

DPO: Differential reinforcement learning with application to optimal configuration search

Chandrajit Bajaj^{1,3} and Minh Nguyen^{2,3}

¹ Department of Computer Science
The University of Texas at Austin

² Department of Mathematics
The University of Texas at Austin

³ Oden Institute for Computational Engineering and Sciences
The University of Texas at Austin

Abstract. Reinforcement learning (RL) with continuous state and action spaces remains one of the most challenging problems within the field. Most current learning methods focus on integral identities such as value functions to derive an optimal strategy for the learning agent. In this paper, we instead study the dual form of the original RL formulation to propose the first differential RL framework that can handle settings with limited training samples and short-length episodes. Our approach introduces Differential Policy Optimization (DPO), a point-wise and stage-wise iteration method that optimizes policies encoded by local-movement operators. We prove a pointwise convergence estimate for DPO and provide a regret bound comparable with current theoretical works. Such pointwise estimate ensures that the learned policy matches the optimal path uniformly across different steps. We then apply DPO to a class of practical RL problems which search for optimal configurations with Lagrangian rewards. DPO is easy to implement, scalable, and shows competitive results on benchmarking experiments against several popular RL methods.

1 Introduction

Reinforcement learning (RL) is a powerful machine learning paradigm with applications in many areas of engineering and biological sciences [2, 10, 12]. In RL, an agent interacts with the outside environment to learn the best action maximizing a certain cumulative reward. Such unsupervised and feedback-based learning approach resembles human thinking, giving promising solutions to challenging tasks without proper labeled data. Nevertheless, RL is not widely applied and popularized as a large-scale industrial solution like other fields such as computer vision or natural language processing. In fact, current RL algorithms are more complex and sample-expensive than their deep learning counterparts. Thus, they offer lower-quality solutions under limited data settings and restricted theoretical sample-efficiency guarantees except in special cases.

In this paper, we aim to make a very first step towards practical RL algorithms, particularly in the case of continuous state and action spaces, for limited data scenarios. Mathematically speaking, an RL agent seeks an optimal policy, the function mapping from the agent’s current state to its optimal action. RL algorithms rely on the cumulative reward objective to assess a particular policy. In applications involving physical systems, the path of the learned actions and states ideally should resemble the optimal path across every time steps to capture the underlying physics. Unfortunately, under the extreme situations of limited training samples, this is met with several challenges. First, since the value function is a sum quantity, the learned policy can put more bias on a particular term and less focus on others to temporarily boost the sum value. Such policy might also search for the best state/action in terms of rewards without concerning about the optimal path as a whole, leading to a sub-optimal path that don’t conform to physics laws. Second, convergence guarantees [18] for RL algorithms mainly look at the value function’s improvement, making it difficult to assess the learned path’s quality across all individual time steps. These weaknesses motivate us to consider an alternative formulation with exclusive focus on the dynamics induced by the optimal policy and individual data points along the learned trajectories (see Section 2). As a direct solution to this new formulation, referred to as differential RL, we develop differential policy optimization (DPO), a simple yet effective policy optimization algorithm. We then prove a theoretical pointwise convergence estimate for DPO that allows policy’s assessment on the whole path. This estimate is also crucial to our derivation of the theoretical sample-efficiency guarantees in the next step.

To theoretically evaluate an RL algorithm’s sample efficiency, researchers often use a quantity called **regret**, which is the gap between the cumulative rewards of the trained policy and the optimal policy over all training samples and episodes. While several research works [4] have addressed the discrete settings with their optimal regret bounds, the case of continuous state and action spaces remains challenging. Under a generic assumption, RL regret bounds can be untractable. Thus, a mild assumption of Lipschitz MDP [15] was applied, and the corresponding regret has a lower bound of $\Omega(K^{\frac{d+1}{d+2}})$ [20], where K is the number of training episodes and d is the dimension of the combined state-action space. Within our framework, we also prove a regret bound for DPO of similar order $O(K^{\frac{2d+3}{2d+4}})$ (see Corollary 4). Nonetheless, these bounds come with some performance deterioration due to their dependence on state-action dimension. Several works impose stronger assumptions to obtain more desirable bounds. Maran et al. [13], for instance, considers cases with smooth transition dynamics and rewards up to order ν with a regret bound $O(K^{\frac{3d/2+\nu+1}{d+2(\nu+1)}})$, which becomes the desirable $O(K^{1/2})$ when $\nu \rightarrow \infty$. Vakili and Olkhovskaya [22] study transition and reward functions in reproducing kernel Hilbert space (RKHS) induced by the Matern covariance function with parameter m . This assumption yields the the regret bound $O(K^{\frac{d+m+1}{d+2m+1}})$ that again becomes $O(K^{1/2})$ when $m \rightarrow \infty$. Our work takes a different approach to achieve more helpful regret bounds. Instead of

considering more specialized problems, we propose a narrower hypothesis class for policy approximation. In particular, we assume that black box environment can provide a hypothesis class where functions differ from the true policy either a d -weakly convex or concave function and linearly bounded (see Definition 3 and Definition 4). DPO algorithm then can attain a regret bound of order $O(K^{5/6})$ (see Corollary 4) that is independent of the state-action dimension.

Contributions. Our main contributions can be summarized as:

1. We propose a novel learning framework with a practical policy optimization algorithm called differential policy optimization (DPO) (see Algorithm 1) to address RL problem in differential form, where the dynamics is evaluated directly without expected cumulative rewards.
2. We provide a rigorous proof for DPO algorithm’s pointwise convergence estimates together with its regret bounds analysis.
3. We demonstrate DPO’s effectiveness in limited data settings through application to complex physics-based and engineering problems with Lagrangian energy rewards.

Organization. In Section 2, we introduce a novel framework of differential reinforcement learning and its learning algorithm called DPO. Here, we also demonstrate the duality between our framework and original RL formulation. In Section 3, we outline lemmas to prepare for the convergence proof and regret analysis for the DPO algorithm in Section 4. In Section 5, we apply DPO algorithm to three RL tasks with Lagrangian energy cost, and perform benchmarkings against TRPO [18], PPO [19] and SAC [8].

2 Differential reinforcement learning

2.1 Problem statement

Differential reinforcement learning’s main goal is to solve the problem $\mathcal{D}$ which aims to find the optimal dynamics $G : \Omega \rightarrow \Omega$ underlying the optimal trajectory:

$$x_0 = X \sim \rho_0, x_1 = G(X), x_2 = G(x_1) = G^{(2)}(X), \dots, x_{H-1} = G^{(H-1)}(X) \quad (1)$$

Here Ω is a compact domain in $\mathbb{R}^d$, and H is the number of steps in an episode. $G^{(k)}(X)$ is the k -times composition function $G(G(\dots G(x) \dots))$. Moreover, ρ_0 is the distribution of the starting point x_0 . In this work, we assume ρ_0 is a continuous function. By interacting with an environment $\mathcal{B}$ (see Definition 1), we learn a policy G_θ parameterized by θ that approximates G .

Definition 1. *An environment $\mathcal{B}$ with respect to an adversarial distribution ρ_0 and a score function g is a black-box system that allows you to check the quality of a policy G_θ . More concretely, $\mathcal{B}$ inputs a policy G_θ and outputs the trajectories $(G_\theta^{(k)}(x))_{k=0}^{H-1}$, and their associated scores $(g(G_\theta^{(k)}(x)))_{k=0}^{H-1}$. Recall that H is the number of steps in an episode.*

The function g allows one to assess the performance of the current policy G_θ and then perform necessary updates to get to the the desired dynamics G . In this paper, we assume that there is a first-order relation between the score function g and the dynamics function G :

$$G = A \cdot dg, \text{ where } A \text{ is a known constant matrix} \quad (2)$$

This relation, which generalizes the assumption under a Hamiltonian system, will be explained in Section 2.4.

2.2 Differential policy optimization (DPO) algorithm

Differential policy optimization (DPO) (see Algorithm 1) is an “time-expansion (Dijkstra-like)” algorithm that iteratively uses appropriate on-policy samples for policy update to ensure an increase in policy quality over each iteration. This algorithm has a similar idea as Trust Region Proximal Policy Optimization. However, thanks to the differential approach that focuses on pointwise estimates, DPO becomes much simpler and easier to implement/optimize in practice compared to its RL counterpart.

