



Efficient Off-Policy Meta-Reinforcement Learning via Probabilistic Context Variables

Kate Rakelly^{1*} Aurick Zhou^{1*} Deirdre Quillen¹ Chelsea Finn¹ Sergey Levine¹

ICML 2019

Introduction

❑ On-policy learning

Only one policy used throughout the system to both explore and select actions.

sample inefficiency

eg: policy gradient

❑ Off-policy learning

Two policies, one for exploring and the other for action selection.

eg: DQN

Introduction

□ Meta-Reinforcement Learning

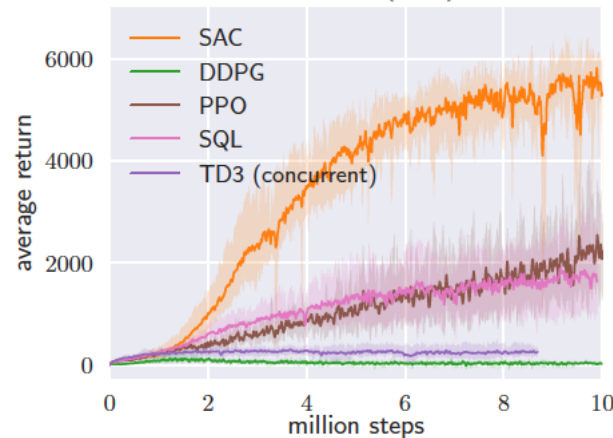
The agent can leverage varied experiences from previous tasks to adapt quickly to the new task at hand.



https://blog.csdn.net/sinat_37422398

Motivation

- Most meta-learning RL systems use on-policy learning. The general problem with on-policy learning is sample inefficiency.

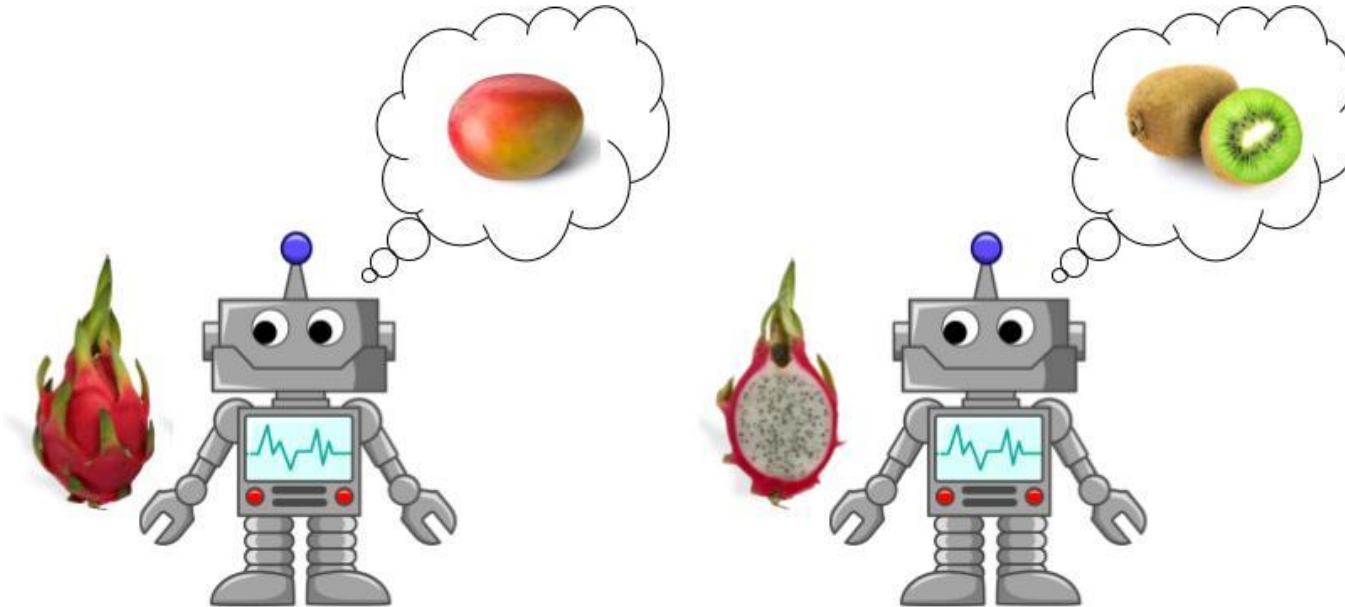


Off-policy RL (SAC) is more sample efficient than on-policy approaches (PPO) by 1-2 orders of magnitude (PPO eventually nears SAC performance), figure from Haarnoja et al. 2018.

Tackle the problem of **efficient off-policy meta-reinforcement learning**.

Method

- ❑ Meta-training : learn a probabilistic **encoder** that accumulates the necessary statistics from past experience into the **context variables**.
- ❑ Meta-testing : the context variables can be sampled and held constant for the duration of an episode, enabling temporally-extended exploration.



Method-Context Variable

A task : $\mathcal{T} = \{p(\mathbf{s}_0), p(\mathbf{s}_{t+1}|\mathbf{s}_t, \mathbf{a}_t), r(\mathbf{s}_t, \mathbf{a}_t)\}$

$\mathbf{c}_n^{\mathcal{T}} = (\mathbf{s}_n, \mathbf{a}_n, r_n, \mathbf{s}'_n)$: one transition in the task $\mathcal{T}$

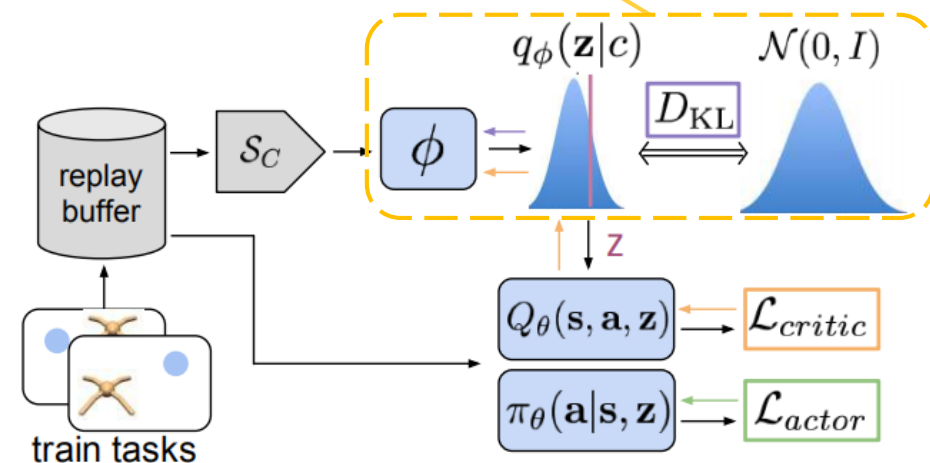
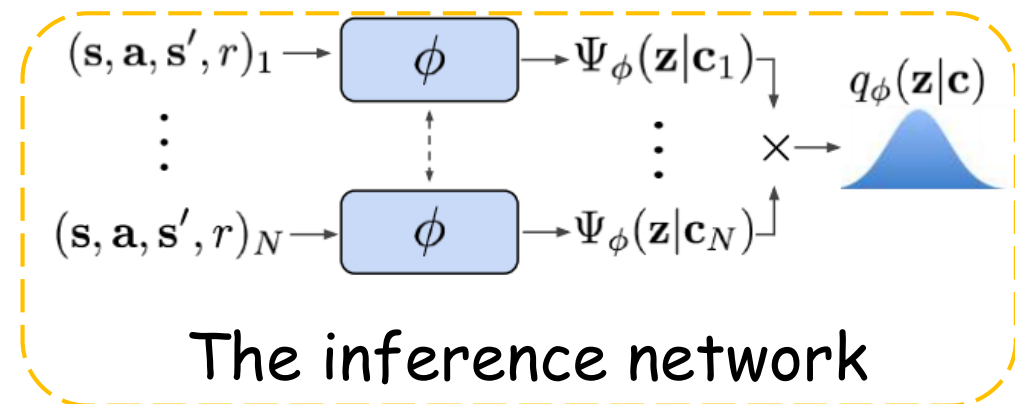
$\mathbf{z}$: latent probabilistic context variable

