

Lecture 5. Model free Control

Prediction: $V_{\pi}(s)$

Control: improve policy $V_{\pi'}(s) \geq V_{\pi}(s)$
 $\pi' \geq \pi$

S.1 Model free

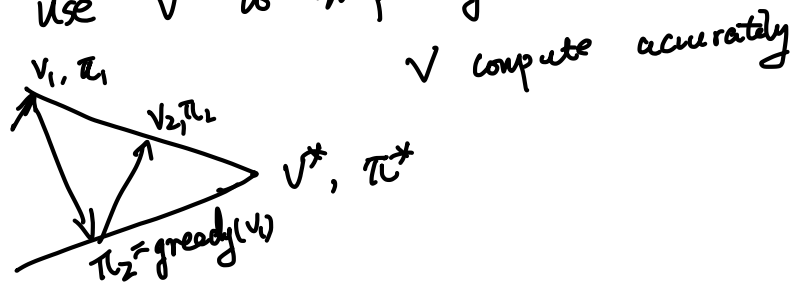
① boat, wave, wind, complicate

② go: know every detail, space too large

S.2 on-policy: learn on job
 learn about π from samples from π

off-policy: learn policy $\pi \neq$ sample policy μ

S.3 use V to improve policy



policy iteration ① MC. to evaluate V

② TD evaluate V

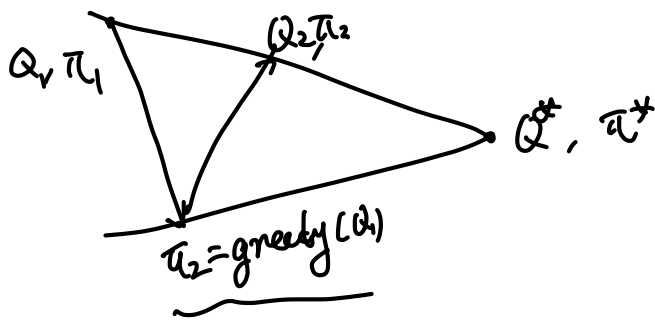
$$\pi'(s) = \arg \max_a R_s^a + \gamma \underbrace{p_{ss'}^a}_{\leftarrow \text{MDP}} V(s')$$

Model free

$$\pi'(s) = \arg \max_a Q(s, a)$$

state-action value function

5.4 Policy iteration with action-value function



① MC

② TD

① 3 4
② 2 2

example



Greedy:

left Reward 0

$$V(\text{left}) = 0$$

right R: +1

$$V(\text{right}) = 1$$

right R: +3

$$V(\text{right}) = \frac{1+3}{2} = 2$$

⋮
right

巨大问题

5.5 ϵ greedy exploration

$$\pi(a|s) = \begin{cases} 1-\epsilon + \frac{\epsilon}{m} & \text{if } a^* = \underset{a}{\operatorname{argmax}} Q(s, a) \\ \frac{\epsilon}{m} & \text{others} \end{cases}$$

$1-\epsilon$ greedy

ϵ uniform

Theorem ϵ -greedy policy improvement

ϵ -greedy π , π' , ϵ -greedy w.r.t q_π

$$\Rightarrow V_{\pi'}(s) \geq V_\pi(s)$$

Pf: $V_{\pi'}(s) \geq q_\pi(s, \pi'(s)) = \sum_a \bar{a}'(a|s) q_\pi(s, a)$

Lect 3 DP

$$\begin{aligned} \boxed{V_\pi(s) \leq q_\pi(s, \pi'(s))} &= \mathbb{E}_{\pi'} [R_{t+1} + \gamma V_\pi(S_{t+1}) \mid S_t = s] \\ &\leq \mathbb{E}_{\pi'} [R_{t+1} + \gamma q_\pi(S_{t+1}, \pi'(S_{t+1})) \mid S_t = s] \\ &\leq \mathbb{E}_{\pi'} [R_{t+1} + \gamma R_{t+2} + \gamma^2 q_\pi(S_{t+2}, \pi'(S_{t+2})) \mid S_t = s] \\ &\leq \mathbb{E}_{\pi'} [R_{t+1} + \gamma R_{t+2} + \dots \mid S_t = s] = \underline{V_{\pi'}(s)} \\ V_\pi(s) &\leq V_{\pi'}(s) \end{aligned}$$

$$V_{\pi'}(s) \geq q_{\bar{\pi}}(s, \pi'(s)) = \sum_a \bar{a}'(a|s) q_{\bar{\pi}}(s, a)$$

$$= \frac{\epsilon}{m} \sum_a q_{\bar{\pi}}(s, a) + (1-\epsilon) \max_a q(s, a)$$

$$\sum_a \bar{a}(a|s) q_{\bar{\pi}}(s, a) \geq \frac{\epsilon}{m} \sum_a q_{\bar{\pi}}(s, a) + (1-\epsilon) \sum_a \frac{\bar{a}(a|s) - \frac{\epsilon}{m}}{1-\epsilon} q_{\bar{\pi}}(s, a)$$

