

Model-based Offline Reinforcement Learning with Lower Expectile Q-Learning

Kwanyoung Park, Youngwoon Lee

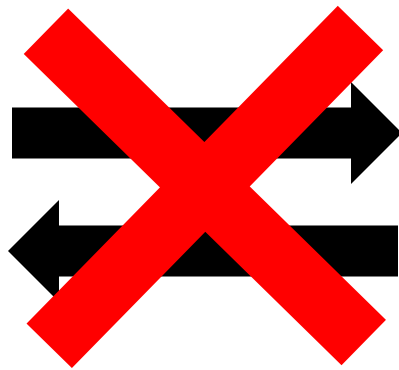
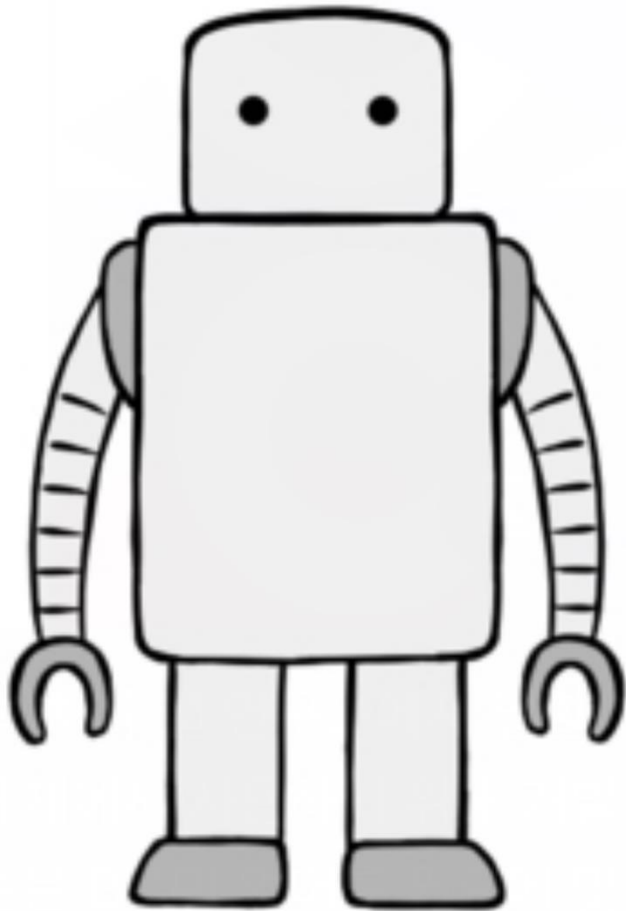


