

Lecture-13

- Today MC $\rightarrow$ TD $\rightarrow$ TD(λ)

* MC estimation of Action-Values

$Q(s,a)$ = average return from the first time action 'a' taken in state 's' in each episode

Problem - What if π never selects action 'a' in states

- Need for estimation of value of all actions at each state, not only action that we currently favor.
- Solution 1: Exploring starts
 - Randomize so δ & A_0 such that every (s,a) has a non-zero probability.
 - Sometimes not possible

• Solution 2: stochastic policy.

— Non-zero prob. to every 'a' in every state

$$\forall s, a \quad \pi(s, a) > 0$$

* MC control with exploring starts

Property — avoid infinite episodes in evaluation step by improving after each episode

— Accumulate Return over all episode

Pseudocode —

Initialize : for all $s \in S, a \in A$
 $q(s, a) \leftarrow$ arbitrary
 $\pi(s) \leftarrow$ arbitrary
 $Returns(s, a) \leftarrow$ empty list

Repeat forever:

- ① Generate an episode using exploring starts & π
- ② for each (s, a) appearing in episode,
 $G \leftarrow$ Returns following first occurrence
 $q(s, a) \leftarrow \text{mean}(Returns(s, a))$

② for each $s \in S$ in the episode

$$\pi(s) \in \arg \max_a q(s, a)$$

* ϵ -Greedy MC control

Pseudo code:

$q(s, a) \leftarrow \text{arbitrary}$

$\pi(s, a) \leftarrow \text{arbitrary}$

Returns(s, a) $\leftarrow$ empty list

Repeat forever:

① Generate episode using π

② for each (s, a)

$G \leftarrow$

③ For each $s \in S$, in episode

$$a^* \leftarrow \arg \max_a q(s, a)$$

$$\pi(s, a) = \begin{cases} 1 - \epsilon + \frac{\epsilon}{|A|} & \text{if } a = a^* \\ \frac{\epsilon}{|A|} & \text{if } a \neq a^* \end{cases}$$

Properties: —

- By variations in policy imp. them
- (SGB: —)

* TD (Temporal difference) learning

Properties — Like MC it learns directly from experiences, i.e. w/o knowledge of P & R

— Like DP, update estimates based on other estimates

• Another MC algorithm

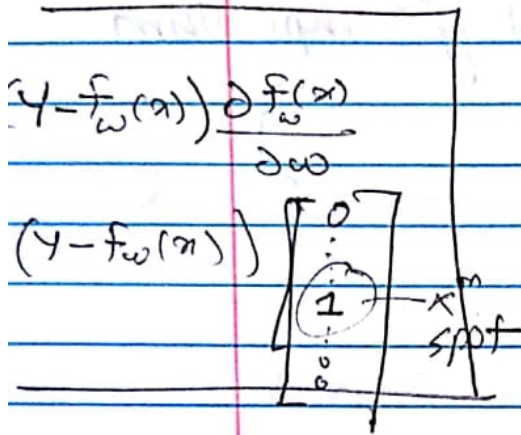
$$f(x) \leftarrow f(x) + \alpha (Y - f(x))$$

↗ R.V. (Target)

α : step size, or learning rate

$$f_w(x), \quad L(w) = \mathbb{E} \left[\frac{1}{2} (Y - f_w(x))^2 \right]$$

$$\omega \leftarrow \omega + \mathbb{E} \left[(Y - f_{\omega}(x)) \cdot \frac{\partial f_{\omega}(x)}{\partial \omega} \right] \times \alpha$$

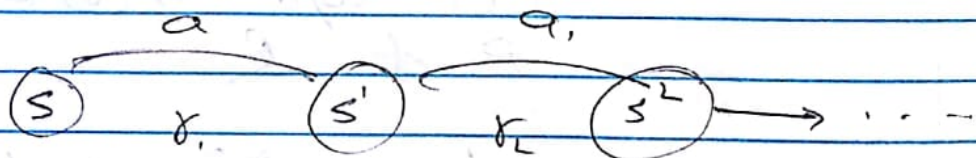


Similarly,

$$V(s_t) \leftarrow V(s_t) + \alpha \cdot (G_t - V(s_t))$$

$$\dots \text{where, } G_t = \sum_{k=0}^{\infty} \gamma^k R_{t+k}$$

- Problem - wait till end of episode
- Que - how can we update before the end of the episode?



Bellman eqⁿ:

$$V(s) = \sum_a \pi(s,a) \sum_{s'} P(s,a,s') \left[R(s,a,s') + \gamma V^{\pi}(s') \right]$$

$$= \mathbb{E} \left[R_t + \gamma V^{\pi}(s_{t+1}) \mid s_t = s \right]$$

TD - target

$$= \mathbb{E} \left[\left(\sum_{k=0}^{\infty} \gamma^k R_{t+k} \right) \mid s_t = s, \pi \right]$$

MC - Target

TD - update (given transition (s, a, s'))

$$V(s) \leftarrow V(s) + \alpha \cdot (r + \gamma \cdot V(s') - V(s))$$

Equivalently

$$V(s_t) \leftarrow V(s_t) + \alpha (R_t + \gamma \cdot V(s_{t+1}) - V(s_t))$$

→ actually

what happened?

→ Expectation

Reward
prediction
error

$$\begin{aligned} \text{TD - error} &= (R_t + \gamma \cdot V(s_{t+1}) - V(s_t)) \\ &= \delta_t \end{aligned}$$

$$\therefore V(s) \leftarrow V(s) + \alpha \delta_t$$

* How much better did things turn out than I expected?

- Is TD method SGD?

$$f(V) = \mathbb{E} \left[\frac{1}{2} (R_t + \gamma \cdot V(s_{t+1}) - V(s_t))^2 \right]$$

$$= \mathbb{E} \left[\frac{1}{2} \delta_t^2 \right]$$

$$\therefore Df(V) = S_t \cdot \frac{\partial S_t}{\partial V}$$

$$= (R_t + \gamma V(s_{t+1}) - V(s_t)) \begin{pmatrix} \frac{\partial V(s_{t+1})}{\partial V} \\ -\frac{\partial V(s_t)}{\partial V} \end{pmatrix}$$

$$= (R_t + \gamma V(s_{t+1}) - V(s_t)) \left[\begin{matrix} \begin{matrix} 0 \\ 0 \\ 0 \\ 0 \end{matrix} \\ \begin{matrix} 1 \\ 0 \\ 0 \\ 0 \end{matrix} \end{matrix} \right]_{s_{t+1}} - \begin{matrix} \begin{matrix} 0 \\ 0 \\ 0 \\ 0 \end{matrix} \\ \begin{matrix} 1 \\ 0 \\ 0 \\ 0 \end{matrix} \end{matrix} \right]_{s_t}$$

$$* V(s_t) \leftarrow V(s_t) + \alpha S_t$$

drop this $\rightarrow V(s_{t+1}) \leftarrow V(s_{t+1}) - \alpha \cdot \gamma \cdot S_t$