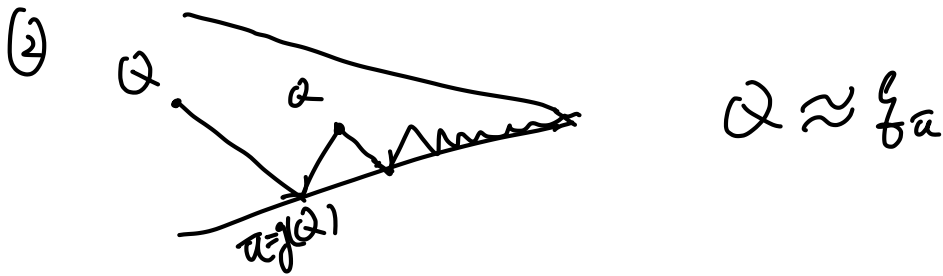
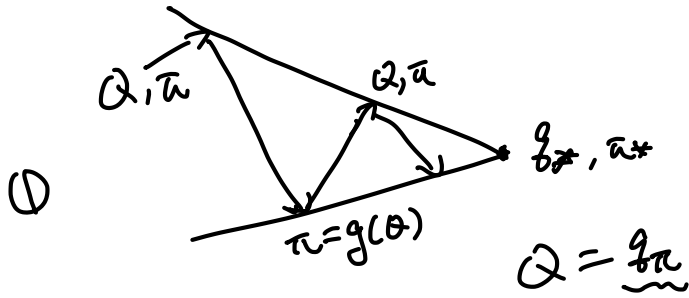
$$\sum_a \frac{\bar{a}(a|s) - \frac{\epsilon}{m}}{1-\epsilon}$$

$$\stackrel{!}{=} \frac{1-\epsilon}{1-\epsilon}$$

$$= 1$$

$$= \sum_a \bar{a}(a|s) q_{\bar{\pi}}(s, a) = V_{\bar{\pi}}(s)$$

S.6. MC. policy iteration



ϵ -greedy

S.6.1 Definition (GLIE)
Greedy in the limit with infinite exploration

① $N_K(S, a) \rightarrow \infty$ as $K \rightarrow \infty$ π K : episode

② $\lim_{K \rightarrow \infty} \pi_K(a|s) = 1$ ($a = \arg \max_{a'} Q_K(s, a')$)

S.6.2 GLIE MC control

① sample k -th episode using $\pi = \{S_1, A_1, R_1, S_2, A_2, R_2, \dots\}$

② S_t, A_t .

$$N(S_t, A_t) \leftarrow N(S_t, A_t) + 1$$

$$Q(S_t, A_t) \leftarrow Q(S_t, A_t) + \frac{1}{N(S_t, A_t)} (G_t - Q(S_t, A_t))$$