ICLR 2025

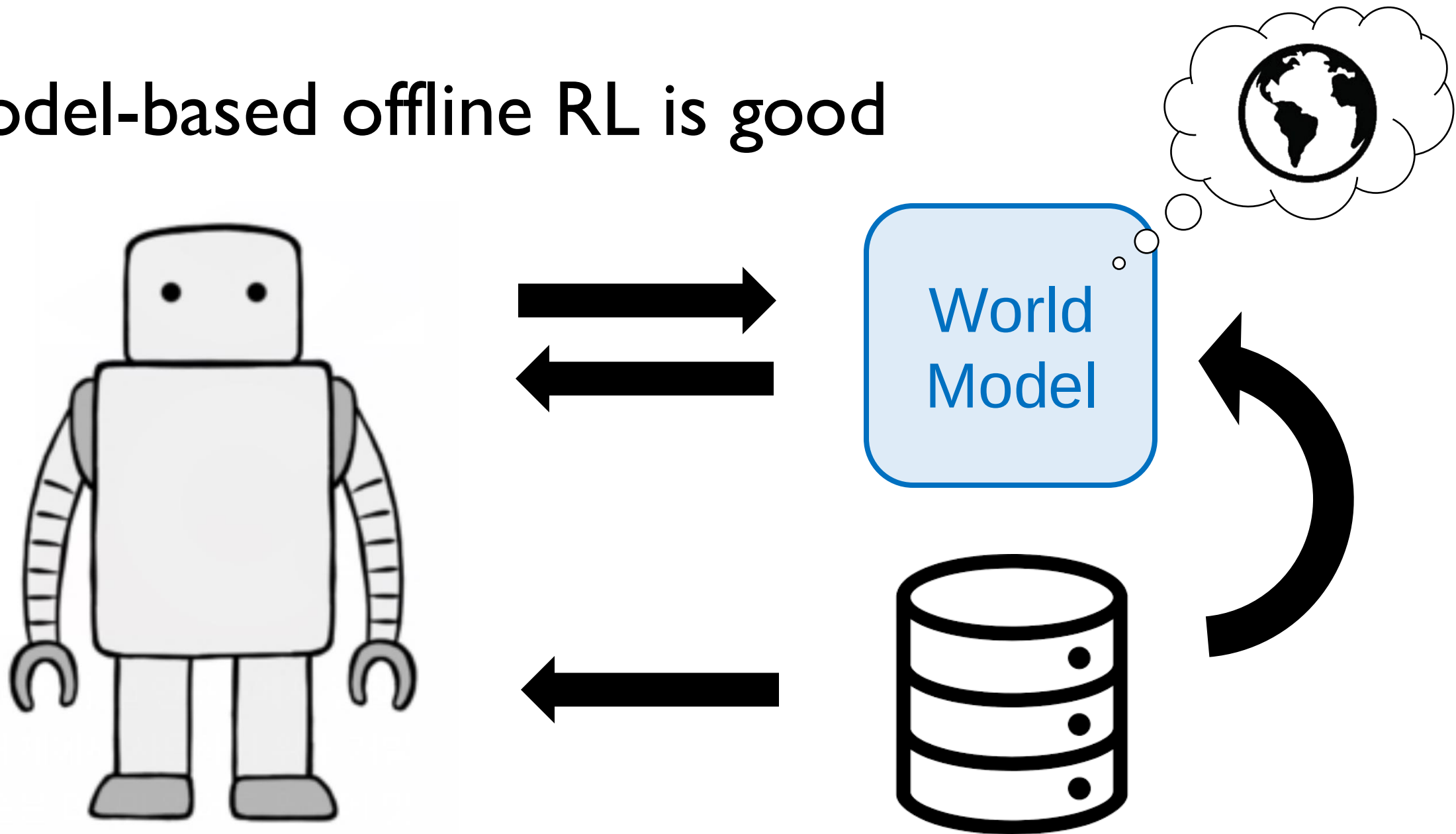
<https://kwanyoungpark.github.io/LEQ/>



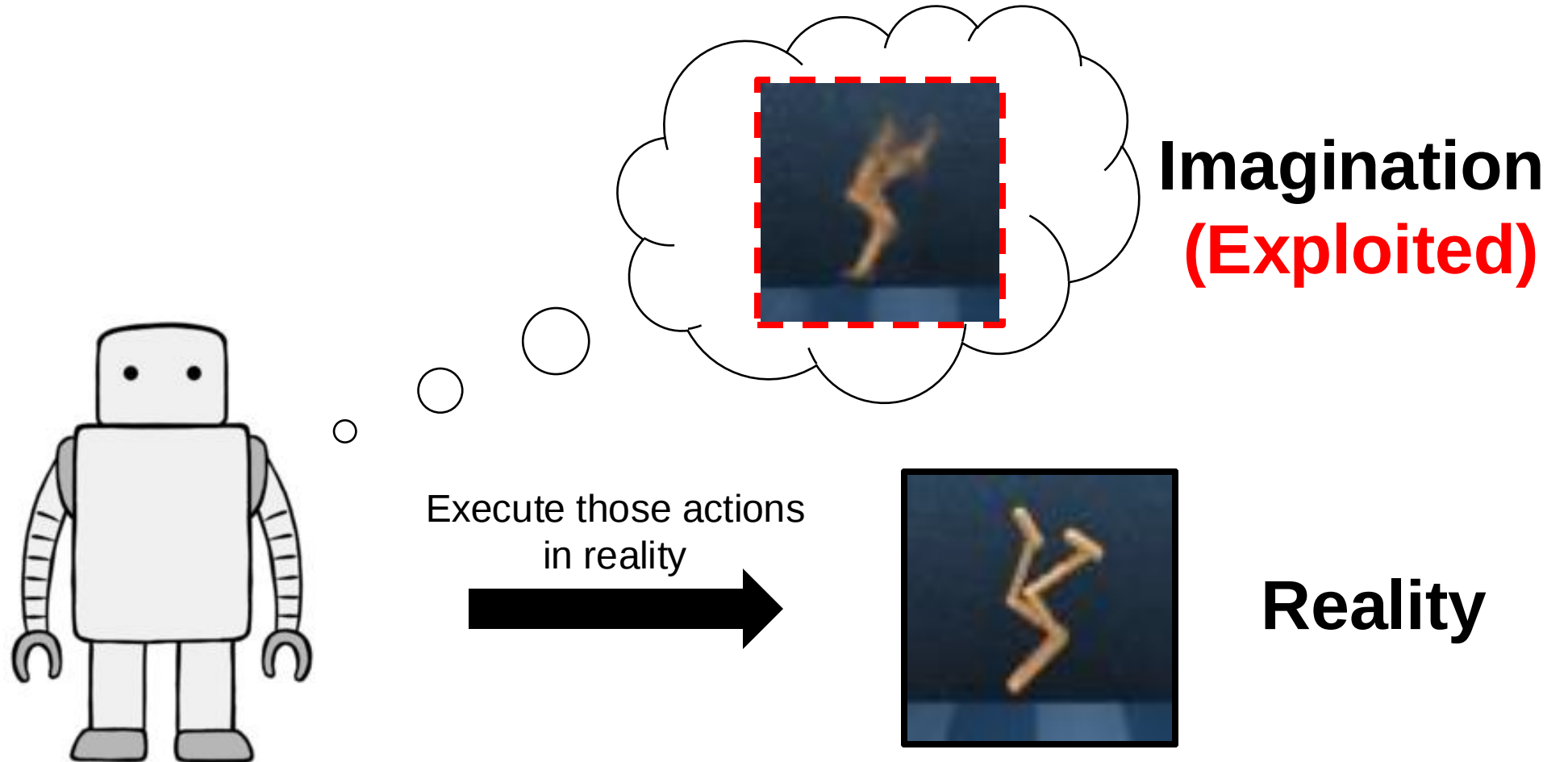
Why offline RL is hard?



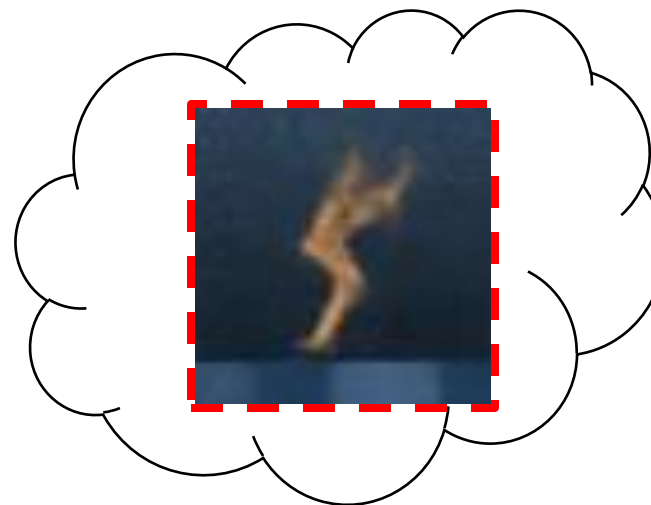
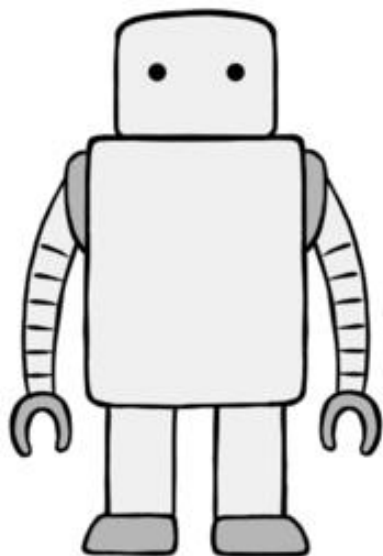
Model-based offline RL is good