2.3 Original RL formulation

We first recall the usual RL mathematical formulation, which has a general idea similar to Section 2.1 but with different details. A finite-horizon discounted Markov decision process (MDP) is defined by the 6-tuple $(\mathcal{S}, \mathcal{A}, \mathbb{P}, r, \rho_0, \gamma)$, where $\mathcal{S}$ and $\mathcal{A}$ be the set of states and of actions, $\mathbb{P} : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}$ is the transition probability distribution, $r : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ is the reward function, $\rho_0 : \mathcal{S} \rightarrow \mathbb{R}$ is the distribution of initial state and $\gamma \in (0, 1)$ is the discount factor. At time t , the agent choose an action $a_t \in \mathcal{A}$ given its current state $s_t \in \mathcal{S}$ based on $\pi(a_t|s_t)$:

$$s_0 \sim \rho_0(s_0), a_t \sim \pi(a_t|s_t), s_{t+1} \sim \mathbb{P}(s_{t+1}|s_t, a_t) \quad (3)$$

The goal is to maximize the expected γ -discounted cumulative reward $\mathcal{J} = \mathbb{E}_\pi \left[\sum_{t=0}^{H-1} \gamma^t r(s_t, a_t) \right]$. For a particular policy π , let V_π denote the following value function obtained by executing π :

$$V_\pi(s_t) = \mathbb{E}_{a, s_1, \dots} \left[\sum_{t=0}^{H-1} \gamma^t r(s_t, a_t) | s_0 = s \right] \quad (4)$$

The value function is then defined as $V(s) := \arg \max_\pi V_\pi(s)$ and is used across several RL algorithms.

Algorithm 1 (Main algorithm) DPO for a generic environment $\mathcal{B}$

Input: a generic environment $\mathcal{B}$, the number of steps per episode H , and the number of samples N_k at stage k with $k \in \overline{1, H-1}$. Here N_k can be chosen based on Theorem 1. We also assume that the hypothesis space for the policy approximator G_{θ_k} in stage k is $\mathcal{H}_k$ for $k \in \overline{1, H}$.

Output: a neural network approximation G_θ that approximates the optimal policy G

- 1: Initialize an empty replay memory queue $\mathcal{M}$.
 - 2: Initialize $k = 1$ as the current stage and a random scoring function g_{θ_0} . Set the initial policy $G_{\theta_0} = A \cdot dg_{\theta_0}$ via automatic differentiation.
 - 3: **repeat**
 - 4: Use N_k starting points $\{X^i\}_{i=1}^{N_k}$ and previous policy $G_{\theta_{k-1}}$ to query $\mathcal{B}$ to get N_k sample trajectories $\left\{G_{\theta_{k-1}}^{(j)}(X^i)\right\}_{j=0}^{H-1}$ together with their scores $\left\{g(G_{\theta_{k-1}}^{(j)}(X^i))\right\}_{j=0}^{H-1}$ for $i \in \overline{1, N_k}$
 - 5: Add the labeled samples $(x, y) = (G_{\theta_{k-1}}^{(k-1)}(X^i), g(G_{\theta_{k-1}}^{(k-1)}(X^i)))$ to $\mathcal{M}$. Also add labeled samples $(x, y) = (G_{\theta_{k-1}}^{(j)}(X^i), g_{\theta_{k-1}}(G_{\theta_{k-1}}^{(j)}(X^i)))$ for $j \in \overline{1, k-2}$ and $i \in \overline{1, N_k}$ to $\mathcal{M}$. The latter addition step is to ensure that the new policy doesn't deviate from the previous policy on samples on which the previous policy already performs well.
 - 6: Train the neural network $g_{\theta_k} \in \mathcal{H}_k$ at stage k using labeled sample from $\mathcal{M}$ with smooth L^1 loss function [7].
 - 7: Set $G_{\theta_k} = A \cdot dg_{\theta_k}$ via automatic differentiation. Update $k \rightarrow k + 1$.
 - 8: **until** $k \geq H$
 - 9: Output $G_{\theta_{H-1}} := A \cdot dg_{\theta_{H-1}}$ via automatic differentiation.
-

2.4 Duality between Standard RL and Differential RL

To illustrate the duality between two approaches, we first consider a more specialized problem with Lagrangian reward (kinetic minus potential):

$$\pi^* = \arg \max_{\pi} \mathbb{E} \left[\sum_{k=0}^{H-1} \gamma^{-k} \left(\frac{a_{t_k}^2}{2} - \mathcal{F}(s_{t_k}) \right) \right]$$

subject to the linear dynamics : $s_{k+1} = s_k + \Delta a_k$, for step size Δ . (5)

Here $\mathcal{F}$ is a potential energy function unknown to the learning agent on the state space $\mathcal{S}$. The action $a_{t_k} = \pi(s_{t_k})$ be the output action from a policy π at discrete time $t_k = k\Delta$. $T = t_H = H\Delta$ is the terminal time, and γ is the magnification factor on the energy. To see the duality, we first look at the continuous-version of Equation (5), which is the optimal control problem of finding the path that

optimizes:

$$J = \int_0^T \beta^{-t} \left(\frac{a^2}{2} - \mathcal{F}(q) \right) dt$$

subject to $\dot{s} = a$

(6)

where $\gamma^{-k} = \beta^{-t_k} = \beta^{-\Delta k}$ so that we have the relation $\gamma = \beta^\Delta$. With this continuous-time format, the duality can be expressed via the following Lemma 1

Lemma 1. *The optimal trajectory/dynamics of Equation (6) in the generalized coordinates has the form:*

$$\begin{aligned} \dot{s} &= \beta^t p \\ \dot{p} &= -\beta^{-t} \nabla \mathcal{F}(s) \end{aligned}$$
(7)

Proof. First of all, the associated Hamiltonian $HF(p, s, a)$ for Equation (6) is $pa - \beta^{-t} \left(\frac{a^2}{2} - \mathcal{F}(q) \right)$. By Pontryagin maximum principle [11], $HF_a = p - \beta^{-t}a = 0$ so that the optimal action $a^* = \beta^t p$. Hence, the reduced Hamiltonian $hf(s, p) := HF(s, p, a^*(s, p))$ has the form:

$$hf(s, p) = \beta^t \|p\|^2 - \beta^t \frac{\|p\|^2}{2} + \beta^{-t} \mathcal{F}(s) = \beta^t \frac{p^2}{2} + \beta^{-t} \mathcal{F}(s)$$
(8)

As a result, other conditions from Pontryagin maximum principle imply that:

$$\begin{aligned} \dot{s} &= hf_p = \beta^t p \\ \dot{p} &= -hf_s = -\beta^{-t} \nabla \mathcal{F}(s) \end{aligned}$$

□

The above dynamics Equation (7) is called **optimal control flow** (OCF). Agents learning this type of dynamics strive to arrive at a stable position with low potential energy $\mathcal{F}$. By considering $x = (s, p)$, our RL problem can be approximated by solving the abstract problem $\mathcal{D}$ in Section 2.1 with the score function $g(x) = g(s, p) = 0.5\|p\|^2 + \mathcal{F}(s)$ and a constant symplectic matrix A . Similar to this specific example, duality between the traditional RL and our differential formulation can be obtained via the reduced Hamiltonian function and is explicitly expressed through the Legendre transform [11]. This type of duality is a generalized version of the duality between Lagrangian mechanics and Hamiltonian mechanics in classical physics. The latter, which is similar to our formulation, imposes a physical condition regarding the Hamiltonian's derivative that leads to a more regularized path:

$$\dot{s} = hf_p \text{ and } \dot{p} = -hf_s$$
(9)

In the differential form Section 2.1, we couple state s and momentum p , which contains information from both state and action, into a single vector $x = (s, p)$ with dimension $d = d_S + d_A$ (sum of state and action spaces' dimensions). The above Hamiltonian condition is a special case of Equation (2).

Remark 1. We also note an interesting property of the optimal solution for Equation (5). From Lemma 1, for $a = \beta^t p$, $\dot{a} = \beta^t \dot{p} + (\log \beta) \beta^t p = -\nabla_s \mathcal{F}(s) + \log \beta a$. Hence, Equation (7) is equivalent to:

$$\begin{aligned}\dot{s} &= a \\ \dot{a} &= (-\log \beta)a + \nabla_s \mathcal{F}_s\end{aligned}$$

or

$$\begin{bmatrix} \dot{s} \\ \dot{a} \end{bmatrix} + \begin{bmatrix} 0 \\ (-\log \beta)a \end{bmatrix} - \begin{bmatrix} 0 & I \\ -I & 0 \end{bmatrix} \begin{bmatrix} \nabla_s h f \\ \nabla_a h f \end{bmatrix} = 0 \quad (10)$$

This is similar to the continuous-version of the Nesterov accelerated method, which has the form of a damped Hamiltonian flow:

$$\begin{bmatrix} \dot{x}_t \\ \dot{p}_t \end{bmatrix} + \begin{bmatrix} 0 \\ \alpha_t p_t \end{bmatrix} - \begin{bmatrix} 0 & I \\ -I & 0 \end{bmatrix} \begin{bmatrix} \nabla_x HF(x_t, p_t) \\ \nabla_p HF(x_t, p_t) \end{bmatrix} = 0 \quad (11)$$

where HF is the associated Hamiltonian flow.

3 Basic notations and supporting lemmas

In this section, we state all supporting lemmas to prove the main estimate for Algorithm 1 (DPO) in Section 4.

