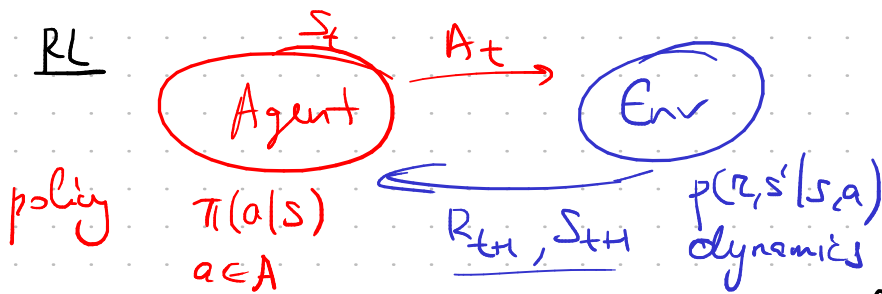


① RL



MDP:

$S_0, A_0, R_1, S_1, A_1, R_2, S_2, A_2, \dots$

Return $G_t = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots = \sum_{k=0}^{\infty} \gamma^k R_{t+k+1}$

$$\begin{cases} V_{\pi}(s) = E_{\pi}[G_t | S_t = s] & V_{*}(s) = \max_{\pi} V_{\pi}(s) \\ Q_{\pi}(s, a) = E_{\pi}[G_t | S_t = s, A_t = a] & Q_{*}(s, a) = \max_{\pi} Q_{\pi}(s, a) \end{cases}$$

Bellman equations

V_{π}, Q_{π} (lin) V_{*}, Q_{*} (max(-))

1) ~~Micro S, max(s, a)~~

2) Micro S, max(s, a) ← Approximate RL

3) He gives $p(r, s' | s, a)$

② Monte-Carlo RL

$$V_{\pi}(s) = E_{\pi}[G_t | S_t = s] \approx \frac{1}{R} \sum_{r=1}^R G_t^{(2)}$$

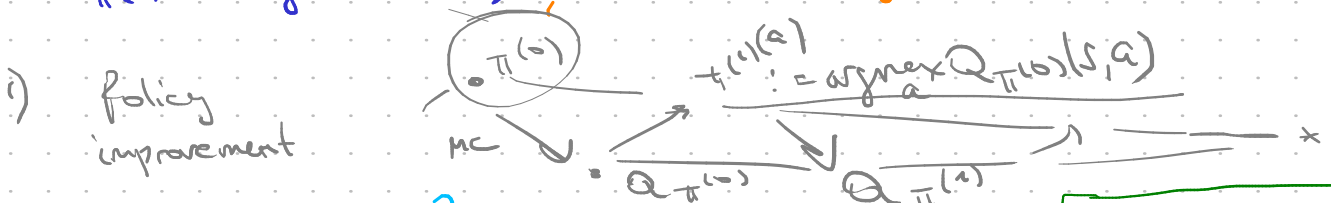
Agent, π :

- 1) $S_{10}, A_{10}, R_{11}, S_{11}, A_{11}, R_{12}, \dots$ — $S_t, A_t, R_{t+1}, S_{t+1}, A_{t+1}, \dots$
- 2) $S_{20}, A_{20}, \dots$

$G_t = R_{t+1} + \gamma V_{t+1} - V_t$
 $G_t | s, a \rightarrow Q(s, a)$

Monte Carlo estimation (π)

- init $V(s)$ with π
- repeat:
 - generate $S_0, A_0, R_1, S_1, A_1, \dots, S_T, A_T, R_{T+1}$ — $R_T, S_T = \text{goal}$
 - for $t = T-1, \dots, 0$:
 - $G := \gamma \cdot G + R_{t+1}$
 - add G to the list $\text{Returns}(S_t) / \text{Returns}(S_t, A_t)$
- $V_{\pi}(s) := \text{Avg}(\text{Returns}(s))$, $Q_{\pi}(s, a) := \text{Avg}(\text{Returns}(s, a))$



2) $\text{Returns}(s, a) \neq \emptyset$?

exploring starts / soft strategies

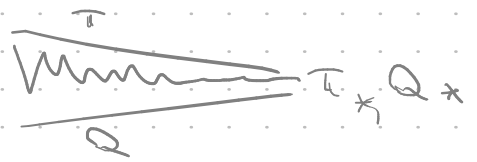
3) first visit vs. every visit MC.