Model can be **exploited**



Reality

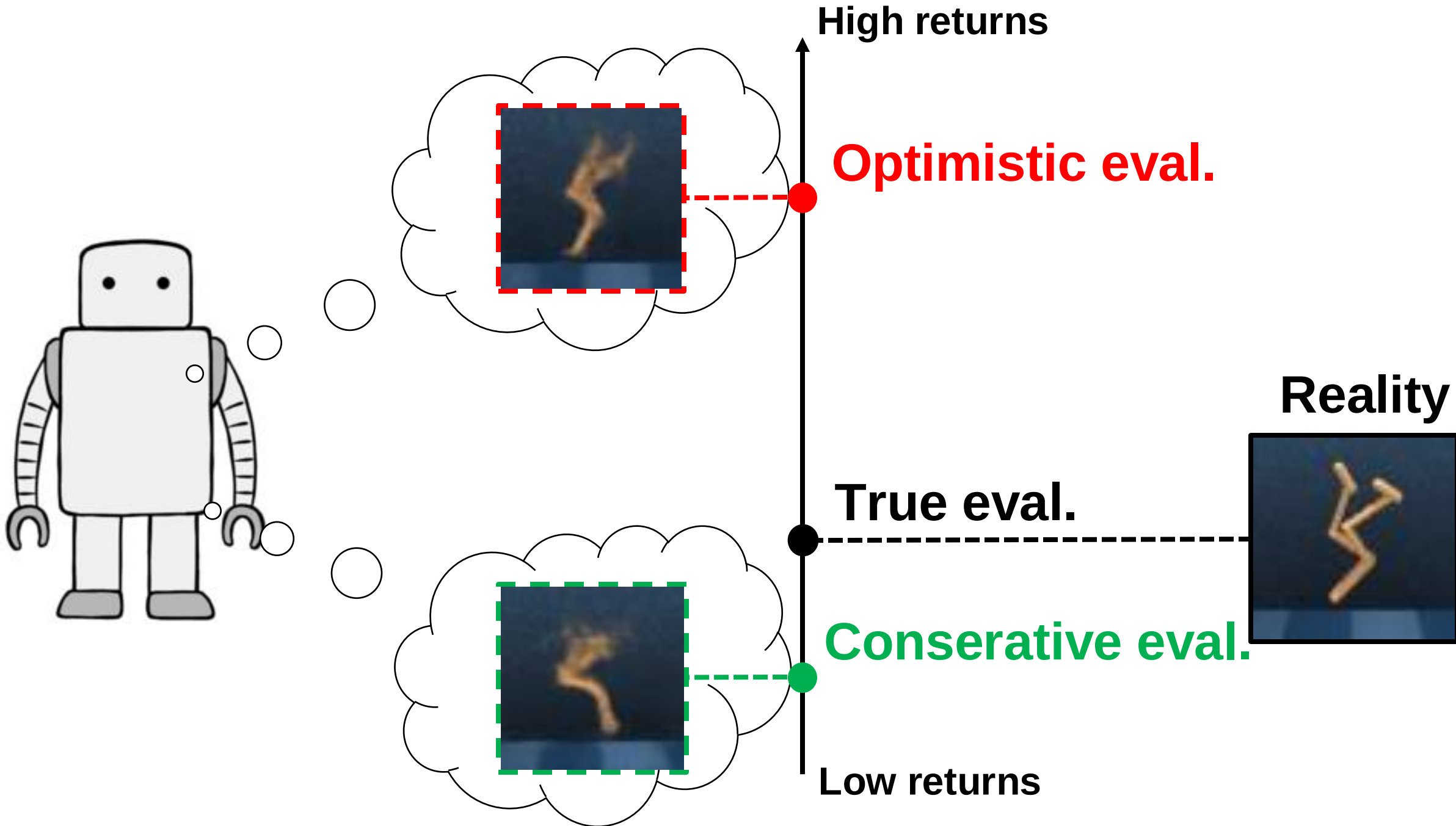


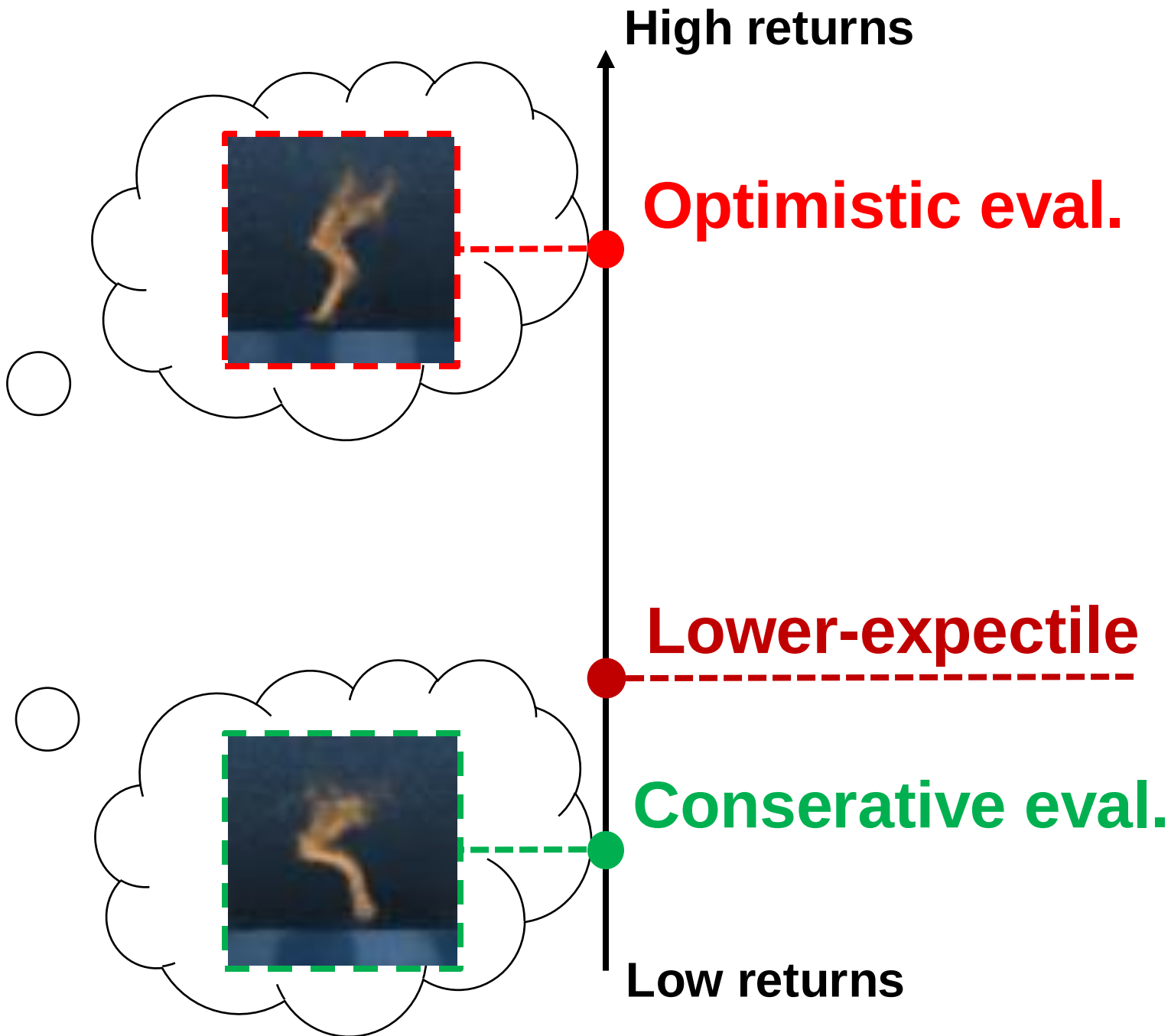
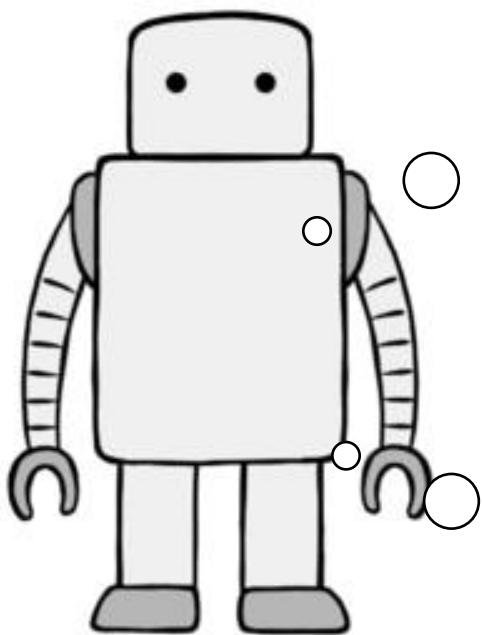
Optimistic



