

MURO: DEPLOYMENT CONSTRAINED REINFORCEMENT LEARNING WITH MODEL-BASED UNCERTAINTY REGULARIZED BATCH OPTIMIZATION

Anonymous authors

Paper under double-blind review

ABSTRACT

In many contemporary applications such as healthcare, finance, robotics, and recommendation systems, continuous deployment of new policies for data collection and online learning is either cost ineffective or impractical. We consider a setting that lies between pure offline reinforcement learning (RL) and pure online RL called deployment constrained RL in which the number of policy deployments for data sampling is limited. To solve this challenging task, we propose a new algorithmic learning framework called Model-based Uncertainty Regularized batch Optimization (MURO). Our framework discovers novel and high quality samples for each deployment to enable efficient data collection. During each offline training session, we bootstrap the policy update by quantifying the amount of uncertainty within our collected data. In the high support region (low uncertainty), we encourage our policy by taking an aggressive update. In the low support region (high uncertainty) when the policy bootstraps into the out-of-distribution region, we downweight it by our estimated uncertainty quantification. Experimental results show that MURO achieves state-of-the-art performance in the deployment constrained RL setting.

1 INTRODUCTION

Recent advances in deep learning have enabled reinforcement learning (RL) to achieve remarkable success in various applications (Silver et al., 2017; OpenAI et al., 2019b; Vinyals et al., 2019; OpenAI et al., 2019a). Despite RL’s success, it suffers many problems. In particular, traditional RL algorithms require the agent to interact with the real world to collect large amounts of online data with the latest learned policy. However, the online deployment of the agent to the real world might be impossible in many real world applications (Matsushima et al., 2020). In the field of robotics or self-driving cars, for example, the cost of deploying the agent to the field may be too risky for the agent itself as well as its surrounding environment. In quantitative finance, trading strategy is usually carefully back-tested and calibrated offline with historical data and simulation. Since the market data has a low signal-to-noise ratio (Chen et al., 2020), online training can easily lead the policy fits to the noise, which is dangerous and can potentially trigger unexpected large monetary loss. In the recommendation system setting (Peska & Vojtas, 2020), after the policy has been trained offline, it will be deployed across different servers for serving many users. In such setting, online training of policy is difficult because the resulting policy might become unstable and/or cause bad user experience. Thus, large chunk of data is collected for the deployed policy. Then, the data is used for offline training. During the next scheduled deployment, the production level (data collection) policy is then updated.

To address this challenge, training an RL agent in an offline fashion seems to offer partial solution. In pure offline RL (Levine et al., 2020) setting, a static dataset is collected by a behavioral (or data collection) policy. Since the RL agent has no access to the online environment, training the RL agent faces many challenges. First, because the learning policy and the behavioral policy have different state visitation frequencies, evaluation of the offline policy is difficult. Second, during the training process, this discrepancy of the distribution might increase over the training time, a phenomenon known as distributional shift. Third, there might be a large extrapolation error when the value function is bootstrapped into out-of-distribution actions (Kumar et al., 2019), a phenomenon that can lead to learning divergence and instability.

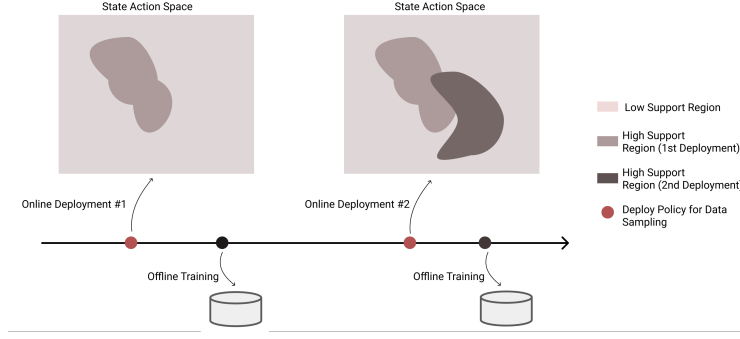


Figure 1: Illustration showing the deployment constrained RL setting. Large batch of data collection only occurs at the deployment point (showing two of them in red). The deployment and offline training occurs in an interleaved fashion. As the number of deployment increases, the state-action space will be explored more and more. Here, thicker-brown colored regions show the area of high support whereas the lighter color shows an under-explored (low support) region.

There has been a large body of literature in solving the pure offline problem, and they can be broadly categorized into the model-based and the model-free approach. In the model-free approach, they are either: 1) enforcing the learning policy to stay close with the behavioral policy as in Fujimoto et al. (2019); Kumar et al. (2019); Wu et al. (2019); Nachum et al. (2019b); Zhang et al. (2020); Nachum et al. (2019a) or 2) using an ensemble of Q values for stabilizing the learning and behaviors as in Ghasemipour et al. (2021); Wu et al. (2019); Nair et al. (2020). On the other hand, the model-based approaches are either: 1) implicitly enforcing learning policy to stay close to the behavioral policy through parameter initialization (Matsushima et al., 2020) or 2) injecting pessimism at the low confidence regions as in Yu et al. (2020); Kidambi et al. (2020).

Despite their empirical success, since the pure offline RL environment is fundamentally different than our deployment setting, we suspect developing novel method for deployment constrained RL can possibly achieve a large headroom of improvement, since pure offline RL algorithms assume zero interaction with the environment.

Our main contribution is the development of the Model-based Uncertainty Regularized batch Optimization (MURO), a novel framework for doing batch policy optimization for the deployment constrained RL setting. To build an accurate model of the environment, our method encourages the data collection policy to *discover high quality and novel data batches during each online deployment*. During the offline training, MURO samples fictitious rollouts from the learned model and performs policy update weighted with *uncertainty regularized coefficient*, a term that quantifies the amount of uncertainty with respect to each state-action pair within the fictitious rollouts and the data. Theoretically, we show that our approach optimizes the lower-bound of the true value function. Intuitively, our optimization process regularizes the update toward the high confidence regions. In areas of low confidence or low data support, MURO discounts the update by the uncertainty regularized coefficient when the policy bootstraps into out-of-distribution states and actions regions. We empirically compare our MURO to other strong baselines such as the state-of-the-art deployment constrained RL algorithms (BREMEN) (Matsushima et al., 2020), and we show that MURO is capable of achieving significant policy improvement while using smaller amounts of data and fewer deployment.

2 BACKGROUND

We consider a discounted, infinite horizon *Markov decision process (MDP)*, and denote $\Omega = (S, A, M, r, \mu_0, \gamma)$, where S is the state space, A is the action space, $M(s'|s, a)$ is the state transition kernel of transiting to a next state s' from state s while taking action a , $r(s, a)$ is the reward function, $\gamma \in (0, 1)$ is the discounting factor, and μ_0 is the initial state distribution.

In RL, the goal is to optimize a policy $\pi(a|s)$ such that the expected discounted return $J_M(\pi)$ is maximized: $J_M(\pi) = \mathbf{E}_{\pi, M, \mu_0} \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t)$. The value function is defined as $V_M^\pi(s) = \mathbf{E}_{\pi, s_{t+1} \sim M(s_t, a_t)} \{\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) | s_0 = s\}$. We denote the optimal policy $\pi^* = \operatorname{argmax}_{\pi} J_M(\pi)$. We

further define ρ_M^π to be the discounted state visited by π on M such that $\rho_M^\pi = (1 - \gamma) \sum_{t=0} \gamma^t p_{S_t^\pi}$, where $p_{S_t^\pi}$ is the distribution of state at time t .

In the **deployment constrained** setting, there is a limitation on the number of times that the policy can be deployed online for data collection. In real world applications such as recommender system or robotic control, due to the associated cost with each deployment, it is desirable to minimize the number of deployments to as small as possible. Let I be the number of times of deployment, and let $|B|$ be the size of a large batch of training data collected when a policy has been taken online (or deployed), then, the total size of data collected throughout the entire process is $I \times |B|$. Further, let $B^{(i)}$ be the batch of data collected during the i_{th} deployment. In each batch B , we store the data transition $\{(s, a, r, s')\}$. Note the difference between the deployment constrained setting and the pure offline RL setting. In the pure offline RL setting, we consider it as a single batch of data with $I = 1$, whereas in the deployment constrained setting, the agent will have the opportunity to interact with the environment by deploying its learning or other behavioral policy for data collection purposes.

