

Stabilizing Q-Learning with Linear Architectures for Provable Efficient Learning

Andrea Zanette, Martin J. Wainwright

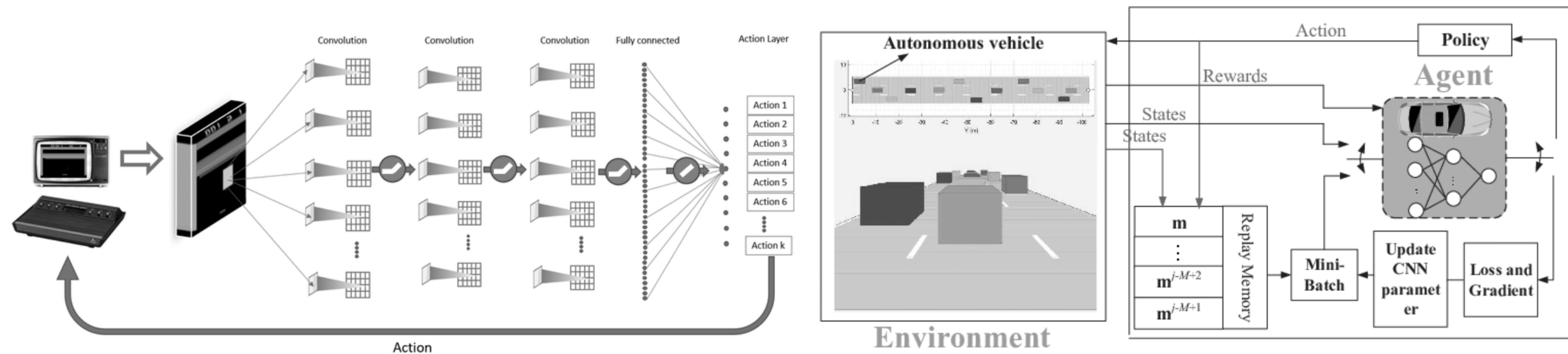


Q-Learning

Core RL algorithm

Based on controlling the projected Bellman error

Practically successful even in large applications



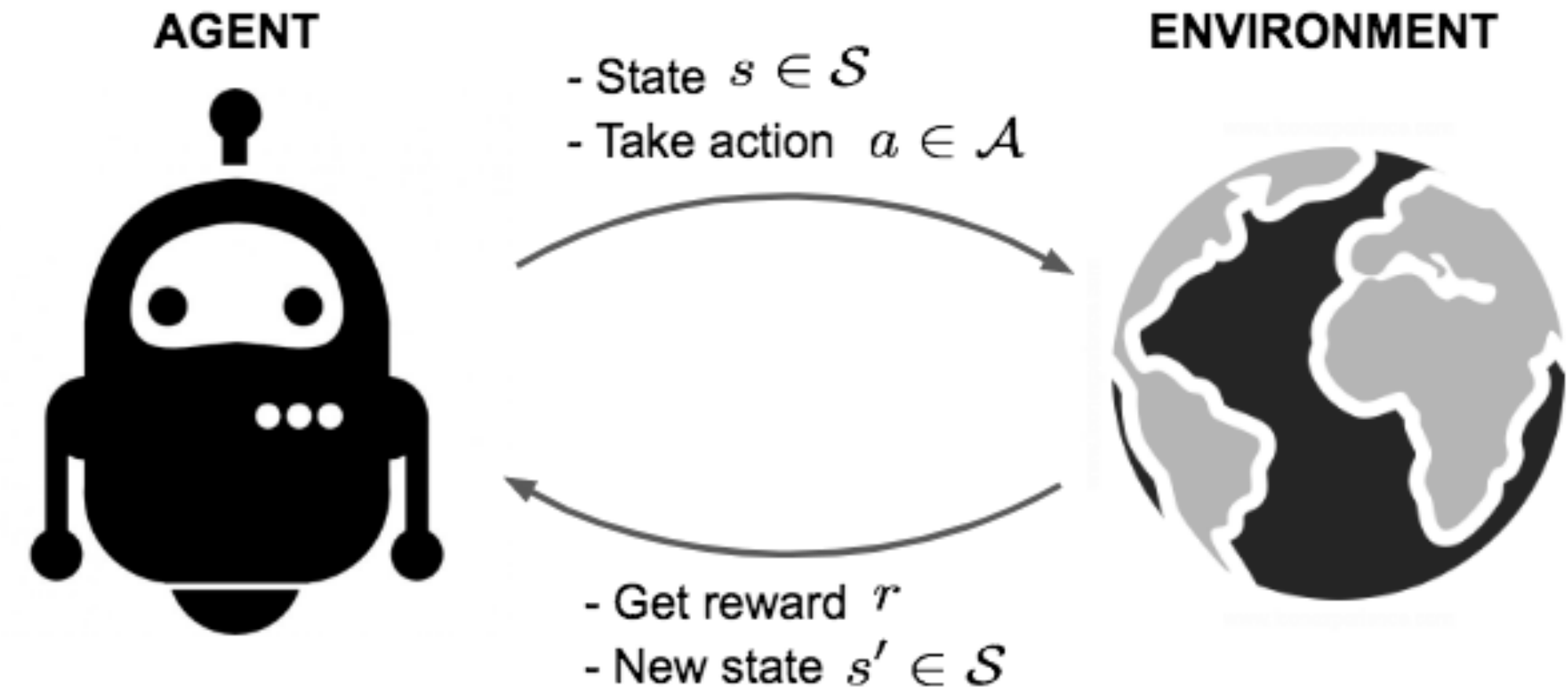
Mnih et al, 2015

Kiran et al, 2020

