
Offline Constrained Multi-Objective Reinforcement Learning via Pessimistic Dual Value Iteration

Runzhe Wu

Shanghai Jiao Tong University
runzhe@sjtu.edu.cn

Yufeng Zhang

Northwestern University
yufengzhang2023@u.northwestern.edu

Zhuoran Yang

Princeton University
zy6@princeton.edu

Zhaoran Wang

Northwestern University
zhaoranwang@gmail.com

Abstract

In constrained multi-objective RL, the goal is to learn a policy that achieves the best performance specified by a multi-objective preference function under a constraint. We focus on the offline setting where the RL agent aims to learn the optimal policy from a given dataset. This scenario is common in real-world applications where interactions with the environment are expensive and the constraint violation is dangerous. For such a setting, we transform the original constrained problem into a primal-dual formulation, which is solved via dual gradient ascent. Moreover, we propose to combine such an approach with pessimism to overcome the uncertainty in offline data, which leads to our Pessimistic Dual Iteration (PEDI). We establish upper bounds on both the suboptimality and constraint violation for the policy learned by PEDI based on an arbitrary dataset, which proves that PEDI is provably sample efficient. We also specialize PEDI to the setting with linear function approximation. To the best of our knowledge, we propose the first provably efficient constrained multi-objective RL algorithm with offline data without any assumption on the coverage of the dataset.

1 Introduction

There has been increased interest in multi-objective RL in recent years. Compared with traditional single-objective RL, the goal of the multi-objective one depends on a preference function, which takes the multiple objectives as input and outputs a scalar. The multi-objective optimization problems are usually constrained, as otherwise, it may cause danger or malfunction in applications. For example, consider a home automation system that helps humans monitor and control home attributes which can be regarded as multi-objectives. Users at different times may value different aspects of its services. Some may think highly of lighting at night, while others may concern the climate. We can formulate the users' preference as a preference function on the multiple objectives. Therefore, the system is dealing with a multi-objective optimization problem. However, the system cannot optimize this problem without any constraints, which might go against human's will, such as generating extreme climate inside a house or performing unacceptably energy-consuming operations.

We formulate the constrained multi-objective Markov decision process (CMOMDP), which is similar to the constrained Markov decision process (CMDP) (Altman, 1999). The difference is that the objectives are multiple and constraints can be nonlinear. We aim to minimize the value of a preference function, which takes multiple objectives as input and outputs a scalar.

Most existing RL methods assume full accessibility of the environment for the agent, which tends to be impractical as the exploration could be expensive (Gottesman et al., 2019) and dangerous (Shalev-Shwartz et al., 2016). Hence, we consider the offline case where the agent has only a historical dataset collected a priori and has no further interactions with the environment. This setting is common and embodied in various scenarios such as healthcare (Chakraborty and Murphy, 2014) and auto-driving (Sun et al., 2020). However, offline RL is less understood theoretically (Levine et al., 2020) than online RL, and how to maximally exploit the dataset remains unclear.

In this paper, we study the offline CMOMDP. The challenges are threefold:

- (i) Different from CMDPs, the nonlinearity of the preference function and constraints of CMOMDPs makes the analysis challenging. Moreover, constraints can be neither convex nor concave in the policy, which means the optimization problems are usually non-convex. It brings great difficulty to designing a provably efficient algorithm.
- (ii) Since further interaction with the environment is prohibited, the dataset’s quality is not guaranteed. The data collected by the experimenter may not sufficiently cover the trajectories induced by the optimal policy. Therefore, the information of the optimal policy could be limited, which probably makes the suboptimality and constraint violation arbitrarily large.
- (iii) As the CMOMDP can be considered a generalization of CMDP (see the reduction in Appendix C), any difficulties emerging in the CMDP will arise here too. For example, due to constraints, the Bellman optimality equation may not hold anymore. Therefore, most existing offline RL algorithms based on dynamic programming are inapplicable.

Challenge (i) is inherent in CMOMDPs, and certain requirements on the preference function and constraints are inevitable. Hence, We suppose they satisfy some conditions. For instance, a geometric analog of Slater’s condition, which is usually assumed in CMDPs, is imposed. To tackle challenge (ii), existing works have put a great effort. One possible principle is pessimism, which applies penalties to ensure pessimistic estimation. This method is successful in ordinary RL, and we extend it to CMOMDPs. We also note that we only impose a minimal assumption on the offline dataset’s compliance in our analysis. For challenge (iii), we apply the convex conjugate and Fenchel’s duality to transform the problem into a primal-dual formulation.

In summary, our work answers the following question:

Is it possible to develop a provably efficient offline algorithm for constrained multi-objective reinforcement learning with minimal assumptions on the dataset?

We propose the Pessimistic Dual Iteration (PEDI) algorithm. Theoretical contributions are as follows:

- (i) By transforming the original constrained optimization problem of CMOMDPs into a primal-dual formulation via convex conjugate and duality, we develop the algorithm for general CMOMDPs with offline dataset, which iterates in a dual gradient ascent manner. Then, we instantiate the algorithm for linear kernel CMOMDPs, which is a large class of CMOMDPs that includes the tabular case.
- (ii) We show in Appendix F the significance of pessimism in offline CMOMDPs. To summarize, we decompose the discrepancy between a value function and the optimal one into spurious correlation, intrinsic uncertainty, and optimization error, among which the spurious correlation is the most difficult to control. However, by maintaining pessimistic estimates of the value functions, PEDI eliminates the spurious correlation.
- (iii) We establish theoretical guarantees for PEDI. Specifically, we demonstrate that two metrics, the suboptimality and the constraint violation, can be bounded from above by the optimization error of $\mathcal{O}(1/\sqrt{K})$ and the intrinsic uncertainty, where K is the number of iterations. When specialized to linear kernel CMOMDPs, the upper bound of suboptimality matches the information-theoretic lower bound up to the optimization error and multiplicative factors of constants related to the CMOMDP, which suggests the near-optimality of PEDI.
- (iv) We show that the error of PEDI is data-dependent, i.e., it depends on how well the dataset covers the trajectories induced by the optimal policy. When the trajectories are assumed further to be sufficiently covered, the suboptimality and constraint violation are $\tilde{\mathcal{O}}(1/\sqrt{N})$ where N is the number of trajectories in the dataset.

To the best of our knowledge, we are the first to propose a provably efficient offline algorithm that considers the constrained multi-objective RL without any assumptions on the coverage of the dataset. As a by-product, our method can be viewed as a highly generalized one as it can easily reduce to certain simpler cases such as CMDPs, which is discussed in Appendix C.

1.1 Related works

Constrained RL. Our work is closely related to CMDPs (Altman, 1999). Several recent works (Efroni et al., 2020; Ding et al., 2021; Qiu et al., 2020; Brantley et al., 2020; Tessler et al., 2018) have studied the MDP with linear constraints and others (Bhatnagar and Lakshmanan, 2012; Chow et al., 2017; Paternain et al., 2019) have considered RL with more complex constraints. Most of them have provided algorithms with both regret and total constraint violation guarantees. However, while these approaches are successful in CMDPs, they cannot be applied to CMOMDPs where objectives are multiple and constraints can be nonlinear. The difference between their methods and ours are (i) they usually assume the Slater’s condition, while we consider its geometric analog, and (ii) most works utilize the Lagrangian multiplier to handle constraints, while in our method, we obtain a primal-dual formulation via conjugate.

Multi-objective RL. Existing studies on multi-objective RL can be divided into two groups: the first maintains a set of Pareto-optimal policies (Barrett and Narayanan, 2008; Castelletti et al., 2011, 2012; Wang and Sebag, 2013), and the second collapses multi-objective vector into a scalar and then apply a standard RL method (Barrett and Narayanan, 2008; Natarajan and Tadepalli, 2005; Van Moffaert et al., 2013; Cheung, 2019). However, these methods are available in unconstrained multi-objective RL but unsuitable for constrained cases. The difficulty is that the constraints invalidate the Bellman optimality principle, which is the cornerstone of these works.