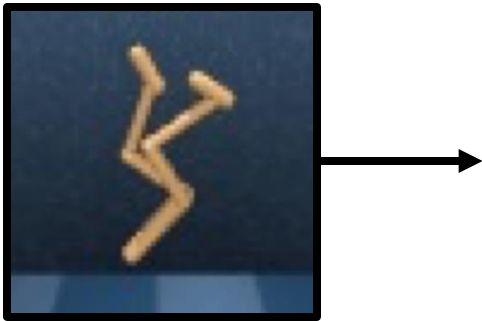
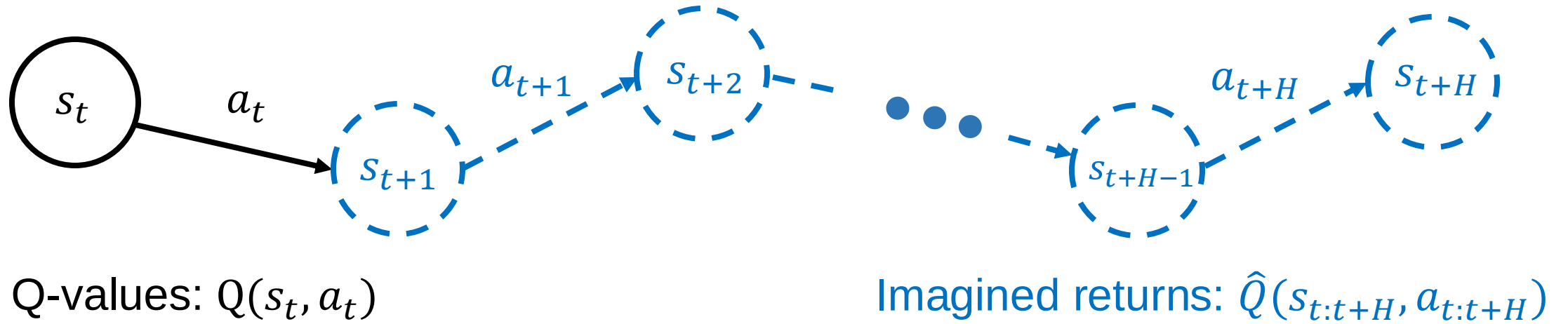
Conservative



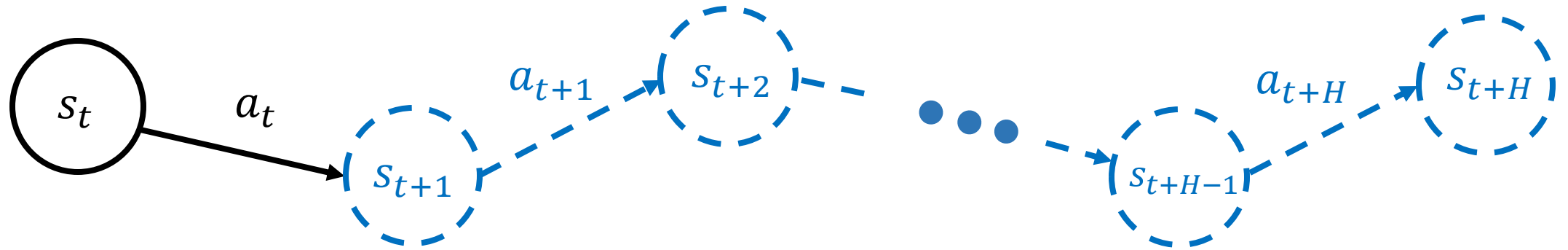




Model-generated trajectory (Imagination)



Model-generated trajectory (Imagination)