Definition 2 defines the number of training samples needed to allow derivative approximation transfer. This definition is used to derive the number of samples needed for Algorithm 1 (DPO).

Definition 2. Recall that p_0 is the a continuous density for starting states. Suppose we are given a function $g : \Omega \rightarrow \Omega$, a hypothesis space $\mathcal{H}$ consists of function $h \in \mathcal{H}$ that approximates g , two positive constants ϵ and δ . We define the function $N(g, \mathcal{H}, \epsilon, \delta)$ to be the number of samples needed so that if we approximate g by $h \in \mathcal{H}$ via $N(g, \mathcal{H}, \epsilon, \delta)$ training samples, then with probability of at least $1 - \delta$, we have the following estimate bound on two function gradients:

$$\|\nabla g(X) - \nabla h(X)\| < \epsilon \quad (12)$$

In other words, we want $N(g, \mathcal{H}, \epsilon, \delta)$ to be large enough so that the original approximation can transfer to the derivative approximation above. If no such $N(g, \mathcal{H}, \epsilon, \delta)$ exists, let $N(g, \mathcal{H}, \epsilon, \delta) = \infty$

Lemma 2 and Lemma 3 below estimates $N(g, \mathcal{H}, \epsilon, \delta)$ (defined in Definition 2) in terms of the required threshold error ϵ .

Lemma 2. Suppose that the hypothesis space $\mathcal{H}$ for approximating the target function $g : \Omega \subset \mathbb{R}^d \rightarrow \mathbb{R}^d$ consists of neural network approximators with bounded weights and biases. Further assume that function g and neural network approximators $h \in \mathcal{H}$ are twice continuously differentiable with bounded first

and second derivative by some constant C . One way this assumption can be satisfied is to choose activation functions that are twice continuously differentiable. Then $N(g, \mathcal{H}, \epsilon, \delta)$ is upper-bounded by $O(\epsilon^{-(2d+4)})$, where we only ignore quadratic factors of δ , polynomial terms of d , and other logarithmic terms.

Definition 3. For $p \in \mathbb{N}$, a function h is called a **p -weakly convex** function if for any $x \in \Omega$, there exists a sufficiently small neighborhood U of x so that:

$$h(y) \geq h(x) + \nabla h(x)(y - x) - C\|y - x\|^p \quad (13)$$

Note that any convex function is p -weakly convex for any $p \in \mathbb{N}$.

Definition 4. A function h is **linearly bounded**, if for some constant $C > 0$, and for any $x \in \Omega$, there exists a sufficiently small neighborhood U of x so that:

$$|h(y) - h(x)| \leq C\|\nabla h(x)\|\|y - x\| \quad (14)$$

Any Lipschitz function on compact domain that has non-zero derivative is linearly bounded. One such example is the function $e^{\alpha\|x\|^2}$ on domains that don't contain 0.

Lemma 3. Suppose that, possibly with knowledge from outside environment or from certain policy experts, the hypothesis space $\mathcal{H}$ is reduced to a smaller family of functions of the form $g + h$, where h is a linearly bounded and p -weakly convex function for some $p \geq d$. Then $N(f, \mathcal{H}, \epsilon, \delta)$ can be upper-bounded by the factor $O(\epsilon^{-6})$ that is independent of the dimension d .

Remark 2. There are many hypothesis spaces consisting of infinitely many elements that satisfy Lemma 3. One such hypothesis space is:

$$\mathcal{H} = \left\{ h(x) := \sum_{i=1}^D a_i e^{b_i \|x\|^2}, \quad 0 \leq a_i \leq A, \quad 0 < b \leq b_i \leq B \right\} \quad (15)$$

for constant $A, B, b > 0$ and for domain Ω doesn't contain 0. Lemma 3 also holds for concave functions and their weak versions.

Finally, we state the supporting Lemma 4 for the Main Theorem 1 which proves the pointwise estimates for DPO algorithm.

Lemma 4. Given L and $\epsilon > 0$, define two sequences $\{\alpha_j\}_{j \geq 0}$ and $\{\epsilon_j\}_{j \geq 0}$ recursively as follows:

$$\alpha_0 = 0 \text{ and } \alpha_j = L\alpha_{j-1} + \epsilon \quad (16)$$

$$\epsilon_1 = \epsilon \text{ and } \epsilon_{k+1} = L\alpha_k + \epsilon + L\epsilon_k \quad (17)$$

Then for each k , we get:

$$\epsilon_k \leq \frac{kL^k \epsilon}{L - 1} \quad (18)$$

4 Theoretical analysis for differential policy optimization

This section shows the convergence of differential policy optimization (DPO, Algorithm 1) based on generalization pointwise estimates. We then use this result to derive suitable regret bounds for DPO.

4.1 Main theorem

Theorem 1. *Suppose we're given a threshold error ϵ , a probability threshold δ , and number of steps per episode H . Assume that $\{N_k\}_{k=1}^{H-1}$ is the sequence of numbers of samples used at each stage Algorithm 1 (DPO) so that:*

$$N_1 = N(g, \mathcal{H}_1, \epsilon, \delta), \text{ and for } k \in \overline{2, H-1} : \quad (19)$$

$$N_k = \max \left\{ N(g_{\theta_{k-1}}, \mathcal{H}_k, \epsilon, \delta_{k-1}/(k-1)), N(g, \mathcal{H}_k, \epsilon, \delta_{k-1}/(k-1)) \right\} \quad (20)$$

We further assume that there exists a Lipschitz constant $L > 0$ such that both the true dynamics G and the policy neural network approximator G_{θ_k} at step k with regularized parameters have their Lipschitz constant at most L for each $k \in \overline{1, H}$. Then, for a general starting point X , with probability at least $1 - \delta$, the following generalization bound for the trained policy G_{θ_k} holds for all $k \in \{1, 2, \dots, H-1\}$:

$$\mathbb{E}_X \|G_{\theta_k}^{(j)}(X) - G^{(j)}(X)\| < \frac{jL^j\epsilon}{L-1} \text{ for all } 1 \leq j \leq k \quad (21)$$

Note that when $N_k \rightarrow \infty$, the errors approach 0 uniformly for all j given a finite terminal time T .

Proof. Define $\delta_k = \delta/3^{m-k} = 3\delta_{k-1}$. Let α_k and ϵ_k be two sequences associated with Lipschitz constant L and threshold error ϵ as in Lemma 4. We prove the generalization bound statement by induction on the stage number k that for probability of at least $1 - \delta_k$,

$$\mathbb{E}_X \|G_{\theta_k}^{(j)}(X) - G^{(j)}(X)\| < \epsilon_j \text{ for all } 1 \leq j \leq k \quad (22)$$

By Lemma 4, proving this statement also proves Theorem 1.

The bound for the base case $k = 1$ is due to the definition of $N_1 = N(g, \mathcal{H}_1, \epsilon, \delta)$ that allows the approximation of g by g_{θ_1} transfers to their derivatives G and G_{θ_1} with error threshold ϵ and probability threshold δ . Assume that the induction hypothesis is true for k . We prove that for a starting (random variable) point X , the following error holds with a probability of at least $1 - \delta_{k+1} = 1 - 3\delta_k$:

$$\mathbb{E}_X \|G_{\theta_{k+1}}^{(j)}(X) - G^{(j)}(X)\| < \epsilon_j \text{ for all } j \leq k+1 \quad (23)$$

First, from induction hypothesis, with probability of at least $1 - \delta_k$:

$$\mathbb{E}_X \|G_{\theta_k}^{(j)}(X) - G^{(j)}(X)\| < \epsilon_j \text{ for all } j \leq k \quad (24)$$

In stage $k + 1$, all previous stages' samples up to stage $k - 1$ is used for $G_{\theta_{k+1}}$. As a result, we can invoke the induction hypothesis on k to yield the same error estimate for $G_{\theta_{k+1}}$ on the first k sample points with probability $1 - \delta_k$:

$$\mathbb{E}_X \|G_{\theta_{k+1}}^{(j)}(X) - G^{(j)}(X)\| < \epsilon_j \text{ for all } j \leq k \quad (25)$$

Recall from Algorithm 1 that $g_{\theta_{k+1}}$ is trained to approximate g_{θ_k} to ensure that the updated policy doesn't deviate too much from current policy. For $j \in \{1, \dots, k-1\}$, $g_{\theta_{k+1}} \in \mathcal{H}_{k+1}$ approximates g_{θ_k} on N_{k+1} samples of the form $G_{\theta_k}^{(j)}(X^i)$ for $i \in \{1, \dots, N_{k+1}\}$. Since $N_{k+1} \geq N(g_{\theta_k}, \mathcal{H}_{k+1}, \epsilon, \delta_k/k)$ allows derivative approximation transfer, for probability of at least $1 - \delta_k/k$, $\mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_k}^{(j)}(X)) - G_{\theta_k}(G_{\theta_k}^{(j)}(X))\| < \epsilon$. Hence, under a probability subspace Γ with probability of at least $(1 - (k-1)\delta_k/k)$, we have:

$$\mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_k}^{(j)}(X)) - G_{\theta_k}(G_{\theta_k}^{(j)}(X))\| < \epsilon \text{ for all } 1 \leq j < k \quad (26)$$