Offline RL. Recent successful offline RL (Lange et al., 2012; Levine et al., 2020) methods (Fujimoto et al., 2019; Laroche et al., 2019; Jaques et al., 2019; Wu et al., 2019; Kumar et al., 2019, 2020; Agarwal et al., 2020; Yu et al., 2020; Kidambi et al., 2020; Siegel et al., 2020; Nair et al., 2020; Liu et al., 2020; Wang et al., 2020a,b; Tosatto et al., 2017; Farahmand et al., 2016, 2010) fall into two categories (Buckman et al., 2020): (i) uncertainty-aware pessimistic algorithms and (ii) proximal pessimistic algorithms. (i) gives a lower estimate for states and actions less covered by the dataset, while (ii) avoids visiting these less visited states and actions by regularization. However, existing methods are largely based on approximate dynamic programming and do not take constraints into account, thus failing to apply. Our proposed method generalizes the ordinary pessimistic algorithm so that it is compatible with multiple objectives and constraints.

In the analysis, we only assume the compliance of the dataset. In comparison with existing works, which assume the sufficient coverage of the dataset such as the finite concentrability coefficients (Antos et al., 2008; Farahmand et al., 2010; Scherrer et al., 2015; Le et al., 2019; Chen and Jiang, 2019; Fu et al., 2020) and uniformly lower bounded densities of visitation measures (Yin et al., 2021), we require no assumptions as they often fail to hold in practice. As Wang et al. (2020a) have studied the influence of insufficient coverage of the dataset, their conclusion is also reflected in our main results. Moreover, we do not impose any constraints on the affinity of the behavior policy and the learned policy (Liu et al., 2020). In addition, our general algorithm uses convex conjugate, a technique used similarly in Yu et al. (2021); Miryoosefi and Jin (2021) but their method cannot apply to offline situations and Yu et al. (2021) do not address the function approximation setting.

1.2 Notations

For simplicity, we write $\langle f, g \rangle_{\mathcal{X}} = \int_{\mathcal{X}} f(x)g(x) dx$ as the inner product on space $\mathcal{X}$ for functions $f, g : \mathcal{X} \rightarrow \mathbb{R}$. We also define the norms $\|f\|_{2,\mathcal{X}} = \sqrt{\langle f, f \rangle_{\mathcal{X}}}$ and $\|f\|_{\infty,\mathcal{X}} = \sup_{x \in \mathcal{X}} |f(x)|$. We denote by $[n]$ the set of positive integers not greater than n , i.e., $[n] = \{i \in \mathbb{Z} | 1 \leq i \leq n\}$. For two vectors $\mathbf{x}, \mathbf{x}' \in \mathbb{R}^D$, we write $\mathbf{x} \geq \mathbf{x}'$ (or $\mathbf{x} \leq \mathbf{x}'$) if $\mathbf{x}$ is not smaller (or larger) than $\mathbf{x}'$ in an elementwise manner. We set $|\mathbf{x}| = (|x_1|, |x_2|, \dots, |x_D|)^\top$ and $\mathbf{x}_+ = (\max\{x_1, 0\}, \max\{x_2, 0\}, \dots, \max\{x_D, 0\})^\top$. We denote by $\mathcal{B}^D$ the D -dimensional unit Euclidean ball, i.e., $\mathcal{B}^D = \{\mathbf{x} \in \mathbb{R}^D : \|\mathbf{x}\|_2 \leq 1\}$. The set of probability distributions on a set $\mathcal{X}$ is denoted by $\Delta(\mathcal{X})$. We denote by $\tilde{O}(\cdot)$ the big O notation ignoring logarithmic factors.

2 Preliminaries

2.1 Constrained multi-objective Markov decision process (CMOMDP)

We consider an episodic CMOMDP given by $\mathcal{M} = (\mathcal{S}, \mathcal{A}, H, \mathcal{P}, c)$ where $\mathcal{S}$ is a state space, $\mathcal{A}$ is an action space, $H \in \mathbb{N}_+$ is the horizon, $\mathcal{P} = \{\mathcal{P}_h\}_{h=1}^H$ is a collection of transition kernels $\mathcal{P}_h : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$, and $c = \{c_h\}_{h=1}^H$ is a collection of cost functions $c_h : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]^D$. Note that by $\mathcal{P}$ we mean both the probability density function and the probability mass function, as they can be unified (see Appendix A for details). We suppose a fixed initial state $\underline{s} \in \mathcal{S}$ for simplicity. However, a stochastic initial state poses no extra difficulty to our analysis. In each episode, given a policy $\pi = \{\pi_h\}_{h=1}^H$ where $\pi_h : \mathcal{S} \rightarrow \Delta(\mathcal{A})$, the agent interacts with the environment as follows. At state s_h , the agent takes action $a_h \in \mathcal{A}$ according to $\pi_h(\cdot | s_h)$, and then receives a stochastic cost $c_h \in [0, 1]^D$, for which we assume $\mathbb{E}[c_h | s_h, a_h] = c_h(s, a)$. The system then transits to the next state s_{h+1} according to $\mathcal{P}_h(\cdot | s_h, a_h)$. The episode terminates at state s_{H+1} where no action is taken and the cost is zero.

Given a policy π , the state-value function $V_h^\pi : \mathcal{S} \rightarrow [0, H]^D$ and the action-value function $Q_h^\pi : \mathcal{S} \times \mathcal{A} \rightarrow [0, H]^D$ at the h -th step are defined as $V_h^\pi(s) = \mathbb{E}_\pi[\sum_{i=h}^H c_i(s_i, a_i) | s_h = s]$ and $Q_h^\pi(s, a) = \mathbb{E}_\pi[\sum_{i=h}^H c_i(s_i, a_i) | s_h = s, a_h = a]$ where the expectation $\mathbb{E}_\pi$ is taken with respect to $a_i \sim \pi_i(\cdot | s_i)$ and $s_{i+1} \sim \mathcal{P}(\cdot | s_i, a_i)$ for $i \in [h : H]$. Note that those value functions are D -dimensional vector-valued, and we use bold letters to represent vectors. We denote their i -th scalar components by $V_h^{i,\pi}$ and $Q_h^{i,\pi}$ respectively. For notational simplicity, we write: $\mathbb{P}_h[V](s, a) = \mathbb{E}_{s' \sim \mathcal{P}_h(\cdot | s, a)}[V(s')]$, $\mathbb{B}_h[V](s, a) = c_h(s, a) + \mathbb{P}_h[V](s, a)$, and $\mathbb{D}_\pi[Q](s) = \mathbb{E}_{a \sim \pi(\cdot | s)}[Q(s, a)]$. Thereby, the Bellman equation for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and $h \in [H]$ can be written as $Q_h^\pi(s, a) = \mathbb{B}_h[V_h^\pi](s, a)$, and $V_h^\pi(s) = \mathbb{D}_\pi[Q_h^\pi](s)$. Our goal is to minimize the value of a preference function subject to some constraints. Formally, we define $g : \mathbb{R}^D \rightarrow \mathbb{R}$ as the preference function and assume that g is 1-Lipschitz and convex and that $g(\mathbf{x}) \geq g(\mathbf{x}')$ holds as long as $\mathbf{x} \geq \mathbf{x}'$. For example, $g(\mathbf{x})$ can be the summation over $\{x_i\}_{i=1}^n$ or the maximum among $\{x_i\}_{i=1}^n$. We aim to find the optimal policy π^* of the following constrained optimization problem,

$$\min_{\pi \in \Delta(\mathcal{A} | \mathcal{S}, H)} g(V_1^\pi(\underline{s})) \quad \text{s.t.} \quad V_1^\pi(\underline{s}) \in \mathcal{W}^*, \quad (1)$$

where $\mathcal{W}^* \subseteq [0, H]^D$ is a target set. We use $V_1^{\pi^*}(\underline{s}) = V_1^{\pi^*}(\underline{s})$ to denote the state-value function of the optimal policy of (1) for simplicity.