Gaussian factors

$$\Psi_{\phi}(\mathbf{z}|\mathbf{c}_n) = \mathcal{N}(f_{\phi}^{\mu}(\mathbf{c}_n), f_{\phi}^{\sigma}(\mathbf{c}_n))$$

$$q_{\phi}(\mathbf{z}|\mathbf{c}_{1:N}) \propto \prod_{n=1}^N \Psi_{\phi}(\mathbf{z}|\mathbf{c}_n)$$

permutation-invariant function of
prior experience



Method-Meta-training

Algorithm 1 PEARL Meta-training

Require: Batch of training tasks $\{\mathcal{T}_i\}_{i=1\dots T}$ from $p(\mathcal{T})$,

learning rates $\alpha_1, \alpha_2, \alpha_3$

1: Initialize replay buffers $\mathcal{B}^i$ for each training task

2: **while** not done **do**

3: **for** each $\mathcal{T}_i$ **do**

4: Initialize context $\mathbf{c}^i = \{\}$

5: **for** $k = 1, \dots, K$ **do**

6: Sample $\mathbf{z} \sim q_\phi(\mathbf{z}|\mathbf{c}^i)$

7: Gather data from $\pi_\theta(\mathbf{a}|\mathbf{s}, \mathbf{z})$ and add to $\mathcal{B}^i$

8: Update $\mathbf{c}^i = \{(\mathbf{s}_j, \mathbf{a}_j, \mathbf{s}'_j, r_j)\}_{j:1\dots N} \sim \mathcal{B}^i$

9: **end for**

10: **end for**

11: **for** step in training steps **do**

12: **for** each $\mathcal{T}_i$ **do**

13: Sample context $\mathbf{c}^i \sim \mathcal{S}_c(\mathcal{B}^i)$ and RL batch $b^i \sim \mathcal{B}^i$

14: Sample $\mathbf{z} \sim q_\phi(\mathbf{z}|\mathbf{c}^i)$

15: $\mathcal{L}_{actor}^i = \mathcal{L}_{actor}(b^i, \mathbf{z})$

16: $\mathcal{L}_{critic}^i = \mathcal{L}_{critic}(b^i, \mathbf{z})$

17: $\mathcal{L}_{KL}^i = \beta D_{KL}(q(\mathbf{z}|\mathbf{c}^i) || r(\mathbf{z}))$

18: **end for**

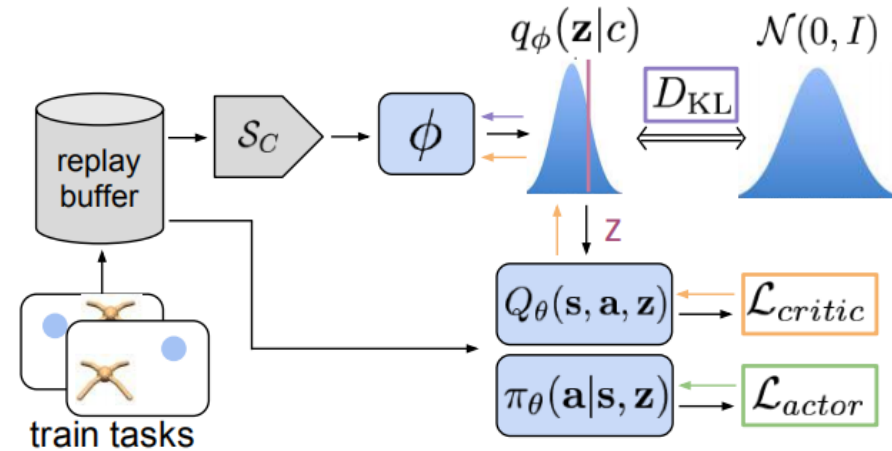
19: $\phi \leftarrow \phi - \alpha_1 \nabla_\phi \sum_i (\mathcal{L}_{critic}^i + \mathcal{L}_{KL}^i)$

20: $\theta_\pi \leftarrow \theta_\pi - \alpha_2 \nabla_\theta \sum_i \mathcal{L}_{actor}^i$

21: $\theta_Q \leftarrow \theta_Q - \alpha_3 \nabla_\theta \sum_i \mathcal{L}_{critic}^i$

22: **end for**

23: **end while**



Meta-training procedure

$$\mathcal{L}_{actor} = \mathbb{E}_{\mathbf{s} \sim \mathcal{B}, \mathbf{a} \sim \pi_\theta, \mathbf{z} \sim q_\phi(\mathbf{z}|\mathbf{c})} \left[D_{KL} \left(\pi_\theta(\mathbf{a}|\mathbf{s}, \bar{\mathbf{z}}) \left\| \frac{\exp(Q_\theta(\mathbf{s}, \mathbf{a}, \bar{\mathbf{z}}))}{Z_\theta(\mathbf{s})} \right\| \right) \right]$$

$$\mathcal{L}_{critic} = \mathbb{E}_{(\mathbf{s}, \mathbf{a}, r, \mathbf{s}') \sim \mathcal{B}, \mathbf{z} \sim q_\phi(\mathbf{z}|\mathbf{c})} [Q_\theta(\mathbf{s}, \mathbf{a}, \mathbf{z}) - (r + \bar{V}(\mathbf{s}', \bar{\mathbf{z}}))]^2$$

Method-Meta-testing/Adaptation

Algorithm 2 PEARL Meta-testing

Require: test task $\mathcal{T} \sim p(\mathcal{T})$

- 1: Initialize context $\mathbf{c}^{\mathcal{T}} = \{\}$
 - 2: **for** $k = 1, \dots, K$ **do**
 - 3: Sample $z \sim q_{\phi}(\mathbf{z}|\mathbf{c}^{\mathcal{T}})$
 - 4: Roll out policy $\pi_{\theta}(\mathbf{a}|\mathbf{s}, \mathbf{z})$ to collect data $D_k^{\mathcal{T}} = \{(\mathbf{s}_j, \mathbf{a}_j, \mathbf{s}'_j, r_j)\}_{j:1\dots N}$
 - 5: Accumulate context $\mathbf{c}^{\mathcal{T}} = \mathbf{c}^{\mathcal{T}} \cup D_k^{\mathcal{T}}$
 - 6: **end for**
-