③

$$\boxed{\epsilon \leftarrow \frac{1}{K}}$$

$$a \leftarrow \epsilon\text{-greedy}(Q)$$

Theorem: GLIE MC control converges to the optimal action-value function

$$Q(s, a) \rightarrow q_*(s, a)$$

§. 7. MC vs TD control

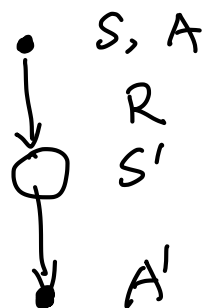
TD: lower variance

: online

: incomplete sequences

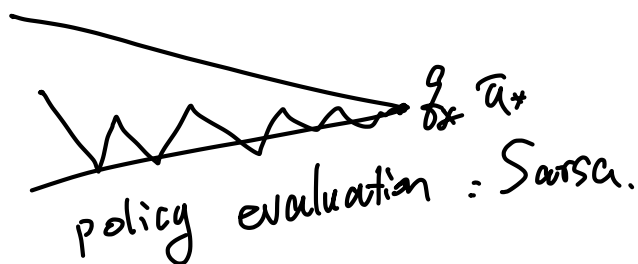
use TD to replace MC

§. 7.1



SARSA

$$Q(s, a) \leftarrow Q(s, a) + \alpha (R + \gamma Q(s', a') - Q(s, a))$$



Sarsa Algorithm for On-Policy Control

Initialize $Q(s, a), \forall s \in \mathcal{S}, a \in \mathcal{A}(s)$, arbitrarily, and $Q(\text{terminal-state}, \cdot) = 0$
 Repeat (for each episode):
 Initialize S
 Choose A from S using policy derived from Q (e.g., ϵ -greedy)
 Repeat (for each step of episode):
 Take action A , observe R, S'
 Choose A' from S' using policy derived from Q (e.g., ϵ -greedy)
 $Q(S, A) \leftarrow Q(S, A) + \alpha [R + \gamma Q(S', A') - Q(S, A)]$
 $S \leftarrow S'; A \leftarrow A'$;
 until S is terminal

Theorem convergence of Sarsa

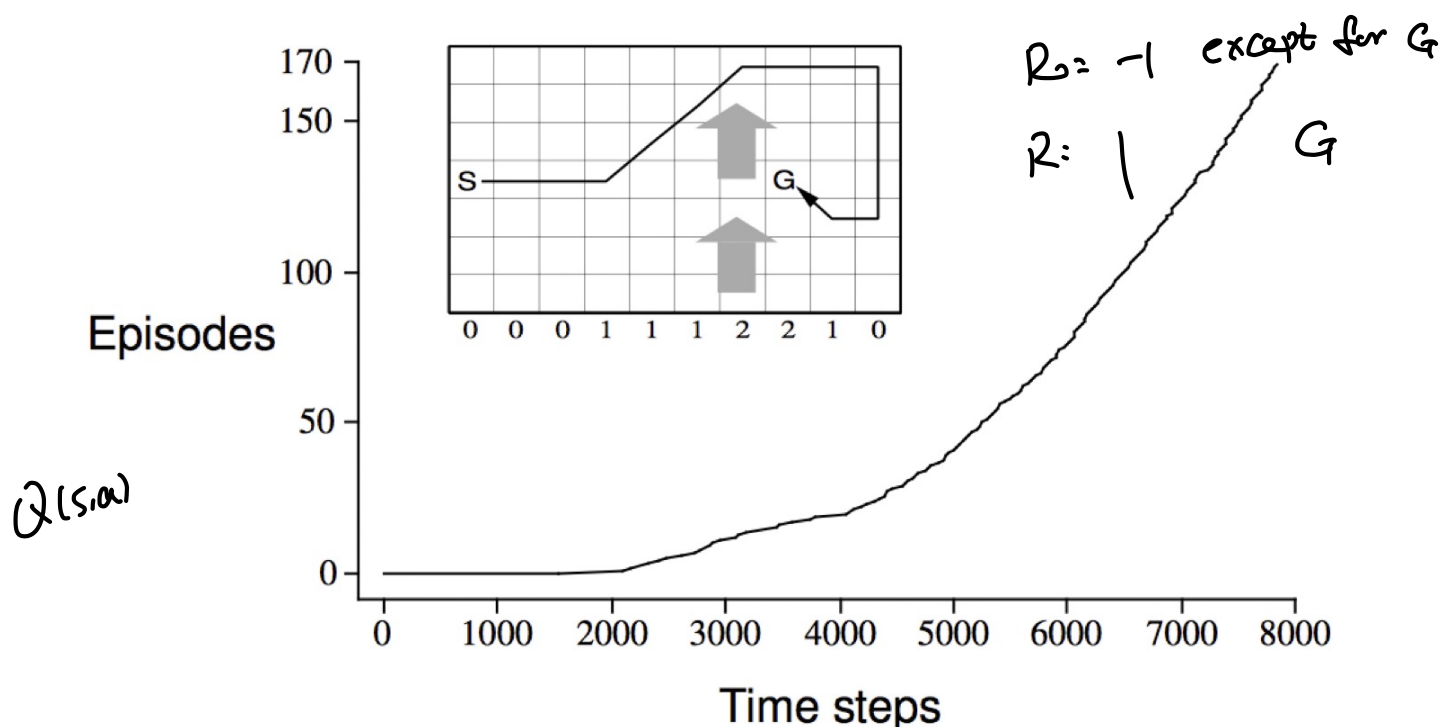
Condition:

① GLIE

② $\sum_{t=1}^{\infty} \alpha_t = +\infty$

$\sum_{t=1}^{\infty} \alpha_t^2 < +\infty$

Sarsa on the Windy Gridworld



5.8. n-step Sarsa.