Sutton, Barto
 RL: an introduction

1998

2018

Monte Carlo control w/ expl. starts



- init π, Q
- repeat:

- sample S_0, A_0
- generate $S_0, A_0, R_1, S_1, A_1, R_2, S_2, A_2, \dots, S_T = \mathcal{B}$
- for $t = T-1, \dots, 0$:
 - $G := R_t + \gamma V(S_{t+1})$
 - add G to Returns (S_t, A_t)
 - $Q(S_t, A_t) := \text{Avg}(\text{Ret}(S_t, A_t))$
 - $\pi(S_t) := \underset{a}{\text{argmax}} Q(S_t, a)$

③ On-policy vs. off-policy MC control

when no π'
overl. Q_π

when no π'
overl. $Q_{\pi'}$

on-policy MC control

$$\pi(S_t) := \begin{cases} \underset{a}{\text{argmax}} Q(S_t, a), & 1-\epsilon \\ \text{Unif}(a), & \epsilon \end{cases}$$

$$\pi(a|S_t) = \frac{\epsilon}{|A|}, a \neq a_* \quad , \quad (1-\epsilon) \frac{\epsilon}{|A|}, a = a_*$$

then $\forall s \quad Q_\pi(s, \pi'(s)) \geq V_\pi(s)$ if $\pi, \pi' = \epsilon$ -soft

$$Q_\pi(s, \pi'(s)) = \sum_a \pi'(a|s) Q_\pi(s, a) =$$

$$= \frac{\epsilon}{|A|} \sum_a Q_\pi(s, a) + (1-\epsilon) \cdot \max_a Q_\pi(s, a)$$

$$V_\pi(s) = \sum_a \pi(a|s) Q_\pi(s, a) = \frac{\epsilon}{|A|} \sum_a Q_\pi(s, a) + \sum_a \left(\pi(a|s) - \frac{\epsilon}{|A|} \right) Q_\pi(s, a)$$

$$(1-\epsilon) \cdot \sum_a \frac{\pi(a|s) - \epsilon/|A|}{1-\epsilon} \cdot Q_\pi(s, a)$$

$$\sum_a = 1 \leq \max_a Q_\pi(s, a)$$

- off-policy MC control

- generate episodes with π' , $\pi' - \epsilon$ -soft
- learn Q_π for $\pi \neq \pi'$, e.g., π - greedy

importance sampling: $E_p[f(\bar{x})] = \int \frac{p(\bar{x})}{q(\bar{x})} f(\bar{x}) d\bar{x} = E_{q(\bar{x})} \left[f(\bar{x}) \frac{p(\bar{x})}{q(\bar{x})} \right]$

$$p(\text{Traj}(\pi) | S_t) = p(A_t, R_{t+1}, S_{t+1}, A_{t+1}, \dots | S_t, \pi) =$$

$$= \underbrace{p(A_t | S_t, \pi)}_{\pi(a|s)} \underbrace{p(R_{t+1}, S_{t+1} | S_t, A_t)}_{p(r, s' | s, a)} p(A_{t+1} | S_{t+1}, \pi) p(R_{t+2}, S_{t+2} | S_{t+1}, A_{t+1}) \dots$$

$$E_\pi[G_t | S_t] = E_{\pi'} \left[G_t \cdot \frac{p(\text{Traj}(\pi))}{p(\text{Traj}(\pi'))} \right] = E \left[G_t \cdot \frac{\pi(A_t | S_t)}{\pi'(A_t | S_t)} \right]$$

$$= E_{\pi'} \left[G_t \cdot \prod_{k=t}^{T-1} \frac{\pi(A_k | S_k)}{\pi'(A_k | S_k)} \right]$$