Experiments-Sample Efficiency and Performance

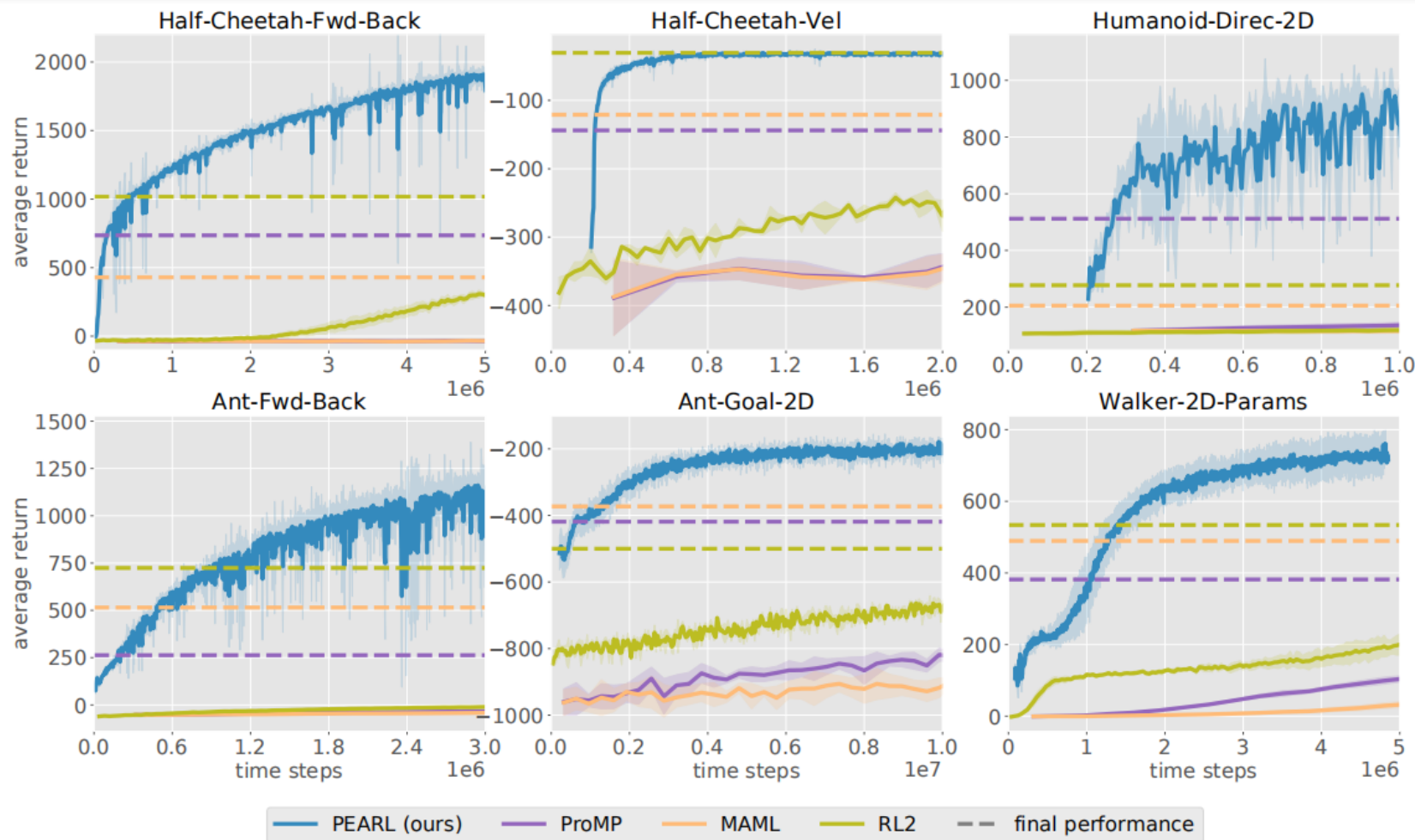


Figure 3. **Meta-learning continuous control.** Test-task performance vs. samples collected during *meta-training*. Our approach PEARL outperforms previous meta-RL methods both in terms of asymptotic performance and meta-training sample efficiency across six benchmark tasks. **Dashed lines correspond to the maximum return achieved by each baseline after 1e8 steps.** By leveraging off-policy data during meta-training, PEARL is 20-100x more sample efficient than the baselines, and achieves consistently better or equal final performance compared to the best performing prior method in each environment. See Appendix A for the full timescale version of this plot.