Q-Learning with function approximation uses
experience replay, target networks

Reinforcement Learning Exploration

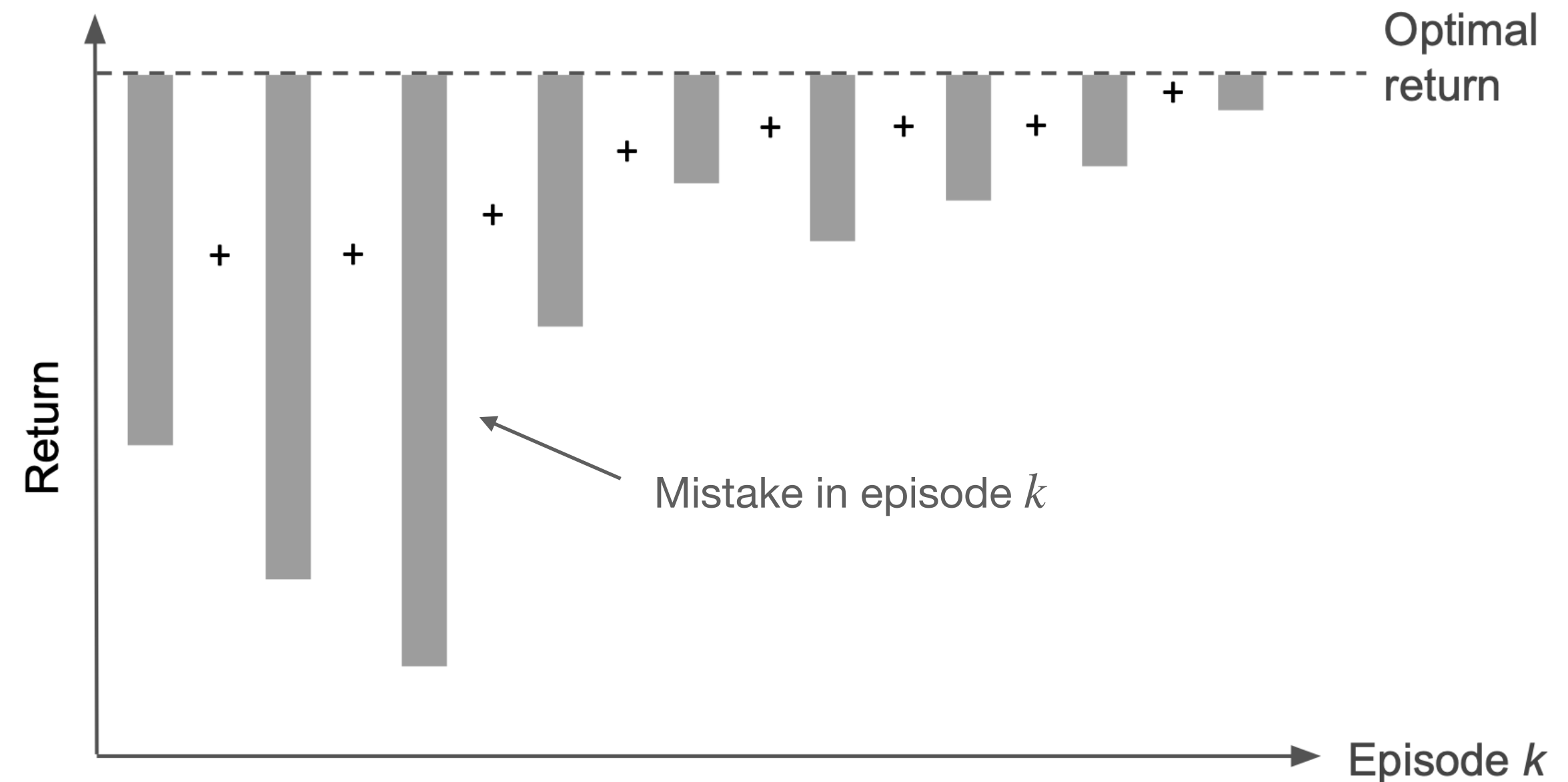
Setting



Goal: Minimize Cumulative Regret

$$\text{Regret} = \sum_{k=1}^K V^{\star} - V^{\pi_k}$$

Optimal return Agent's return



Q-Learning with Linear Function Approximation

Parametric Approximation

$$Q_w(s, a) = \phi(s, a)^\top w$$

Learning rate

Update

$$w^+ = w - \alpha \underbrace{[Q_w(s, a) - r - \gamma \max_{a'} Q_w(s', a')]}_{\text{TD error}} \underbrace{\nabla Q_w(s, a)}_{\text{Gradient of Q-function}}$$

Current parameter

TD error

Gradient of Q-function

Issue: not a stochastic gradient update

1. 'Gradient' is not really a gradient of any loss
2. Non-stationary (s, a, r, s') data distribution