$$n=1 \text{ (Sarsa)} \quad g_t^{(1)} = R_{t+1} + \gamma Q(S_{t+1})$$

$$n=2 \quad g_t^{(2)} = R_{t+1} + \gamma R_{t+2} + \gamma^2 Q(S_{t+2})$$

$\vdots$

$$n=\infty \text{ MC} \quad g_t^{(\infty)} = R_{t+1} + \gamma R_{t+2} + \dots + \gamma^{T-1} R_T$$

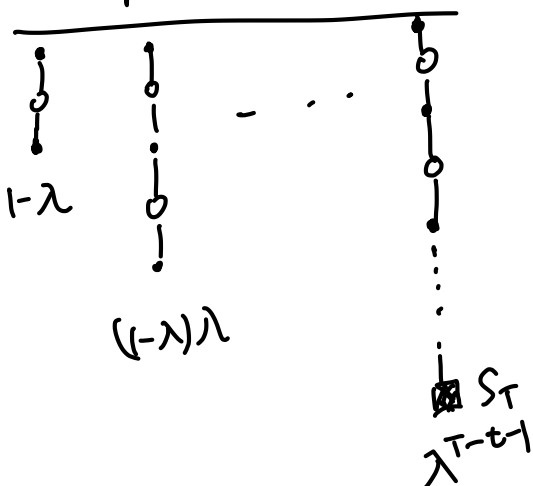
Define n-step Q-return

$$g_t^{(n)} = R_{t+1} + \gamma R_{t+2} + \dots + \gamma^{n-1} R_{t+n} + \gamma^n Q(S_{t+n})$$

n-step Sarsa

$$Q(S_t, A_t) \leftarrow Q(S_t, A_t) + \alpha (g_t^{(n)} - Q(S_t, A_t))$$

5.9. Forward view Sarsa(λ)



$$1 = \sum_{n=1}^{\infty} (1-\lambda) \lambda^{n-1} \quad \frac{1-\lambda^{\infty}}{1-\lambda} = 1-\lambda$$

$$Q(S_t, A_t) \leftarrow Q(S_t, A_t) + \alpha (g_t^{\lambda} - Q(S_t, A_t))$$

$$g_t^{\lambda} = (1-\lambda) \sum_{n=1}^{\infty} \lambda^{n-1} g_t^{(n)}$$

5.10 Backward View Sarsa (λ)

Sarsa(λ) Algorithm

eligibility
trace

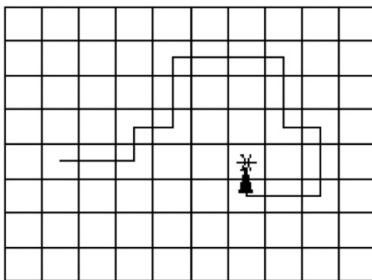
```

Initialize  $Q(s, a)$  arbitrarily, for all  $s \in \mathcal{S}, a \in \mathcal{A}(s)$ 
Repeat (for each episode):
   $E(s, a) = 0$ , for all  $s \in \mathcal{S}, a \in \mathcal{A}(s)$ 
  Initialize  $S, A$ 
  Repeat (for each step of episode):
    Take action  $A$ , observe  $R, S'$ 
    Choose  $A'$  from  $S'$  using policy derived from  $Q$  (e.g.,  $\epsilon$ -greedy)
     $\delta \leftarrow R + \gamma Q(S', A') - Q(S, A)$ 
     $E(S, A) \leftarrow E(S, A) + 1$ 
    For all  $s \in \mathcal{S}, a \in \mathcal{A}(s)$ :
       $Q(s, a) \leftarrow Q(s, a) + \alpha \delta E(s, a)$ 
       $E(s, a) \leftarrow \gamma \lambda E(s, a)$ 
     $S \leftarrow S'; A \leftarrow A'$ 
  until  $S$  is terminal
  