Experiments-Posterior Sampling For Exploration

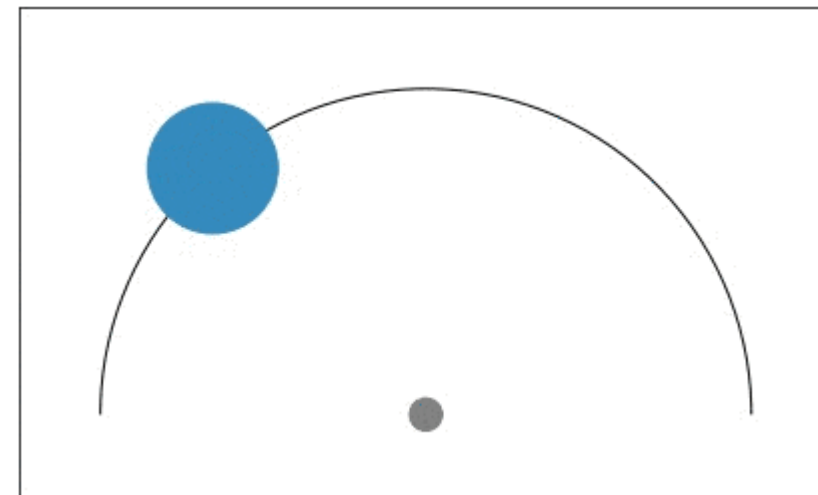
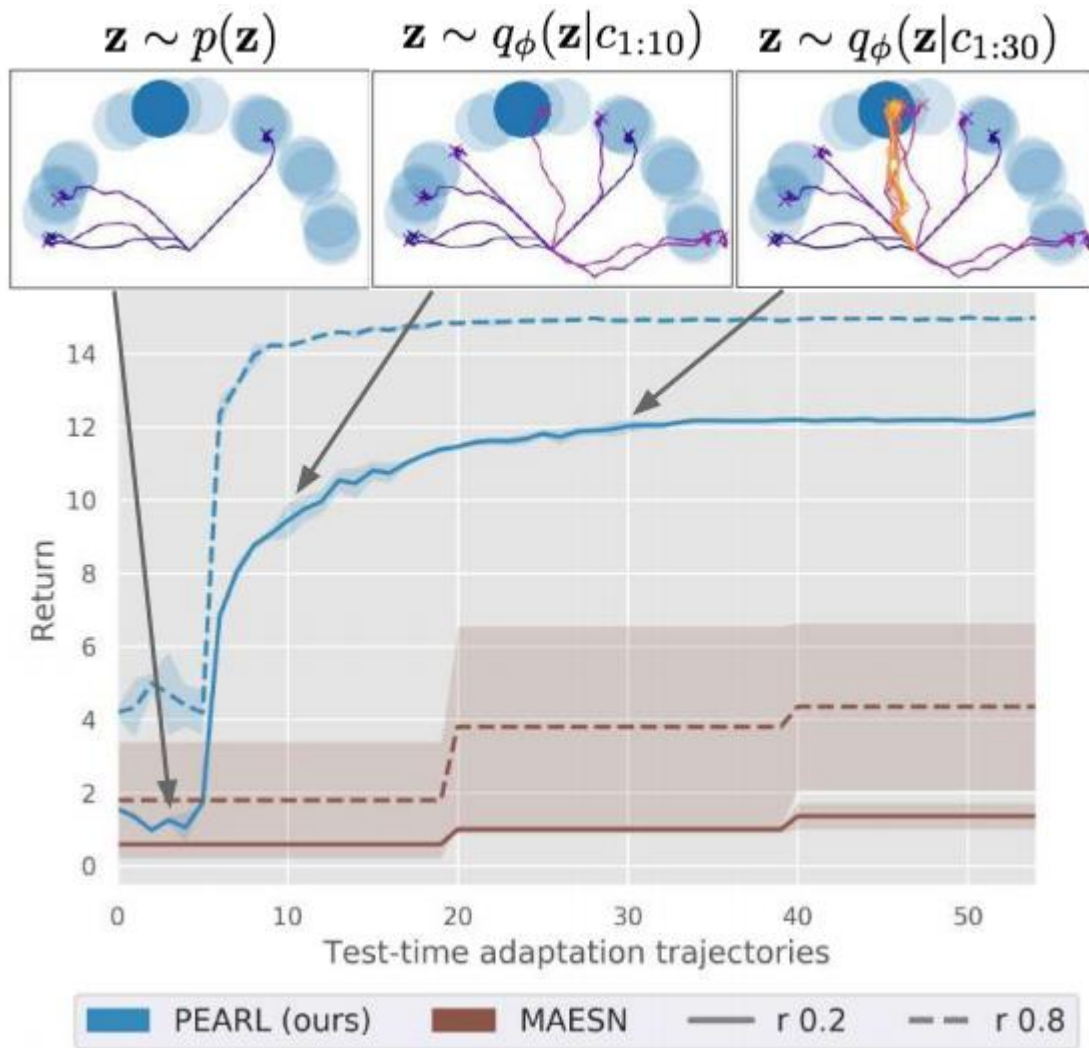
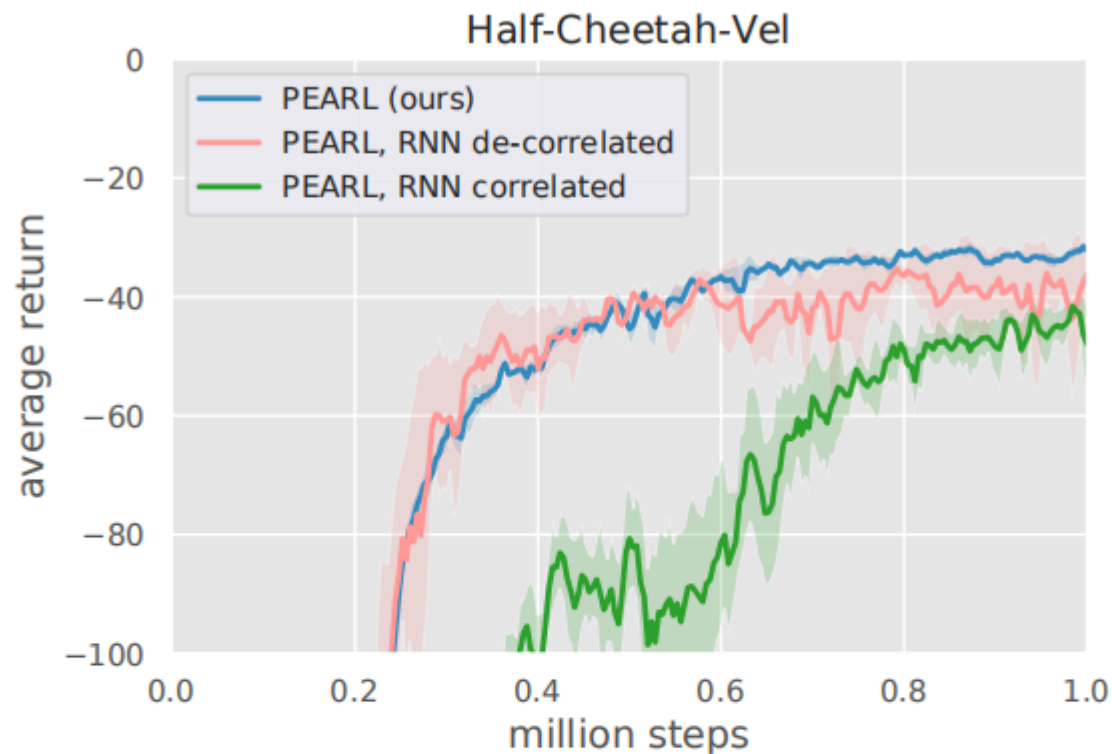
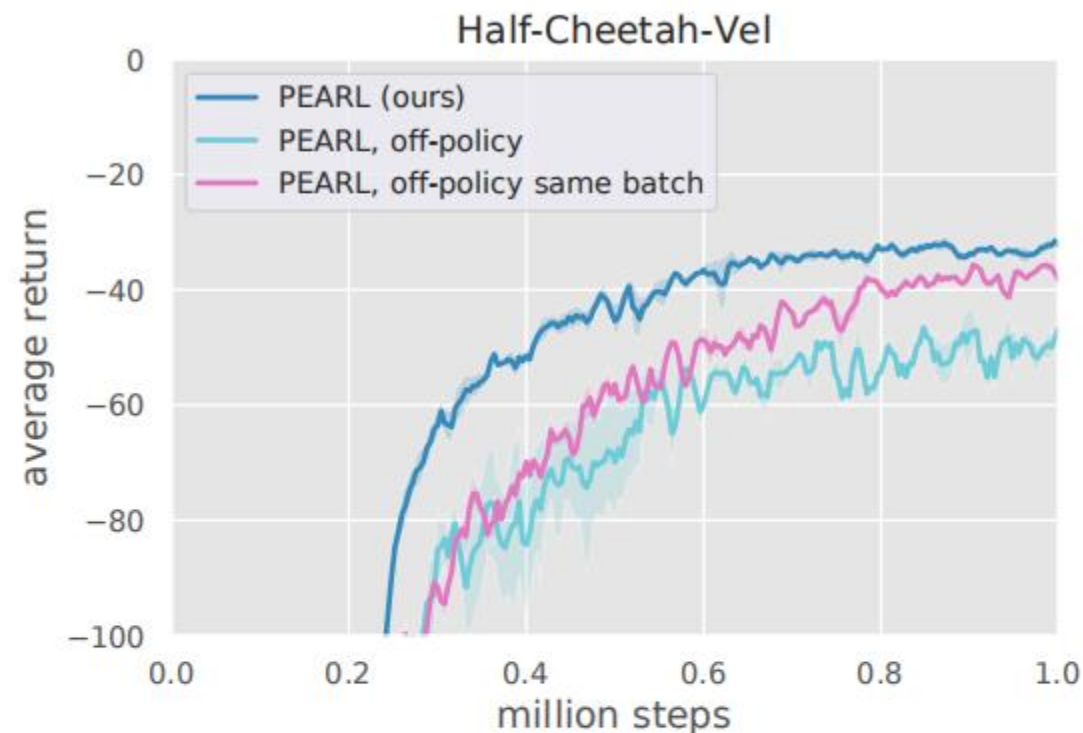