Solving (1) is impossible when we have no geometric requirements on the target set $\mathcal{W}^*$. Corresponding to the Slater's condition¹ commonly imposed in CMDPs, we establish its geometric analogue for CMOMDPs, as studied in Yu et al. (2021). Let $\mathcal{V}$ be the set of achievable values, i.e., $\mathcal{V} = \{V_1^\pi(\underline{s}) : \text{any policy } \pi\}$, and $\mathcal{W} = \mathcal{V} \cap \mathcal{W}^*$ be the achievable values within the target set. We suppose $\mathcal{W}$ is nonempty. We denote by $\partial\mathcal{W}^*$ and $\partial\mathcal{V}$ the boundaries of $\mathcal{W}^*$ and $\mathcal{V}$, respectively. With a little abuse of notations, we set $\partial\mathcal{W} = \partial\mathcal{W}^* \cap \partial\mathcal{V}$ as the intersection of boundaries. If $\partial\mathcal{W}$ is nonempty, then for each $\mathbf{W} \in \partial\mathcal{W}$, we define the angle between support vectors² at $\mathbf{W}$ as

$$\gamma(\mathbf{W}) = \min \{ \angle(\mathbf{a}, \mathbf{b}) \mid \mathbf{a}, \mathbf{b} \text{ are support vectors of } \mathcal{W}^* \text{ and } \mathcal{V} \text{ at } \mathbf{W}, \text{ respectively} \}$$

where $\angle(\mathbf{a}, \mathbf{b})$ denotes the angle between $\mathbf{a}$ and $\mathbf{b}$. We impose geometric requirements on the target set $\mathcal{W}^*$ as stated below.

Assumption 1. *We assume the target set $\mathcal{W}^*$ is closed and convex, and there exists an upper bound $\gamma_{\max} \in [\frac{\pi}{2}, \pi)$ such that $\max_{\mathbf{W} \in \partial\mathcal{W}} \gamma(\mathbf{W}) < \gamma_{\max}$. Moreover, we assume $\mathcal{W}^*$ is a lower set, i.e., for any $\mathbf{W} \in \mathcal{W}^*$, it holds that $\mathbf{W}' \in \mathcal{W}^*$ for any $\mathbf{0} \leq \mathbf{W}' \leq \mathbf{W}$.*

We write $\rho = 2/\sin(\gamma_{\max})$ for notational simplicity. Explanations and justifications of Assumption 1 are provided in Appendix B where we show that this assumption is reasonable and closely related to Slater's condition. However, the optimization problem (1) is still hard even with these assumptions. Instead, we solve it by considering the suboptimality and constraint violation as performance metrics, which are defined as

$$\text{SubOpt}(\pi) = g(V_1^\pi(\underline{s})) - g(V_1^{\pi^*}(\underline{s})), \quad \text{Violation}(\pi) = \text{dist}(V_1^\pi(\underline{s}), \mathcal{W}^*), \quad (2)$$

¹For CMDPs with constraints in the form of $\mathbf{V}^\pi \leq \mathbf{b}$, the Slater's condition assumes the existence of a policy π such that $\mathbf{V}^\pi < \mathbf{b}$.

²The support vectors at a point are normals to the support hyperplanes at this point.

where we define $\text{dist}(\mathbf{V}_1^\pi(\underline{s}), \mathcal{W}^*) = \min_{\mathbf{W} \in \mathcal{W}^*} \|\mathbf{V}_1^\pi(\underline{s}) - \mathbf{W}\|_2$.

To show that the formulation of CMOMDPs is reasonable, we state the relationship between CMOMDPs and CMDPs in the following proposition, which suggests that the CMOMDP is a generalization of the CMDP. The proof is deferred to Appendix C.1.

Proposition 1. *The CMDP is a special case of the CMOMDP, and the Slater’s condition assumed in CMDPs corresponds to the existence of $\gamma_{\max}$ in Assumption 1 in CMOMDPs.*

2.2 Offline learning

We consider the offline setting where we only have access to a dataset $\mathcal{D} = \{(s_h^\tau, a_h^\tau, c_h^\tau)\}_{h,\tau=1}^{H,N}$ with N trajectories collected a priori by an experimenter. We only make the following assumption on the data collecting process.

Assumption 2 (Compliance of dataset). *We assume that the dataset $\mathcal{D}$ is compliant with the CMOMDP $\mathcal{M}$, i.e., for any $\mathbf{c} \in [0, 1]^D$, $s' \in \mathcal{S}$, $h \in [H]$, and $\tau \in [N]$,*

$$\Pr_{\mathcal{D}}(s_{h+1}^\tau = s', c_h^\tau = \mathbf{c} \mid \mathcal{F}_h^\tau) = \Pr(s_{h+1} = s', c_h = \mathbf{c} \mid s_h = s_h^\tau, a_h = a_h^\tau), \quad (3)$$

where $\mathcal{F}_h^\tau$ is the σ -algebra generated via $\mathcal{F}_h^\tau = \sigma(\{(s_i^n, a_i^n, c_i^n) \mid (n-1)H + i \leq (\tau-1)H + h\})$. The probability on the left-hand side of (3) is with respect to the joint distribution of the data collecting process, and that of the right-hand side is with respect to the underlying CMOMDP.

Assumption 2 is adapted from Jin et al. (2020b). It implies that the dataset $\mathcal{D}$ is generated from the CMOMDP and possesses the Markov property. Such an assumption holds when the experimenter collects the dataset by interacting with the CMOMDP. In particular, we do not require that the data collecting process well explores the state-action space. In other words, we do not impose any assumption on the coverage of the dataset.

3 Pessimistic Dual Iteration (PEDI)

In this section, we first motivate our method to solve CMOMDPs in Section 3.1. Then, we develop our algorithm, Pessimistic Dual Iteration (PEDI), for general CMOMDPs in Section 3.2 and introduce an instantiation for linear kernel CMOMDPs in Section 3.3.

3.1 Primal-dual formulation

To solve problems like (1), recent works usually apply Lagrangian multipliers to handle the constraints. However, we seek to use another dual method. Specifically, we consider the following problem:

$$p^* = \min_{\pi} (\text{SubOpt}(\pi) + \nu \text{Violation}(\pi)), \quad (4)$$

where ν is a scaling constant that serves as a trade-off between suboptimality and constraint violation. In Appendix D, we show that when $\nu > 1$, (1) and (4) share the solution, and thus they are equivalent. We note that the formulation of (4) is different from existing works on CMDPs where ν usually plays the role of the Lagrangian multiplier. Intuitively, a large ν lies more emphasis on satisfying the constraints, while a small ν focus on reducing the suboptimality.

However, p^* is hard to solve in (4), since the relationship between the objective and the policy π is complex. Therefore, we substitute the policy π with the state-value function in the following way,

$$\begin{aligned} p^* &= \min_{\pi} \left(g(\mathbf{V}_1^\pi(\underline{s})) - g(\mathbf{V}_1^*(\underline{s})) + \nu \text{dist}(\mathbf{V}_1^\pi(\underline{s}), \mathcal{W}^*) \right) \\ &= \min_{\mathbf{V} \in \mathcal{V}} \left(g(\mathbf{V}) - g(\mathbf{V}_1^*(\underline{s})) + \nu \text{dist}(\mathbf{V}, \mathcal{W}^*) \right) \end{aligned} \quad (5)$$