- init $S_{t-3} \rightarrow S_{t-2} \rightarrow S_{t-1} \rightarrow S_t$

- repeat: $G_t = 0, w := 1$

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

- $G := \gamma \cdot G + R_{t+1}$

- $w := w \cdot \frac{\pi(A_t | S_t)}{\pi'(A_t | S_t)}$

- add $w \cdot G$ to $\text{Returns}(S_t, A_t)$, update Q

- $\pi'(a | S_t) := \begin{cases} a_x, & 1-\epsilon \\ \text{rand}, & \epsilon \end{cases}, \pi(a | S_t) = a_x$

$$\text{Avg}(\text{Returns}(\dots)) = \frac{1}{|\text{Ret}|} \sum_t G_t \cdot \frac{\pi(A_t | S_t)}{\pi'(A_t | S_t)}$$

$$= \frac{\sum_t G_t \cdot w_t^{\pi, \pi'}}{\sum_t w_t^{\pi, \pi'}}$$

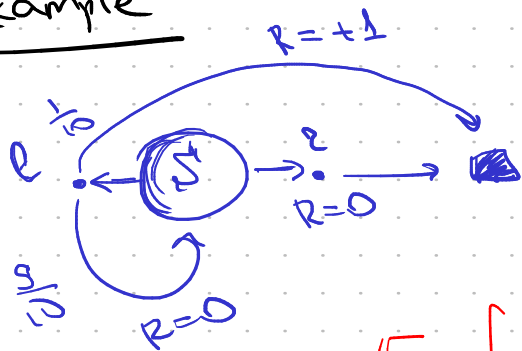
greedy: $\text{can } \pi(a|s) > 0$
 no $\pi'(a|s) > 0$

$$S_0 \xrightarrow{A_0} S_1 \xrightarrow{A_1} S_2 \xrightarrow{A_2} \dots \xrightarrow{A_{t-1}} S_t \xrightarrow{A_t} S_{t+1}$$

$$\frac{\pi \cdot \pi_1 \cdot \pi_2 \cdot \dots \cdot 1 - \pi}{\pi' \cdot \pi'_1 \cdot \pi'_2 \cdot \dots \cdot 1}$$

④ Example

$\gamma=1$



$$Q(s, \text{right}) = 0$$

$$Q_+(s, \text{left}) = 1$$

$$\pi = \text{left}$$

$$\pi' = \left(\frac{1}{2}, \frac{1}{2}\right)$$

$$E_{\pi'} [G_t \cdot w_t^{\pi, \pi'}] = 1$$

$$\text{Var}[\dots] = E[\dots^2] - E[\dots]^2 = E[\dots^2] - 1$$

$$E_{\pi'} \left[G_0 \prod_{t=0}^{T-1} \frac{\pi(A_t|S_t)}{\pi'(A_t|S_t)} \right] = \sum_{\text{Traj}} p(\text{Traj}|\pi') \cdot G_0^2 \left(\prod_{t=0}^{T-1} \dots \right)^2 =$$

$$= \sum_{\text{Traj}: G=1} \left(w_0^{\pi, \pi'} \right)^2 \cdot p(\text{Traj}|\pi') =$$

$$= \frac{1}{2} \cdot \left(\frac{1}{10}\right) \cdot \left(\frac{1}{4/2}\right)^2 + \frac{1}{2} \cdot \frac{9}{10} \cdot \frac{1}{2} \cdot \left(\frac{1}{10}\right) \cdot \left(\frac{1}{1/2 \cdot 4/2}\right)^2 +$$

$$+ \frac{1}{2} \cdot \frac{9}{10} \cdot \frac{1}{2} \cdot \frac{9}{10} \cdot \frac{1}{2} \cdot \left(\frac{1}{10}\right) \cdot \left(\frac{1}{1/2 \cdot 1/2 \cdot 4/2}\right)^2 + \dots$$

$$= \frac{1}{10} \sum_{t=0}^{\infty} \left(\frac{1}{2}\right)^{t+1} \cdot \left(\frac{9}{10}\right)^t \cdot 2^{2(t+1)} = \frac{1}{5} \cdot \sum_{t=0}^{\infty} \left(\frac{9}{5}\right)^t = \infty$$

$$\begin{bmatrix} \mu_1 \\ \vdots \\ \mu_n \end{bmatrix} \begin{bmatrix} a_1 \\ \vdots \\ a_n \end{bmatrix} \begin{bmatrix} \mu_1 \\ \vdots \\ \mu_n \end{bmatrix}$$