Figure 4. Sparse 2D navigation. The agent must navigate to a previously unseen goal (dark blue, other test goals in light blue) with reward given only when inside the goal radius – radius of 0.2 (illustrated) and 0.8 are tested here. The agent is trained to navigate to a training set of goals, then tested on a distinct set of unseen test goals. By using posterior sampling to explore efficiently, PEARL is able to start adapting to the task after collecting on average only 5 trajectories, outperforming MAESN (Gupta et al., 2018).

Experiments-Ablations

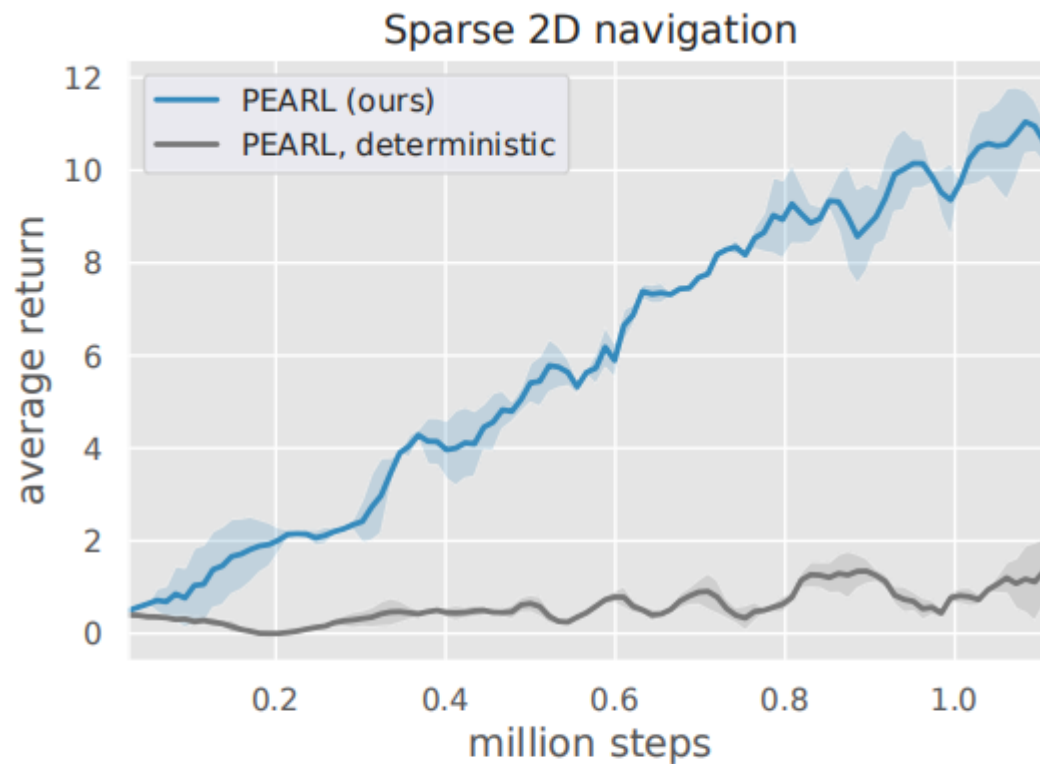


*Figure 5. **Recurrent encoder ablation.*** We compare our encoder architecture to a recurrent network. For the RNN, we sample context as trajectories rather than unordered transitions. Sampling the actor-critic batch as de-correlated transitions (“RNN de-correlated”) fares much better than sampling trajectories (“RNN correlated”). In addition to better performance, our encoder has a computational advantage.



*Figure 6. **Context sampling ablation.*** PEARL samples context batches of recently collected transitions de-correlated with the batches sampled for training the actor-critic. We compare to sampling context from the entire history (“off-policy context”), as well as using the same sampled batch for the context and the actor-critic batch (“off-policy same batch”).

Experiments-Ablations



*Figure 7. **Deterministic latent context.*** We compare PEARL to a variant with deterministic latent context on the sparse reward 2D navigation domain. As expected, without a mechanism for reasoning about uncertainty over tasks, this approach is unable to explore effectively and performs poorly.



Meta-Q-Learning

Rasool Fakoor¹, Pratik Chaudhari^{2*}, Stefano Soatto¹, Alexander Smola¹

¹ Amazon Web Services

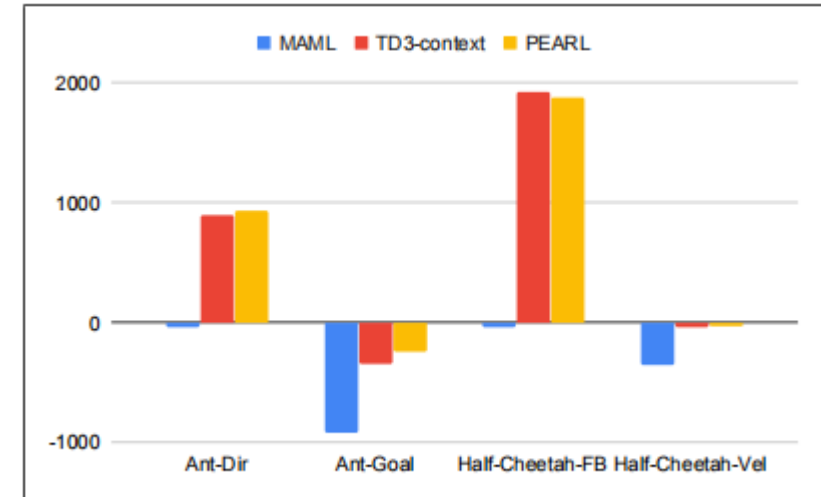
² University of Pennsylvania

Email: {fakoor, soattos, smola}@amazon.com, pratikac@seas.upenn.edu

ICLR 2020

Contribution

- ❑ Q-learning is **competitive** with state-of-the-art meta-RL algorithms if given access to a **context variable** that is a representation of the past trajectory.
- ❑ A **multi-task objective** to maximize the average reward across the training tasks is an effective method to meta-train RL policies.
- ❑ Past data from the meta-training replay buffer can be **recycled to adapt** the policy on a new task using off-policy updates.