Nevertheless, solving this problem is still hard since the objective is nonlinear in $\mathbf{V}$ and thus standard RL cannot apply. Therefore, we utilize the convex conjugate to “linearize” the objective. This technique is widely used in related works (Miryoosefi and Jin, 2021; Yu et al., 2021). Specifically, by convex conjugate, we have

$$g(\mathbf{V}) = \max_{\beta \in \mathcal{B}^D} (\beta^\top \mathbf{V} - g^*(\beta)), \quad \text{dist}(\mathbf{V}, \mathcal{W}^*) = \max_{\alpha \in \mathcal{B}^D} (\alpha^\top \mathbf{V} - \max_{\mathbf{x} \in \mathcal{W}^*} \alpha^\top \mathbf{x}) \quad (6)$$

for any $\mathbf{V} \in \mathcal{V}$. The effective domains, $\mathcal{B}^D$, of both α and β are obtained by the 1-Lipschitzness and Corollary 13.3.3 in Rockafellar (1970). To see why the second equality of (6) holds, we notice that

$$\text{dist}(\mathbf{V}, \mathcal{W}^*) = \max_{\alpha \in \mathcal{B}^D} \left(\alpha^\top \mathbf{V} - \underbrace{\max_{\mathbf{x} \in \mathbb{R}^D} (\alpha^\top \mathbf{x} - \text{dist}(\mathbf{x}, \mathcal{W}^*))}_{(*)} \right),$$

and $(*)$ attains the maximum when $\mathbf{x} \in \mathcal{W}^*$ since $\alpha \in \mathcal{B}^D$. We call α and β the dual variables throughout this paper. By plugging (6) into (5), we have

$$\mathbf{p}^* = \min_{\mathbf{V} \in \mathcal{V}} \max_{\alpha, \beta \in \mathcal{B}^D} \mathcal{L}(\mathbf{V}; \alpha, \beta) = \beta^\top \mathbf{V} - g^*(\beta) - g(\mathbf{V}_1^*(s)) + \nu \alpha^\top \mathbf{V} - \nu \max_{\mathbf{x} \in \mathcal{W}^*} \alpha^\top \mathbf{x}. \quad (7)$$

Its dual problem can be written as $\mathbf{d}^* = \max_{\alpha, \beta \in \mathcal{B}^D} D(\alpha, \beta)$, where $D(\alpha, \beta)$ is the dual function defined as $D(\alpha, \beta) = \min_{\mathbf{V} \in \mathcal{V}} \mathcal{L}(\mathbf{V}; \alpha, \beta)$. We notice that $\mathcal{L}(\mathbf{V}; \alpha, \beta)$ is convex in $\mathbf{V} \in \mathcal{V}$ and concave in $\alpha, \beta \in \mathcal{B}^D$, which implies the strong duality, $\mathbf{p}^* = \mathbf{d}^*$. Hence, dual methods can apply.

When α and β are fixed, solving $D(\alpha, \beta)$ is a planning problem. When the model is known, it suffices to conduct value iteration on the model. However, what we have is a dataset collected a priori in offline RL. Therefore, we develop the offline planning algorithm that is a pessimistic variant of the value iteration (Jin et al., 2020b), which applies penalties to ensure pessimism. We explain why pessimism is crucial to offline CMOMDP in Appendix F. We now describe the high-level intuition of our approach. Consider a meta-algorithm that constructs the empirical transition kernel $\hat{\mathcal{P}}$ and the empirical cost function $\hat{c}$ based on the dataset $\mathcal{D}$. As defined below, we introduce the uncertainty quantifier that bounds the estimation error from above with a confidence parameter $\xi \in (0, 1)$.

Definition 1 (ξ -uncertainty quantifier). *We call $\{(\Gamma_h^{\mathcal{P}}, \Gamma_h^c)\}_{h=1}^H$ a ξ -uncertainty quantifier if the following event*

$$\mathcal{E} = \left\{ \left| \hat{\mathcal{P}}_h(s'|s, a) - \mathcal{P}_h(s'|s, a) \right| \leq \Gamma_h^{\mathcal{P}}(s, a, s'), \left| \hat{c}_h^i(s, a) - c_h^i(s, a) \right| \leq \Gamma_h^c(s, a) \right. \\ \left. \text{for all } (s, a, s') \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}, h \in [H], i \in [D] \right\} \quad (8)$$

holds with probability at least $1 - \xi$. Here we write $\Gamma_h^c = (\Gamma_h^{c^1}, \Gamma_h^{c^2}, \dots, \Gamma_h^{c^D})$.

Then, we can construct the empirical counterparts of $\mathbb{P}_h$ and $\mathbb{B}_h$ by $\hat{\mathbb{P}}_h[\mathbf{V}](s, a) = \langle \mathbf{V}(\cdot), \hat{\mathcal{P}}_h(\cdot | s, a) \rangle_{\mathcal{S}}$ and $\hat{\mathbb{B}}_h[\mathbf{V}](s, a) = \hat{c}_h(s, a) + \hat{\mathbb{P}}_h[\mathbf{V}](s, a)$. Note that for any function $V : \mathcal{S} \rightarrow [0, H]$, under the event $\mathcal{E}$, it holds for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and any $h \in [H]$ that

$$\left| (\hat{\mathbb{P}}_h - \mathbb{P}_h)[V](s, a) \right| = \left| \int_{\mathcal{S}} (\hat{\mathcal{P}}_h(s'|s, a) - \mathcal{P}_h(s'|s, a)) V(s') ds' \right| \leq H \int_{\mathcal{S}} \Gamma_h^{\mathcal{P}}(s, a, s') ds'.$$

Therefore, we define $\Gamma_h(s, a) = H \int_{\mathcal{S}} \Gamma_h^{\mathcal{P}}(s, a, s') ds'$, which quantifies the uncertainty arising from the transition kernel. We use the ξ -uncertainty quantifier as the penalty function for pessimistic planning, which leads to a conservative estimation of the value function. We formally describe our planning method in Algorithm 1, where θ denotes a projection vector to be specified later. By adding the pessimistic penalty $\Gamma_h \cdot \mathbf{1} + \Gamma_h^c$ to the estimated value, we ensure a conservative estimation. We also truncate $\mathbf{Q}_h^k$ and $\mathbf{V}_h^k$ in $[0, H - h + 1]^D$ to improve the estimation.

It remains to solve $\max_{\alpha, \beta \in \mathcal{B}^D} D(\alpha, \beta)$. We utilize the projected subgradient method which produces a sequence of dual variables $\{\alpha^k\}_{k=1}^K$ and $\{\beta^k\}_{k=1}^K$ that approximately solves the dual problem. A detailed description of projected subgradient method is provided in Appendix H.6.2. For a brief description, we update the dual variables via

$$\alpha^{k+1} \leftarrow \Pi_{\mathcal{B}^D} \left\{ \alpha^k + \eta^k (\mathbf{V}^k - \arg \max_{\mathbf{x} \in \mathcal{W}^*} (\alpha^k)^\top \mathbf{x}) \right\}, \\ \beta^{k+1} \leftarrow \Pi_{\mathcal{B}^D} \left\{ \beta^k + \eta^k (\mathbf{V}^k - \partial g^*(\beta^k)) \right\}. \quad (9)$$

where $\mathbf{V}^k \in \arg \min_{\mathbf{V} \in \mathcal{V}} \mathcal{L}(\mathbf{V}; \alpha^k, \beta^k)$ and η^k is the step length at the k -th iteration.

As claimed in Proposition 1, the CMDP is a special case of the CMOMDP. Moreover, we show in Appendix C.2 that when reducing to the CMDP, the objective (7) reduces to the Lagrangian formulation of the CMDP problem. Hence, our proposed method has good generalization ability.