In **Model-Based RL** (MBRL), the world dynamics are learned from the sampled data. We estimate the ground truth environment transition dynamics M^* by learning a forward next state transition dynamics model $\hat{M} : S \times A \rightarrow S$. Let $\hat{M}^{(i)}$ be the estimated dynamics bootstrapped from data batch $B^{(1)} \cup B^{(2)} \cup \dots \cup B^{(i)}$, in which we assume the model is updated with the aggregated data batches up to the i_{th} deployment. We denote by $\hat{M}_\theta^{(i)}$ the estimated dynamics learned by neural networks parameterized by θ .

3 ALGORITHMIC FRAMEWORK

For the purpose of exposition, we start by presenting an idealized version of MURO algorithmic framework. Then, we describe a practical version that we have implemented and perform quite well. The detailed algorithm is summarized in Algorithm 1.

The heart of the algorithm relies on uncertainty estimation. We postulate that by having more diversified samples (data coverage), we can learn a more accurate representation of the environment, which is the key ingredient of MBRL. To get the most out of each deployment, our algorithm tries to explore the under-represented region (the low support area of our collected data) so as to minimize the amount of uncertainty or maximize the information gain. After building an accurate model, we can use it to generate fictitious rollouts for policy training. The idea is that we want to weight the contribution of each policy update by taking the uncertainty quantification on the fictitious rollouts (each state action pair) by **uncertainty regularized coefficient** ($U_{\hat{M}^{(i)}, M^*}(s, a)$) which characterize the level of uncertainty results from model estimation error (eq. (2)). For fictitious rollouts that we have low data support, we discount the policy update by $U_{\hat{M}^{(i)}, M^*}(s, a)$ and discourage the policy from bootstrapping into unknown regions. In doing this, we are *regularizing* the policy update into high support/confidence regions.

Theoretical Results. All the proofs in this section are deferred to the appendix. In the deployment constrained setting, we are interested in iterative improvement after each deployment. Let π_{ref} be a reference policy that we will deploy to collect the data, $d(\cdot, \cdot)$ be the closeness of the two policies, $D(\cdot, \cdot)$ be the discrepancy between the M^* and $\hat{M}$, we make the followings four assumptions:

$$V_{M^*}^\pi \geq V_{\hat{M}}^\pi - D_{\pi_{\text{ref}}, \delta}(\hat{M}, \pi), \text{ s.t. } d(\pi, \pi_{\text{ref}}) \leq \delta \quad (\text{A1}) \quad \hat{M} = M^* \implies D_{\pi_{\text{ref}}}(\hat{M}, \pi) = 0 \quad (\text{A2})$$

$$\text{L-Lipschitz in } V_{\hat{M}}^\pi \text{ w.r.t to some norm such that } |V_{\hat{M}}^\pi(s) - V_{\hat{M}}^\pi(s')| \leq L \|s - s'\|, \forall s, s' \quad (\text{A3})$$

$$D_{\pi_{\text{ref}}}(\hat{M}, \pi) \text{ is given by the form of } E_{\tau \sim \pi_{\text{ref}}, M^*} \{f(\hat{M}, \pi, \tau)\} \quad (\text{A4})$$

Let $\hat{M}^{(i)}$ be the model estimated from the data collected up to the i_{th} deployment. Adapted from Schulman et al. (2015); Luo et al. (2018):

Lemma 1. *The difference of the value functions in-between each deployment is bounded by:*

$$|V_{\hat{M}^{(i+1)}}^\pi - V_{\hat{M}^{(i)}}^\pi| \leq \kappa L \mathbf{E}_{(s, a) \sim \rho_{\hat{M}^{(i+1)}}^\pi} (\|\hat{M}^{(i+1)}(s, a) - \hat{M}^{(i)}(s, a)\|) \quad (1)$$

We define the following function $U_{\hat{M}^{(i)}, M^*}(s, a) : S \times A \rightarrow \mathbb{R}$ as an **uncertainty regularized coefficient**:

$$\mathbf{E}_{(s,a) \sim \rho_{\hat{M}^{(i)}}^\pi} \{U_{\hat{M}^{(i)}, M^*}(s, a)\} \triangleq g\left(\kappa \mathbf{E}_{(s,a) \sim \rho_{\hat{M}^{(i)}}^\pi} \|M^*(s, a) - \hat{M}^{(i)}(s, a)\|\right) \quad (2)$$

where $g(x) \triangleq (V_{\hat{M}^{(i)}}^\pi - x)(V_{\hat{M}^{(i)}}^\pi)^{-1}$, and $\kappa \triangleq (1 - \gamma)^{-1}\gamma$.

Proposition 1. Let $U_{\hat{M}^{(i)}, M^*}(s, a)$ be the uncertainty regularized coefficient defined by eq. (2), then the performance of π on the ground-truth M^* is lower bounded by that of the estimated $\hat{M}^{(i)}$, weighted by $U_{\hat{M}^{(i)}, M^*}(s, a)$:

$$V_{M^*}^\pi \geq V_{\hat{M}^{(i)}}^\pi \mathbf{E}_{(s,a) \sim \rho_{M^*}^\pi} \{U_{\hat{M}^{(i)}, M^*}(s, a)\}. \quad (3)$$

Proposition.1 says that in optimizing the lower bound (the RHS of eq. (3)) with $V_{\hat{M}^{(i)}}^\pi$, we can maximize the overall performance of π on the real dynamics M^* .

Interpretation of $U_{\hat{M}^{(i)}, M^*}(s, a)$ (uncertainty regularized coefficient). Here, the term can be interpreted as an uncertainty quantification measure as a result of *model estimation error* between $\hat{M}^{(i)}$ and M^* , with $\hat{M}^{(i)}$ being learned from the collected (and limited amount of) data up to the i_{th} deployment. On the regions where we have high data support, $\hat{M}^{(i)}(s, a)$ is close to $M^*(s, a)$, thus their discrepancy decreases and $U_{\hat{M}^{(i)}, M^*}(s, a)$ approaches to unity. On the regions that are less certain, the discrepancy increases and thus $U_{\hat{M}^{(i)}, M^*}(s, a)$ decreases.

Remark 1. If our estimated model $\hat{M}^{(i)}$ is reasonably good, so that the model estimation error $\kappa \|M^*(s, a) - \hat{M}^{(i)}(s, a)\|$ is smaller than $V_{\hat{M}^{(i)}}^\pi$, and that $V_{\hat{M}^{(i)}}^\pi$ is positive, then $U_{\hat{M}^{(i)}, M^*}(s, a)$ is a scalar between 0 and 1, and is approximately proportional to the accuracy of the model estimation.

Thus to optimize $V_{M^*}^\pi$, we instead optimize its lower bound, which is the value function at the estimated model $V_{\hat{M}^{(i)}}^\pi$ weighted by the uncertainty regularized coefficient. We assign a higher weight at the regions of high confidence, and a lower weight at the regions of low confidence. Next, we utilize our result from eq. (3) for establishing our MURO lower bound optimization algorithm.

Algorithm 1 MURO: Model-based Uncertainty Regularized batch Optimization

Given: Size of data batches B , Number of deployments I , $D_{all} \leftarrow \{\}$

for $i = 1, 2, 3, \dots, I$ of deployments **do**

- 1 deploy π online for data collection
- 2 $D_{all} \leftarrow D_{all} \cup \text{Collect-Data-Low-Support-Region}(\pi)$
- 3 $\hat{M}_\theta \leftarrow \text{Learn approximate transition dynamics model}$
- 4 Train Uncertainty-Labeler(D_{all})
- 4 $\pi \leftarrow \text{Train with MBRL(Uncertainty-Labeler, } \hat{M}_\theta, D_{all})$ as in section 3.2

return π

3.1 PRACTICAL IMPLEMENTATION

In this section, we explain the practical implementation of algorithm 1 in detail. For readers convenience, we have labeled the line number in the algorithm, and we will refer to the line number as we explain below.