We prove by induction on j that under this probability subspace Γ , we have:

$$\mathbb{E}_X \|G_{\theta_{k+1}}^{(j)}(X) - G_{\theta_k}^{(j)}(X)\| < \alpha_j \text{ for all } 1 \leq j \leq k \quad (27)$$

In fact, for the induction step, one get:

$$\begin{aligned} \mathbb{E}_X \|G_{\theta_{k+1}}^{(j)}(X) - G_{\theta_k}^{(j)}(X)\| &\leq \mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_{k+1}}^{(j-1)}(X)) - G_{\theta_{k+1}}(G_{\theta_k}^{(j-1)}(X))\| \\ &\quad + \mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_k}^{(j-1)}(X)) - G_{\theta_k}(G_{\theta_k}^{(j-1)}(X))\| \\ &\leq L\mathbb{E}_X \|G_{\theta_{k+1}}^{(j-1)}(X) - G_{\theta_k}^{(j-1)}(X)\| + \epsilon \leq L\alpha_{j-1} + \epsilon = \alpha_j \end{aligned}$$

Finally, we look at the approximation of g by $g_{\theta_{k+1}} \in \mathcal{H}_{k+1}$ on the specific sample points $\{G_{\theta_k}^{(k)}(X^i)\}_{i=1}^{N_{k+1}}$. Definition of $N_{k+1} \geq N(g, \mathcal{H}_{k+1}, \epsilon, \delta_k/k)$ again allow derivative approximation transfer so that with probability at least $1 - \delta_k/k$:

$$\mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_k}^{(k)}(X)) - G(G_{\theta_k}^{(k)}(X))\| < \epsilon \quad (28)$$

Now consider the probability subspace $\mathcal{S}$ under which 3 inequalities Equation (24), Equation (27), and Equation (28) hold. The subspace $\mathcal{S}$ has the probability measure of at least $1 - (\delta_k + (k-1)\delta_k/k + \delta_k/k) = 1 - 2\delta_k = 1 - (\delta_{k+1} - \delta_k)$, and under $\mathcal{S}$, we have:

$$\begin{aligned} \mathbb{E}_X \|G_{\theta_{k+1}}^{(k+1)}(X) - G^{(k+1)}(X)\| &\leq \mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_{k+1}}^{(k)}(X)) - G_{\theta_{k+1}}(G_{\theta_k}^{(k)}(X))\| \\ &\quad + \mathbb{E}_X \|G_{\theta_{k+1}}(G_{\theta_k}^{(k)}(X)) - G(G_{\theta_k}^{(k)}(X))\| + \mathbb{E}_X \|G(G_{\theta_k}^{(k)}(X)) - G(G^{(k)}(X))\| \\ &\leq L\mathbb{E}_X \|G_{\theta_{k+1}}^{(k)}(X) - G_{\theta_k}^{(k)}(X)\| + \epsilon + L\mathbb{E}_X \|G_{\theta_k}^{(k)}(X) - G^{(k)}(X)\| \\ &\leq L\alpha_k + \epsilon + L\epsilon_k = \epsilon_{k+1} \end{aligned}$$

Merging this inequality with probability subspace where the inequality in Equation (25) holds leads to the estimate on the final step for stage $k + 1$ for the induction step. $\square$

Together with the general pointwise estimates for DPO algorithm in Theorem 1, Lemma 2 and Lemma 3 allow us to explicitly state the number of training episodes required for two scenarios considered in this work: one work with general neural network approximators and the other with more restricted (weakly convex and linear bounded) difference functions:

Corollary 1. *In Algorithm 1 (DPO), suppose we are given fixed step size and fixed number of steps per episode H . Further assume that for all $k \in \overline{1, H-1}$, $\mathcal{H}_k$ is the same everywhere and is the hypothesis space $\mathcal{H}$ consisting of neural network approximators with bounded weights and biases. Then with the sequence of numbers of training episodes $N_k = O(\epsilon^{-(2d+4)})$, the pointwise estimates Equation (21) hold.*

Corollary 2. *Again, in Algorithm 1 (DPO), suppose we are given fixed step size and fixed number of steps per episode H . Suppose $\mathcal{H}_k$ is a hypothesis subspace consisting of $h \in \mathcal{H}_k$ so that $h - g$ and $h - g_{\theta_{k-1}}$ are both p -weakly convex and linearly bounded. Then with the sequence of numbers of training episodes $N_k = O(\epsilon^{-6})$, we obtain the pointwise estimates Equation (21).*

Corollary 3. *In order to justify the Lipschitz constant L assumption, we note that the dynamics function G and its approximator often have the discretized dynamics form: $G(x) = x + \Delta F(X)$, where Δ is the step size, and F is a bounded function. In this case, even when the step size is infinitely small, Algorithm 1 (DPO)'s training converges with reasonable numbers of training episodes.*

Proof. The Lipschitz constant L of G in this case is bounded by $1 + C\Delta$ for some constant $C > 0$. By scaling, WLOG, assume that for finite-time horizon problem, the terminal time $T = 1$ so that the number of steps is $m = 1/\Delta$. Hence, for $N_k = O(1/\Delta^{2p})$, $\epsilon = O(\Delta^p)$ for $p > 2$, the error bounds in Theorem 1 are upper-bounded by:

$$\|G_{\theta_k}^{(k)}(X) - G^{(k)}(X)\| \leq \frac{kL^k\epsilon}{L-1} \leq \frac{1}{\Delta} \left(1 + \frac{C}{N_k}\right)^{N_k} O(\Delta^p) \frac{1}{\Delta} \leq e^C O(\Delta^{p-2}) \rightarrow 0 \quad (29)$$

for $k \leq H-1$, as $\Delta \rightarrow 0$. $\square$

4.2 Regret bound analysis

In this section, we define explicitly the regret mentioned in Section 1 and then derive two regret bounds for DPO algorithm (Algorithm 1). Suppose K episodes are used during the training process and suppose a policy π^k is applied at the beginning of the k -th episode with the starting state s^k (sampled from adversary distribution ρ_0) for $k \in \overline{1, K}$. If we focus on the total number of training episodes used, and assuming that number of steps H is fixed. Then, the quantity **Regret** is defined as the following function of the number of episodes K :

$$\text{Regret}(K) = \sum_{k=1}^K (V(s^k) - V_{\pi^k}(s^k)) \quad (30)$$

We now derive an upper bound on $\mathbf{Regret}(K)$ via Equation (30):

Corollary 4. *Suppose that number of steps per episode H is fixed and relatively small. If in Algorithm 1 (DPO), the number of training samples N_k have the scale of $O(\epsilon^{-\mu})$, regret bound for is upper-bounded by $O(K^{(\mu-1)/\mu})$. As a result, for the general case in Corollary 1, we obtain a regret bound of $O(K^{(2d+3)/(2d+4)})$. For the special cases with restricted hypothesis space in Corollary 2 the regret bound is approximately $O(K^{5/6})$.*

Proof. For a fixed H , Equation (21) gives us the estimate $\mathbb{E}_X[G_{\theta_k}^{(j)}(X) - G(X)] \approx O(\epsilon)$ for the state-action pair X at step j between $(N_1 + \dots + N_{k-1} + 1)^{th}$ and $(N_1 + \dots + N_k)^{th}$ episodes during stage k . Assuming a mild Lipschitz condition on reward function, the gap between the optimal reward and the reward obtained from the learned policy at step j during these episodes is also approximately $O(\epsilon)$. Summing these uniform bounds over all j and over all episodes gives:

$$\mathbf{Regret}(K) \leq H(N_1 + \dots + N_{H-1})C\epsilon = CHK\epsilon \quad (31)$$

Since $N_k = O(\epsilon^{-\mu})$, $K = HO(\epsilon^{-\mu})$. With a fixed H , $\epsilon = O(K^{-1/\mu})$. Hence, Equation (31) leads to $\mathbf{Regret}(K) \leq KO(K^{-1/\mu}) = O(K^{(\mu-1)/\mu})$ as desired. $\square$

5 Experiments on Lagrangian rewards problems

We consider 3 specific problems under the form of Equation (5): materials deformation, topological materials deformation, and molecular dynamics.