Stabilizing Q-Learning with Linear Architectures

1) **Exploration Bonus**

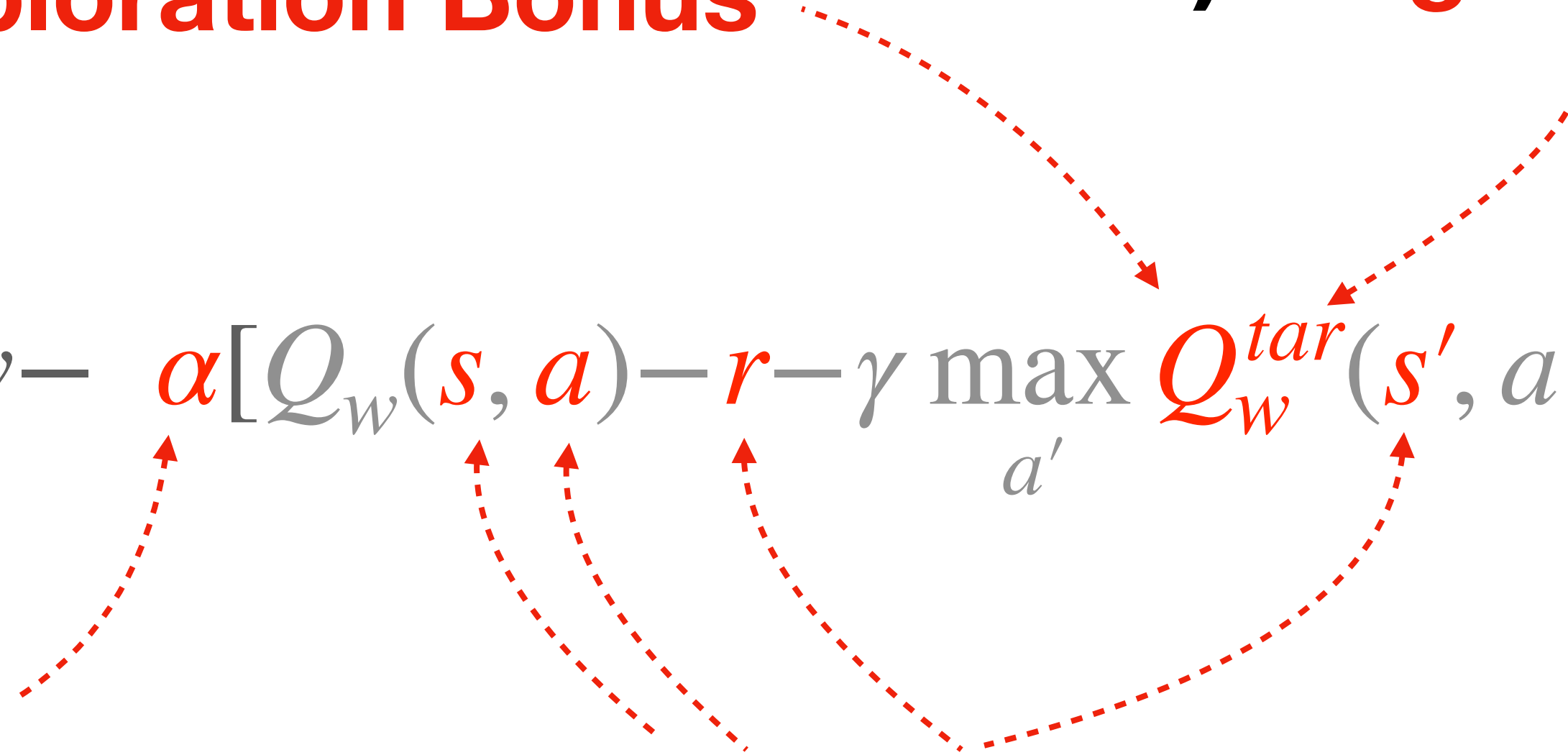
2) **Target Network**

Update

$$w^+ = w - \alpha [Q_w(s, a) - r - \gamma \max_{a'} Q_w^{tar}(s', a')] \nabla Q_w(s, a)$$

4) **Learning Rate**

3) **Replay Mechanism**



Target Network

Target Network

Ensures Q-Learning
Minimizes a Loss

Periodic Updates
(slower timescale)

$$\mathcal{L}(w) = \sum_{s,a,r,s'} [Q_w(s,a) - r - \gamma \max_{a'} Q_w^{tar}(s',a')]^2$$