Learning the transition dynamics (line 2): For estimating the transition dynamics $\hat{M}_\theta$, we use an ensemble of N deterministic dynamics models with multi-layer fully-connected perceptrons as in Mishra et al. (2017); Kurutach et al. (2018) to reduce model bias. They are trained to predict the next state with L_2 loss:

$$\min_{\theta} \frac{1}{|D_{all}|} \sum_{(s_t, a_t, s_{t+1}) \in D_{all}} \|s_{t+1} - \hat{M}_\theta(s_t, a_t)\|_2^2. \quad (4)$$

Since we use the data collected up to i_{th} deployment, we drop the (i) superscript on $\hat{M}_\theta^{(i)}$ hereafter.

Uncertainty-Labeler (line 3). This module is responsible for characterizing the levels of uncertainty in the collected data and approximated $U_{\hat{M}^{(i)}, M^*}(s, a)$. Specifically, we want to identify the regions in our data that have high support or low support with respect to uncertainty. Since the oracle is unavailable to us, here we can only approximate the uncertainty by state-actions visitation frequency.

Shall the state-actions visited frequently, more pairs of them will show up in the data. If we were to train models on the batches of data to predict the next state transition dynamics, then, our prediction will be more accurate on the states that we have seen (in the collected data), and less accurate on the unfamiliar state. Thus, we have established that the uncertainty quantification measure as the next state prediction error.

In the actual implementation, we have used K ensembles of Probabilistic Neural Networks (PNN) (Chua et al., 2018) to *capture the uncertainty* within the data. Let $\hat{P}_\phi : S \times A \rightarrow S$ be the PNN with parameter ϕ . In PNN, the network has its output neurons parameterized by a Gaussian distribution in the effort of capturing the uncertainty. As explained above, we use $\hat{P}_\phi$ to quantify the uncertainty by prediction error. The lower the prediction error in relation to the ground truth, the higher the support. Our uncertainty-labeler approximates the uncertainty regularized coefficient $U_{\hat{M}^{(i)}, M^*}(s, a)$ by $\hat{U}(a, s)$ (define below shortly), and we have dropped the dependency on subscript since $\hat{U}$ is trained with the data collected up to the i_{th} deployment.

The actual implementation of $\hat{U}(a, s)$ is motivated by remark 1. Let $\hat{\tau}$ be the fictitious trajectory generated by $\hat{M}_\theta$ (the learned dynamics), we use $\hat{P}_\phi$ to quantify the uncertainty. Since the ground truth state is unavailable to us in the fictitious rollouts, to calculate the prediction error, we further approximate it by the intra-discrepancy error within the ensembles. We randomly sample two models (a, b) from K Uncertainty-Labeler ensemble, for each state-action pairs in $\hat{\tau}$, we label them :

$$\hat{U}(s_t, a_t) = \exp(-\alpha \|\hat{s}_{t+1}^{(a)} - \hat{s}_{t+1}^{(b)}\|_1), \text{ for } (s_t, a_t) \in \hat{\tau} \quad (5)$$

where $\hat{s}_{t+1}^{(a)} = \hat{P}_\phi^{(a)}(s_t, a_t)$ is the next state predicted by the sampled a_{th} model from the Uncertainty-Labeler (and similarly for b index), and $\alpha > 0$ is a temperature parameter. Here, we have assumed that our estimated models $\hat{P}_\phi$ are operating on the "reasonably good" regime (as by remark 1). To verify how good our estimated $\hat{U}(s_t, a_t)$ in approximating the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$, we perform additional ablation study on Appendix G.

Note that here, we have used different sets and different types of networks for estimating the model dynamics $\hat{M}_\theta$ and the uncertainty $\hat{P}_\phi$. We use $\hat{M}_\theta$ (deterministic networks) to generate fictitious rollouts, and we use $\hat{P}_\phi$ (PNN, a specialized uncertainty probabilistic network) as the "judge" to quantify the uncertainty for the generated fictitious rollouts. We separate out these two networks intentionally to avoid possible error propagation between the two modules. $\hat{P}_\phi$ is trained by (eq. (17)). For detail discussion, see Appendix F.

Collect Data from Low Support Region (line 1). To maximize the benefit out of each deployment and to build an accurate representation $\hat{M}_\theta$ of the environment, we emphasize that we want to achieve the maximal data coverage as well as exploring the under-explored region (the low support area).

Exploration in the deployment constrained setting is non-trivial because it is hard to evaluate which action-state pairs will fill the low support region and lead to high data coverage. During online data sampling, for each state that our agent encounters, we identify the amount of uncertainty by taking the actions which lead to maximal uncertainty (prediction error) between our predicted next state ($\hat{s}'$) and the ground truth next state (as by eqn.16 in the Appendix E). Following a trajectory of maximal prediction error leads to novel experience discovery and reduce the number of unknown regions within the data. In the ablation study (Fig.4 c. and d.), we show that this strategy leads to more accurate learned dynamics $\hat{M}_\theta$. (For detailed implementation, please refer to Appendix E).

3.2 OFFLINE TRAINING WITH MBRL METHOD

Next, we define our offline model-based training method (line 4). In the function below, we provide a practical instantiation of MURO offline model-based method for learning a policy.

Our method utilizes trust region policy optimization (TRPO) (Schulman et al., 2015) with fictitious trajectory generated with learned ensembles of transition dynamics $\hat{M}_\theta$ weighted by Uncertainty-Labeler (or uncertainty regularized coefficient).

Fictitious trajectory generated from learned dynamics (line 7 to line 8). After each deployment, the environment dynamics transitions are estimated by an ensemble of $\{\hat{M}_\theta\}_1^N$ trained using D_{all} . To generate a fictitious trajectory, we first randomly sample a model $j \in (1, 2, \dots, N)$, and then we rollout

Function MURO MBRL Training ($\hat{M}_\theta, D_{all}, \text{Uncertainty-Labeler}$) :

```

5 | Initialize  $\pi_{\text{init}}$  with the data collection policy.
6 | for training iterations do
7 |   Randomly sample a dynamics model from  $N$  ensembles of  $\hat{M}_\theta$ 
8 |   for optimization steps do
9 |      $\hat{\tau} \leftarrow$  sample fictitious trajectory from  $\hat{M}_\theta$ 
10 |     $\hat{\tau}_{\text{labeled}} \leftarrow$  label fictitious trajectory with Uncertainty-Labeler ( $\hat{\tau}$ ) as in eq.(5)
    train with uncertainty regularized TRPO with  $\hat{\tau}_{\text{labeled}}, \pi_{\text{init}}$  as in eq.(6)

```

the trajectory $\hat{\tau}$ by running the learning policy π with the next state as $\hat{s}_{t+1} = \hat{M}_{\theta_j}(\hat{s}_t, a_t = \pi(\hat{s}_t))$, for $t \in 1, \dots, T$.

TRPO training with Uncertainty (line 9 to line 10). Next, we train the policy with uncertainty regularized TRPO. In the offline setting, TRPO is trained using imaginary rollouts generated from the learned dynamics $\hat{M}_\theta$. Utilizing Prop 1, we use TRPO to optimize the $V_{\hat{M}^*}^\pi$ by improving its lower-bound (the RHS). Here, we replace $V_{\hat{M}^{(i)}}^\pi(s)$ by $A^{\pi_{\theta_k}}(s, a)$, and we approximated $U_{\hat{M}^{(i)}, \hat{M}^*}(s, a)$ by $\hat{U}(s, a)$ (eqn.5), with TRPO:

$$\begin{aligned}
& \underset{\vartheta}{\operatorname{argmax}} E_{\pi_\vartheta(s), s, \hat{P}_\phi^{(a,b)}, \hat{M}_\theta} \left\{ \frac{\pi_\vartheta(a|s)}{\pi_{\vartheta_k}(a|s)} A^{\pi_{\vartheta_k}}(s, a) \hat{U}(a, s) \right\} \\
& \text{s.t. } E_{a \sim \pi_\vartheta(s), s, \hat{P}_\phi^{(a,b)}, \hat{M}_\theta} \{ D_{KL}(\pi_\vartheta(\cdot|s) \| \pi_{\vartheta_k}(\cdot|s)) \} \leq \delta
\end{aligned} \tag{6}$$

where we set $\pi_{\vartheta_0} = \pi_{\text{init}}$ as the initial policy of the TRPO at the first iteration (as a way to impose the constraint on π_{ref}). Here, $A^{\pi_{\vartheta_k}}(s, a)$ is the advantage function of policy π_{ϑ_k} following a fictitious trajectory generated by $\hat{M}_\theta$. Here, when the policy bootstraps into out-of-distribution states and actions regions (low support area), we discount the update. In doing this, our optimization process regularizes the update toward the high confidence regions.

