

# From Bandits to PPO

## I. Bandit

\*

action value

$$q_{\pi}(a) = \mathbb{E}(R_t | A_t = a)$$

Hedge UCB

\*

softmax policy

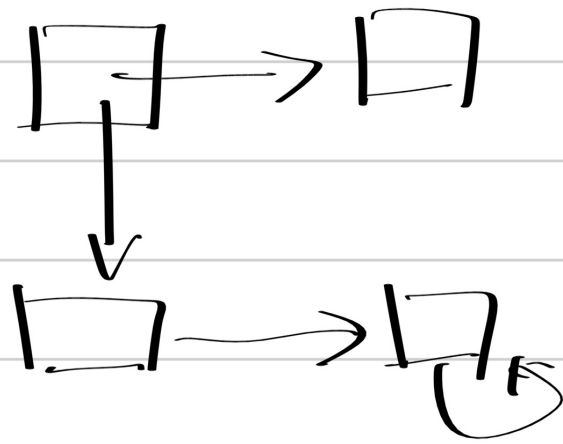
$$\pi_t(a) = \frac{\exp(H_t(a))}{\sum_b \exp(H_t(b))}$$

\* Preference Update Rule

$$J(\pi) = \mathbb{E}[R_t] = \sum \pi_t(a) q_{\pi}(a)$$

{ contextual bandit

$\pi$  state  
action  
policy



(immediate) reward

{ discounted return  
average return

$$V(s) \leftarrow V(s) + \sum_{s', r} P(s', r | s, a) (\gamma V(s') + r - V(s))$$

$$V = \gamma P V + r$$

$$(I_n - \gamma P) V = r$$

inverse -

\* contraction mapping

### III policy iteration value iteration

#### \* Bellman Optimality Equation

$$V^*(s) = \max_a \sum_{s', r} p(s', r | s, a) (\gamma V^*(s') + r)$$

\* Essence : DP

+

Numerical Linear Algebra

You have to know  
"probability"

Model - Based.

Curse of Dimensionality

↳ To update "s", we have to  
[Full Sweep] state space.

# IV Monte Carlo & TD learning

- Model free
- Finite states
- "Tabular"

How to evaluate

\* Monte Carlo  
sample a complete trajectory

$$V(s_t) \leftarrow V(s_t) + \alpha [G_t - V(s_t)]$$

\* Temporal Difference learning

$$V(s_t) \leftarrow V(s_t) + \alpha [R_{t+1} + \gamma V(s_{t+1}) - V(s_t)]$$





# V Q-Learning

$$Q(s,a) \leftarrow Q(s,a) + \alpha [R_{t+1} + \gamma \max_{a'} Q(s',a') - Q(s,a)]$$

$$\gamma \max_{a'} Q(s',a')$$

so have to be tabular

X DQN

Couple "Sample" with  
"Function Approximation"

$$Loss(\theta) = \mathbb{E}_{s,a,r,s'} |r + \gamma \max_{a'} Q(s',a') - Q(s,a)|^2$$

For stability

# VI, Policy Gradient

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{s_t, a_t \sim \pi_{\theta}} Q^{\pi}(s_t, a_t) \nabla_{\theta} \log \pi(a_t | s_t)$$

$$J_{\theta} = \int_{\mathcal{P}} R(p) P(p | \theta) dp$$

$$P(p | \theta) = p(s_0) \prod_{t=0}^{\infty} P(s_{t+1} | s_t, a_t) \pi_{\theta}(a_t | s_t)$$

$$\nabla_{\theta} J(\theta) = \nabla_{\theta} \int_{\mathcal{P}} R(p) P(p | \theta) dp = \int_{\mathcal{P}} \nabla_{\theta} R(p) P(p | \theta) dp$$

$$= \int_{\mathcal{P}} R(p) (\nabla_{\theta} P(p | \theta)) dp$$

$$= \int_{\mathcal{P}} R(p) (\nabla_{\theta} \ln P(p | \theta)) P(p | \theta) dp$$

$$= \mathbb{E}_{p \sim \pi_{\theta}} R(p) \nabla_{\theta} (\ln P(p | \theta)) = \mathbb{E}_{p \sim \pi_{\theta}} R(p) \nabla_{\theta} \left( \sum_{t=0}^{\infty} P(s_{t+1} | s_t, a_t) \pi(a_t | s_t) \right) p(s_0) +$$

What to optimize

→ value-based.

→ policy-based

VII, Reinforce.

use Monte Carlo to  
estimate  $Q$

Monte Carlo Return