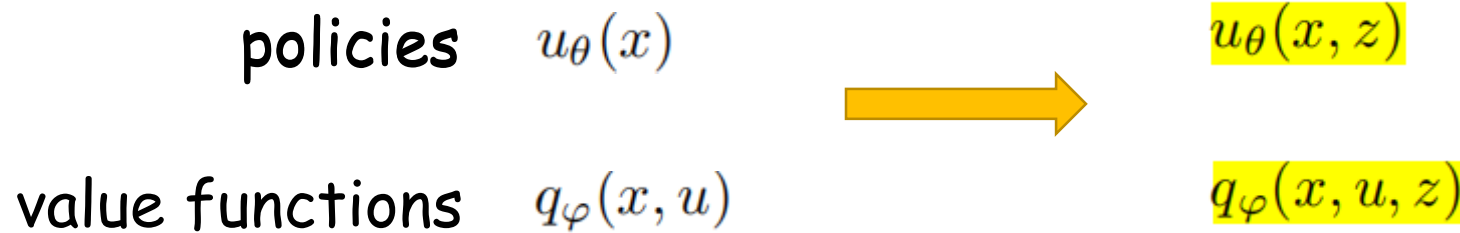
$$\hat{\theta}_{\text{meta}} = \arg \max_{\theta} \frac{1}{n} \sum_{k=1}^n \mathbb{E}_{\tau \sim D^k} [\ell^k(\theta)]$$

Method-Context Variable

□ recurrent context variable z_t

Set z_t to the hidden state at time t of a Gated Recurrent Unit (GRU) model.

z_t depends on $\{(x_i, u_i, r_i)\}_{i < t}$.



Method-Meta-training

Algorithm 1: MQL - Meta-training

Input: Set of training tasks $\mathcal{D}_{\text{meta}}$

- 1 Initialize the replay buffer
- 2 Initialize parameters θ of an off-policy method, e.g., TD3
- 3 **while** *not done* **do**
 - 4 // Rollout and update policy
 - 5 Sample a task $D \sim \mathcal{D}_{\text{meta}}$
 - 6 Gather data from task D using policy π_{θ} while feeding transitions through context GRU. Add trajectory to the replay buffer.
 - 7 $\mathcal{B} \leftarrow$ Sample mini-batch from buffer
 - 8 Update parameters θ using mini-batch $\mathcal{B}$ and Eqn. (15)
- 9 $\theta_{\text{meta}} \leftarrow \theta$
- 10 **return** θ_{meta} , replay buffer

$$\hat{\theta}_{\text{meta}} = \arg \min_{\theta} \frac{1}{n} \sum_{k=1}^n \mathbb{E}_{\tau \sim D^k} [\text{TD}^2(\theta)];$$

Method-Meta-training

□ Objective

$$\hat{\theta}_{\text{meta}} = \arg \max_{\theta} \frac{1}{n} \sum_{k=1}^n \mathbb{E}_{\tau \sim D^k} [\ell^k(\theta)] \xrightarrow{\text{TD3}} \hat{\theta}_{\text{meta}} = \arg \min_{\theta} \frac{1}{n} \sum_{k=1}^n \mathbb{E}_{\tau \sim D^k} [\text{TD}^2(\theta)];$$
$$\text{TD}^2(\theta) = (q_{\varphi}(x, u) - r^k - \gamma q_{\varphi}(x', u_{\theta}(x')))^2$$

□ Analysis

MAML : $\ell_{\text{meta}}^k(\theta) = \ell^k(\theta + \alpha \nabla_{\theta} \ell^k(\theta))$

Taylor series expansion & Gradient:

$$\nabla \ell_{\text{meta}}^k(\theta) = \nabla \ell^k(\theta) + 2\alpha(m-1) (\nabla^2 \ell^k(\theta)) \nabla \ell^k(\theta) + \mathcal{O}(\alpha^2).$$



gradient

$$\ell^k(\theta) + \alpha(m-1) \|\nabla \ell^k(\theta)\|_2^2$$

Method-Meta-testing/Adaptation

- Update the policy using the new data

$$\arg \max_{\theta} \left\{ \mathbb{E}_{\tau \sim D^{\text{new}}} [\ell^{\text{new}}(\theta)] - \frac{\lambda}{2} \|\theta - \hat{\theta}_{\text{meta}}\|_2^2 \right\}. \quad (18)$$

- Exploits the meta-training replay buffer

$$\arg \max_{\theta} \left\{ \mathbb{E}_{\tau \sim \mathcal{D}_{\text{meta}}} [\beta(\tau; D^{\text{new}}, \mathcal{D}_{\text{meta}}) \ell^{\text{new}}(\theta)] - \frac{\lambda}{2} \|\theta - \hat{\theta}_{\text{meta}}\|_2^2 \right\}. \quad (19)$$

$$\beta(\tau; D^{\text{new}}, \mathcal{D}_{\text{meta}}) = \frac{P(\tau \in D^{\text{new}})}{P(\tau \in D_{\text{meta}})}$$

- Effective Sample Size

Normalized Effective Sample Size ($\widehat{\text{ESS}}$): A related quantity to $\beta(x)$ is the normalized Effective Sample Size ($\widehat{\text{ESS}}$) which we define as the relative number of samples from the target distribution $p(x)$ required to obtain an estimator with performance (say, variance) equal to that of the importance sampling estimator (10). It is not possible to compute the ESS without knowing both densities

$$\widehat{\text{ESS}} = \frac{1}{m} \frac{(\sum_{k=1}^m \beta(x_k))^2}{\sum_{k=1}^m \beta(x_k)^2} \in [0, 1]$$
$$\lambda = 1 - \widehat{\text{ESS}}$$

$$\arg \max_{\theta} \left\{ \mathbb{E}_{\tau \sim D^{\text{new}}} [\ell^{\text{new}}(\theta)] + \mathbb{E}_{\tau \sim \mathcal{D}_{\text{meta}}} [\beta(\tau; D^{\text{new}}, \mathcal{D}_{\text{meta}}) \ell^{\text{new}}(\theta)] - \left(1 - \widehat{\text{ESS}}\right) \|\theta - \hat{\theta}_{\text{meta}}\|_2^2 \right\}$$