4 EMPIRICAL RESULTS

We empirically evaluate our proposed MURO algorithm with five continuous control benchmarks using the MuJoCo¹ physics simulator: Walker2d, Hopper, Half-Cheetah, Ant, and Cheetah-Run. **Baseline.** We compare our MURO with BREMEN, a state-of-the-art algorithm² that is designed specifically for deployment constrained setting. For BREMEN, we used the exact same hyperparameters as what was originally proposed in Matsushima et al. (2020). As for comparison, we also adapted MOPO (Yu et al., 2020), a pure-offline algorithm, to deployment-constrained setting. We used the latest learned policy for data collection.

Evaluation Setup. Following Matsushima et al. (2020), our evaluation set up consisted of 5 deployments for Half-Cheetah, Ant, Cheetah-Run and 10 deployments for Walker2D and Hopper. To see the trade-off between the number of deployment versus data sample size, we perform our empirical tests on two separate cumulative data sizes of 500k and 250k.

For the settings of 500k / (250k) sample size experiments, the environments Half-Cheetah, Ant and Cheetah-Run have a per deployment data collection batch size $|B| = 100k/(50k)$ (with 5 deployments in total). For the environments of Walker2d and Hopper, they have a per deployment data collection batch size of $|B| = 50k/(25k)$ (with 10 deployments in total).

Deployment Setup. In the deployment constrained setting, the data collection only happens during the deployment. No training happens during this period. Only until a data batch size of $|B|$ has been collected, the agent will be taken offline for training and policy update. At the first (initial) deployment, a random policy is used for data collection. After that, the data collection will be replaced by the updated policy and will be launched to deploy again. We plot our 500k data size results in Fig.2, and we plot our 250k data size results on Fig.3. For all experimental results, we averaged over 5 random seeds.

¹<http://www.mujoco.org/>

²BREMEN has been compared against to SAC, Model-Ensembles-TRPO, BCQ, and BRAC (all adapted to deployment-constrained setting) and shows state-of-the-art performance.

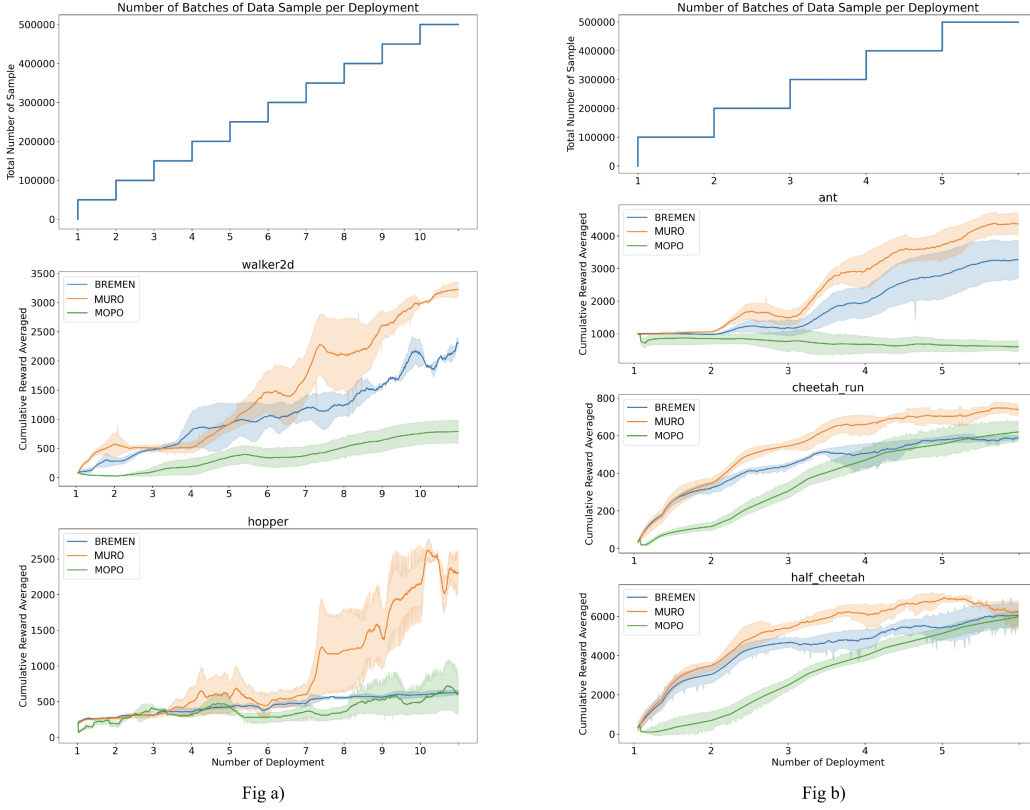


Figure 2: Empirical Evaluations of a) Walker2D and Hopper for $I = 10$ deployments, $|B| = 50k$, and b) Ant, Half_Cheetah and Cheetah_Run for $I = 5$ deployments, $|B| = 100k$. Total data consumption in both settings is $I \times |B| = 500k$. The x-axis is aligned showing the number of deployment.

Deployment	2nd	3rd	4th	5th
MURO	0.995	0.959	0.945	0.941
Baseline(BREMEN)	0.972	0.903	0.891	0.858

Table 1: Table showing the amount of novelty of each data batches collected between each deployment for cheetah_run. Novelty is measured as cosine distance between each observation (state) in $B^{(i)}$ versus the previous aggregated batches ($B^{(1)} \dots \cup B^{(i-1)}$), averaged over the number of transitions.

Empirical Details: 500k data size. The top(first) figure of Fig.2 shows the total sample size of each deployment on the y-axis, and on the x-axis, we show the number of deployments. We align along the x-axis for all figures. Our method works the best on the Walker2D and Hopper, in which our MURO approach achieves significant cumulative rewards especially in the longer deployments (after 6_{th}). In the Hopper environment, we see that the baselines (BREMEN and MOPO) show incapable of learning a meaningful policy while our MURO significantly outperforms. On the other hand, in the Ant, Cheetah-Run and Half-Cheetah environment (as shown in part b. of Fig.2), our MURO is capable of achieving a higher performance using a smaller amount of data. For instance, in the Cheetah-Run environment, MURO already achieved 550 points in the second deployment but the baselines take four to five deployments to achieve the same score.

Empirical Details: 250k data size. Different than the previous set of experiments, here, we reduce the total data size by half to 250k. In general, when we reduce the data size, the performances of all algorithms will be reduced. Despite this, our MURO algorithm performs quite well even in this setting. Comparing to the 500k data size experiment, we also observe similar trends. In Fig.3 a), we plot the results for Walker2d and Hopper. For the former, we see a significant winning margin in Walker2d whereas in Hopper, the winning only happens in the $8_{th} \sim 9_{th}$ deployments. In Fig.3 b), we plot the results for Ant, Cheetah-Run and Half-Cheetah. We observe that our MURO algorithm achieves faster learning and stronger performance.