- Greedy, $\mu_i^{(0)} = 10$

⑤ TD-learning (temporal difference)

$$V_{\pi}(s) = E_{\pi} [G_t | S_t = s] = E_{\pi} [R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots | S_t = s]$$

$$= E_{\pi} [R_{t+1} | S_t = s] + \gamma E_{\pi} [R_{t+2} + \gamma^2 R_{t+3} + \dots | S_t = s]$$

$\approx R_{t+1}$ $\approx G_{t+1}$

$S_0, A_0, R_1, S_1, A_1, \dots, S_t, A_t, R_{t+1}, S_{t+1}$

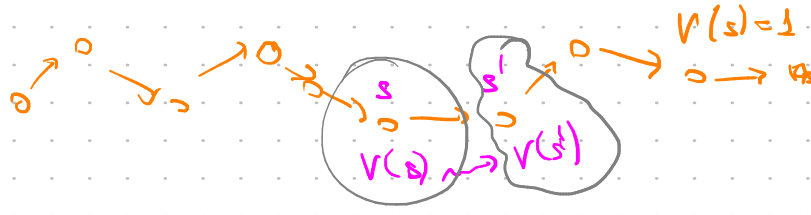
bootstrap

Experience tuples:

(s, a, r, s')

$$V_{\pi}(s') \approx R_{t+1} + \gamma \cdot V_{\pi}(s_{t+1})$$

$\sim p(R, s' | s, a)$



TD estimation

$$V_{\pi}(s) := V_{\pi}(s) + \alpha \cdot (\underbrace{r + \gamma \cdot V_{\pi}(s')}_{\text{target}} - \underbrace{V_{\pi}(s)}_{\text{old value}})$$

$$Q_{\pi}(s, a) := Q_{\pi}(s, a) + \alpha (r + \gamma \cdot Q_{\pi}(s', a') - Q_{\pi}(s, a))$$



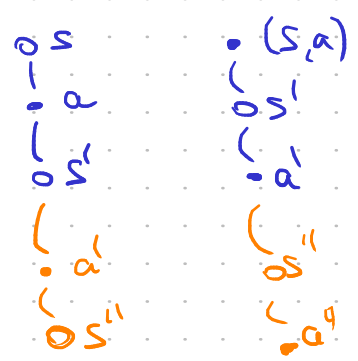
$$D = \left\{ \begin{array}{l} (s' \rightarrow \text{terminal}, +1) \\ (s' \rightarrow \text{terminal}, +1) \\ (s' \rightarrow \text{terminal}, 0) \\ (s' \rightarrow \text{terminal}, +1) \\ (s \rightarrow s' \rightarrow \text{terminal}, 0) \end{array} \right\}$$

TD:
 $V(s') \approx 3/5$
 $V(s) \approx 3/5$

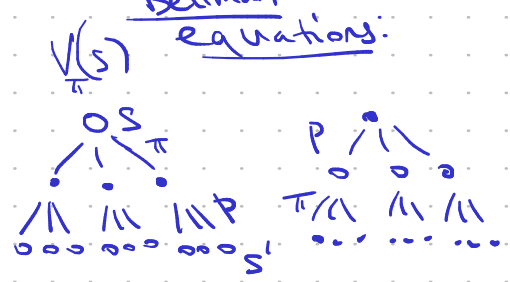
MC: $V(s') = 3/5$
 $V(s) = 0$

Backup diagrams

TD:



Bellman Equations:



MC



n-step TD:

$$Q_{\pi}(S_t, A_t) = E[R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots | S_t, A_t]$$

$$Q_{\pi}(S_t, A_t) := Q_{\pi}(S_t, A_t) + \alpha (R_{t+1} + \gamma R_{t+2} + \gamma^2 Q_{\pi}(S_{t+2}, A_{t+2}) - Q_{\pi}(S_t, A_t))$$

TD: 1-step $\rightarrow$ 2-step $\rightarrow$ 3-step $\rightarrow \dots \rightarrow$ MC

⑥ On-policy TD control / Sarsa

Sarsa:

exp. tuple

(s, a, r, s')