$$Q_w^{tar} \leftarrow Q_w$$

Fixed Bellman backup



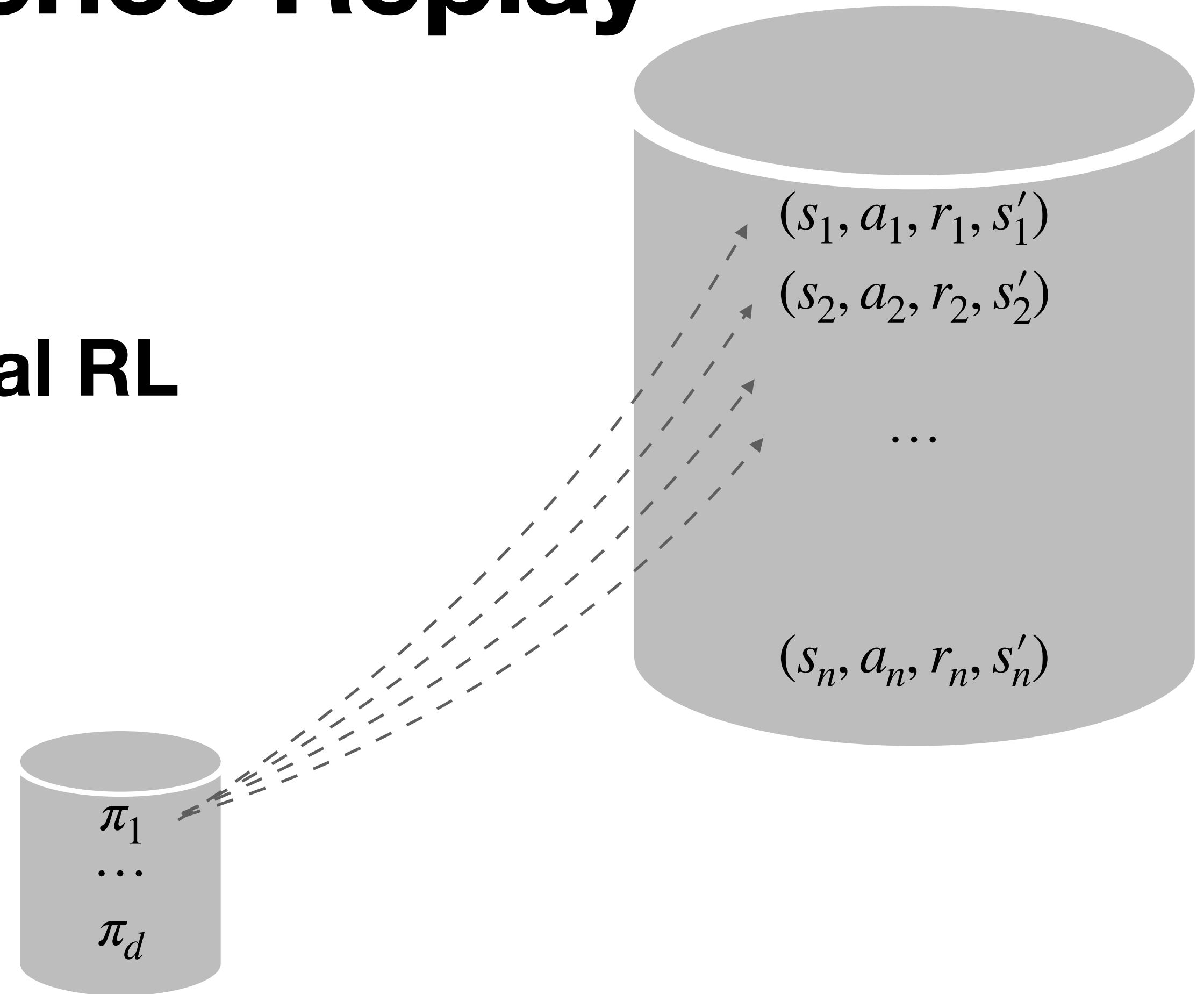
Policy Replay for Experience Replay

Experience replay

essential for state-of-the-art empirical RL
(e.g., Mnih et al., 2015)

Policy replay

same effect but limited memory



Main Result

Q-Learning with Linear Function Approximation is provably efficient

Cumulative Regret

Horizon

Feature dimension

$$\widetilde{O}(H^3 d^{3/2} \sqrt{K})$$

Episodes

Detailed description: This diagram shows the asymptotic bound for cumulative regret, $\widetilde{O}(H^3 d^{3/2} \sqrt{K})$. Dashed arrows point from the labels 'Horizon' and 'Feature dimension' to the terms H and d respectively. Another dashed arrow points from the label '# Episodes' to the term $\sqrt{K}$.

Total Memory Used

$$\widetilde{O}(d^3 H^2)$$

Per-Step Computations

$$\widetilde{O}(d^2 A)$$

Actions

Detailed description: This diagram shows the asymptotic bound for per-step computations, $\widetilde{O}(d^2 A)$. A dashed arrow points from the label '# Actions' to the term A .

Thank you for your Attention!