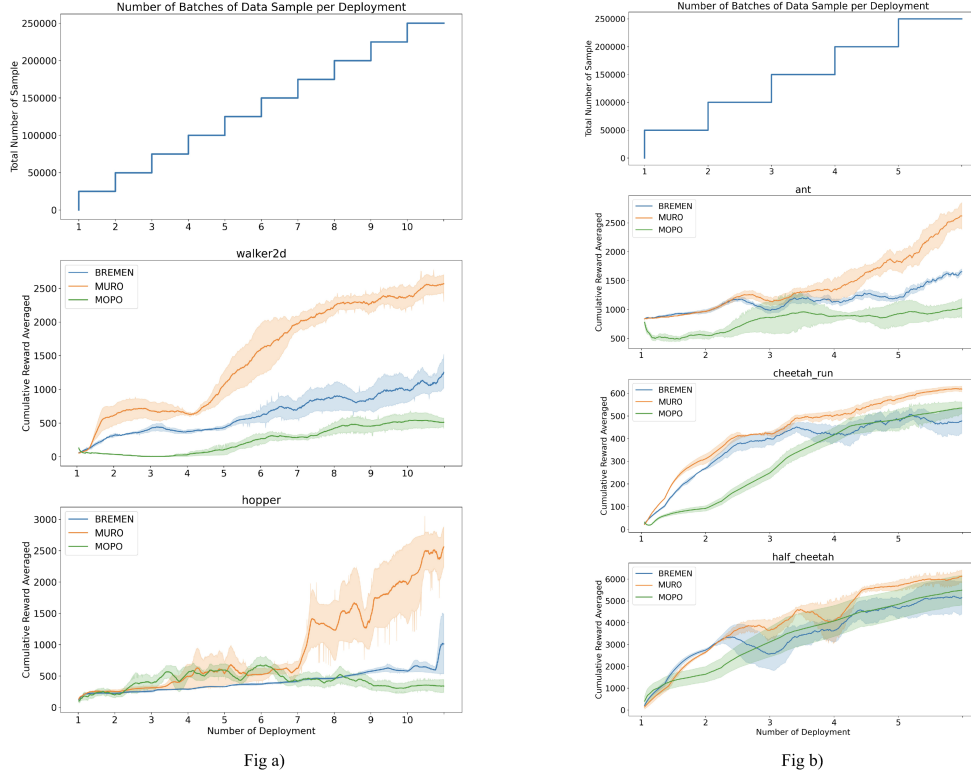


Figure 3: Empirical Evaluations of a) Walker2D and Hopper for $I = 10$ deployments, $|B| = 25k$ and b) Ant, Half_Cheetah and Cheetah_Run for $I = 5$ deployments, $|B| = 50k$. Total data consumption in both settings is $I \times |B| = 250k$. The x-axis is aligned showing the number of deployment.

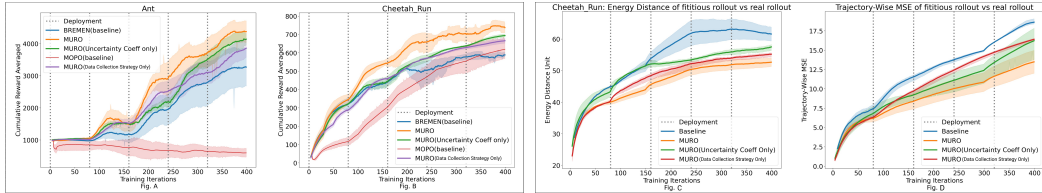


Figure 4: Ablation study on $I \times |B| = 500k$ setting: isolating the effects of having either uncertainty coeff only or data collection strategy only (as in eqn.16) on MURO. We plot the cumulative rewards on fig. a) and b) for ant and cheetah_run. We see that the effect of uncertainty coeff is stronger than the effect of data collection strategy. On subplots c) and d) we show quality of the learned model dynamics. We compare the fictitious rollouts against real rollouts in terms of energy distance (the **lower** the better) and trajectory-wise MSE (the **lower** the better) for the cheetah_run. We see that MURO (data collection strategy) gives stronger effect. The combined effects of two components (which becomes MURO) gives the best performance not only on the cumulative rewards (subplots a and b), but also on the learned dynamics (subplots c and d).

5 ABLATION STUDY

In this section, we examine the MURO algorithm in terms of the following three aspects: 1) discovery of high quality and novel data batches during each deployment; 2) whether this will lead to more accurate learning of model dynamics; 3) isolation effect of having uncertainty coefficient only or data-collection strategy only. To assess the novelty of data-batches during each deployment, we calculate the average cosine distance between the current batch ($B^{(i)}$) versus the aggregation of all of the previous batches ($B^1 \cup B^2 \cup \dots B^{(i-1)}$) and then show the result in Table.1. Note that deployment 1 is not shown because the initial data collection policy is a random policy. Our result shows that our MURO algorithm leads to much higher novel transitions discovery.

In Fig.4 subplots a and b, we isolate each component and show the effects of including either 1) only the uncertainty coefficient or 2) only our data collection strategy as in eqn.16. In terms of cumulative rewards, the former (coeff only) has a stronger effect than latter (data collection strategy only). Combining both gives the optimal effect on the cumulative rewards and leads learning an accurate model (subplots c and d).

Next, in Fig.4 subplots c and d, we compare the performance of the fictitious rollouts from the learned model dynamics versus the rollouts from the real environment. We first plot the energy distance in c) and we show the mean square error (MSE) of trajectory-wise rollouts in d). As the number of deployment increases, our method is capable of discovering high quality transitions which result in a better estimated model. In terms of the learned model dynamics, the data collection strategy only has a stronger effect in learning a more accurate model.

Lastly, in Appendix G, we show how good is the approximated $\hat{U}(s, a)$ in estimating the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$.

6 RELATED WORKS

Deployment-constrained RL. To the best of our knowledge, Matsushima et al. (2020) is the first work that proposed the concept of deployment-constrained efficiency. In their paper, the authors proposed the algorithm Behavior-Regularized Model-ENsemble (BREMEN) that enforces implicit KL-divergence between the learning policy and the behavioral policy by parameter initialization. In contrast, our work optimizes the lower bound of the true value function with the uncertainty regularized coefficient, an approach with theoretical support. In addition, we show how to make use of the uncertainty measure to maximize the benefit of each online deployment, and to connect between the online (data collection) and offline (policy learning) setting. Empirically, this leads to higher performance while using fewer number of deployment.

On the other hand, there also some related research works such as Bai et al. (2019); Pan et al. (2020) and in semi-batch RL such as Ernst et al. (2005); Lange et al. (2012); Jaakkola et al. (1999); Chu & Kitani (2020).

Model-Based RL. In model-based approach, the world representation is learned first and then is used for generating imaginary rollouts. Related research works are Chua et al. (2018); Janner et al. (2019); Luo et al. (2018); Munos & Szepesvári (2008). However, direct application of MBRL methods into the offline setting can be challenging due to distribution shifts. Nevertheless, the closely related research is Kurutach et al. (2018), in which an ensemble of estimated model dynamics is used for generating fictitious rollouts for stabilizing effects. Similar approaches have also been investigated by Zhang et al. (2019); Kaiser et al. (2019); Veerapaneni et al. (2020); Feinberg et al. (2018).

Offline RL. Unlike deployment-constrained setting, offline RL assumes no interaction with the environment. Thus, the learning policy needs to reason about the behavioral policy and make policy update base on that. The two most related literatures are Yu et al. (2020) and Kidambi et al. (2020). In Yu et al. (2020), the authors proposed an uncertainty penalized MDP for explicitly penalizing the uncertainty. On the other hand, the authors from Kidambi et al. (2020) proposed a pessimistic MDP that divides the environment into two regions: known or unknown. When the agent is entering into the unknown region, the MDP undergoes a halt state (or absorbing state) for penalizing this action. Here instead, we show that in the deployment setting, we have developed the uncertainty regularized coefficient which directly applies to the learning value function. We show that optimizing this uncertainty regularized value function in-between each deployment is the same as optimizing the lower bound of the true value function. Intuitively, our approach regularizes the optimization toward the higher confidence regions.

Other model-free offline RL works include Fujimoto et al. (2019); Kumar et al. (2019); Wu et al. (2019); Nachum et al. (2019b); Zhang et al. (2020); Nachum et al. (2019a) and Ghasemipour et al. (2021); Wu et al. (2019); Nair et al. (2020).

7 CONCLUSION

In this paper, we have proposed the algorithmic framework MURO for optimizing the policy learning under the deployment-constrained setting.