Algorithm 1 Pessimistic planning.

```

1: function PESSPLANNING( $\theta, \mathcal{D}, \{(\Gamma_h^{\mathcal{P}}, \Gamma_h^{\mathcal{C}})\}_{h=1}^H$ )
2:   for step  $h = H, H-1, \dots, 1$  do
3:      $\bar{Q}_h(\cdot, \cdot) \leftarrow \hat{c}_h(\cdot, \cdot) + \hat{\mathbb{P}}_h[\mathbf{V}_{h+1}](\cdot, \cdot) + \Gamma_h \cdot \mathbf{1} + \Gamma_h^{\mathcal{C}}$ 
4:      $Q_h(\cdot, \cdot) \leftarrow \min \{ \bar{Q}_h(\cdot, \cdot), (H-h+1) \cdot \mathbf{1} \}_+$ 
5:      $\pi_h(\cdot) \leftarrow \arg \min_{a \in \mathcal{A}} Q_h(\cdot, a) \cdot \theta$ 
6:      $V_h(\cdot) \leftarrow \mathbb{D}_{\pi_h}[Q_h](\cdot)$ 
7:   end for
8:   return  $\pi = \{\pi_h\}_{h=1}^H, Q = \{Q_h\}_{h=1}^H, V = \{V_h\}_{h=1}^H$ 
9: end function

```

3.2 General CMOMDPs

Based on the above analysis, we now develop the Pessimistic Dual Iteration (PEDI) (Algorithm 2) for general CMOMDPs, which is motivated by the dual gradient method in solving the minimax optimization problem. Specifically, PEDI solves $D(\alpha, \beta)$ via PESSPLANNING (Algorithm 1) and updates dual variables via (9), alternately. The output policy $\hat{\pi}$ is a mixed policy executed by randomly selecting a policy π^k from $\{\pi^k\}_{k=1}^K$ with equal probability beforehand and then exclusively following π^k thereafter. Hence, it holds that, for the mixed policy $\hat{\pi}$, $V_1^{\hat{\pi}}(s) = \frac{1}{K} \sum_{k=1}^K V_1^{\pi^k}(s)$. Note that the mixed policy may not be a Markov policy (Altman, 1999). However, it would be useful to extend the definition of a policy $\pi = \{\pi_h\}_{h=1}^H$ so as to allow π_h to depend not only on h but also on some initial randomizing mechanism. That is, we may have a set of policies $\mathcal{U}$ and a distribution $q \in \Delta(\mathcal{U})$. Then we define the mixed policy of $\mathcal{U}$, $\pi_{\mathcal{U}}$, as a policy to be executed by first using q to choose some policy $\pi \in \mathcal{U}$ and then proceeding with only that policy. The optima of (1) over the mixed policy is still attained by the Markov policy π^* . Thus, $\text{SubOpt}(\hat{\pi})$ remains a reasonable performance metric of our algorithm.

Algorithm 2 Pessimistic Dual Iteration (PEDI).

Input: offline dataset $\mathcal{D}$, step length $\{\eta^k\}_{k=1}^K$, scaling parameter ν
Output: the mixed policy $\hat{\pi} = \{\hat{\pi}_h\}_{h=1}^H$ of $\{\pi^k\}_{k=1}^K$

```

1: randomly initialize dual variables  $\alpha^1, \beta^1 \in \mathcal{B}^D$ , and set  $\theta^1 = \nu\alpha^1 + \beta^1$ 
2: construct  $\{\hat{c}_h\}_{h=1}^H$  and  $\{\hat{\mathcal{P}}_h\}_{h=1}^H$  based on  $\mathcal{D}$ 
3: construct the  $\xi$ -uncertainty quantifier  $\{(\Gamma_h^{\mathcal{P}}, \Gamma_h^{\mathcal{C}})\}_{h=1}^H$  based on  $\mathcal{D}$ 
4: for iteration  $k = 1, 2, \dots, K$  do
5:    $\pi^k, Q^k, V^k \leftarrow \text{PESSPLANNING}(\theta^k, \mathcal{D}, \{(\Gamma_h^{\mathcal{P}}, \Gamma_h^{\mathcal{C}})\}_{h=1}^H)$ 
6:   update  $\alpha^k, \beta^k$  to  $\alpha^{k+1}, \beta^{k+1}$  via (9)
7:    $\theta^{k+1} \leftarrow \nu\alpha^{k+1} + \beta^{k+1}$ 
8: end for

```

3.3 Linear kernel CMOMDPs

In this section, we study PEDI on settings with linear function approximation such as linear kernel CMOMDPs. To that end, we first introduce the linear kernel CMOMDP, which is a generalization of the linear kernel MDP (Yang and Wang, 2019; Jin et al., 2020a; Cai et al., 2020). Then, we propose an instantiation of Algorithm 2 for linear kernel CMOMDPs.

Definition 2 (Linear kernel CMOMDP). *The CMOMDP $(\mathcal{S}, \mathcal{A}, H, \mathcal{P}, c)$ is a linear kernel CMOMDP with a known kernel feature map $\psi : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}^{d_1}$ and a known value feature map $\varphi : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}^{d_2}$, if for any $h \in [H]$ there exist unknown vectors $\theta_h \in \mathbb{R}^{d_1}$ and $\theta_h^i \in \mathbb{R}^{d_2}$ for any $i \in [D]$ such that for any $(s, a, s') \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}$,*

$$\mathcal{P}_h(s' | s, a) = \psi(s, a, s')^\top \theta_h, \quad c_h^i(s, a) = \varphi(s, a)^\top \theta_h^i.$$

Let $d = \max\{d_1, d_2\}$. We assume there exists a constant $R > 0$ such that

$$R^{-2} \max_{s' \in \mathcal{S}} (y^\top \psi(s, a, s'))^2 \leq \int_{\mathcal{S}} (y^\top \psi(s, a, s'))^2 ds' \leq d$$

for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and $y : \|y\|_2 \leq 1$. Moreover, we assume $\|\theta_h\|_2 \leq \sqrt{d}$ and $\|\theta_h^c\|_2 \leq \sqrt{d}$ for $i \in [d]$.

The existence of constant R naturally holds when $\psi(s, a, \cdot)$ is upper bounded and Lipschitz continuous for all $(s, a) \in \mathcal{S} \times \mathcal{A}$. Note that the tabular CMOMDP is a special case of the linear kernel CMOMDP with $R = 1$ (see Appendix G for a detailed discussion).

An instantiation of PEDI for linear kernel CMOMDPs is deferred to Appendix E where we propose a method to estimate the empirical transition kernel $\hat{\mathcal{P}}$ and cost function $\hat{c}$ and thereby construct a ξ -uncertainty quantifier.

4 Theoretical results

In this section, we first show upper bounds for suboptimality and constraint violation for PEDI on general CMOMDPs in Section 4.1. In Section 4.2, we show that PEDI achieves the minimax optimality up to the optimization error and multiplicative factors of constants related to the CMOMDP when specified to linear kernel CMOMDPs. Moreover, when the dataset has sufficient coverage over the trajectories induced by the optimal policy, the intrinsic uncertainty is guaranteed to be $\tilde{\mathcal{O}}(1/\sqrt{N})$ where N is the number of trajectories in the dataset.

4.1 Results of general CMOMDPs

By using pessimistic planning (Algorithm 1), we ensure a pessimistic policy $\hat{\pi}$. We formally state this notion in the following lemma.

Lemma 1 (Pessimism). *Suppose that Assumptions 1 and 2 hold. By Algorithm 2, under event $\mathcal{E}$ defined in (8), for any $(s, a) \in \mathcal{S} \times \mathcal{A}$ and $h \in [H]$, it holds that $\mathbf{V}_h^k(s) \geq \mathbf{V}_h^{\pi^k}(s)$, $\mathbf{Q}_h^k(s, a) \geq \mathbf{Q}_h^{\pi^k}(s, a)$, and $\text{dist}(\mathbf{V}_1^k(\underline{s}), \mathcal{W}^*) \geq \text{dist}(\mathbf{V}_1^{\pi^k}(\underline{s}), \mathcal{W}^*)$. Furthermore, it holds that*

$$g(\hat{\mathbf{V}}_1(\underline{s})) \geq g(\mathbf{V}_1^{\hat{\pi}}(\underline{s})) \quad \text{and} \quad \text{dist}(\hat{\mathbf{V}}_1(\underline{s}), \mathcal{W}^*) \geq \text{dist}(\mathbf{V}_1^*(\underline{s}), \mathcal{W}^*),$$

where $\hat{\mathbf{V}}_1(\underline{s}) = \frac{1}{K} \sum_{k=1}^K \mathbf{V}_1^k(\underline{s})$ and $\hat{\pi}$ is the output of Algorithm 2.

Proof of Lemma 1. See Appendix H.1 for a detailed proof. $\square$

We are now ready to establish the following theorem, which characterizes the suboptimality and constraint violation of PEDI for general offline CMOMDPs.