Method-Meta-testing/Adaptation

Algorithm 2: MQL - Adaptation

Input: Test task D , meta-training replay buffer, meta-trained policy θ_{meta}

- 1 Initialize temporary buffer buf
- 2 $\theta \leftarrow \theta_{\text{meta}}$
- 3 buf $\leftarrow$ Gather data from D using $\pi_{\theta_{\text{meta}}}$
- 4 Update Eqn. (18) using buf
- 5 Fit $\beta(D)$ using buf and meta-training replay buffer using Eqn. (12)
- 6 Estimate $\widehat{\text{ESS}}$ using $\beta(D)$ using Eqn. (13)
- 7 **for** $i \leq n$ **do**
- 8 $\bar{\theta} \leftarrow$ sample mini-batch from meta-training replay buffer
- 9 Calculate β for $\bar{\theta}$
- 10 Update θ using Eqn. (19)
- 11 Evaluate θ on a new rollout from task D
- 12 **return** θ

$$\arg \max_{\theta} \left\{ \mathbb{E}_{\tau \sim D^{\text{new}}} [\ell^{\text{new}}(\theta)] - \frac{\lambda}{2} \|\theta - \hat{\theta}_{\text{meta}}\|_2^2 \right\}. \quad (18)$$

$$\beta(\tau; D^{\text{new}}, \mathcal{D}_{\text{meta}}) = \frac{P(\tau \in D^{\text{new}})}{P(\tau \in \mathcal{D}_{\text{meta}})}$$

$$\widehat{\text{ESS}} = \frac{1}{m} \frac{(\sum_{k=1}^m \beta(x_k))^2}{\sum_{k=1}^m \beta(x_k)^2} \in [0, 1] \quad (13)$$

$$\arg \max_{\theta} \left\{ \mathbb{E}_{\tau \sim \mathcal{D}_{\text{meta}}} [\beta(\tau; D^{\text{new}}, \mathcal{D}_{\text{meta}}) \ell^{\text{new}}(\theta)] - \frac{\lambda}{2} \|\theta - \hat{\theta}_{\text{meta}}\|_2^2 \right\}. \quad (19)$$

Experiments

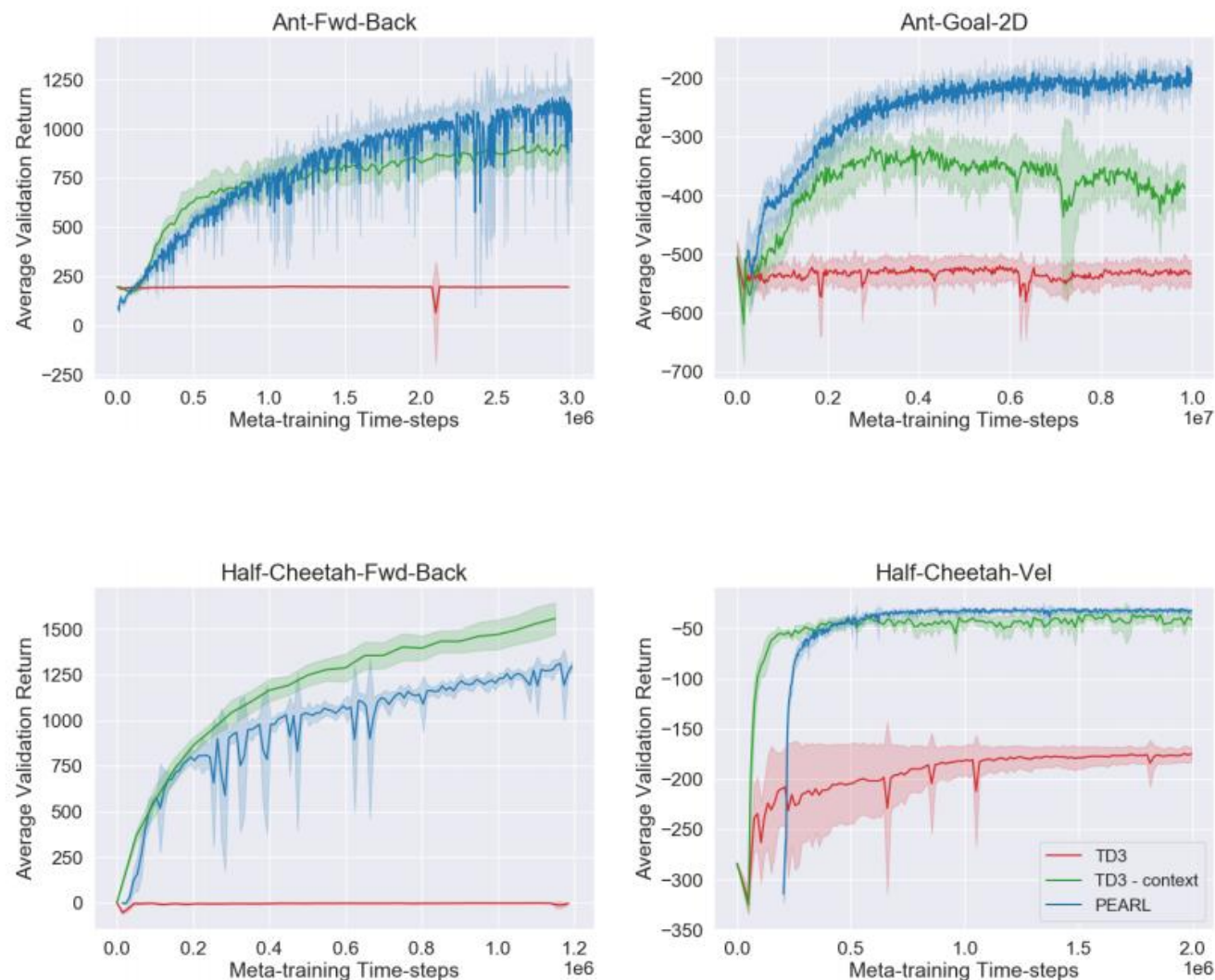


Figure 2: **Average undiscounted return of TD3 and TD3-context compared with PEARL for validation tasks from four meta-RL environments.** The agent fails to learn if the policy is conditioned only on the state. In contrast, everything else remaining same, if TD3 is provided access to context, the rewards are much higher. In spite of not adapting on the validation tasks, TD3-context is comparable to PEARL.