REFERENCES

- Yu Bai, Tengyang Xie, Nan Jiang, and Yu-Xiang Wang. Provably efficient q-learning with low switching cost. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- Luyang Chen, Markus Pelger, and Jason Zhu. Deep learning in asset pricing. *Available at SSRN 3350138*, 2020.
- Wen-Hsuan Chu and Kris M. Kitani. Neural batch sampling with reinforcement learning for semi-supervised anomaly detection. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm (eds.), *Computer Vision – ECCV 2020*, pp. 751–766, Cham, 2020. Springer International Publishing. ISBN 978-3-030-58574-7.
- Kurtland Chua, Roberto Calandra, Rowan McAllister, and Sergey Levine. Deep reinforcement learning in a handful of trials using probabilistic dynamics models. *Advances in Neural Information Processing Systems*, 31, 2018.
- Damien Ernst, Pierre Geurts, and Louis Wehenkel. Tree-based batch mode reinforcement learning. *Journal of Machine Learning Research*, 6:503–556, 04 2005.
- Vladimir Feinberg, Alvin Wan, Ion Stoica, Michael I Jordan, Joseph E Gonzalez, and Sergey Levine. Model-based value estimation for efficient model-free reinforcement learning. *arXiv preprint arXiv:1803.00101*, 2018.
- Scott Fujimoto, David Meger, and Doina Precup. Off-policy deep reinforcement learning without exploration. In *International Conference on Machine Learning*, pp. 2052–2062. PMLR, 2019.
- Seyed Kamyar Seyed Ghasemipour, Dale Schuurmans, and Shixiang Shane Gu. Emaq: Expected-max q-learning operator for simple yet effective offline and online rl. In *International Conference on Machine Learning*, pp. 3682–3691. PMLR, 2021.
- Tommi Jaakkola, Satinder Singh, and Michael Jordan. Reinforcement learning algorithm for partially observable markov decision problems. *Advances in Neural Information Processing Systems*, 7, 11 1999.
- Michael Janner, Justin Fu, Marvin Zhang, and Sergey Levine. When to trust your model: Model-based policy optimization. *Advances in Neural Information Processing Systems*, 32:12519–12530, 2019.
- Lukasz Kaiser, Mohammad Babaeizadeh, Piotr Milos, Blazej Osinski, Roy H Campbell, Konrad Czechowski, Dumitru Erhan, Chelsea Finn, Piotr Kozakowski, Sergey Levine, et al. Model-based reinforcement learning for atari. *arXiv preprint arXiv:1903.00374*, 2019.
- Rahul Kidambi, Aravind Rajeswaran, Praneeth Netrapalli, and Thorsten Joachims. Morel: Model-based offline reinforcement learning. In H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 21810–21823. Curran Associates, Inc., 2020.
- Aviral Kumar, Justin Fu, Matthew Soh, George Tucker, and Sergey Levine. Stabilizing off-policy q-learning via bootstrapping error reduction. *Advances in Neural Information Processing Systems*, 32:11784–11794, 2019.
- Thanard Kurutach, Ignasi Clavera, Yan Duan, Aviv Tamar, and Pieter Abbeel. Model-based trust-region policy optimization. In *International Conference on Learning Representations*, 2018.
- Sascha Lange, Thomas Gabel, and Martin Riedmiller. Batch reinforcement learning. *Reinforcement Learning: State of the Art*, 01 2012. doi: 10.1007/978-3-642-27645-3_2.
- Sergey Levine, Aviral Kumar, George Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems, 2020.
- Yuping Luo, Huazhe Xu, Yuezhi Li, Yuandong Tian, Trevor Darrell, and Tengyu Ma. Algorithmic framework for model-based deep reinforcement learning with theoretical guarantees. In *International Conference on Learning Representations*, 2018.

- Tatsuya Matsushima, Hiroki Furuta, Yutaka Matsuo, Ofir Nachum, and Shixiang Gu. Deployment-efficient reinforcement learning via model-based offline optimization. In *International Conference on Learning Representations*, 2020.
- Nikhil Mishra, Pieter Abbeel, and Igor Mordatch. Prediction and control with temporal segment models. In *International Conference on Machine Learning*, pp. 2459–2468. PMLR, 2017.
- Rémi Munos and Csaba Szepesvári. Finite-time bounds for fitted value iteration. *Journal of Machine Learning Research*, 9(27):815–857, 2008. URL <http://jmlr.org/papers/v9/munos08a.html>.
- Ofir Nachum, Yinlam Chow, Bo Dai, and Lihong Li. Dualdice: Behavior-agnostic estimation of discounted stationary distribution corrections. *Advances in Neural Information Processing Systems*, 32:2318–2328, 2019a.
- Ofir Nachum, Bo Dai, Ilya Kostrikov, Yinlam Chow, Lihong Li, and Dale Schuurmans. Algaedice: Policy gradient from arbitrary experience. *arXiv preprint arXiv:1912.02074*, 2019b.
- Ashvin Nair, Murtaza Dalal, Abhishek Gupta, and Sergey Levine. Accelerating online reinforcement learning with offline datasets, 2020.
- OpenAI, :, Christopher Berner, Greg Brockman, Brooke Chan, Vicki Cheung, Przemyslaw Debiak, Christy Dennison, David Farhi, Quirin Fischer, Shariq Hashme, Chris Hesse, Rafal Jozefowicz, Scott Gray, Catherine Olsson, Jakub Pachocki, Michael Petrov, Henrique Ponde de Oliveira Pinto, Jonathan Raiman, Tim Salimans, Jeremy Schlatter, Jonas Schneider, Szymon Sidor, Ilya Sutskever, Jie Tang, Filip Wolski, and Susan Zhang. Dota 2 with large scale deep reinforcement learning, 2019a.
- OpenAI, Ilge Akkaya, Marcin Andrychowicz, Maciek Chociej, Mateusz Litwin, Bob McGrew, Arthur Petron, Alex Paino, Matthias Plappert, Glenn Powell, Raphael Ribas, Jonas Schneider, Nikolas Tezak, Jerry Tworek, Peter Welinder, Lilian Weng, Qiming Yuan, Wojciech Zaremba, and Lei Zhang. Solving rubik’s cube with a robot hand, 2019b.
- Feiyang Pan, Jia He, Dandan Tu, and Qing He. Trust the model when it is confident: Masked model-based actor-critic. In H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 10537–10546. Curran Associates, Inc., 2020. URL <https://proceedings.neurips.cc/paper/2020/file/77133be2e96a577bd4794928976d2ae2-Paper.pdf>.
- Ladislav Peska and Peter Vojtas. Off-line vs. on-line evaluation of recommender systems in small e-commerce. In *Proceedings of the 31st ACM Conference on Hypertext and Social Media*, HT ’20, pp. 291–300, New York, NY, USA, 2020. Association for Computing Machinery. ISBN 9781450370981. doi: 10.1145/3372923.3404781. URL <https://doi.org/10.1145/3372923.3404781>.
- John Schulman, Sergey Levine, Pieter Abbeel, Michael Jordan, and Philipp Moritz. Trust region policy optimization. In *International conference on machine learning*, pp. 1889–1897. PMLR, 2015.
- David Silver, Thomas Hubert, Julian Schrittwieser, Ioannis Antonoglou, Matthew Lai, Arthur Guez, Marc Lanctot, Laurent Sifre, Dhharshan Kumaran, Thore Graepel, Timothy Lillicrap, Karen Simonyan, and Demis Hassabis. Mastering chess and shogi by self-play with a general reinforcement learning algorithm, 2017.
- Rishi Veerapaneni, John D. Co-Reyes, Michael Chang, Michael Janner, Chelsea Finn, Jiajun Wu, Joshua Tenenbaum, and Sergey Levine. Entity abstraction in visual model-based reinforcement learning. In Leslie Pack Kaelbling, Danica Kragic, and Komei Sugiura (eds.), *Proceedings of the Conference on Robot Learning*, volume 100 of *Proceedings of Machine Learning Research*, pp. 1439–1456. PMLR, 30 Oct–01 Nov 2020. URL <https://proceedings.mlr.press/v100/veerapaneni20a.html>.