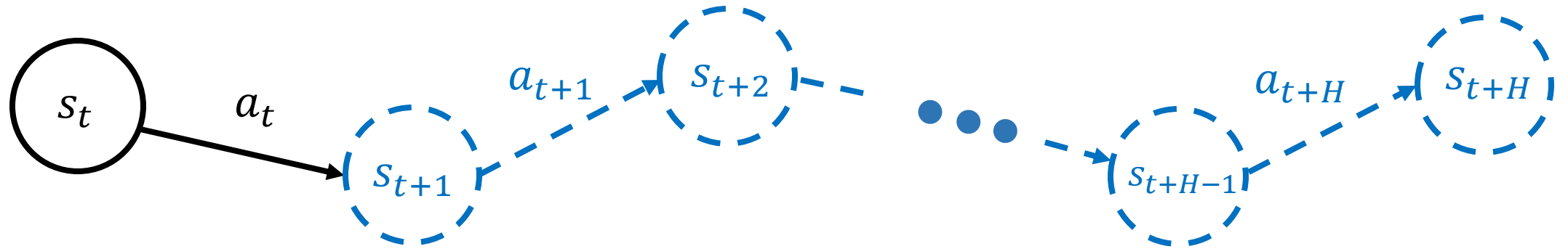
Q-values: $Q(s_t, a_t)$

Imagined returns: $\hat{Q}(s_{t:t+H}, a_{t:t+H})$



Optimistic
($Q < \hat{Q}$)

Model-generated trajectory (Imagination)



Q-values: $Q(s_t, a_t)$

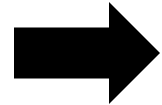
Imagined returns: $\hat{Q}(s_{t:t+H}, a_{t:t+H})$



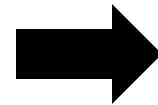
Conservative

$(Q > \hat{Q})$

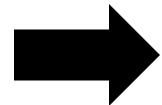
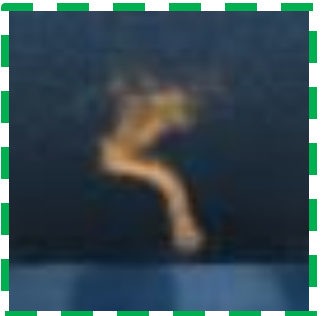
Lower-Expectile Q-Learning (LEQ)



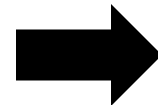
Optimistic



$$0.1 \cdot \frac{(Q_\phi(s_t, a_t) - \hat{Q})^2}{\text{Bellman loss}}$$

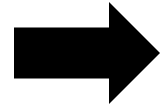


Conservative

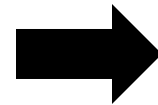


$$0.9 \cdot \frac{(Q_\phi(s_t, a_t) - \hat{Q})^2}{\text{Bellman loss}}$$

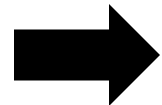
Lower-Expectile Q-Learning (LEQ)



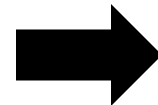
Optimistic



$$0.1 \cdot \frac{\nabla_{\theta} \hat{Q}(s_{t:t+H}, a_{t:t+H})}{\text{Policy Gradient}}$$



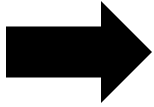
Conservative



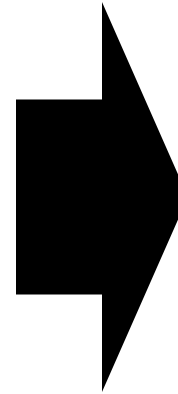
$$0.9 \cdot \frac{\nabla_{\theta} \hat{Q}(s_{t:t+H}, a_{t:t+H})}{\text{Policy Gradient}}$$

Lower-Expectile Q-Learning (LEQ)

Optimistic

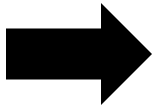


0.1 · Bellman loss
0.1 · Policy gradient



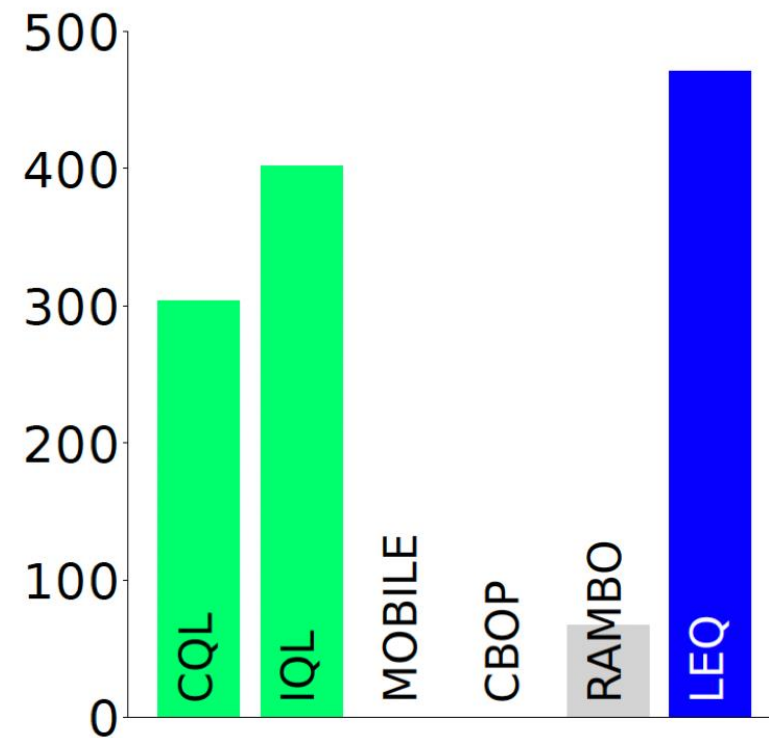
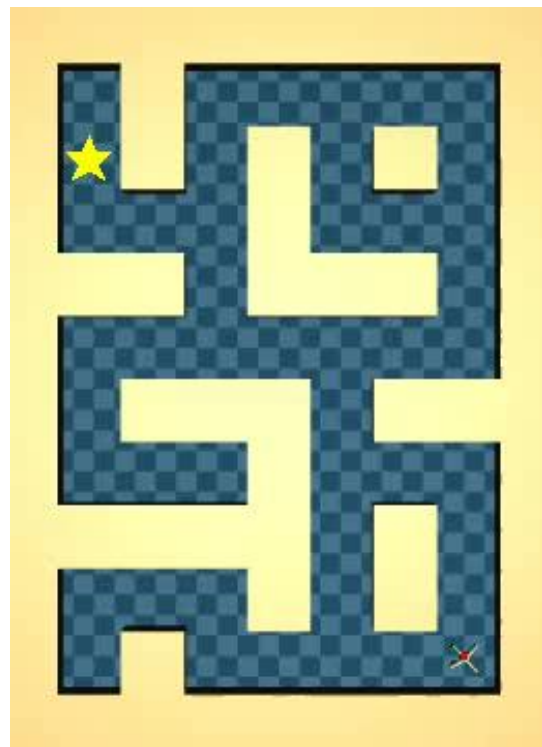
$$Q_{\phi}(s_t, a_t) = \text{Lower-expectile}$$
$$\nabla_{\theta} L_{\pi}(\theta) \approx \nabla_{\theta} (\text{Lower-expectile})$$