```

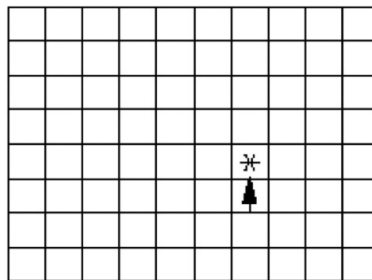
Sarsa(λ) Gridworld Example

$R=0$

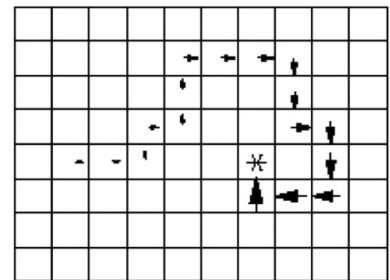
Path taken



Action values increased
by one-step Sarsa



Action values increased
by Sarsa(λ) with $\lambda=0.9$



S.11 off-policy learning

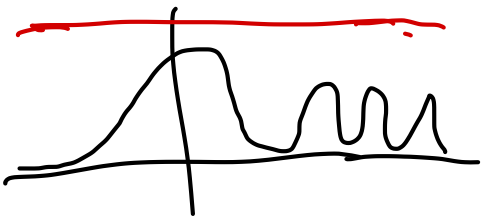
target policy $\pi(a|s) \rightarrow V_a(s)$ or $Q_a(s, a)$

$$\{S_1, A_1, R_1, S_2, A_2, \dots, S_T\} \sim \mathcal{M}$$

- ① learn from observing other agents
- ② Re-use experience
- ③ learn about optimal policy while following exploratory policy
- ④ Train multiple policies

S.12. Importance sampling

$$\begin{aligned} E_{x \sim p} [f(x)] &= \sum p(x) f(x) \\ &= \sum f(x) \frac{p(x)}{q(x)} q(x) \\ &= E_{x \sim q(x)} \left[f(x) \frac{p(x)}{q(x)} \right] \end{aligned}$$



S.13. Imp. sampling for off-policy

$$G_t \quad \mu(A_t|S_t) \quad \mu(A_{t+1}|S_{t+1}) \quad \dots$$

$$G_t^{\pi/\mu} = \frac{\pi(A_t|S_t)}{\mu(A_t|S_t)} \frac{\pi(A_{t+1}|S_{t+1})}{\mu(A_{t+1}|S_{t+1})} \dots \frac{\pi(A_T|S_T)}{\mu(A_T|S_T)} G_t$$

$$MC \quad V(S_t) \leftarrow V(S_t) + \alpha (G_t^{\pi/\mu} - V(S_t))$$

$$TD \quad V(S_t) \leftarrow V(S_t) + \alpha \left(\frac{\pi(A_t|S_t)}{\sum_a \pi(a|S_t)} (R_{t+1} + \gamma V(S_{t+1}) - V(S_t)) \right)$$

$$5.14 \quad \pi(S_{t+1}) = \arg \max_{a'} Q(S_{t+1}, a')$$

$$R_{t+1} + \gamma Q(S_{t+1}, A')$$

$$= R_{t+1} + \gamma Q(S_{t+1}, \arg \max_{a'} Q(S_{t+1}, a'))$$

$$= R_{t+1} + \gamma \max_{a'} Q(S_{t+1}, a')$$

Q-learning Control: (no policy explicitly)

$$Q(S, A) \leftarrow Q(S, A) + \alpha [R + \gamma \max_{a'} Q(S', a') - Q(S, A)]$$

Q-Learning Algorithm for Off-Policy Control

Initialize $Q(s, a), \forall s \in S, a \in \mathcal{A}(s)$, arbitrarily, and $Q(\text{terminal-state}, \cdot) = 0$

Repeat (for each episode):

Initialize S

Repeat (for each step of episode):

Choose A from S using policy derived from Q (e.g., ϵ -greedy)

Take action A , observe R, S'

$$Q(S, A) \leftarrow Q(S, A) + \alpha [R + \gamma \max_a Q(S', a) - Q(S, A)]$$

$S \leftarrow S'$;

until S is terminal

Relationship Between DP and TD (2)

Full Backup (DP)	Sample Backup (TD)
Iterative Policy Evaluation <i>value</i> $V(s) \leftarrow \mathbb{E}[R + \gamma V(S') \mid s]$	TD Learning <i>value</i> $V(S) \stackrel{\alpha}{\leftarrow} R + \gamma V(S')$
Q-Policy Iteration <i>采样</i> $Q(s, a) \leftarrow \mathbb{E}[R + \gamma Q(S', A') \mid s, a]$	Sarsa <i>采样</i> $Q(S, A) \stackrel{\alpha}{\leftarrow} R + \gamma Q(S', A')$
Q-Value Iteration $Q(s, a) \leftarrow \mathbb{E}\left[R + \gamma \max_{a' \in \mathcal{A}} Q(S', a') \mid s, a\right]$	Q-Learning $Q(S, A) \stackrel{\alpha}{\leftarrow} R + \gamma \max_{a' \in \mathcal{A}} Q(S', a')$ <i>终极.</i>

where $x \stackrel{\alpha}{\leftarrow} y \equiv x \leftarrow x + \alpha(y - x)$