5.1 Materials deformation

In materials engineering [3], the process of deforming a raw figure of a particular material such as wood, metal, or plastic into a desired shape is essential. As each deformation is costly in terms of energy usage, we demand a minimum number of deformations. In abstract terms, given a particular shape Γ_0 in the space of shapes $\mathcal{S}(\text{shape})$, one wants to deform Γ_0 as fast as possible to get to a desired shape Γ^* , where the desired property is characterized by minimizing a functional cost $\mathcal{C} : \mathcal{S}(\text{shape}) \rightarrow \mathbb{R}$. Assuming that the deformation/hamming machine can only incrementally deform the current object's outer part, we proceed to formulate this problem in terms of Equation (5) as follows: the state s for a 2D object Γ is the vector of discrete points on its boundary. The object $\Gamma = \Gamma(s)$ is then constructed from s by the 1D cubic spline approximation so that $\mathcal{F}(s) = \mathcal{C}(\Gamma(s))$. The action a is defined as the hamming machine's modification at the boundary's discrete points. For experiments, we choose the functional cost to be the dimensional-less quantity: $\mathcal{C}(\Gamma) = \frac{\mathbf{Peri}(\Gamma)}{\sqrt{\mathbf{Area}(\Gamma)}}$, where $\mathbf{Peri}$ and $\mathbf{Area}$ are perimeter and area of Γ . The distribution ρ_0 gives a random polygon at the start of each episode.

5.2 Topological materials deformation

Rather than restricting the deformation to the boundary only, one allows the hamming machine to modify any point inside and even change the topological structure of the object. For this new setting, the state s of Γ is now a discretization vector of Γ 's level-set function on a fixed grid. The original Γ can then be constructed via a bicubic spline approximation [1]. The action a is the modification of the level-set function at grid points. The starting distribution ρ_0 is an uniform distribution on grid points.

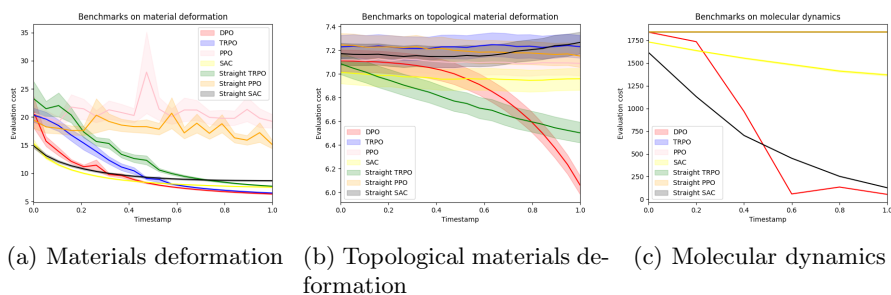


Fig. 1: Evaluation costs along episodes of 7 algorithms on 3 Lagrangian reward-type problems. DPO moves faster towards smaller evaluation costs and closer to the optimal configurations than other methods.

5.3 Molecular dynamics

Molecule's configuration is the 3D spatial arrangement between its atoms and is essential to the molecule's functioning. For a given molecule, we're interested in capturing its movement from a random configuration to a stable configuration, where the potential energy is minimized. With respect to the formulation Equation (5), the state s of each molecular configuration is a vector of its dihedral angles denoted by $(\phi_j, \psi_j)_j$, where j are the angle indices. Given these angles, we can recover the atom coordinates $(x_i)_{i=1}^N = X((\phi_j, \psi_j)_j)$ as a function of dihedral angles. Such coordinates allow us to calculate the functional $\mathcal{F}(s) = \mathcal{F}((\phi_j, \psi_j)_j) = U(X((\phi_j, \psi_j)_j)) = U((x_i)_{i=1}^N)$ where U is the energy function on the atom coordinates given by PyRosetta Python package [5]. The action a is the modification of the molecule's dihedral angles. The choice of the dihedral angles is due to the rigidity property of molecule. For this experiment, we consider the octa-alanine with the primary sequence consisting of eights amino acids A . The starting distribution ρ_0 is purposely chosen to be an uniform distribution over a very small compact intervals around zero to see how the agent performs with limited exploration from the start.

5.4 Benchmark results

Models: We run 4 models including Algorithm 1 (DPO), Trust Region Policy Optimization (TRPO) [18], Proximal policy optimization (PPO) [19], and Soft Actor-Critic (SAC) [8] on 3 problems with RL formulation in Equation (5) as outlined above. We additionally run 3 benchmark models TRPO, PPO, SAC on these 3 problems with the straightforward negative energy reward $r(s, a) = -\mathcal{F}(s)$. We call the 3 trained models on this type of straightforward rewards: Straight TRPO, Straight PPO, and Straight SAC, bringing in a total of 6 benchmark models. Implementations of benchmark models come from Stable-Baselines3 package [16].

Parameters: For each case, we use 200 test episodes with normalized time horizon $[0, 1]$ ($T = 1$). The number of episodes, the number of time steps, and the magnification factor γ are outlined in Table 2. Here γ is specific to each problem and is independent of the learning methods. All training and testing have been conducted on Nvidia A100 GPU.

Table 1: Final evaluation costs from 7 different algorithms on materials deformation, topological materials deformation and molecular dynamics.

	Materials deformation	Topological materials deformation	Molecular dynamics
DPO	6.323	6.061	53.340
TRPO	6.503	7.230	1842.299
PPO	19.229	7.089	1842.296
SAC	7.528	6.959	1369.605
Straight TRPO	7.709	6.502	1842.272
Straight PPO	15.117	7.151	1842.316
Straight SAC	8.686	7.267	126.449

Metrics: The main metrics are the evaluation costs, which are the values of $\mathcal{F}$ evaluated along the test episodes. The goal is to arrive at the minimum $\mathcal{F}$ -values in a fast and smooth manner.

Results: From Table 1 and Figure 1, DPO (Algorithm 1) performs the best across 3 problems. Figure 1 showing evaluation costs with uncertainty at each time step further demonstrates that DPO yields more desirable optimal paths.

Related model-based reinforcement learning algorithms. Several model-based RL algorithms [23] are not applicable here because:

1. Policy search with back-propagation through time requires more information than what the system $\mathcal{B}$ could provide. For instance, SVG [9] and iLQR [21]

Table 2: Problem-specific parameters

	Materials deformation	Topological materials deformation	Molecular dynamics
Number of episodes	5000	5000	800
Number of steps	20	20	6
Magnification factor γ	0.99	0.81	0.0067

require explicit forms of $\mathcal{F}$'s derivatives, while PILCO [6] needs the convolution of $\mathcal{F}$ with a Gaussian kernel. For simple problems such as **Acrobot** or **Ant-walk**, such information can be hand-calculated. But this is impossible for the 3 complex problems considered in this section. For the first two, $\mathcal{F}$ is a combination of a geometric functional together with cubic spline approximation operator. For the third problem, the energy functional contains molecular statistics unknown constants.

2. Shooting algorithms [14] require the agent to perform replanning at every time step. They need to execute the action sequence starting from an arbitrary time t , while Section 2.1 only allows $t = 0$.

6 Conclusion

In this paper, we develop a new approach called differential reinforcement learning focusing on the quality of the individual data points along the learned path. The duality between our problem and the original RL formulation is demonstrated as a generalization of the classical duality between Lagrangian and Hamiltonian mechanics. We propose a simple yet effective algorithm called differential policy optimization (DPO, see Algorithm 1) as a concrete solution to our approach. Theoretically, DPO not only possesses pointwise convergence guarantees but also achieves regret bounds comparable to current literature. Experimentally, DPO demonstrates promising performance on 3 physics-based problems with complex reward systems where the learning opportunities are limited under a small number of trajectories and number of time steps per episode. In the future, we aim to take the next few steps in promoting practical RL algorithms for continuous state and action spaces. These include: improving solution algorithms for differential RL both theoretically and empirically, and drawing more connections between our new approach to current state-of-the-art RL methods.

Author acknowledgement: This research was supported in part by a grant from the NIH DK129979, in part from the Peter O'Donnell Foundation, the Michael J. Fox Foundation, Jim Holland-Backcountry Foundation, and in part from a grant from the Army Research Office accomplished under Cooperative

Agreement Number W911NF-19-2-0333.

Author Contributions: This research work (Differential RL and DPO) is developed by Minh Nguyen and improved by Chandrajit Bajaj and Minh Nguyen. Alphabetical order is currently used to order author names. The implementation is done by Minh Nguyen and is available at <https://github.com/mpnguyen2/dpo>.