- Oriol Vinyals, I. Babuschkin, W. Czarnecki, Michaël Mathieu, Andrew Dudzik, J. Chung, D. Choi, R. Powell, Timo Ewalds, P. Georgiev, Junhyuk Oh, Dan Horgan, Manuel Kroiss, Ivo Danihelka, Aja Huang, L. Sifre, Trevor Cai, John P. Agapiou, Max Jaderberg, A. S. Vezhnevets, Rémi Leblond, Tobias Pohlen, Valentin Dalibard, D. Budden, Yury Sulsky, James Molloy, T. L. Paine, Caglar Gulcehre, Ziyu Wang, T. Pfaff, Yuhuai Wu, Roman Ring, Dani Yogatama, Dario Wünsch, Katrina McKinney, O. Smith, T. Schaul, T. Lillicrap, K. Kavukcuoglu, Demis Hassabis, Chris Apps, and D. Silver. Grandmaster level in starcraft ii using multi-agent reinforcement learning. *Nature*, pp. 1–5, 2019.
- Tingwu Wang, Xuchan Bao, Ignasi Clavera, Jerrick Hoang, Yeming Wen, Eric Langlois, Shunshi Zhang, Guodong Zhang, Pieter Abbeel, and Jimmy Ba. Benchmarking model-based reinforcement learning. *arXiv preprint arXiv:1907.02057*, 2019.
- Yifan Wu, George Tucker, and Ofir Nachum. Behavior regularized offline reinforcement learning. *arXiv preprint arXiv:1911.11361*, 2019.
- Tianhe Yu, Garrett Thomas, Lantao Yu, Stefano Ermon, James Y Zou, Sergey Levine, Chelsea Finn, and Tengyu Ma. Mopo: Model-based offline policy optimization. *Advances in Neural Information Processing Systems*, 33:14129–14142, 2020.
- Marvin Zhang, Sharad Vikram, Laura Smith, Pieter Abbeel, Matthew Johnson, and Sergey Levine. Solar: Deep structured representations for model-based reinforcement learning. In *International Conference on Machine Learning*, pp. 7444–7453. PMLR, 2019.
- Shangdong Zhang, Bo Liu, and Shimon Whiteson. GradientDICE: Rethinking generalized offline estimation of stationary values. In Hal Daumé III and Aarti Singh (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 11194–11203. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/zhang20r.html>.

A PROOF OF LEMMA 1

Adapted from Schulman et al. (2015); Luo et al. (2018), we let W_j be the expected return when executing $\hat{M}^{(i+1)}$ for the first j steps, and then switch to $\hat{M}^{(i)}$ for the remaining steps.

Proof.

$$W_j = \mathbf{E}_{\substack{a_t \sim \pi(s_t) \\ \forall j > t \geq 0, s_{t+1} \sim \hat{M}^{(i+1)}(\cdot | s_t, a_t) \forall t \geq j, s_{t+1} \sim \hat{M}^{(i)}(\cdot | s_t, a_t)}} \left\{ \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s \right\} \quad (7)$$

Thus, we have $W_0 = V_{\hat{M}^{(i)}}^{\pi}$, and $W_{\infty} = V_{\hat{M}^{(i+1)}}^{\pi}$. Next, we write:

$$V_{\hat{M}^{(i+1)}}^{\pi} - V_{\hat{M}^{(i)}}^{\pi} = \sum_{j=0}^{\infty} (W_{j+1} - W_j) \quad (8)$$

We expand W_j and W_{j+1} so that we can cancel the shared terms:

$$\begin{aligned} W_j &= R_j + \mathbf{E}_{s_j, a_j \sim \pi, \hat{M}^{(i+1)}} \left\{ \mathbf{E}_{s_{j+1} \sim \hat{M}^{(i)}(\cdot | s_t, a_t)} \{ \gamma^{j+1} V_{\hat{M}^{(i)}}^{\pi}(s_{j+1}) \} \right\} \\ W_{j+1} &= R_j + \mathbf{E}_{s_j, a_j \sim \pi, \hat{M}^{(i+1)}} \left\{ \mathbf{E}_{s_{j+1} \sim \hat{M}^{(i+1)}(\cdot | s_t, a_t)} \{ \gamma^{j+1} V_{\hat{M}^{(i)}}^{\pi}(s_{j+1}) \} \right\} \end{aligned}$$

where R_j is the expected return of the first j time step. Next, we cancel the share terms so that:

$$W_{j+1} - W_j = \gamma^{j+1} \mathbf{E}_{s_j, a_j \sim \pi, \hat{M}^{(i+1)}} \left\{ \mathbf{E}_{s' \sim \hat{M}^{(i+1)}(\cdot | s_j, a_j)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \} - \mathbf{E}_{s' \sim \hat{M}^{(i)}(\cdot | s_j, a_j)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \} \right\} \quad (9)$$

Thus, based on eq. (9), we have:

$$V_{\hat{M}^{(i+1)}}^{\pi} - V_{\hat{M}^{(i)}}^{\pi} = \kappa \mathbf{E}_{s_j, a_j \sim \pi, \hat{M}^{(i+1)}} \left\{ \mathbf{E}_{s' \sim \hat{M}^{(i+1)}(\cdot | s_j, a_j)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \} - \mathbf{E}_{s' \sim \hat{M}^{(i)}(\cdot | s_j, a_j)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \} \right\} \quad (10)$$

Let $\nu(s, a) = \mathbf{E}_{s' \sim \hat{M}^{(i+1)}(\cdot | s, a)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \} - \mathbf{E}_{s' \sim \hat{M}^{(i)}(\cdot | s, a)} \{ V_{\hat{M}^{(i)}}^{\pi}(s') \}$, since we have assumed the Lipschitzness of $V_{\hat{M}}^{\pi}$, we can bound $|\nu(s, a)| \leq L |\hat{M}^{(i+1)}(s, a) - \hat{M}^{(i)}(s, a)|$, then combine with triangle inequality, we have:

$$|V_{\hat{M}^{(i+1)}}^{\pi} - V_{\hat{M}^{(i)}}^{\pi}| \leq \kappa L \mathbf{E}_{(s, a) \sim \rho_{\hat{M}^{(i+1)}}^{\pi}} (\|\hat{M}^{(i+1)}(s, a) - \hat{M}^{(i)}(s, a)\|) \quad (11)$$

□

B PROOF OF PROPOSITION 1

Proof. Starting from eqn.10, we substitute with $V_{M^*}^{\pi}$ and $V_{\hat{M}^{(i)}}^{\pi}$, and multiple both side with -1. Due to the Lipschitzness assumption, we can then bound $|\nu(s, a)|$ (with the corresponding substitution) on eqn.10 and by the definition of $U_{\hat{M}^{(i)}, M^*}(s, a)$, thus we have:

$$V_{M^*}^{\pi} \geq V_{\hat{M}^{(i)}}^{\pi} \mathbf{E}_{(s, a) \sim \rho_{M^*}^{\pi}} \{ U_{\hat{M}^{(i)}, M^*}(s, a) \}. \quad (12)$$

□

E EXPLORATION TO COLLECT DATA FROM LOW SUPPORT REGION

In order to build an accurate representation of the world model, during each deployment, we want to collect the data from the low support (or un-visited) regions. we make use of the uncertainty labeler to guide exploration to the un-visited regions for novel data discovery. This is achieved by injecting the $\zeta(a, s)$ as an exploration noise with the zero-mean normal distribution: $\mathcal{N}(0, \sigma = \zeta(a, s))$, where ζ is:

$$\zeta(a, s) = \max_{i \in \{\hat{P}_\phi\}_i^K} (\|s_{t+1} - \hat{s}_{t+1}\|_1) \quad (13)$$

where $\hat{s}_{t+1}$ is the predicted next state and we take the maximum prediction error of the model within $\{\hat{P}_\phi\}_i^K$ ensemble. In the Ablation Study (Section.5), we show that this exploration strategy leads to novel data discovery, and also contribute to better learned model dynamics (Fig.4 subplots c. and d.).

Specifically, the action will be parameterized by a stochastic Gaussian policy (with parameter μ_θ) as:

$$a_t = \tanh(\mu_\theta(s_t)) + \epsilon_{\text{const}} + \epsilon_\zeta \quad (16)$$

where ϵ_{const} and ϵ_ζ are:

$$\begin{aligned} \epsilon_{\text{const}} &\sim \mathcal{N}(\mu = 0, \sigma = 0.01) \\ \epsilon_\zeta &\sim \mathcal{N}\left(\mu = 0, \sigma = \zeta(a = \tanh(\mu_\theta(s_t)) + \epsilon_{\text{const}}, s = s_t)\right) \end{aligned}$$

Here, ϵ_{const} is an additive noise with a constant variance of 0.01, and on top of this, we also added another Gaussian noise with variance equal to $\zeta(a, s)$ from the uncertainty labeler to guide exploration.