Conservative



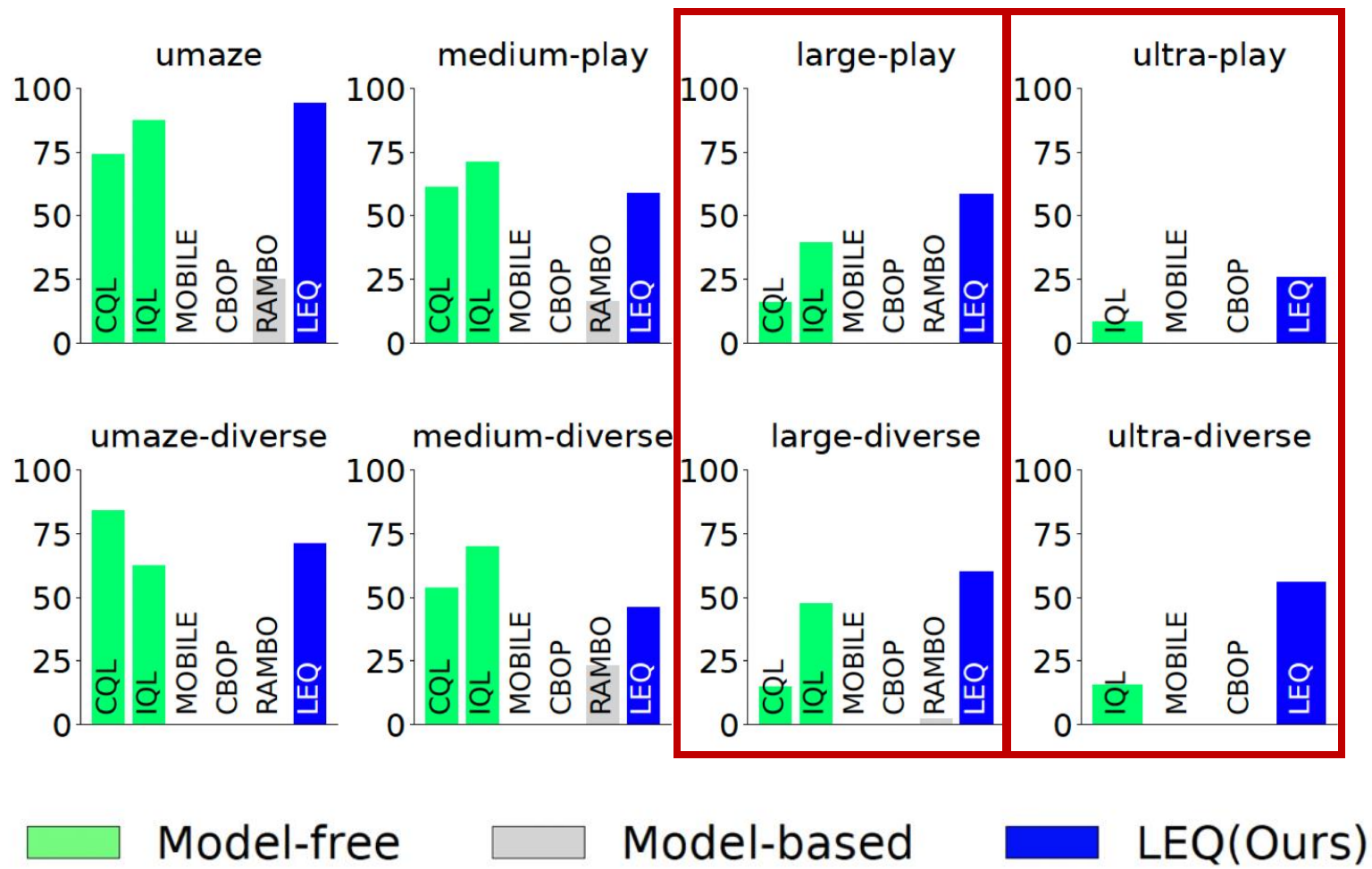
0.9 · Bellman loss
0.9 · Policy gradient

AntMaze

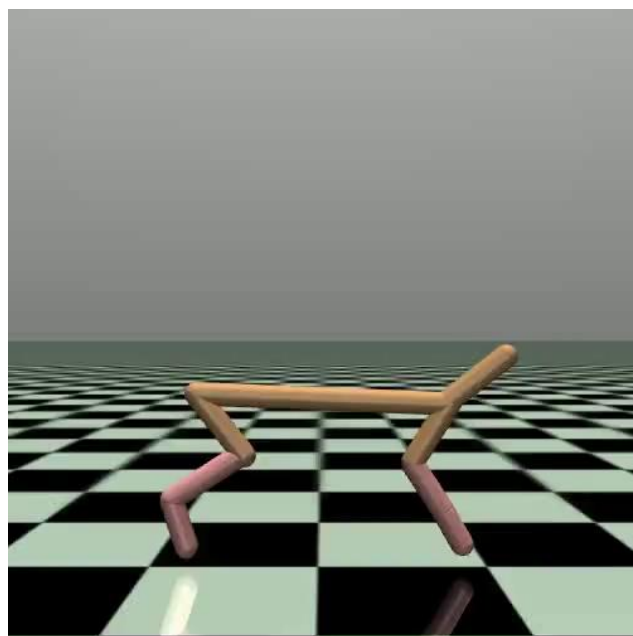


 Model-free  Model-based  LEQ(Ours)