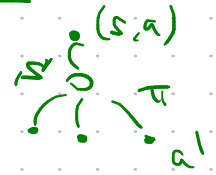
- bootstrap a no $\pi(a|s) = \begin{cases} \text{argmax}_a Q(s, a), & 1-\epsilon \\ \text{unif}, & \epsilon \end{cases}$
- no boot $r, s',$ but a'
- $Q(s, a) := Q(s, a) + \alpha (r + \gamma \cdot Q(s', a') - Q(s, a))$



start

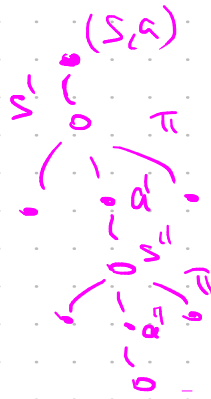
Q-learning

finish



Expected Sarsa

$$Q(s, a) := Q(s, a) + \alpha (r + \gamma \cdot \underbrace{\sum_{a'} \pi(a'|s') Q(s', a')}_{\text{target}} - Q(s, a))$$

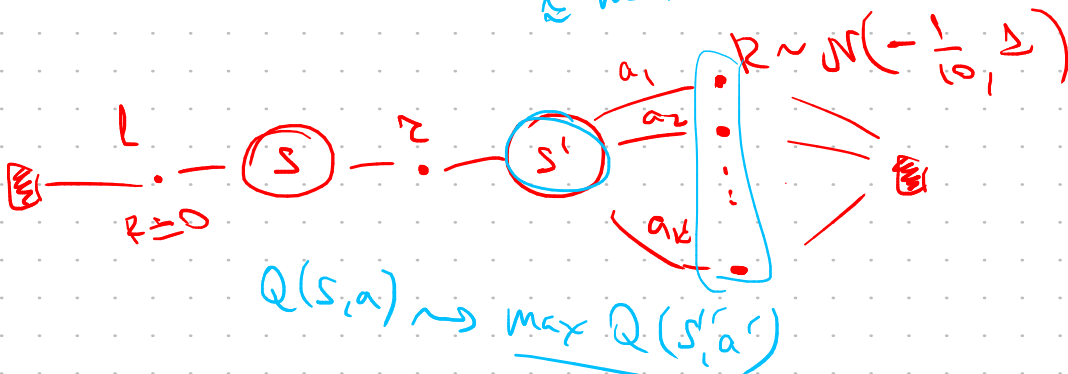
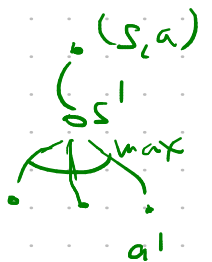


$$\text{Target} = r + \gamma \cdot \sum_{a_1 \neq a'} \pi(a_1|s') Q(s', a_1) + \gamma \cdot \pi(a'|s') (r' + \gamma \cdot \sum_{a_2} \pi(a_2|s'') Q(s'', a_2))$$

⑦ Q-learning / off-policy TD control (Watkins, 1989)