Theorem 1 (Suboptimality and constraint violation for general CMOMDPs). *Suppose that Assumptions 1 and 2 hold. Under event $\mathcal{E}$ defined in (8), for the output policy $\hat{\pi}$ of Algorithm 2, when we set $\nu = \rho$ and $\eta^k = 2G^{-1}\sqrt{D/k}$ (or $2G^{-1}\sqrt{D/K}$ if K is predefined), where $G = 2(1 + \nu)H\sqrt{D}$, it holds that*

$$\text{SubOpt}(\hat{\pi}) \leq \epsilon_K + \text{IntUncert}_{\mathcal{D}}^{\pi^*}, \quad \text{Violation}(\hat{\pi}) \leq \frac{2}{\rho}(\epsilon_K + \text{IntUncert}_{\mathcal{D}}^{\pi^*}), \quad (10)$$

where we let C denote an absolute constant and define

$$\epsilon_K = C(1+\rho)\sqrt{\frac{DH^2}{K}}, \quad \text{IntUncert}_{\mathcal{D}}^{\pi^*} = 2(1+\rho)\sqrt{D} \sum_{h=1}^H \mathbb{E}_{\pi^*} [\Gamma_h(s_h, a_h) + \|\Gamma_h^c(s_h, a_h)\|_\infty \mid s_1 = \underline{s}].$$

Proof of Theorem 1. See Appendix H.2 for a detailed proof. Here we provide a proof sketch, which consists of two steps.

First step: upper bounding $\text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi})$, which is our objective (4). Here $\hat{\pi}$ is the output of PEDI (Algorithm 2). By pessimism, we have $\text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi}) \leq$

$g(\widehat{\mathbf{V}}_1(\underline{s})) - g(\mathbf{V}_1^*(\underline{s})) + \nu \text{dist}(\widehat{\mathbf{V}}_1(\underline{s}), \mathcal{W}^*)$ where $\widehat{\mathbf{V}}_1(\underline{s}) = \frac{1}{K} \sum_{k=1}^K \mathbf{V}_1^k(\underline{s})$. We do linearization via convex conjugate (6) and then get

$$\begin{aligned} \text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi}) &\leq K^{-1} \max_{\|\beta\| \leq 1} \left\{ \beta \cdot \sum_{k=1}^K \mathbf{V}_1^k(\underline{s}) - \sum_{k=1}^K g^*(\beta) \right\} - g(\mathbf{V}_1^*(\underline{s})) \\ &\quad + K^{-1} \nu \max_{\|\alpha\| \leq 1} \left\{ \alpha \cdot \sum_{k=1}^K \mathbf{V}_1^k(\underline{s}) - \sum_{k=1}^K \max_{\mathbf{x} \in \mathcal{W}^*} \alpha \cdot \mathbf{x} \right\}. \end{aligned}$$

Recall that PEDI uses the projected subgradient method, which approximates the max operator above by a sequence of variables at the expense of a sublinear approximation error, which leads to

$$\begin{aligned} \text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi}) &\leq K^{-1} \sum_{k=1}^K \left\{ (\beta^k)^\top \mathbf{V}_1^k(\underline{s}) - g^*(\beta^k) \right\} - g(\mathbf{V}_1^*(\underline{s})) \\ &\quad + K^{-1} \nu \sum_{k=1}^K \left\{ (\alpha^k)^\top \mathbf{V}_1^k(\underline{s}) - \max_{\mathbf{x} \in \mathcal{W}^*} (\alpha^k)^\top \mathbf{x} \right\} \\ &\quad + C(1 + \nu) \sqrt{DH^2/K}, \end{aligned}$$

where C is a constant. We rearrange the right-hand side above with some tricks and then obtain

$$\text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi}) \leq K^{-1} \sum_{k=1}^K \left[(\theta^k)^\top (\mathbf{V}_1^k(\underline{s}) - \mathbf{V}_1^*(\underline{s})) \right] + C(1 + \nu) \sqrt{DH^2/K}.$$

The last term in the right hand side above is exactly ϵ_K . Intuitively, the term $(\theta^k)^\top (\mathbf{V}_1^k(\underline{s}) - \mathbf{V}_1^*(\underline{s}))$ indicates the projected difference of state-value functions along θ , which motivates the PESSPLANNING (Algorithm 1). We apply results from this pessimistic approach and then conclude

$$\text{SubOpt}(\hat{\pi}) + \nu \text{Violation}(\hat{\pi}) = g(\widehat{\mathbf{V}}_1(\underline{s})) - g(\mathbf{V}_1^*(\underline{s})) + \nu \text{dist}(\widehat{\mathbf{V}}_1(\underline{s}), \mathcal{W}^*) \leq \epsilon_K + \text{IntUncert}_{\mathcal{D}}^{\pi^*}.$$

Second step: upper bounding $\text{SubOpt}(\hat{\pi})$ and $\text{Violation}(\hat{\pi})$ individually. This is done by two facts: (i) $\text{dist}(\cdot, \cdot) \geq 0$, and (ii) $g(\widehat{\mathbf{V}}_1(\underline{s})) - g(\mathbf{V}_1^*(\underline{s})) \geq -\nu \text{dist}(\widehat{\mathbf{V}}_1(\underline{s}), \mathcal{W}^*)/2$. We plug them into the resulting inequality in the first step and then complete the proof. $\square$

As shown in Theorem 1, our pessimistic algorithm bounds the suboptimality and constraint violation from above by the sum of the optimization error $\epsilon_K \in \mathcal{O}(1/\sqrt{K})$ and the intrinsic uncertainty $\text{IntUncert}_{\mathcal{D}}^{\pi^*}$. The optimization error ϵ_K vanishes as K goes into infinity. Later, we will show that the intrinsic uncertainty, $\text{IntUncert}_{\mathcal{D}}^{\pi^*}$, is impossible to eliminate through an analysis of PEDI on linear kernel CMOMDPs in Section 4.2.

4.2 Results of linear kernel CMOMDPs

The following theorem characterizes the suboptimality and constraint violation of PEDI for linear kernel CMOMDPs. Recall that we instantiate PEDI to linear kernel CMOMDPs in Appendix E.

Theorem 2 (Suboptimality and constraint violation for linear kernel CMOMDPs). *We set $\lambda = 1$ and $\kappa = CR\sqrt{d \log(dN)} + \log(DH/\xi)$ where $C > 0$ is an absolute constant. Then, under Assumptions 1 and 2, $\{\Gamma_h^P, \Gamma_h^r\}_{h=1}^H$ defined in (24) (Appendix E) is a ξ -uncertainty quantifier. Hence, under the same settings of ν and η^k in Theorem 1, (10) holds with probability at least $1 - \xi$ for the output $\hat{\pi}$, and the intrinsic uncertainty is specified to*

$$\text{IntUncert}_{\mathcal{D}}^{\pi^*} = 2(1 + \rho) \sqrt{D} \sum_{h=1}^H \mathbb{E}_{\pi^*} \left[H \int_{\mathcal{S}} \kappa \cdot \|\psi(s_h, a_h, s')\|_{\Lambda_h^{-1}} ds' + \kappa \cdot \|\varphi(s_h, a_h)\|_{\Lambda_{\varphi, h}^{-1}} |s_1 = \underline{s} \right].$$

Proof of Theorem 2. See Appendix H.3 for a detailed proof. $\square$

In what follows, we show that when the dataset sufficiently covers the trajectories of the optimal policy, the intrinsic uncertainty is $\mathcal{O}(1/\sqrt{N})$ where N is the number of trajectories in the dataset.

