

Policy gradient

~~$V(s), Q(s,a)$~~
value functions

~~$s \rightarrow \boxed{\hat{Q}} = \hat{Q}(s,a)$~~

$$s \rightarrow \boxed{\pi} \xrightarrow{\bar{\theta}} \pi(a|s, \theta) = \Pr[A_t = a | S_t = s]$$

$\pi(a|s, \theta)$
 $V(s, \theta)$ } actor-critic methods

① M.S. represent models π

$$s, a \rightarrow \boxed{h} \xrightarrow{\bar{\theta}} \frac{h(s, a, \bar{\theta})}{\text{score}} \rightarrow \boxed{\text{softmax}} \rightarrow \pi(a|s) = \frac{e^{h(s, a, \bar{\theta})}}{\sum_{a'} e^{h(s, a', \bar{\theta})}}$$

$$h(s, a, \bar{\theta}) = \bar{\theta}^T \bar{x}(s, a)$$

② Monte Carlo $\pi(a|s)$ - reg. by π

③ $\bar{\theta}_{t+1} = \bar{\theta}_t + \alpha \nabla J(\bar{\theta}_t)$

$$J(\bar{\theta}) = \underbrace{V_{\pi_{\bar{\theta}}}(s_0)}_{\bar{\theta} \rightarrow \max}$$

Policy gradient theorem

$$\nabla_{\theta} V_{\pi}(s) = \nabla_{\theta} \left[\sum_a \underbrace{\pi(a|s)}_{\theta} Q_{\pi}(s, a) \right] =$$

$$= \sum_a \left[\left(\nabla_{\theta} \pi(a|s) \right) \cdot Q_{\pi}(s, a) + \pi(a|s) \cdot \nabla_{\theta} Q_{\pi}(s, a) \right] =$$

$$= \sum_a \left[Q_{\pi}(s, a) \nabla_{\theta} \pi(a|s) + \pi(a|s) \nabla_{\theta} \left[\sum_{s', r} p(s', r | s, a) \cdot \cancel{(r + V_{\pi}(s'))} \right] \right] =$$

$$= \sum_a \left[Q_{\pi}(s, a) \nabla_{\theta} \pi(a|s) + \pi(a|s) \sum_{s'} p(s' | s, a) \nabla_{\theta} V_{\pi}(s') \right] =$$

$$= \sum_a \left[\underbrace{Q_{\pi}(s, a) \nabla_{\theta} \pi(a|s)} + \pi(a|s) \sum_{s'} p(s' | s, a) \left(\underbrace{Q_{\pi}(s', a') \nabla_{\theta} \pi(a' | s')} + \pi(a' | s') \sum_{s''} p(s'' | s', a') \nabla_{\theta} V_{\pi}(s'') \right) \right]$$

$$= \sum_{s'} \sum_{k=0}^{\infty} \Pr[y s \rightarrow s' \text{ for } k \text{ times} \mid \text{no op } \pi] \cdot \sum_a Q_{\pi}(s', a) \nabla_{\theta} \pi(a|s')$$

$$\nabla_{\theta} J(\theta) = \nabla_{\theta} V_{\pi}(s_0) = \sum_s \underbrace{\left(\sum_{k=0}^{\infty} \Pr[s_0 \rightarrow s \mid k] \right)}_{\mathbb{E}[\# \text{ nonag. b } s]} \underbrace{\left(\sum_a Q_{\pi}(s, a) \nabla_{\theta} \pi(a|s) \right)}_{\mu(s) = \Pr[s]} =$$

$$= \sum_s \mathbb{E}[\# \text{ nonag. b } s] \sum_a \left(\right)$$

$$\nabla_{\theta} J(\theta) \propto \sum_s \mu(s) \sum_a Q_{\pi}(s, a) \nabla_{\theta} \pi(a|s)$$

policy
gradient
theorem

① All-actions method (actor-critic)

$$\bar{\theta}_{t+1} := \bar{\theta}_t + \alpha \sum_a \hat{Q}(s, a | \bar{w}) \nabla_{\bar{\theta}} \pi(a | s, \bar{\theta})$$