References

1. Bajaj, C.L., Xu, G.L., Zhang, Q.: Bio-molecule surfaces construction via a higher-order level-set method. *J. Comput. Sci. Technol.* **23**(6), 1026–1036 (2008)
2. Biagioni, D., Zhang, X., Wald, D., Vaidhynathan, D., Chintala, R., King, J., Zamzam, A.S.: Powergridworld: a framework for multi-agent reinforcement learning in power systems. In: *Proceedings of the Thirteenth ACM International Conference on Future Energy Systems*. p. 565–570. e-Energy '22, Association for Computing Machinery, New York, NY, USA (2022)
3. Brown, N.K., Garland, A.P., Fadel, G.M., Li, G.: Deep reinforcement learning for the rapid on-demand design of mechanical metamaterials with targeted nonlinear deformation responses. *Engineering Applications of Artificial Intelligence* **126**, 106998 (2023)
4. Cai, Q., Yang, Z., Jin, C., Wang, Z.: Provably efficient exploration in policy optimization. In: *Proceedings of the 37th International Conference on Machine Learning*. *Proceedings of Machine Learning Research*, vol. 119, pp. 1283–1294. PMLR (13–18 Jul 2020)
5. Chaudhury, S., Lyskov, S., Gray, J.J.: PyRosetta: a script-based interface for implementing molecular modeling algorithms using rosetta. *Bioinformatics* **26**(5), 689–691 (2010)
6. Deisenroth, M., Rasmussen, C.E.: Pilco: A model-based and data-efficient approach to policy search. In: *Proceedings of the 28th International Conference on machine learning (ICML-11)*. pp. 465–472. Citeseer (2011)
7. Girshick, R.: Fast r-cnn. In: *International Conference on Computer Vision (ICCV)* (2015)
8. Haarnoja, T., Zhou, A., Abbeel, P., Levine, S.: Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor (2018)
9. Heess, N., Wayne, G., Silver, D., Lillicrap, T., Tassa, Y., Erez, T.: Learning continuous control policies by stochastic value gradients. In: *Proceedings of the 28th International Conference on Neural Information Processing Systems - Volume 2*. p. 2944–2952. NIPS'15, MIT Press, Cambridge, MA, USA (2015)
10. Kalashnikov, D., Irpan, A., Pastor, P., Ibarz, J., Herzog, A., Jang, E., Quillen, D., Holly, E., Kalakrishnan, M., Vanhoucke, V., Levine, S.: Scalable deep reinforcement learning for vision-based robotic manipulation. In: *Proceedings of The 2nd Conference on Robot Learning*. *Proceedings of Machine Learning Research*, vol. 87, pp. 651–673. PMLR (29–31 Oct 2018)
11. Kirk, D.E.: *Optimal Control Theory: An Introduction*. Prentice-Hall, London, England (1971)
12. Lutz, I.D., Wang, S., Norn, C., Courbet, A., Borst, A.J., Zhao, Y.T., Dosey, A., Cao, L., Xu, J., Leaf, E.M., Treichel, C., Litvicov, P., Li, Z., Goodson, A.D., Rivera-Sánchez, P., Bratovianu, A.M., Baek, M., King, N.P., Ruohola-Baker, H., Baker,

- D.: Top-down design of protein architectures with reinforcement learning. *Science* **380**(6642), 266–273 (2023)
13. Maran, D., Metelli, A.M., Papini, M., Restell, M.: No-regret reinforcement learning in smooth mdps (2024)
 14. Nagabandi, A., Kahn, G., Fearing, R.S., Levine, S.: Neural network dynamics for model-based deep reinforcement learning with model-free fine-tuning. In: 2018 IEEE International Conference on Robotics and Automation (ICRA). IEEE (2018)
 15. Rachelson, E., Lagoudakis, M.G.: On the locality of action domination in sequential decision making. In: International Symposium on Artificial Intelligence and Mathematics (2010)
 16. Raffin, A., Hill, A., Gleave, A., Kanervisto, A., Ernestus, M., Dormann, N.: Stable-baselines3: Reliable reinforcement learning implementations. *Journal of Machine Learning Research* **22**(268), 1–8 (2021)
 17. Rebeschini, P.: Algorithmic foundations of learning (2022), <https://www.stats.ox.ac.uk/~rebeschini/teaching/AFoL/22/>
 18. Schulman, J., Levine, S., Abbeel, P., Jordan, M., Moritz, P.: Trust region policy optimization. In: Proceedings of the 32nd International Conference on Machine Learning. Proceedings of Machine Learning Research, vol. 37, pp. 1889–1897. PMLR, Lille, France (07–09 Jul 2015)
 19. Schulman, J., Wolski, F., Dhariwal, P., Radford, A., Klimov, O.: Proximal policy optimization algorithms (2017)
 20. Slivkins, A.: Introduction to multi-armed bandits (2022)
 21. Tassa, Y., Erez, T., Todorov, E.: Synthesis and stabilization of complex behaviors through online trajectory optimization. In: 2012 IEEE/RSJ International Conference on Intelligent Robots and Systems. IEEE (2012)
 22. Vakili, S., Olkhovskaya, J.: Kernelized reinforcement learning with order optimal regret bounds (2024)
 23. Wang, T., Bao, X., Clavera, I., Hoang, J., Wen, Y., Langlois, E., Zhang, S., Zhang, G., Abbeel, P., Ba, J.: Benchmarking model-based reinforcement learning (2019)

A Basic algorithmic learning theory under i.i.d setting

We first introduce basic learning theory result on independent identically distributed (i.i.d) samples. Proofs for all lemmas in this section can be found at [17] and can also be found in any modern learning theory literature.

Notations: Throughout this section, $\mathcal{X}$ be the feature space, $\mathcal{Y}$ be the label space, $\mathcal{Z} = \mathcal{X} \times \mathcal{Y}$. Let $\mathcal{H}$ be a hypothesis space consisting of hypothesis $h : \mathcal{X} \rightarrow \mathcal{Y}$. Let $l : \mathcal{H} \times \mathcal{Z} \rightarrow \mathbb{R}_+$ be a loss function on labeled sample $z = (x, y)$ for $x \in \mathcal{X}$ and $y \in \mathcal{Y}$. Our loss function will have the following particular form: $l(h, (x, y)) = \phi(h(x), y)$ for a given associated function $\phi : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}_+$. We use capital letters to represent random variables. We also use the following notations to define the sets: $\mathcal{H} \circ \{Z_1, \dots, Z_n\} := \{(h(Z_1), \dots, h(Z_n)) \in \mathbb{R}^n, h \in \mathcal{H}\}$ and $\mathcal{L} \circ \{Z_1, \dots, Z_n\} := \{l(h, Z_1), \dots, l(h, Z_n) \in \mathbb{R}^n, h \in \mathcal{H}\}$. Also define $e(h)$ to be the average loss over new test data $e(h) := \mathbb{E}_Z[l(h, Z)]$, and $E(h)$ be the empirical loss over $n \in \mathbb{Z}_+$ i.i.d samples $Z_1, \dots, Z_n$: $E(h) := \frac{1}{n} \sum_{i=1}^n l(h, Z_i)$.

Definition 5. The Rademacher complexity of a set $\mathcal{T}$ is defined as:

$$\mathbf{Rad}(\mathcal{T}) = \mathbb{E} \left[\sup_{t \in \mathcal{T}} \frac{1}{n} \sum_{i=1}^n B_i t_i \right] \quad (32)$$

for Bernoulli random variables $B_i \in \{-1, 1\}$.

Lemma 5. Suppose $\phi(\cdot, y)$ is γ -Lipschitz for any $y \in \mathcal{Y}$ for some $\gamma > 0$. Then:

$$\mathbb{E} \left[\sup_{h \in \mathcal{H}} \{e(h) - E(h)\} \right] \leq 2\mathbb{E} \left[\mathbf{Rad}(\mathcal{L} \circ \{Z_1, \dots, Z_n\}) \right] \quad (33)$$

$$\leq 2\gamma \mathbb{E} \left[\mathbf{Rad}(\mathcal{H} \circ \{Z_1, \dots, Z_n\}) \right] \quad (34)$$

Lemma 6. For a hypothesis space $\mathcal{H}$, a loss function l bounded in the interval $[0, c]$, and for n i.i.d labeled samples $Z_1, \dots, Z_n$, with probability of at least $1 - \delta$, the following bound holds:

$$\sup_{h \in \mathcal{H}} (e(h) - E(h)) < 4\mathbf{Rad}(\mathcal{L} \circ \{Z_1, \dots, Z_n\}) + c \sqrt{\frac{2 \log(1/\delta)}{n}} \quad (35)$$

Lemma 7. For hypothesis space $\mathcal{H}$ consisting of (regularized) neural network approximators with weights and biases bounded by a constant, there exists a certain constant $C_1, C_2 > 0$ so that for a set of n random variables $Z_1, \dots, Z_n$:

$$\mathbf{Rad}(\mathcal{H} \circ \{Z_1, \dots, Z_n\}) \leq \frac{1}{\sqrt{n}} (C_1 + C_2 \sqrt{\log d}) \quad (36)$$

Throughout this paper, we assume that the optimization error can be reduced to nearly 0, so that $E(h) \approx 0$ if the hypothesis space $\mathcal{H}$ contains function to be learned. From Lemma 6, the average estimation error generally scales with $c \sqrt{\frac{2 \log(1/\delta)}{n}}$

B Proofs of lemmas in Section 3

First of all, note that the continuous density ρ_0 for starting states on Ω is bounded below by a constant.

