

Accelerating Reinforcement Learning with Learned Skill Priors

Feng Xu
2021/1/19

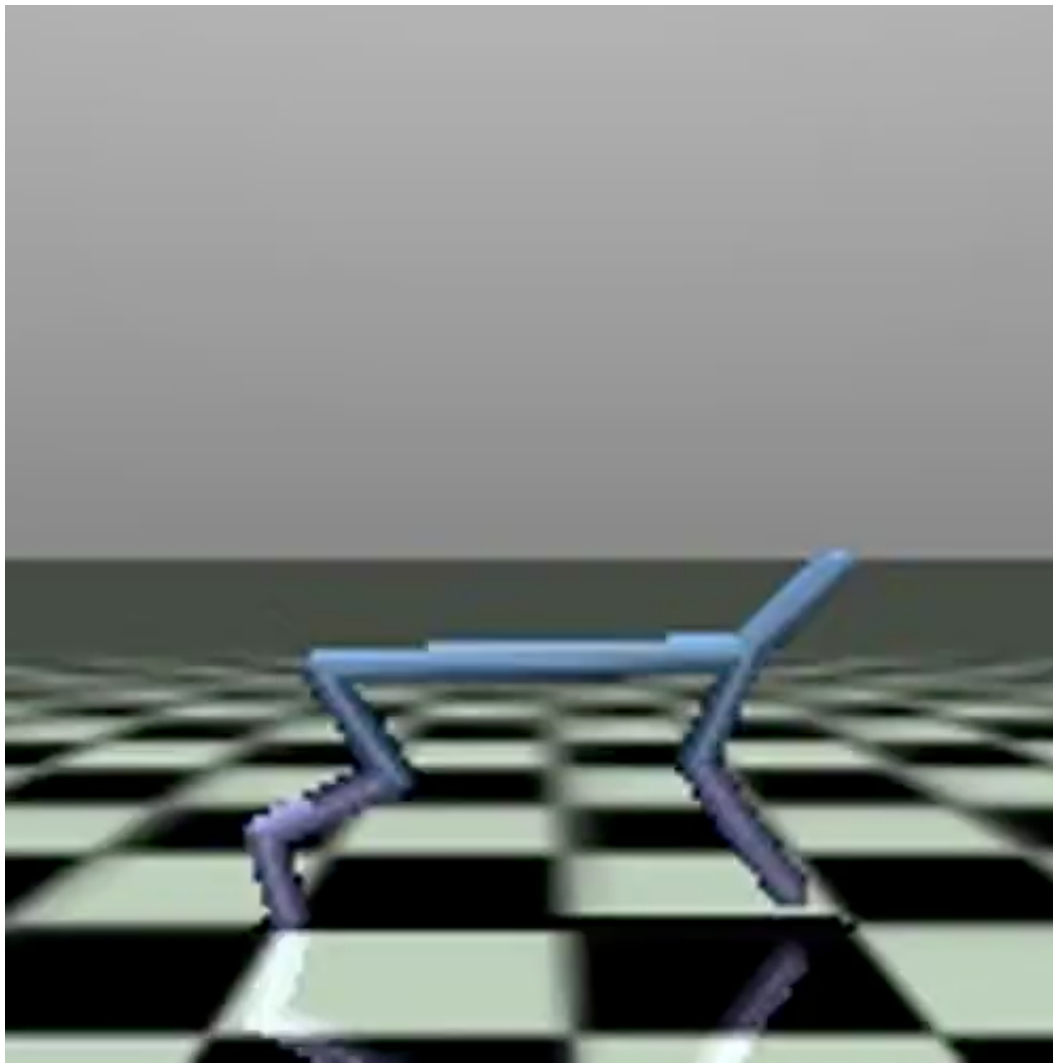


Outline

- Background
- Algorithm
- Experiment

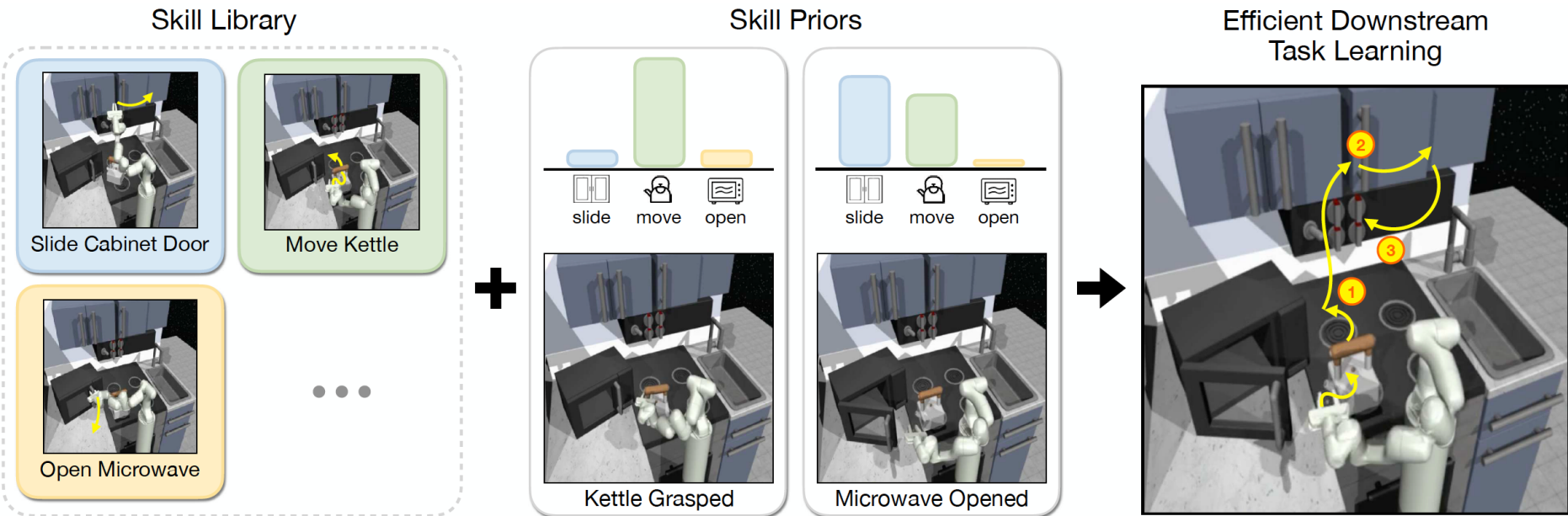
- Learning to Learn
 - Leverage prior experience
- Previous methods:
 - Black Box based: Pearl, RL²
 - Optimization based: MAML and its variants
- Instant adaptation

What's wrong with conventional RL algorithms



Motivation

- Leverage prior experience to accelerate the learning process of downstream tasks



Prior-Guided Skill Transfer

- Extract skill embedding and prior from offline data
 - Learning continuous skill embedding and skill prior
- Prior guided learning of downstream tasks with a hierarchical policy
 - Skill prior regularized reinforcement learning