② REINFORCE (Williams, 1992)

$$\nabla J(\theta) \propto \mathbb{E}_{\pi} \left[\sum_a Q_{\pi}(s_t, a) \nabla_{\theta} \pi(a | s_t) \right] =$$

$$= \mathbb{E}_{\pi} \left[\sum_a \pi(a | s_t) Q_{\pi}(s_t, a) \frac{\nabla_{\theta} \pi(a | s_t)}{\pi(a | s_t)} \right] =$$

$$= \mathbb{E}_{\pi} \left[\underbrace{Q_{\pi}(s_t, A_t)}_{\mathbb{E}_{\pi}[G_t | s_t, A_t]} \cdot \frac{\nabla_{\theta} \pi(A_t | s_t)}{\pi(A_t | s_t)} \right] =$$

$$= \mathbb{E}_{\pi} \left[G_t \cdot \frac{\nabla_{\theta} \pi(A_t | s_t)}{\pi(A_t | s_t)} \right]$$

$$\bar{\theta}_{t+1} := \bar{\theta}_t + \alpha G_t \frac{\nabla_{\theta} \pi(A_t | S_t, \bar{\theta}_t)}{\pi(A_t | S_t, \bar{\theta}_t)}$$

$$\bar{\theta}_{t+1} := \bar{\theta}_t + \alpha G_t \cdot \nabla_{\theta} \log \pi(A_t | S_t, \bar{\theta}_t)$$

③ REINFORCE with baseline

$$\nabla J(\theta) \propto \sum_s \mu(s) \sum_a \underbrace{Q_{\pi}(s, a) - b(s)}_{-b(s)} \nabla_{\theta} \pi(a | s, \theta)$$

$$\begin{aligned} \sum_s \mu(s) \sum_a (Q_{\pi}(s, a) - b(s)) \nabla_{\theta} \pi(a | s, \theta) &= \\ &= \sum_s \mu(s) \sum_a Q_{\pi}(s, a) \nabla_{\theta} \pi - \sum_s \mu(s) b(s) \nabla_{\theta} \left[\sum_a \pi(a | s, \theta) \right] \end{aligned}$$

$\overset{\text{"1"}}{\sum_a \pi(a | s, \theta)}$
 $\underset{\text{"0"}}{\sum_a \pi(a | s, \theta)}$

$$\nabla J(\theta) \propto \sum_s \mu(s) \sum_a \underbrace{(Q_{\pi}(s, a) - b(s))}_{\text{baseline}} \nabla_{\theta} \pi(a | s, \theta)$$

REINFORCE w/ baseline:

$$\bar{\theta}_{t+1} = \bar{\theta}_t + \alpha \underbrace{(G_t - \underbrace{b(s_t)}_{\hat{V}(s_t | \bar{w})})}_{\hat{V}(s_t | \bar{w})} \nabla_{\theta} \log \pi(A_t | S_t, \bar{\theta}_t)$$

- $\alpha_{\theta}, \alpha_w, \bar{\theta}, \bar{w}$

- loop:

- gen. episode $S_0, A_0, R_1, S_1, A_1, \dots, S_{T-1}, A_{T-1}, R_T$ no $\pi, G=0$

- for $t = T-1, \dots, 1, 0$:

- $\underline{G} := R_t + \gamma G$

- $\bar{w} := \bar{w} + \alpha_w \cdot (G - \hat{V}(s_t | \bar{w})) \nabla_{\bar{w}} \hat{V}(s_t | \bar{w})$

- $\bar{\theta} := \bar{\theta} + \alpha_{\theta} (G - \hat{V}(s_t | \bar{w})) \cdot \nabla_{\theta} \log \pi(A_t | S_t, \bar{\theta})$

One-step actor-critic

$$\bar{\theta}_{t+1} = \bar{\theta}_t + \alpha_{\theta} \left(R_{t+1} + \gamma \hat{V}(s_{t+1} | \bar{w}) - \hat{V}(s_t | \bar{w}) \right) \nabla_{\theta} \log \pi(A_t | s_t, \bar{\theta})$$

④ Continuous actions

$a \in \mathbb{R}$

~~$s \sim p(s|a)$~~ $s = [a]$

$$\pi(a | s, \theta) = \frac{1}{\sqrt{2\pi} \cdot \underbrace{\sigma(s|\theta)}_{\text{" } \bar{\theta}_{\sigma}^T \bar{x}(s)}} \cdot e^{-\frac{1}{2\sigma(s|\theta)^2} (a - \underbrace{\mu(s|\theta)}_{\text{" } \bar{\theta}_{\mu}^T \bar{x}(s)})^2}$$

$$\nabla_{\bar{\theta}_{\mu}, \bar{\theta}_{\sigma}} \log \pi(a | s, \theta) = \text{gnp}$$