Lemma 2. *Suppose that the hypothesis space $\mathcal{H}$ for approximating the target function $g : \Omega \subset \mathbb{R}^d \rightarrow \mathbb{R}^d$ consists of neural network approximators with bounded weights and biases. Further assume that function g and neural network approximators $h \in \mathcal{H}$ are twice continuously differentiable with bounded first and second derivative by some constant C . One way this assumption can be satisfied is to choose activation functions that are twice continuously differentiable. Then $N(g, \mathcal{H}, \epsilon, \delta)$ is upper-bounded by $O(\epsilon^{-(2d+4)})$, where we only ignore quadratic factors of δ , polynomial terms of d , and other logarithmic terms.*

Proof. Suppose we're given $\epsilon > 0$. Take $\epsilon_1 > 0$ so that $16C\epsilon_1 < \epsilon/(2d)$, and $\epsilon_2 = 0.5$. Now take $n \in \mathbb{N}$ large enough so that $C_1 \sqrt{\frac{\log(1/(\delta/(3d)))}{n}} < \epsilon/(2d)$ for an appropriate constant C_1 . Then take $m \in \mathbb{N}$ large enough so that $(1 - c_1\epsilon_2\epsilon_1^d)^m < \delta/(3n)$, where c_1 is an appropriate geometric constants (see paragraphs below). Let $M = m + n$.

Train a neural network function $h \in \mathcal{H}$ to approximate the target function g on N samples, where N is given by:

$$N = \frac{\log((6M)/\delta)}{(C\epsilon_1^2\delta/(6M))^2} \text{ or } \sqrt{\frac{\log(1/(\delta/(6M)))}{N}} = C\epsilon_1^2 \frac{\delta}{6M} = N(g, \mathcal{H}, \epsilon, \delta) \quad (37)$$

Before going to the main proof, we dissect N to get its asymptotic rate in terms of ϵ and d . First of all $\epsilon_1 = O(\epsilon/(2d))$. Next, $n \approx \log(1/\delta)/(\epsilon/(2d))^2$, and $m \approx \log(3n/\delta)\epsilon_1^{-d}$. Hence by ignoring log terms, polynomial terms in d , and quadratic factor of δ , $M = m + n$ is approximately $O(\epsilon^{-d})$. As a result, $N \approx \epsilon^{-(2d+4)}$.

Choose a set $S = \{X_1, \dots, X_m, Y_1, \dots, Y_n\}$ consisting of $M = m + n$ random samples from distribution ρ_0 which are independent with g : m samples $X_1, \dots, X_m$ and n sample $Y_1, \dots, Y_n$.

We apply Lemma 6 to $\mathcal{H}$ with i.i.d labeled samples $Z = (X, g(X))$ and loss function l with the associated $\phi(y, \hat{y}) = |y - \hat{y}|$, which is 1-Lipschitz for a fixed $\hat{y}$. In this case, from Lemma 6, for probability of at least $1 - \delta/(6M)$, $\mathbb{E}_U[|h(U) - g(U)|] < C\epsilon_1^2\delta/(6M)$ for the random variable $U \in S$. As a result, by Markov inequality, with probability at least $1 - \delta/(6M) - \delta/(6M) = 1 - \delta/(3M)$, $|h(U) - g(U)| < C\epsilon_1^2$ for each $U \in \{X_1, \dots, X_m, Y_1, \dots, Y_n\}$. Hence there exists a probability subspace Γ with probability at least $1 - \delta/(3M) * M = 1 - \delta/3$ so that $|h(U) - g(U)| < C\epsilon_1^2$ for all $U \in S$.

The probability that a particular sample (random variable) X_i is in the hyper-cone with an conic angle difference of ϵ_2 surrounding the direction $(\nabla h(Y) - \nabla g(Y))$ of the hyper-spherical ϵ_1 -circular neighborhood of Y is at least $c_1 \epsilon_2 \epsilon_1^d$. Here a ϵ_1 -circular neighborhood here include points with radius sizes between $\epsilon_1/2$ and ϵ_1 . The probability that no m samples is in this cone is at most $(1 - c_1 \epsilon_2 \epsilon_1^d)^m < \delta/(3n)$. As a result, on Γ except a subspace with probability less than $\delta/(3n)$, there exists k (depends on both Y and $X_1, \dots, X_m$) so that:

$$\epsilon_1/2 < \|X_k - Y\| < \epsilon_1 \quad (38)$$

$$(\nabla h(Y) - \nabla g(Y)) \cdot (X_k - Y) > (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| \|X_k - Y\| \quad (39)$$

Second-order Taylor expansion for f and g at each Y yields:

$$h(X_k) = h(Y) + \nabla h(Y)(X_k - Y) + \frac{1}{2} \|X_k - Y\|^2 U_h(X_k, Y) \quad (40)$$

$$g(X_k) = g(Y) + \nabla g(Y)(X_k - Y) + \frac{1}{2} \|X_k - Y\|^2 U_g(X_k, Y) \quad (41)$$

where U_h and U_g are the second derivative terms of h and g in respectively, and are bounded by C . As a result:

$$\begin{aligned} & (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| (\epsilon_1/2) \\ & < (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| \|X_k - Y\| \\ & < (\nabla h(Y) - \nabla g(Y)) \cdot (X_k - Y) \\ & < |h(Y) - g(Y)| + |h(X_k) - g(X_k)| + 2C \|X_k - Y\|^2 \\ & < 2C \epsilon_1^2 + 2C \|X_k - Y\|^2 < 4C \epsilon_1^2 \end{aligned}$$

Thus, $\|\nabla h(Y) - \nabla g(Y)\| < 8C \epsilon_1 / (1 - \epsilon_2) = 16C \epsilon_1$. Therefore, on the subspace Γ_0 with probability of at least $1 - \delta/3 - n * (\delta/(3n)) = 1 - (2\delta)/3$, $\|\nabla h(Y) - \nabla g(Y)\| < 16C \epsilon_1$ for all samples $Y \in \{Y_1, \dots, Y_n\}$.

In order to prove that $\mathbb{E}_X \|\nabla h(X) - \nabla g(X)\| < \epsilon$ and thus finishing the proof for $N(g, \mathcal{H}, \epsilon, \delta) = O(\epsilon^{-(2d+4)})$, we only need to prove bounds on individual components of $\|\nabla h(X) - \nabla g(X)\|$:

$$\mathbb{E}_X \|(\nabla h(X) - \nabla g(X))_s\| < \epsilon/d$$

where x_s is the s^{th} component of a vector $x \in \mathbb{R}^d$. To this end, if we have a bound of order $O(1/\sqrt{n})$ for the Rademacher complexity of the hypothesis space $\mathcal{H}_s = \{\partial_s h, h \in \mathcal{H}\}$ for $s \in \overline{1, d}$, then we can invoke the corollary Lemma 6 on $\mathcal{H}_s$ to get:

$$\begin{aligned} \mathbb{E}_X \|(\nabla h(X) - \nabla g(X))_s\| & < \frac{1}{n} \sum_{i=1}^n \|(\nabla h(Y_i) - \nabla g(Y_i))_s\| + \epsilon/(2d) \\ & < \epsilon/(2d) + \epsilon/(2d) = \epsilon/d \end{aligned}$$

on T_0 except a set of probability of at most $\delta/(3d)$. Then we finish the proof of Lemma 2 by summing all inequalities over d components.

For the final argument to get a Rademacher complexity bound on $\mathcal{H}_s$, we use the chain rule to express each element of $\mathcal{H}_s$ as $h_1 \cdots h_R$, where R is the fixed number of layers in $\mathcal{H}$'s neural network architecture. Here each h_i can be expressed as the composition of Lipschitz (activation) functions and linear functions alternatively with bounded weight and bias. Those elements h_i then form another neural network hypothesis space. By invoking Lemma 7, we obtain a bound of order $O(1/\sqrt{n})$ on individual h_i 's. To connect these h_i 's, we express the product $h_1 \cdots h_R$ as:

$$\begin{aligned} h_1 \cdots h_R &= \prod_{i=1}^R (h_i + D - D) = \sum_{W \subseteq [R]} (-D)^{R-|W|} \prod_{i \in W} (h_i + D) \\ &= \sum_{W \subseteq [R]} (-D)^{R-|W|} \exp \left(\sum_{i \in W} \log(h_i + D) \right) \end{aligned}$$

Here D is a constant large enough to make the log function well-defined and to make Lipschitz constant of the log function bounded above by another constant. Now we use the simple bounds on $\mathbf{Rad}(\mathcal{T} + \mathcal{T}')$ and $\mathbf{Rad}(f \circ \mathcal{T})$ for some sets $\mathcal{T}$ and $\mathcal{T}'$ and some Lipschitz function f to derive the Rademacher complexity bound of the same order $O(1/\sqrt{n})$.