$$Q(s, a) := Q(s, a) + \alpha (r + \gamma - \max_{a'} Q(s', a') - Q(s, a))$$

$$Q(s, a) = E_{\pi} [G_t | s, a] \quad E[\max] \neq \max E[\dots]$$



Winner's curse

Agreement

$$\hat{x}_1, \hat{x}_2, \dots, \hat{x}_k$$

$$\hat{x}_k \sim \mathcal{N}(\hat{x}_k | x, \sigma^2)$$

x

max

Double Q-learning

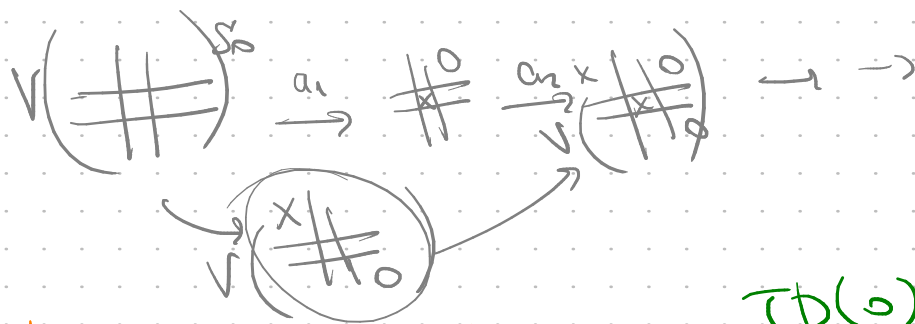
$Q_1(s,a)$

$Q_2(s,a)$

$$(s,a,r,s') \quad p=\frac{1}{2}: \quad Q_1(s,a) := Q_1(s,a) + \alpha [r + \gamma \cdot Q_2(s', \arg\max_{a'} Q_1(s',a')) - Q_1(s,a)]$$

$$p=\frac{1}{2}: \quad Q_2(s,a) := -\infty - Q_1(s', \arg\max_{a'} Q_2(s',a')) \leftarrow -\infty$$

After states



as_1



$\tau b(0)$

$\tau b(x)$



⑧ Approximate RL

$$s \rightarrow \boxed{\quad} \rightarrow \hat{V}(s, \bar{w}) \approx V(s)$$

$$s \rightarrow \boxed{\quad} \rightarrow \hat{Q}(s, a, \bar{w}) \approx Q(s, a)$$

$$L(\bar{w}) = \sum_s (V(s) - \hat{V}(s, \bar{w}))^2 \xrightarrow{\bar{w}} \min$$

GD: $\bar{w} := \bar{w} - \alpha \cdot \nabla_{\bar{w}} L(\bar{w})$

$$L(\bar{w}) = \mathbb{E}_{\mathcal{S}} [L(\mathcal{S}, \bar{w})]$$

SGD: $\bar{w} := \bar{w} + \alpha \cdot \sum_{i=1}^m \underbrace{2(V(s_i) - \hat{V}(s_i, \bar{w}))}_{\text{error}} \cdot \nabla_{\bar{w}} \hat{V}(s_i, \bar{w})$

$$\checkmark L(\bar{w}) = \sum_{s \in \mathcal{S}} \mu(s) (V(s) - \hat{V}(s, \bar{w}))^2 \xrightarrow{\bar{w}} \min$$

$$\mu(s) = \Pr_{\tau} [S_t = s] = \text{"how often, how often"}$$

Gradient MC estimation

$$\bar{w} := \bar{w} + \alpha \left(G_t - \hat{V}(s_t, \bar{w}) \right) \cdot \nabla_{\bar{w}} \hat{V}(s_t, \bar{w})$$

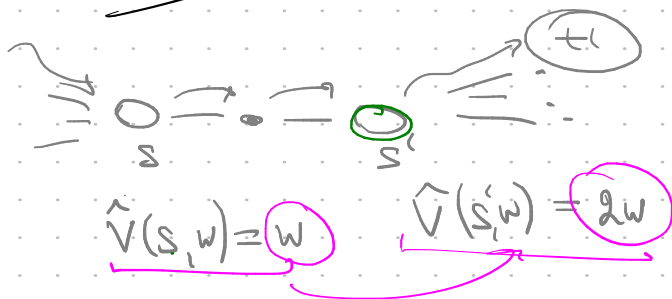
Semi-gradient TD estimation

$$\bar{w} := \bar{w} + \alpha \left(R_{t+1} + \gamma \hat{V}(s_{t+1}, \bar{w}) - \hat{V}(s_t, \bar{w}) \right) \nabla_{\bar{w}} \hat{V}(s_t, \bar{w})$$

Semi-gradient Sarsa / on-policy TD control

$$\bar{w} := \bar{w} + \alpha \left(R_{t+1} + \gamma \hat{Q}(s_{t+1}, A_{t+1}, \bar{w}) - \hat{Q}(s_t, A_t, \bar{w}) \right) \nabla_{\bar{w}} \hat{Q}(s_t, A_t, \bar{w})$$

~~Off-policy semi-gradient TD control~~



Deadly triad

- bootstrap / TD
- off-policy
- approximation / semi-gradient

Double DQN deep Q-network

experience replay

$$\{(s, a, r, s')\}$$