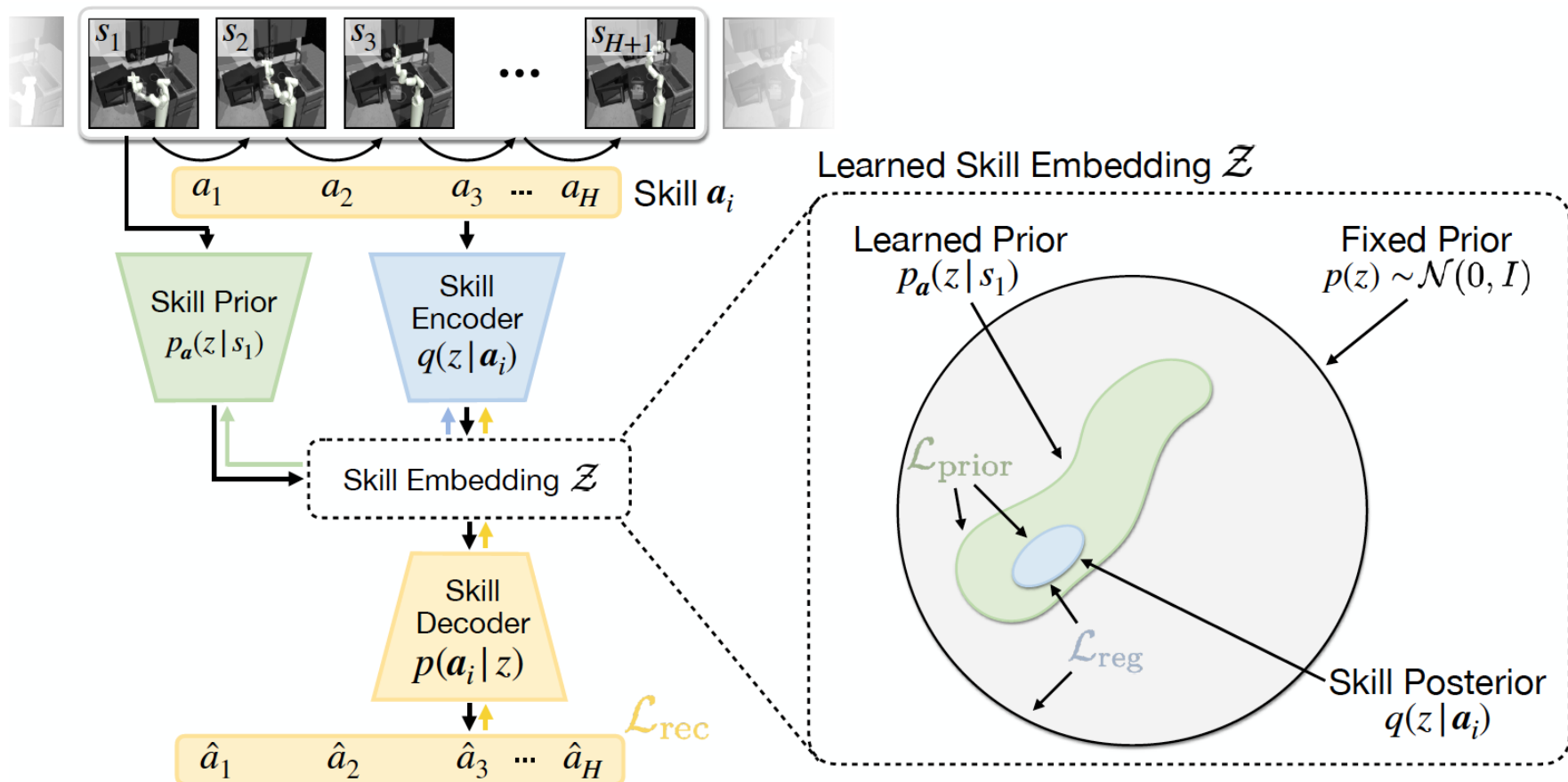
- MDP $\{S, \mathcal{A}, \mathcal{T}, R, \rho, \gamma\}$
- Pre-recorded experience of trajectories

$$\tau_i = \{(s_0, a_0), \dots, (s_{T_i}, a_{T_i})\}.$$

- Learn a policy that maximize expected return of T horizon

$$J(\theta) = \mathbb{E}_{\pi} \left[\sum_{t=0}^{T-1} \gamma^t r_t \right]$$

Overall Structure



- A skill is defined as a sequence of actions with fixed horizon H $\{a_t^i, \dots, a_{t+H-1}^i\}$
- Optimization objection(ELBO)

$$\log p(\mathbf{a}_i) \geq \mathbb{E}_q \left[\underbrace{\log p(\mathbf{a}_i|z)}_{\text{reconstruction}} - \beta \underbrace{(\log q(z|\mathbf{a}_i) - \log p(z))}_{\text{regularization}} \right]$$

- Output: skill encoder $q(z|\mathbf{a}_i)$ and decoder $p(\mathbf{a}_i|z)$

- To adjust to environment and task -> skill prior

$$p_{\mathbf{a}}(z|\cdot) \quad p_{\mathbf{a}}(z|s_t)$$

- Optimization objective:
 - $\max \mathbb{E}_{(s, \mathbf{a}_i) \sim \mathcal{D}} D_{\text{KL}}(q(z|\mathbf{a}_i), p_{\mathbf{a}}(z|s_t))$