Another simpler way to achieve a similar result is to upper-bound $\|\nabla h(X) - \nabla g(X)\|$ by $\|\nabla h(Y_i) - \nabla g(Y_i)\| + 2C\|X - Y_i\|$ and use the probability subspace in which one of the Y_i is close enough to X . However, we need a similar argument for Lemma 3, so we moved forward with the current approach. $\square$

Lemma 3. *Suppose that, possibly with knowledge from outside environment or from certain policy experts, the hypothesis space $\mathcal{H}$ is reduced to a smaller family of functions of the form $g + h$, where h is a linearly bounded and p -weakly convex function for some $p \geq d$. Then $N(f, \mathcal{H}, \epsilon, \delta)$ can be upper-bounded by the factor $O(\epsilon^{-6})$ that is independent of the dimension d .*

Proof. Suppose we're given $\epsilon > 0$. Take $\epsilon_1 = \epsilon^\alpha$ with $\alpha = 1/d$, and set $\epsilon_2 = 1/2$. Now choose $m \in \mathbb{N}$ large enough so that $(1 - c_1 \epsilon_2 \epsilon_1^d)^m < \epsilon^2$, where c_1 is an appropriate geometric constant. From here, we can see that $m \approx O(\epsilon_1^{-d}) = O(\epsilon)$.

Choose $N \approx O(\epsilon^{-6}) \in \mathbb{N}$:

$$N = O\left(\frac{\log(d/\delta)}{\epsilon^6}\right) \text{ so that } \sqrt{\frac{\log(1/(\delta/d))}{N}} \approx O(\epsilon^3) \quad (42)$$

Train a neural network function g to approximate f on N samples $Y_1, \dots, Y_N$ so that $g(Y_i) \approx f(Y_i)$ for $i \in \overline{1, N}$. We prove that this N is large enough to allow derivative approximation transfer.

Choose $\beta_k = k/d$ and $\delta_k = k/d$ for $k \in \overline{0, d}$. We prove by induction on $k \in \overline{0, d}$ that there exists a subspace of probability at least $1 - \delta_k$ so that $\|\nabla f(Y) - \nabla g(Y)\| < C_0 \epsilon^{\beta_k}$ for all $Y \in \{Y_1, \dots, Y_n\}$ and for some constant C_0 .

For $k = 0$, the bound is trivial. For the induction step from k to $k + 1$, we first consider the loss function l_k of the form $l_k(h, (x, y)) = l_k(h, (x, g(x))) = \phi_k(h(x), g(x))$, where $\phi_k(y, \hat{y}) = \text{clip}(|y - \hat{y}|, 0, C\epsilon^{\alpha+\beta_k})^d$. Here clip is to denote a clip function, and, in this case, is obviously a Lipschitz function with Lipschitz constant 1. First, the loss function l_k is bounded by $(C\epsilon^{\alpha+\beta_k})^d$. Now note the following simple inequality:

$$|a^d - b^d| = |a - b| \left| \sum_{k=0}^{d-1} a^k b^{d-1-k} \right| < |a - b| d (C\epsilon^{\alpha+\beta_k})^{d-1}$$

for $a, b < C\epsilon^{\alpha+\beta_k}$. The inequality shows that ϕ_k has the Lipschitz constant bounded by $d(C\epsilon^{\alpha+\beta_k})^{d-1}$. We are now ready to go the main step of the induction step.

Condition on Y , we repeat the same argument in Lemma 2's proof to show that except for a subspace with probability less than ϵ^2 , there exists $j \in \overline{1, m}$ (depends on both Y and $X_1, \dots, X_m$) so that:

$$\epsilon_1/2 < \|X_j - Y\| < \epsilon_1 = \epsilon^\alpha \quad (43)$$

$$(\nabla h(Y) - \nabla g(Y)) \cdot (X_j - Y) > (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| \|X_j - Y\| \quad (44)$$

Under this subspace, because $h - g$ is linearly bounded,

$$|h(X_j) - g(X_j)| \leq C \|\nabla h(Y) - \nabla g(Y)\| \|Y - X_j\| < C\epsilon^{\alpha+\beta_k}$$

As a result, $|h(X_j) - g(X_j)|^d = \phi_k(h(X_j), g(X_j))$. Next, since $h - g$ is p -weakly convex and hence d -weakly convex, we have:

$$\begin{aligned} & (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| (\epsilon_1/2) \\ & < (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\| \|X_j - Y\| \\ & \leq (\nabla h(Y) - \nabla g(Y)) \cdot (X_j - Y) \\ & \leq |h(X_j) - g(X_j)| + |h(Y) - g(Y)| + C_1 \|X_j - Y\|^d \\ & < |h(Y) - g(Y)| + |h(X_j) - g(X_j)| + 2C_1 \epsilon^{\alpha d} \\ & = |h(X_j) - g(X_j)| + 2C_1 \epsilon^2 \\ & = \phi_k(h(X_j), g(X_j))^{1/d} + 2C_1 \epsilon^2 \\ & \leq \left(\sum_{i=1}^m \phi_k(h(X_i), g(X_i)) \right)^{1/d} + 2C_1 \epsilon^2 \end{aligned}$$

By taking expectation over $X_1, \dots, X_m$, we obtain

$$\begin{aligned}
& (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\|(\epsilon_1/2) \\
& < (1 - \epsilon^2) \left(\left(\mathbb{E}_{X_1, \dots, X_m} \left[\sum_{i=1}^m \phi_k(h(X_i), g(X_i)) \right] \right)^{1/d} + 2C_1\epsilon^2 \right) + C_2\epsilon^2 \\
& < \left(m \mathbb{E}_X [\phi_k(h(X), g(X))] \right)^{1/d} + C_3\epsilon^2 \\
& = \left(m \mathbb{E}_X [l_k(h(X), g(X))] \right)^{1/d} + C_3\epsilon^2
\end{aligned}$$

for appropriate constants $C_2, C_3 > 0$.

By Lemma 5 and Lemma 6 on the Lipschitz loss function l_k bounded by $(C\epsilon^{\alpha+\beta_k})^d$ with Lipschitz constant $d(C\epsilon^{\alpha+\beta_k})^{d-1}$, with probability of at least $1 - \delta/d$, we can continue the sequence of upper-bounds:

$$\begin{aligned}
& (1 - \epsilon_2) \|\nabla h(Y) - \nabla g(Y)\|(\epsilon_1/2) \\
& \leq C_4(\epsilon^{-1}(\epsilon^{\alpha+\beta_k})^{d-1}\epsilon^3)^{1/d} + C_3\epsilon^2 \\
& \leq C_5\epsilon^{\alpha+\beta_k+1/d} = C_5\epsilon^{\alpha+\beta_{k+1}}
\end{aligned}$$

for appropriate constants $C_4, C_5 > 0$.

Hence, we finish the induction step because except on a subspace with probability at most $1 - \delta_k - \delta/d = \delta_{k+1}$, the inequality for induction hypothesis at $(k+1)$ holds. For $\beta_d = 1$, we obtain $\|\nabla h(Y) - \nabla g(Y)\| < \epsilon$ for all $Y \in \{Y_1, \dots, Y_N\}$. By repeating the argument in Lemma 2's proof, we obtain the expected bound

$$\mathbb{E}_X [\|\nabla h(Y) - \nabla g(Y)\|] < C_6\epsilon \quad (45)$$

for an appropriate constant $C_6 > 0$. By rescaling by a constant, we showed that $N(g, \mathcal{H}, \epsilon, \delta) \approx O(\epsilon^{-6})$ $\square$

Lemma 4. *Given L and $\epsilon > 0$, define two sequences $\{\alpha_j\}_{j \geq 0}$ and $\{\epsilon_j\}_{j \geq 0}$ recursively as follows:*

$$\alpha_0 = 0 \text{ and } \alpha_j = L\alpha_{j-1} + \epsilon \quad (16)$$

$$\epsilon_1 = \epsilon \text{ and } \epsilon_{k+1} = L\alpha_k + \epsilon + L\epsilon_k \quad (17)$$

Then for each k , we get:

$$\epsilon_k \leq \frac{kL^k\epsilon}{L-1} \quad (18)$$

Proof. First, $\alpha_j = (L^{j-1} + \dots + 1)\epsilon$. Hence,

$$\begin{aligned}
\frac{\epsilon_k}{L^k} & \leq \frac{L(L^{k-2} + \dots + 1) + 1}{L^k} \epsilon + \frac{\epsilon_{k-1}}{L^{k-1}} \\
& = \frac{L^k - 1}{L^k(L-1)} \epsilon + \frac{\epsilon_{k-1}}{L^{k-1}} < \frac{\epsilon}{L-1} + \frac{\epsilon_{k-1}}{L^{k-1}}
\end{aligned}$$

Hence, by simple induction, $\frac{\epsilon_k}{L^k} < \frac{(k-1)\epsilon}{L-1} + \frac{\epsilon_1}{L} < \frac{k\epsilon}{L-1}$. As a result, $\epsilon_k < \frac{kL^k\epsilon}{L-1}$ as desired. $\square$