Corollary 1 (Dataset with sufficient coverage). *Suppose that the τ -th trajectory $\{(s_h^\tau, a_h^\tau, c_h^\tau)\}_{h=1}^H$ of the dataset $\mathcal{D}$ is generated by a behavior policy $\pi^{b,\tau}$ for $\tau \in [N]$. Meanwhile, we assume that there exists a constant $\varsigma > 0$ such that $\mu_h^*(s, a)/\mu_h^{b,\tau}(s, a) \leq \varsigma$ for any $h \in [H], \tau \in [N]$, and $(s, a) \in \mathcal{S} \times \mathcal{A}$. Here $\{\mu_h^*\}_{h=1}^H$ and $\{\mu_h^{b,\tau}\}_{h=1}^H$ are the visitation measures of π^* and $\pi^{b,\tau}$. Then, we have with probability at least $1 - \xi$ that*

$$\text{IntUncert}_{\mathcal{D}}^{\pi^*} \in \tilde{\mathcal{O}}(C_{\varsigma\kappa}(1 + \rho)H^2\sqrt{d|\mathcal{S}|/N}).$$

Proof of Corollary 1. See Appendix H.5 for a detailed proof. $\square$

Here $|\mathcal{S}|$ denotes a certain measure on $\mathcal{S}$. For instance, when $\mathcal{S} = [0, 1]$ and we take the Lebesgue measure, we have $|\mathcal{S}| = 1$. Note that Corollary 1 holds under a weaker condition than the uniform coverage (Wang et al., 2020a; Rashidinejad et al., 2021), where the condition is $\max_{\pi} \mu^{\pi}/\mu^b \leq C$ for a certain constant C . However, our algorithm is still able to achieve the intrinsic uncertainty of $\tilde{\mathcal{O}}(\sqrt{1/N})$. In particular, our condition, $\mu^*/\mu^b \leq \varsigma$, can hold when we have expert demonstration in the dataset $\mathcal{D}$ (Buckman et al., 2020).

In Theorem 2, we notice that when the number of iteration K goes into infinity, the optimization error ϵ_K vanishes, and the suboptimality and constraint violation are bounded from above by the intrinsic uncertainty $\text{IntUncert}_{\mathcal{D}}^{\pi^*}$ only. However, this is the best effort that we can put. The following theorem indicates that the term $\text{IntUncert}_{\mathcal{D}}^{\pi^*}$ is impossible to eliminate since it arises from the information-theoretic lower bound.

Theorem 3 (Information-theoretic lower bound). *For the output $\hat{\pi}$ of any algorithm of offline CMOMDPs, there exists a linear kernel CMOMDP $\mathcal{M}$ with initial state $\underline{s} \in \mathcal{S}$ and a dataset $\mathcal{D}$ compliant with $\mathcal{M}$, such that*

$$\mathbb{E}_{\mathcal{D}} \left[\frac{\text{SubOpt}(\hat{\pi})}{\sum_{h=1}^H \mathbb{E}_{\pi^*} \left[\|\varphi(s_h, a_h)\|_{\Lambda_{\varphi,h}^{-1}} + |\mathcal{S}|^{-1} \int_{\mathcal{S}} \|\psi(s_h, a_h, s')\|_{\Lambda_h^{-1}} ds' \mid s_1 = \underline{s} \right]} \right] \geq c,$$

where the expectation is taken with respect to the data collecting process, and $c > 0$ is an absolute constant.

Proof of Theorem 3. See Appendix H.4 for a detailed proof. $\square$

By Theorem 2 and Theorem 3, we find that the upper bound of suboptimality of Algorithm 2 for linear kernel CMOMDP matches the information-theoretic lower bound up to the optimization error $\mathcal{O}(1/\sqrt{K})$ and some multiplicative factors of horizons, dimensions, and geometric properties of the target set, which means that PEDI is near-optimal for linear kernel CMOMDP.

5 Experimental Results

We also conduct numerical experiments to verify our theory. Please see Appendix I for details. The results show that the proposed algorithm is not only provably efficient but also applicable.

6 Conclusion

In this paper, we propose a general algorithm, Pessimistic Dual Iteration (PEDI), for constrained multi-objective reinforcement learning with offline data. We show our algorithm delivers a data-dependent upper bound for suboptimality and constraint violation. The upper bound is composed of the optimization error and the intrinsic uncertainty. Moreover, we show that the bound of suboptimality is minimax optimal for linear kernel CMOMDPs up to the optimization error and multiplicative factors of constants related to the CMOMDP. When the trajectories of optimal policy are sufficiently covered, the intrinsic uncertainty is proved to be $\tilde{\mathcal{O}}(1/\sqrt{N})$ where N is the number of trajectories in the dataset.

Acknowledgments and Disclosure of Funding

We thank all anonymous reviewers for their useful comments. Zhaoran Wang acknowledges National Science Foundation (Awards 2048075, 2008827, 2015568, 1934931), Simons Institute (Theory of Reinforcement Learning), Amazon, J.P. Morgan, and Two Sigma for their supports. Zhuoran Yang acknowledges Simons Institute (Theory of Reinforcement Learning).

References

- Abbasi-Yadkori, Y., Pál, D., and Szepesvári, C. (2011). Improved algorithms for linear stochastic bandits. In *NIPS*, volume 11, pages 2312–2320.
- Agarwal, R., Schuurmans, D., and Norouzi, M. (2020). An optimistic perspective on offline reinforcement learning. In *International Conference on Machine Learning*, pages 104–114. PMLR.
- Altman, E. (1999). *Constrained Markov decision processes*, volume 7. CRC Press.
- Antos, A., Szepesvári, C., and Munos, R. (2008). Learning near-optimal policies with bellman-residual minimization based fitted policy iteration and a single sample path. *Machine Learning*, 71(1):89–129.
- Barrett, L. and Narayanan, S. (2008). Learning all optimal policies with multiple criteria. In *Proceedings of the 25th international conference on Machine learning*, pages 41–47.
- Bäuerle, N. and Rieder, U. (2010). Markov decision processes. *Jahresbericht der Deutschen Mathematiker-Vereinigung*, 112(4):217–243.
- Bhatnagar, S. and Lakshmanan, K. (2012). An online actor-critic algorithm with function approximation for constrained markov decision processes. *Journal of Optimization Theory and Applications*, 153(3):688–708.
- Brantley, K., Dudik, M., Lykouris, T., Miryoosefi, S., Simchowitz, M., Slivkins, A., and Sun, W. (2020). Constrained episodic reinforcement learning in concave-convex and knapsack settings. *arXiv preprint arXiv:2006.05051*.
- Buckman, J., Gelada, C., and Bellemare, M. G. (2020). The importance of pessimism in fixed-dataset policy optimization. *arXiv preprint arXiv:2009.06799*.
- Cai, Q., Yang, Z., Jin, C., and Wang, Z. (2020). Provably efficient exploration in policy optimization. In *International Conference on Machine Learning*, pages 1283–1294. PMLR.
- Castelletti, A., Pianosi, F., and Restelli, M. (2011). Multi-objective fitted q-iteration: Pareto frontier approximation in one single run. In *2011 International Conference on Networking, Sensing and Control*, pages 260–265. IEEE.
- Castelletti, A., Pianosi, F., and Restelli, M. (2012). Tree-based fitted q-iteration for multi-objective markov decision problems. In *The 2012 international joint conference on neural networks (IJCNN)*, pages 1–8. IEEE.
- Chakraborty, B. and Murphy, S. A. (2014). Dynamic treatment regimes. *Annual review of statistics and its application*, 1:447–464.
- Chen, J. and Jiang, N. (2019). Information-theoretic considerations in batch reinforcement learning. In *International Conference on Machine Learning*, pages 1042–1051. PMLR.
- Cheung, W. C. (2019). Regret minimization for reinforcement learning with vectorial feedback and complex objectives. *Advances in Neural Information Processing Systems*, 32:726–736.
- Chow, Y., Ghavamzadeh, M., Janson, L., and Pavone, M. (2017). Risk-constrained reinforcement learning with percentile risk criteria. *The Journal of Machine Learning Research*, 18(1):6070–6120.
- Ding, D., Wei, X., Yang, Z., Wang, Z., and Jovanovic, M. (2021). Provably efficient safe exploration via primal-dual policy optimization. In *International Conference on Artificial Intelligence and Statistics*, pages 3304–3312. PMLR.