F DETAILED IMPLEMENTATION

We used ADAM as the optimizer with a learning rate of 1e-3 for the model dynamics $\hat{T}$ with ensembles size of $N = 5$. For the uncertainty-labeler, we use ensembles of PNN with $K = 3$ and a learning rate of 1e-3. For the behavioral cloning, we used a learning rate of 5e-4. For all the collected data, we divided them into 85% for training, and 15% for validation (for model validation). The $\hat{T}, \hat{P}$ are trained with early stopping when their performance on the validation set no longer improves after consecutive 3 episodes. We used this same setting for the behavioral cloning as well.

Model Architecture For our policy network, we parameterized it by two layers of fully-connected neural network with hidden units of 200. For the $\hat{T}$ model dynamics, we parameterized it by two layers of fully-connected neural network with hidden units of 1024. We used this same configuration with $\hat{P}$ uncertainty-labeler and implemented with PNN.

Training Time The overall training time differs for each environment. We train all models on Nvidia T4 GPU. For the 500k data size experiment, the entire training duration (1 run) for Walker2D and Hopper environments are 18 hours and 26 hours respectively. For the Half-Cheetah, Ant, and Cheetah-Run environment, it is a lot more faster. It takes 7 hours, 12 hours, and 8 hours per run, respectively. For the 250k data size experiments, the training time is about 25 minutes faster than the 500k experiments.

Since we are applying Dyna-style update with neural network based dynamics models, following Wang et al. (2019) Matsushima et al. (2020), we used the following reward functions for our dynamics models (for model-based training only) as:

- Walker2d, Hopper: $\dot{x}_t - 0.001\|a_t\|_2^2 + 1$
- CheetahRun: $\max(0, \min(\dot{x}_t/10, 1))$
- Ant: $\dot{x}_t - 0.1\|a_t\|_2^2 - 3.0(z_t - 0.57)^2 + 1$
- HalfCheetah: $\dot{x}_t - 0.1\|a_t\|_2^2$

We enabled termination in rollouts for the Hopper and Walker2D environments, and disabled that of the Ant, HalfCheetah, and CheetahRun environments (with a maximum step of 1000 for each episode). For the CheetahRun task, we adopt it from the DM control suit³.

³https://github.com/deepmind/dm_control

	L	Rollouts Length	TRPO's δ
Ant	2,000	250	0.05
HalfCheetah	2,000	250	0.1
Hopper	6,000	1,000	0.05
Walker2d	2,000	1,000	0.05
CheetahRun	2,000	250	0.05

Table 2: Hyper-parameters for MUSBO Algorithm

Other Hyper-parameters We searched α over the set of $\{0.28, 0.028, 0.0028\}$ and we used the same $\alpha = 0.028$ (the temperature parameter for eq.(5)) for all environments. Similarly, for the action parameterization eq. (16), we used the same constant $\sigma = 0.01$ (variance term of ϵ_{const}) for all environments. The σ of ϵ_{const} is searched over the set of $\{0.01, 0.05, 0.1\}$.

We searched the rollout length on $\{250, 1000\}$, and the δ on $\{0.01, 0.05, 0.1\}$. We summarized these three parameters (L , Rollouts Length, δ) as above in table 2.

For discount factor γ and GAE λ , we used the same set of hyperparameters as in Wang et al. (2019). Specifically, we used the same $\gamma = 0.99$ for all environment. Also, we used GAE $\lambda = 0.95$ for all environment except for Ant which has a GAE $\lambda = 0.97$.

Hyper-parameters for Baseline. For the BREMEN, we used the exact parameters as Matsushima et al. (2020). For MOPO, we adapted it to the deployment setting. We used the latest learned policy for deployment, and then launched it to collect data batch of size $|B|$. For the CheetahRun task, we used the same set of parameters of HalfCheetah. On all environments, we tried hyper-parameters search on the ranges as originally proposed by the paper Yu et al. (2020), we didn't find any improvement over the same set of parameters as originally proposed. Thus, we used the same set of parameters as originally proposed.

Training of the Probabilistic Neural Networks (PNN). In PNN (Chua et al., 2018), the network has its output neurons parameterized by a Gaussian distribution in the effort of capturing the uncertainty. This module is trained to minimize the following loss:

$$\text{loss}_{PNN}(\phi) = \sum_{n=1}^K \{ \mu_{\phi}(s_n, a_n) - s_{n+1} \}^{\top} \Sigma_{\phi}^{-1}(s_n, a_n) (\mu_{\phi}(s_n, a_n) - s_{n+1}) + \log \det \Sigma_{\phi}(s_n, a_n) \quad (17)$$

where ϕ is the neural network learning parameters, μ_{ϕ} and Σ_{ϕ} are the mean and variance of the Gaussian distribution.

G ABLATION: HOW GOOD IS THE APPROXIMATED $\hat{U}(s, a)$ IN ESTIMATING THE ACTUAL $U_{\hat{M}^{(i)}, M^*}(s, a)$?

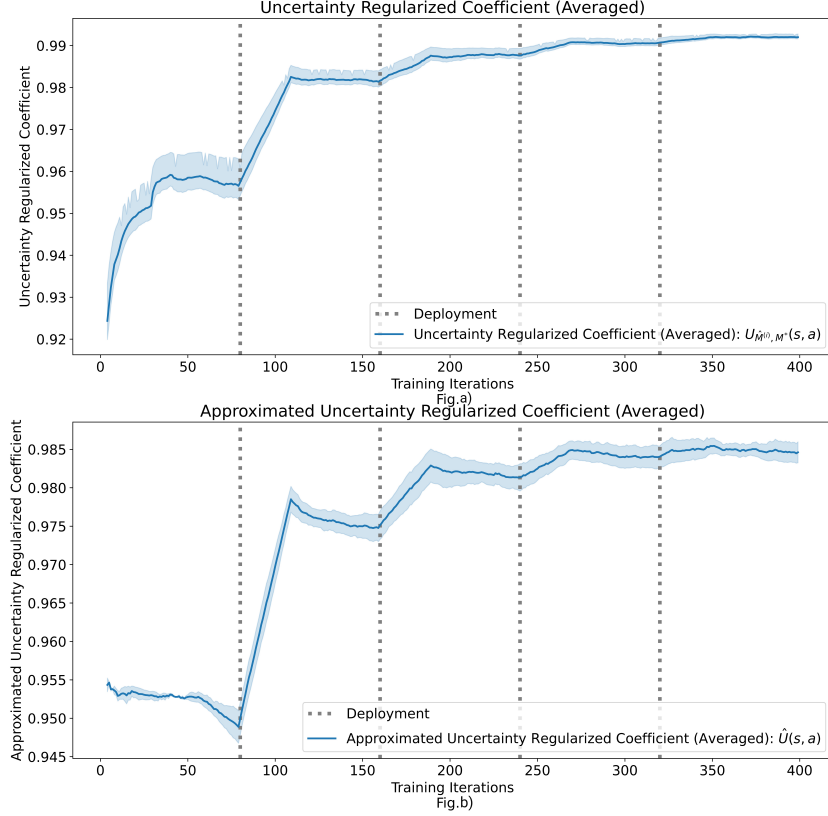


Figure 5: Ablation study on $I \times |B| = 500k$ setting for the ant environment. we plot the approximated $\hat{U}(s, a)$ on the bottom graph and we plot the true $U_{\hat{M}^{(i)}, M^*}(s, a)$ on the top graph. The approximated $\hat{U}(s, a)$ is very good at estimating the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$.

In this section, we study the empirical values of the approximated $\hat{U}(s, a)$ and the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$ to verify how good our approximation in estimating the real one. The reason for not training with the exact definition of the $U_{\hat{M}^{(i)}, M^*}(s, a)$ directly (but instead approximating it with model estimation error) is because of the instability effect that we observed empirically. Having the value function at the denominator can sometimes cause instability issues especially at the early training stage. On Fig.5, we show the empirical values of the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$ at the top (subplot a) and we show the approximated $\hat{U}(s, a)$ (subplot b) at the bottom. As is shown on the plot, the estimated $\hat{U}(s, a)$ is very good in approximating the actual $U_{\hat{M}^{(i)}, M^*}(s, a)$.