- A hierarchical policy learning scheme:
 - Skill embedding space as action space
 - Policy $\pi_{\theta}(z|s_t) \rightarrow \{a_t^i, \dots, a_{t+H-1}^i\} \sim p(\mathbf{a}_i|z)$
- Equivalent to solving a MDP: $\{\mathcal{S}, \mathcal{A}, \mathcal{T}, R, \rho, \gamma\}$
 - Replace the action space with skill space
 - H-step rewards
 - H-step transitions

Guided Learning: Incorporate with Maximum Entropy RL

- Naïve implementation

$$J(\theta) = \mathbb{E}_{\pi} \left[\sum_{t=1}^T \gamma^t r(s_t, a_t) + \alpha \mathcal{H}(\pi(a_t|s_t)) \right]$$

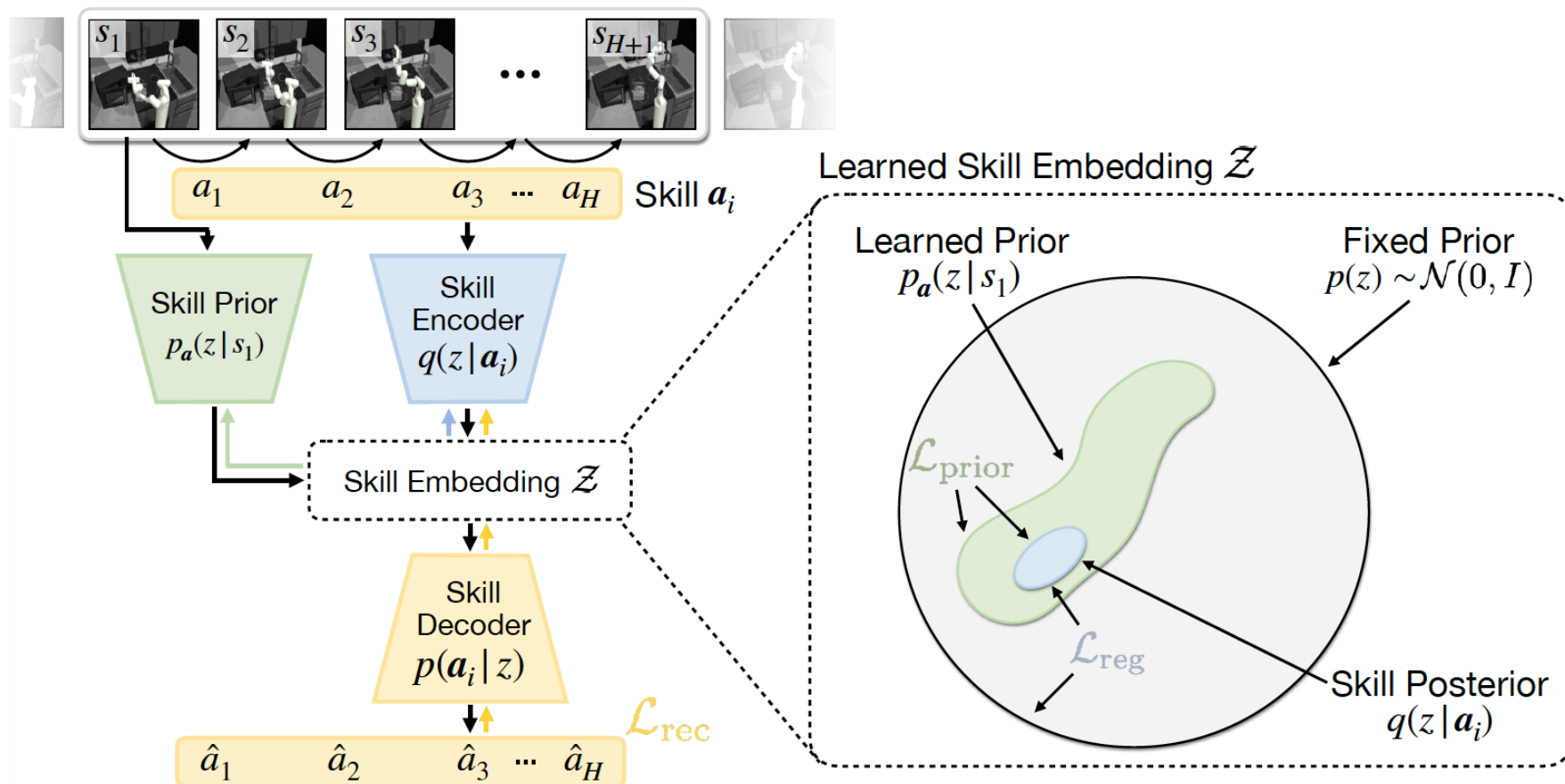
- Equivalent to the negated KL divergence between the policy and a uniform action prior:

$$U(\hat{a}_t): \mathcal{H}(\pi(a_t|s_t)) = -\mathbb{E}_{\pi} [\log \tilde{\pi}(a_t|s_t)] \propto -D_{\text{KL}}(\pi(a_t|s_t), U(\hat{a}_t))$$

- Replace the entropy term with negated KL divergence from prior

$$J(\theta) = \mathbb{E}_{\pi} \left[\sum_{t=1}^T \tilde{r}(s_t, z_t) - \alpha D_{\text{KL}}(\pi(z_t|s_t), p_{\mathbf{a}}(z_t|s_t)) \right]$$

Overall Structure

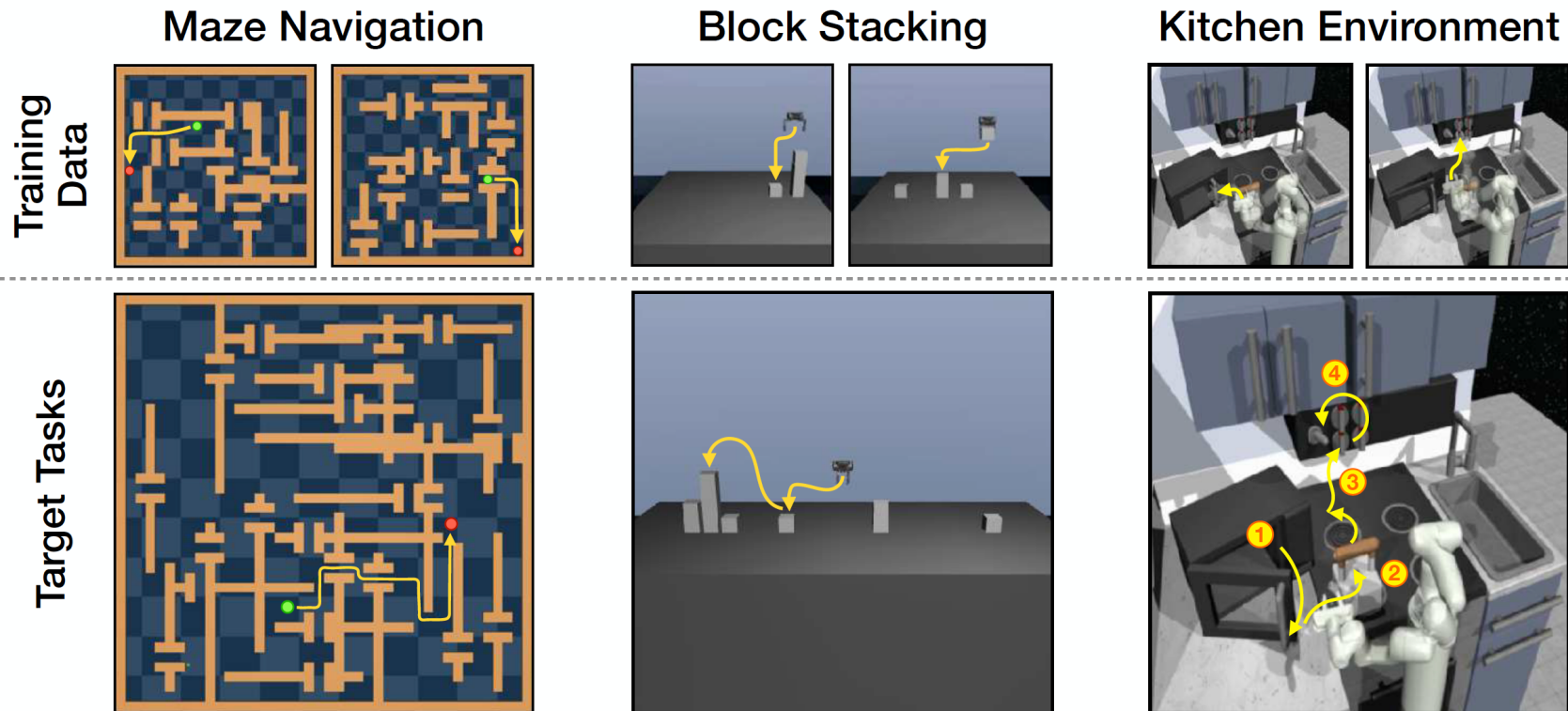


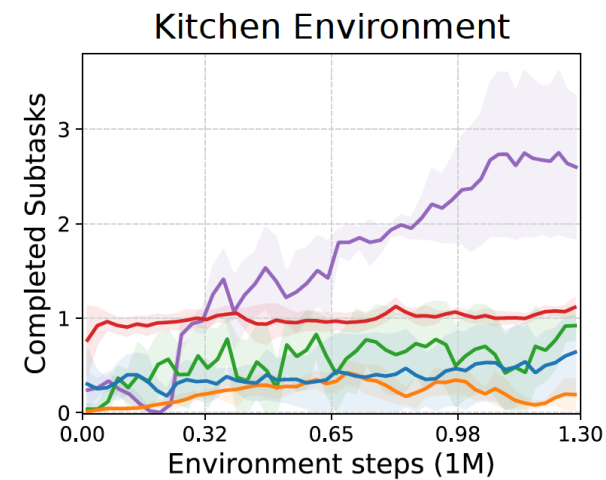
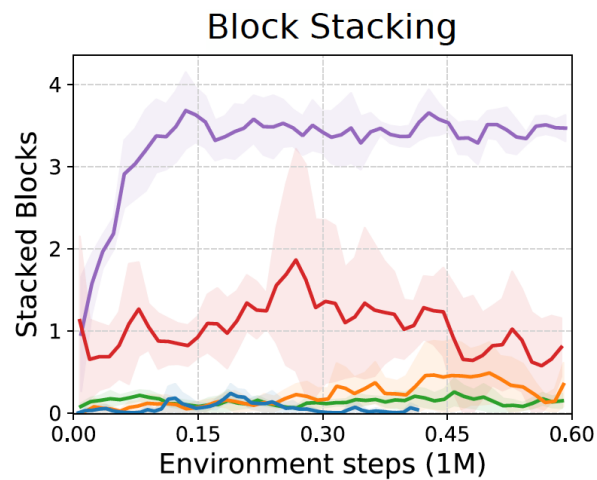
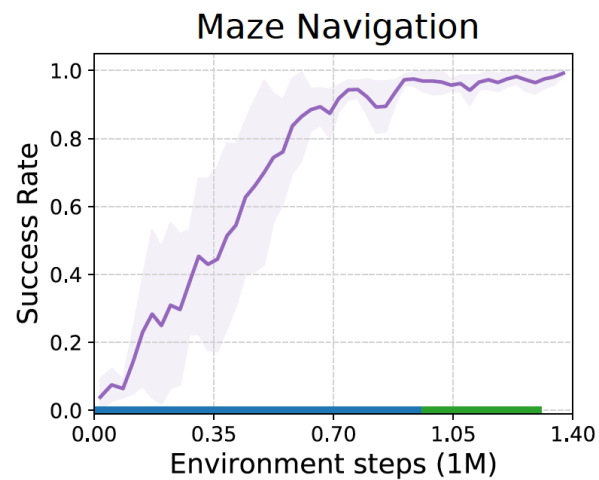
Overall Training Process