- Efroni, Y., Mannor, S., and Pirodda, M. (2020). Exploration-exploitation in constrained mdps. *arXiv preprint arXiv:2003.02189*.
- Farahmand, A.-m., Ghavamzadeh, M., Szepesvári, C., and Mannor, S. (2016). Regularized policy iteration with nonparametric function spaces. *The Journal of Machine Learning Research*, 17(1):4809–4874.
- Farahmand, A. M., Munos, R., and Szepesvári, C. (2010). Error propagation for approximate policy and value iteration. In *Advances in Neural Information Processing Systems*.
- Fu, Z., Yang, Z., and Wang, Z. (2020). Single-timescale actor-critic provably finds globally optimal policy. *arXiv preprint arXiv:2008.00483*.
- Fujimoto, S., Meger, D., and Precup, D. (2019). Off-policy deep reinforcement learning without exploration. In *International Conference on Machine Learning*, pages 2052–2062. PMLR.
- Gottesman, O., Johansson, F., Komorowski, M., Faisal, A., Sontag, D., Doshi-Velez, F., and Celi, L. A. (2019). Guidelines for reinforcement learning in healthcare. *Nature medicine*, 25(1):16–18.
- Jaques, N., Ghandeharioun, A., Shen, J. H., Ferguson, C., Lapedriza, A., Jones, N., Gu, S., and Picard, R. (2019). Way off-policy batch deep reinforcement learning of implicit human preferences in dialog. *arXiv preprint arXiv:1907.00456*.
- Jin, C., Yang, Z., Wang, Z., and Jordan, M. I. (2020a). Provably efficient reinforcement learning with linear function approximation. In *Conference on Learning Theory*, pages 2137–2143. PMLR.
- Jin, Y., Yang, Z., and Wang, Z. (2020b). Is pessimism provably efficient for offline rl? *arXiv preprint arXiv:2012.15085*.
- Kidambi, R., Rajeswaran, A., Netrapalli, P., and Joachims, T. (2020). Morel: Model-based offline reinforcement learning. *arXiv preprint arXiv:2005.05951*.
- Kumar, A., Fu, J., Tucker, G., and Levine, S. (2019). Stabilizing off-policy q-learning via bootstrapping error reduction. *arXiv preprint arXiv:1906.00949*.
- Kumar, A., Zhou, A., Tucker, G., and Levine, S. (2020). Conservative q-learning for offline reinforcement learning. *arXiv preprint arXiv:2006.04779*.
- Lange, S., Gabel, T., and Riedmiller, M. (2012). Batch reinforcement learning. In *Reinforcement learning*, pages 45–73. Springer.
- Laroche, R., Trichelair, P., and Des Combes, R. T. (2019). Safe policy improvement with baseline bootstrapping. In *International Conference on Machine Learning*, pages 3652–3661. PMLR.
- Le, H., Voloshin, C., and Yue, Y. (2019). Batch policy learning under constraints. In *International Conference on Machine Learning*, pages 3703–3712. PMLR.
- Levine, S., Kumar, A., Tucker, G., and Fu, J. (2020). Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*.
- Liu, Y., Swaminathan, A., Agarwal, A., and Brunskill, E. (2020). Provably good batch reinforcement learning without great exploration. *arXiv preprint arXiv:2007.08202*.
- Miryoosefi, S. and Jin, C. (2021). A simple reward-free approach to constrained reinforcement learning. *arXiv preprint arXiv:2107.05216*.
- Nair, A., Dalal, M., Gupta, A., and Levine, S. (2020). Accelerating online reinforcement learning with offline datasets. *arXiv preprint arXiv:2006.09359*.
- Natarajan, S. and Tadepalli, P. (2005). Dynamic preferences in multi-criteria reinforcement learning. In *Proceedings of the 22nd International Conference on Machine learning*, pages 601–608.
- Paternain, S., Calvo-Fullana, M., Chamon, L. F., and Ribeiro, A. (2019). Safe policies for reinforcement learning via primal-dual methods. *arXiv preprint arXiv:1911.09101*.

- Qiu, S., Wei, X., Yang, Z., Ye, J., and Wang, Z. (2020). Upper confidence primal-dual reinforcement learning for cmdp with adversarial loss. *arXiv e-prints*, pages arXiv–2003.
- Rashidinejad, P., Zhu, B., Ma, C., Jiao, J., and Russell, S. (2021). Bridging offline reinforcement learning and imitation learning: A tale of pessimism. *arXiv preprint arXiv:2103.12021*.
- Rockafellar, R. T. (1970). *Convex analysis*, volume 36. Princeton university press.
- Scherrer, B., Ghavamzadeh, M., Gabillon, V., Lesner, B., and Geist, M. (2015). Approximate modified policy iteration and its application to the game of tetris. *J. Mach. Learn. Res.*, 16:1629–1676.
- Shalev-Shwartz, S., Shammah, S., and Shashua, A. (2016). Safe, multi-agent, reinforcement learning for autonomous driving. *arXiv preprint arXiv:1610.03295*.
- Siegel, N. Y., Springenberg, J. T., Berkenkamp, F., Abdolmaleki, A., Neunert, M., Lampe, T., Hafner, R., Heess, N., and Riedmiller, M. (2020). Keep doing what worked: Behavioral modelling priors for offline reinforcement learning. *arXiv preprint arXiv:2002.08396*.
- Sun, P., Kretschmar, H., Dotiwalla, X., Chouard, A., Patnaik, V., Tsui, P., Guo, J., Zhou, Y., Chai, Y., Caine, B., et al. (2020). Scalability in perception for autonomous driving: Waymo open dataset. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 2446–2454.
- Tessler, C., Mankowitz, D. J., and Mannor, S. (2018). Reward constrained policy optimization. *arXiv preprint arXiv:1805.11074*.
- Tosatto, S., Pirotta, M., d’Eramo, C., and Restelli, M. (2017). Boosted fitted q-iteration. In *International Conference on Machine Learning*, pages 3434–3443. PMLR.
- Van Moffaert, K., Drugan, M. M., and Nowé, A. (2013). Scalarized multi-objective reinforcement learning: Novel design techniques. In *2013 IEEE Symposium on Adaptive Dynamic Programming and Reinforcement Learning (ADPRL)*, pages 191–199. IEEE.
- Wang, R., Foster, D. P., and Kakade, S. M. (2020a). What are the statistical limits of offline rl with linear function approximation? *arXiv preprint arXiv:2010.11895*.
- Wang, W. and Sebag, M. (2013). Hypervolume indicator and dominance reward based multi-objective monte-carlo tree search. *Machine learning*, 92(2-3):403–429.
- Wang, Z., Novikov, A., Żołna, K., Springenberg, J. T., Reed, S., Shahriari, B., Siegel, N., Merel, J., Gulcehre, C., Heess, N., et al. (2020b). Critic regularized regression. *arXiv preprint arXiv:2006.15134*.
- Wu, Y., Tucker, G., and Nachum, O. (2019). Behavior regularized offline reinforcement learning. *arXiv preprint arXiv:1911.11361*.
- Yang, L. and Wang, M. (2019). Sample-optimal parametric q-learning using linearly additive features. In *International Conference on Machine Learning*, pages 6995–7004. PMLR.
- Yin, M., Bai, Y., and Wang, Y.-X. (2021). Near-optimal provable uniform convergence in offline policy evaluation for reinforcement learning. In *International Conference on Artificial Intelligence and Statistics*, pages 1567–1575. PMLR.
- Yu, T., Thomas, G., Yu, L., Ermon, S., Zou, J., Levine, S., Finn, C., and Ma, T. (2020). Mopo: Model-based offline policy optimization. *arXiv preprint arXiv:2005.13239*.
- Yu, T., Tian, Y., Zhang, J., and Sra, S. (2021). Provably efficient algorithms for multi-objective competitive rl. *arXiv preprint arXiv:2102.03192*.
- Zinkevich, M. (2003). Online convex programming and generalized infinitesimal gradient ascent. In *Proceedings of the 20th international conference on machine learning (icml-03)*, pages 928–936.