning Dynamics (Section 5.3)

As we mentioned in the main text, the optimization problem (P4) can be transformed into a tractable convex problem with linear constraints – this is shown in Appendix G, where we used such an attack to obtain the upper bound of Theorem 2. We refer the reader to Section G.1, and in particular, the optimization problems (P7) and (P8) for more details.

Proof of Theorem 4 By Lemma 9, the solution $\widehat{P}$ for the problem (P4) can be efficiently calculated with a convex program with linear constraints. As (P4) is the same as (P2) with just an additional constraint, this $\widehat{P}$ is a feasible solution for (P2) and by Theorem 2, $\pi_\dagger$ is ϵ -robust optimal in $\widehat{M} = (S, A, \overline{R}, \widehat{P})$ which is what the learner observes. Lemma 2 gives the bound for expected average mismatches:

$$\mathbb{E} [\text{AvgMiss}(T)] \leq \frac{K(T, \widehat{M})}{T}$$

To show the upper bound we consider the attack (16) we used to prove the upper bound in Theorem 2:

$$\widehat{P}'(s, a, s') = (1 - \Lambda(s, a)) \cdot \overline{P}(s, a, s') + \Lambda(s, a) \cdot P_s(s, a, s')$$

Since we are using the optimal solution of (P4), and by equation (18) we have

$$\|\widehat{P} - \overline{P}\|_p \leq \|\widehat{P}' - \overline{P}\|_p \leq 2 \cdot \|\Lambda\|_p$$

Using this we can write

$$\begin{aligned}
 \mathbb{E} \left[\sum_{t=1}^T \left\| \hat{P}_t(s_t, a_t, \cdot) - \bar{P}(s_t, a_t, \cdot) \right\|_1^p \right] &= \mathbb{E} \left[\sum_{t=1}^T \left\| \hat{P}(s_t, a_t, \cdot) - \bar{P}(s_t, a_t, \cdot) \right\|_1^p \right] \\
 &= \mathbb{E} \left[\sum_{t=1}^T \mathbb{1}[a_t \neq \pi_{\dagger}(s_t)] \left\| \hat{P}(s_t, a_t, \cdot) - \bar{P}(s_t, a_t, \cdot) \right\|_1^p \right] \\
 &\leq \left\| \hat{P} - \bar{P} \right\|_{\infty}^p T \mathbb{E} [\text{AvgMiss}(T)] \\
 &\leq (2 \cdot \|\Lambda\|_{\infty})^p K(T, \widehat{M})
 \end{aligned}$$

As $p \geq 1$ similar to proof of Theorem 3 by jensen inequality we have

$$\begin{aligned}
 \mathbb{E} [\text{AvgCost}(T)] &= \frac{1}{T} \mathbb{E} \left[\left(\sum_{t=1}^T \left\| \hat{P}_t(s_t, a_t, \cdot) - \bar{P}(s_t, a_t, \cdot) \right\|_1^p \right)^{1/p} \right] \\
 &\leq \frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \left\| \hat{P}_t(s_t, a_t, \cdot) - \bar{P}(s_t, a_t, \cdot) \right\|_1^p \right]^{1/p} \\
 &\leq \frac{2}{T} \|\Lambda\|_{\infty} K(T, \widehat{M})^{1/p}
 \end{aligned}$$