Algorithm 1 SPiRL: Skill-Prior RL

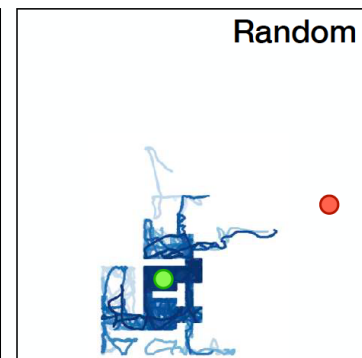
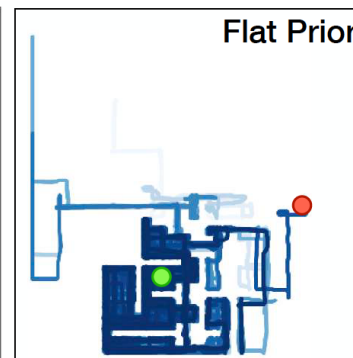
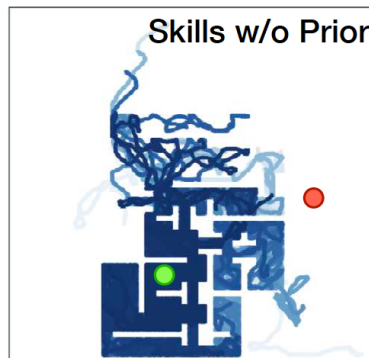
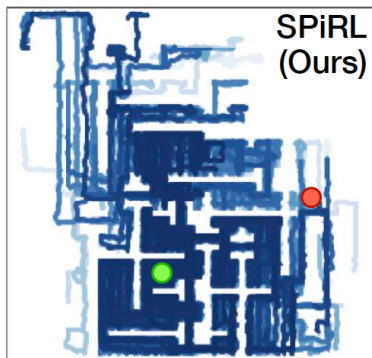
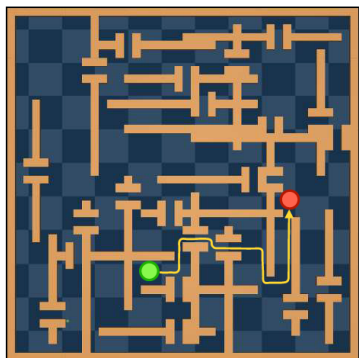
- 1: **Inputs:** H -step reward function $\tilde{r}(s_t, z_t)$, discount γ , target divergence δ , learning rates $\lambda_\pi, \lambda_Q, \lambda_\alpha$, target update rate τ .
 - 2: Initialize replay buffer $\mathcal{D}$, high-level policy $\pi_\theta(z_t|s_t)$, critic $Q_\phi(s_t, z_t)$, target network $Q_{\bar{\phi}}(s_t, z_t)$
 - 3: **for** each iteration **do**
 - 4: **for** every H environment steps **do**
 - 5: $z_t \sim \pi(z_t|s_t)$ ▷ sample skill from policy
 - 6: $s_{t'} \sim p(s_{t+H}|s_t, z_t)$ ▷ execute skill in environment
 - 7: $\mathcal{D} \leftarrow \mathcal{D} \cup \{s_t, z_t, \tilde{r}(s_t, z_t), s_{t'}\}$ ▷ store transition in replay buffer
 - 8: **for** each gradient step **do**
 - 9: $\bar{Q} = \tilde{r}(s_t, z_t) + \gamma [Q_{\bar{\phi}}(s_{t'}, \pi_\theta(z_{t'}|s_{t'})) - \alpha D_{\text{KL}}(\pi_\theta(z_{t'}|s_{t'}), p_\alpha(z_{t'}|s_{t'}))]$ ▷ compute Q-target
 - 10: $\theta \leftarrow \theta - \lambda_\pi \nabla_\theta [Q_\phi(s_t, \pi_\theta(z_t|s_t)) - \alpha D_{\text{KL}}(\pi_\theta(z_t|s_t), p_\alpha(z_t|s_t))]$ ▷ update policy weights
 - 11: $\phi \leftarrow \phi - \lambda_Q \nabla_\phi [\frac{1}{2} (Q_\phi(s_t, z_t) - \bar{Q})^2]$ ▷ update critic weights
 - 12: $\alpha \leftarrow \alpha - \lambda_\alpha \nabla_\alpha [\alpha \cdot (D_{\text{KL}}(\pi_\theta(z_t|s_t), p_\alpha(z_t|s_t)) - \delta)]$ ▷ update alpha
 - 13: $\bar{\phi} \leftarrow \tau \phi + (1 - \tau) \bar{\phi}$ ▷ update target network weights
 - 14: **return** trained policy $\pi_\theta(z_t|s_t)$
-

Experiment





— SPIRL (Ours)
 — Flat Prior
 — SSP w/o Prior
 — SAC
 — BC + SAC



Thanks !