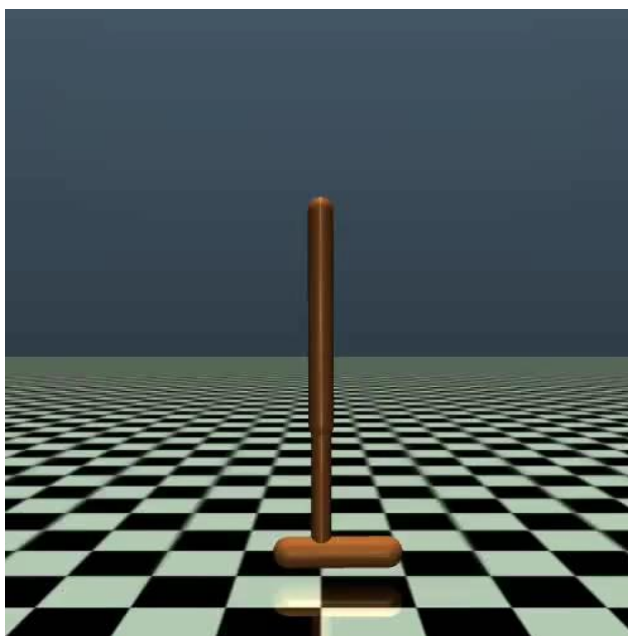
AntMaze



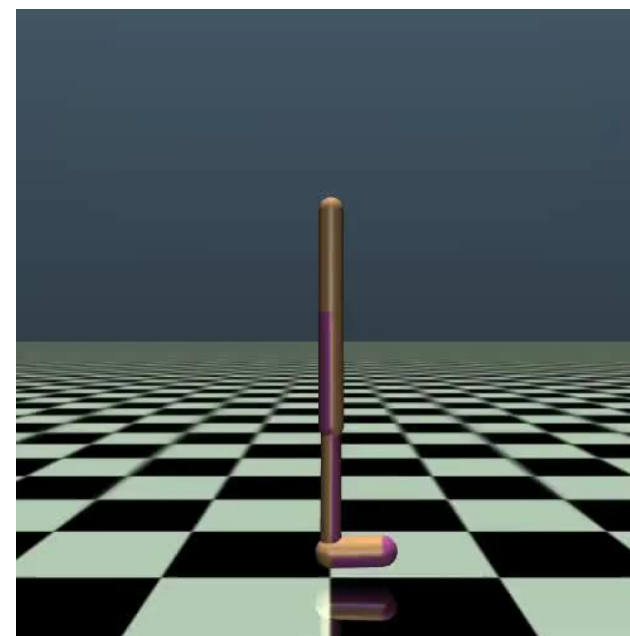
Mujoco Gym (NeoRL, D4RL)



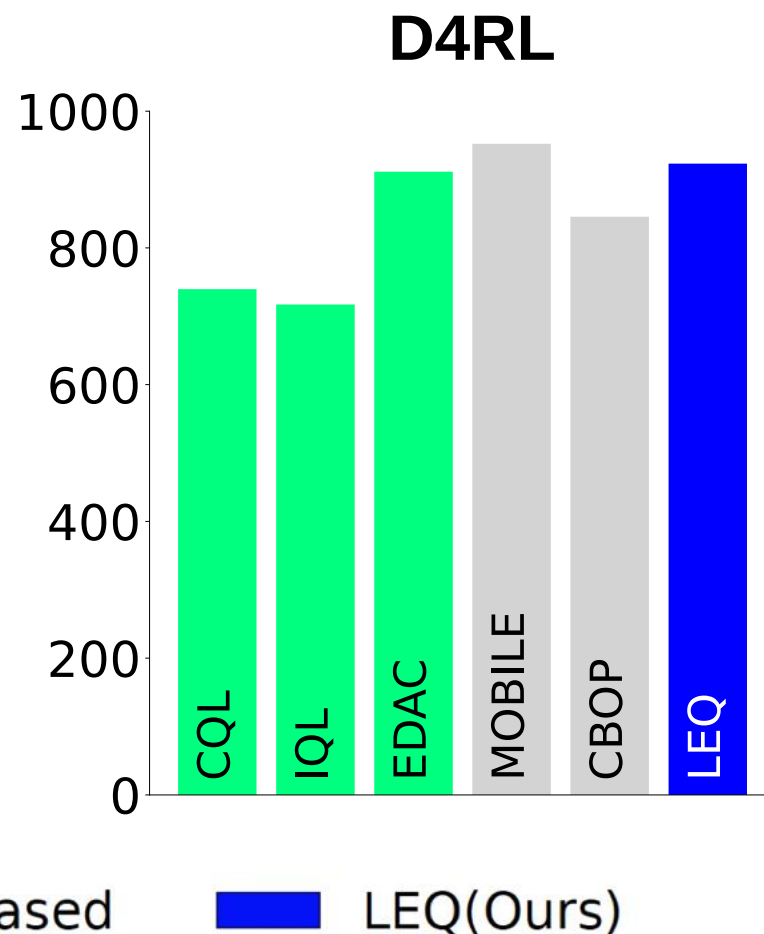
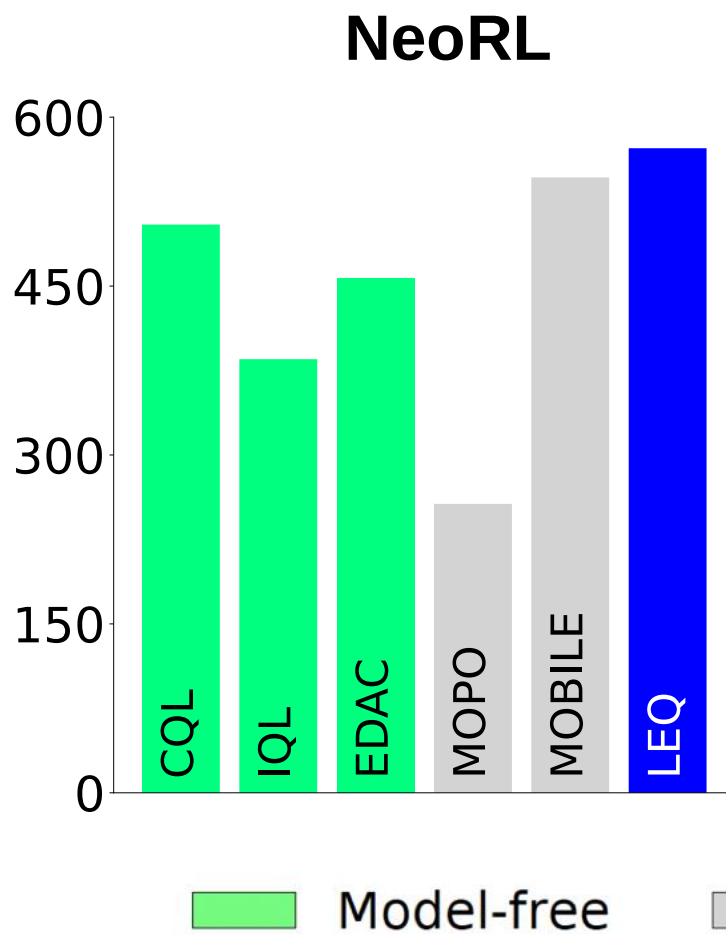
HalfCheetah