Experiments

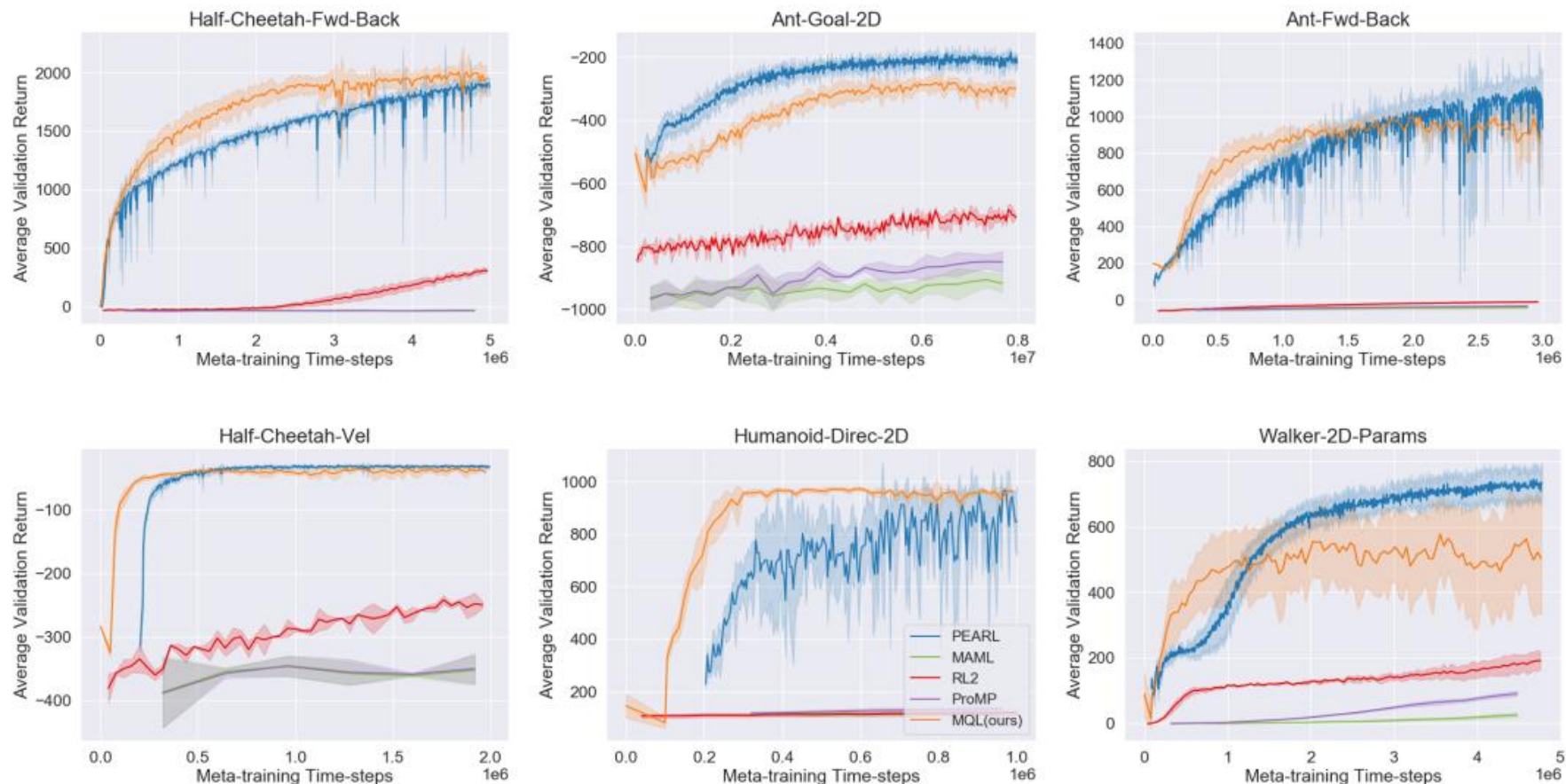


Figure 3: **Comparison of the average undiscounted return of MQL (orange) against existing meta-RL algorithms on continuous-control environments.** We compare against four existing algorithms, namely MAML (green), RL2 (red), PROMP (purple) and PEARL (blue). In all environments except Walker-2D-Params and Ant-Goal-2D, MQL is better or comparable to existing algorithms in terms of both sample complexity and final returns.

Experiments-Ablations

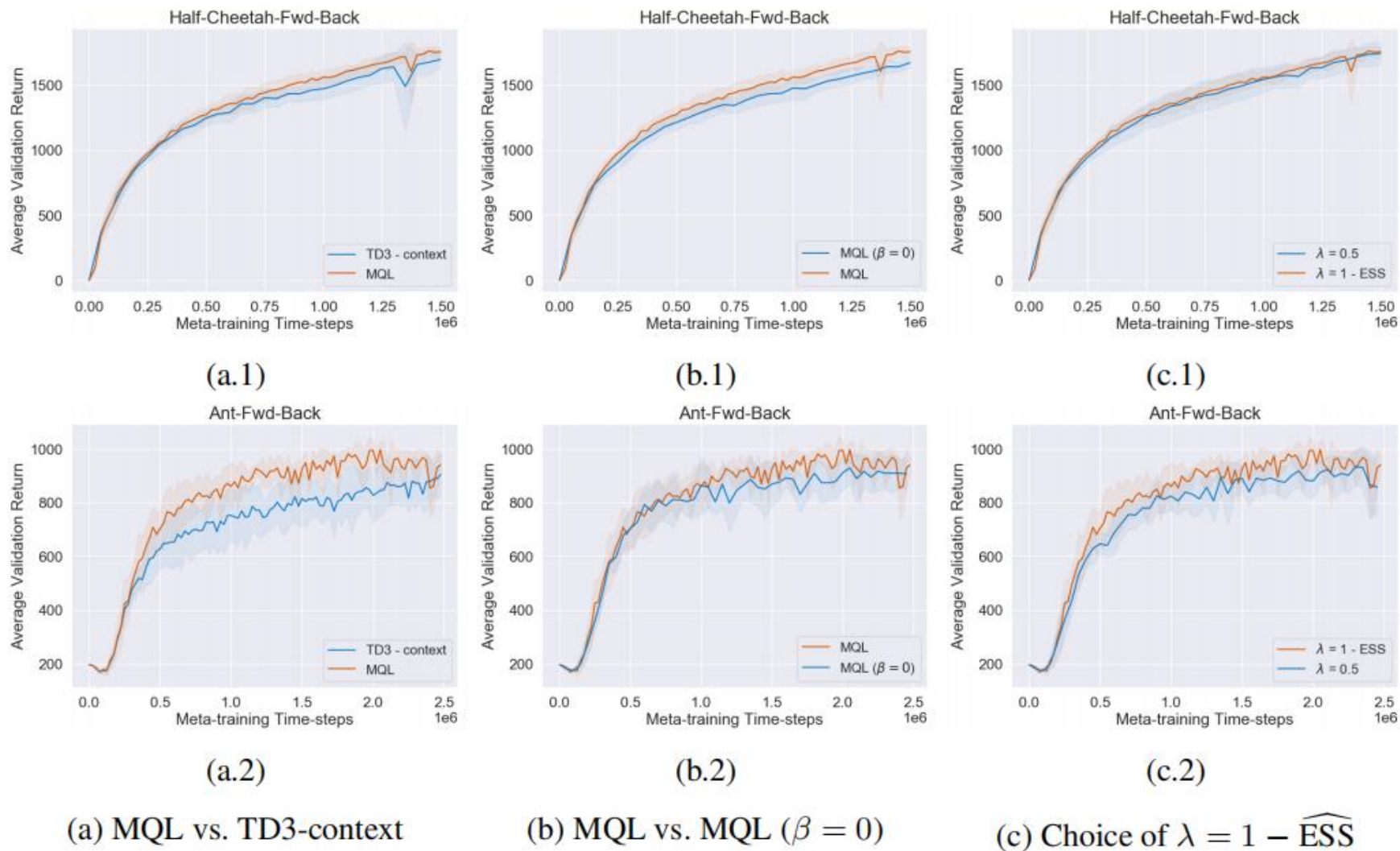


Figure 4: Ablation studies to examine various components of MQL.

Thanks