Hopper



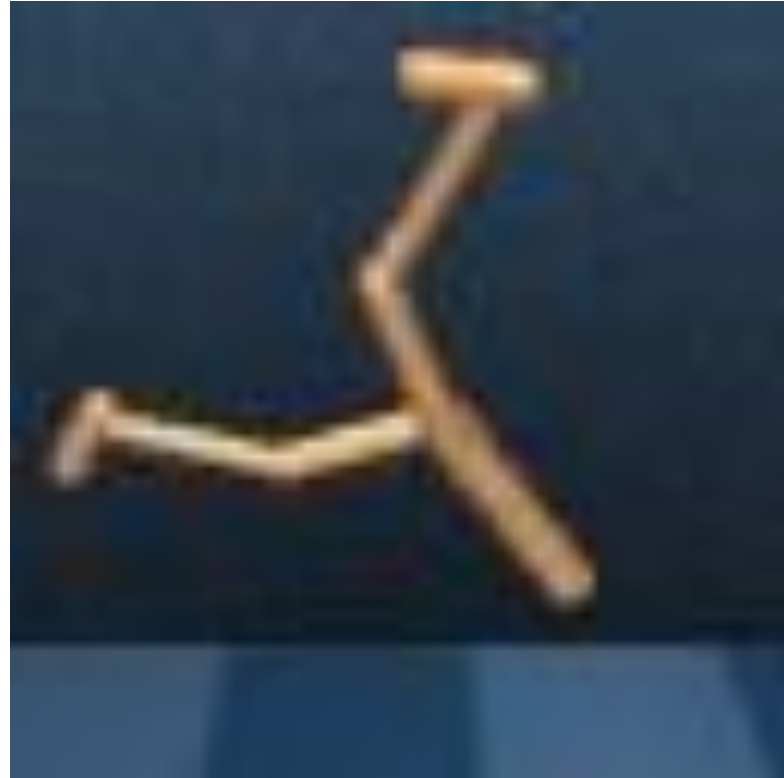
Walker



DMControl (V-D4RL)

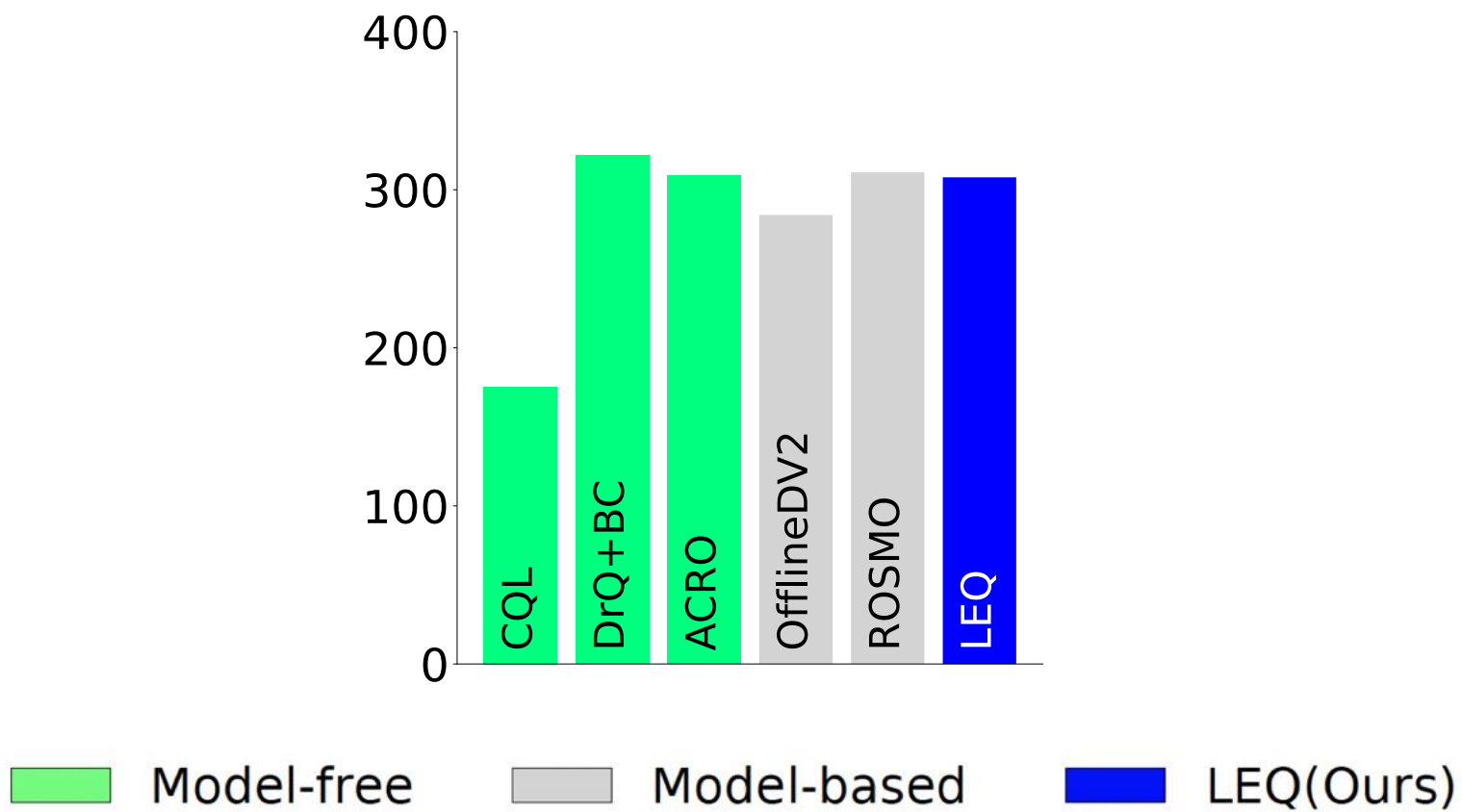


Cheetah-Run



Walker-Walk

V-D4RL



Summary

- LEQ is an offline **model-based** RL algorithm.
- LEQ optimizes the **lower-expectile** of the returns.
 - Assign **larger** weights to **conservative** trajectories
- LEQ works well in various domains.



Project Page