

二. MDPs

Markov process / Markov chain.

时间, 状态都离散的过程 → 马尔科夫链.

马尔科夫过程(链)由元组 (S, P) 组成

S 为有限状态的集合.
 P 为状态转移矩阵 $P_{ss'} = P[S_{t+1}=s' | S_t=s]$
 转移矩阵 P ($|S| \times |S|$ 阶方阵)

(大写的 S_t 为随机状态, 小写的 s 为确定状态),

< episodes (幕) >

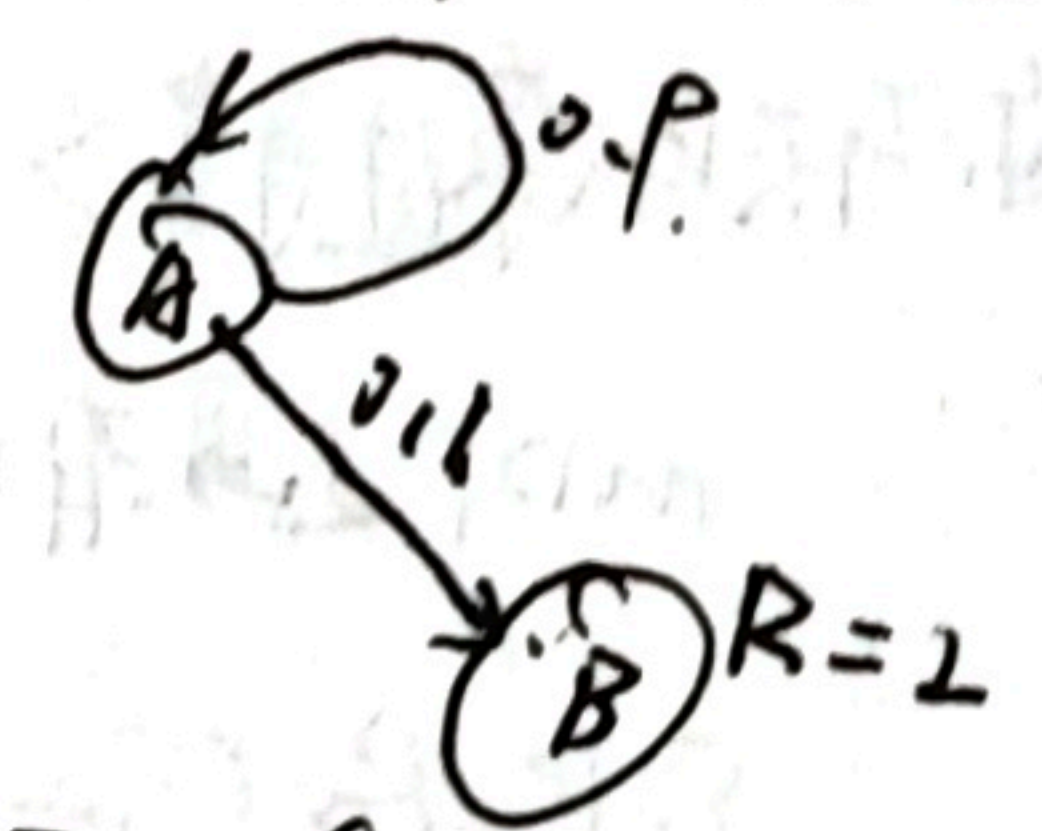
从初始状态 $S_1 = c_1$ 开始, 从马尔科夫链中采样子序列, 并终止于终止状态, 那个子序列就叫做 episode.

< 马尔科夫奖励过程 >

Markov Reward Process (MRP) 由元组 (S, P, R, γ) 组成.

S, P 在 Markov chain 中已有, 额外的 R, γ 为奖励函数, $R_s = E[R_{t+1} | S_t=s]$ 为折扣因子, $\gamma \in [0, 1]$

< 奖励函数 > $R_s = E[R_{t+1} | S_t=s]$ (一般认为到达状态 s 之后, 系统才给出 Reward, 所以称为 $(t+1)$ 的下标, 如果在某 s , 奖励没有随机性, 如 $R_s = R_{t+1} = 2$, 这里方便记为 $R=2$.)



< 回报, return >

从 t 时刻的 S_t 开始, 直至终止状态时, 所有奖励的折现之和 C_t 称为 return.

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

PPT 中称之为 accumulative discounted reward. $\sum_{t=0}^{\infty} \gamma^t r_t$ t 时刻有值之前 s 的 reward 也参与计算.

Return 针对某一条幕说明/计算.

< 价值函数, Value Function > (单一条 episodes 的 return 高并不能说明这个状态好).

价值函数 $V(s)$ 给出状态 s 的长期价值 (long-term value).

input: 状态 output: 价值.

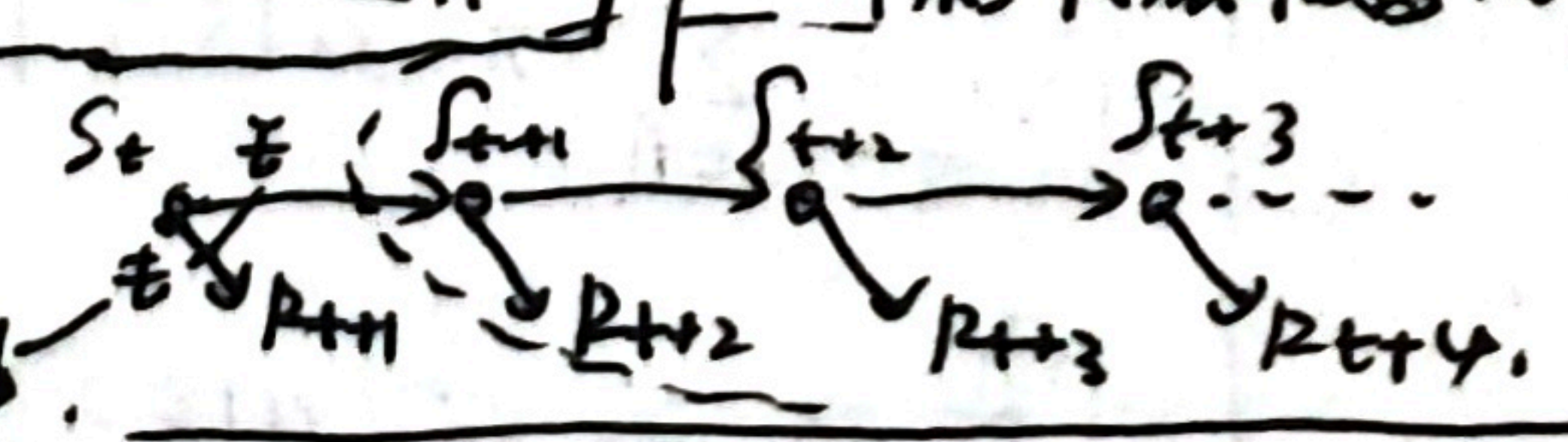
在 MRP 中, 一个 s 的期望回报就是这个 s 的价值函数.

$$V(s) = E[C_t | S_t=s]$$

去除了 Reward 随机性
 去除了状态转移的随机性.

如何计算?

< 贝尔曼方程, Bellman Equation >



根据状态
 运行一步
 即期望.

$$\begin{aligned} V(s) &= E[C_t | S_t=s] \\ &= E[R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots | S_t=s] \\ &= E[R_{t+1} + \gamma (R_{t+2} + \gamma R_{t+3} + \dots) | S_t=s] \\ &= E[R_{t+1} + \gamma C_{t+1} | S_t=s] \end{aligned}$$

$$\Rightarrow V(s) = E[R_{t+1} + \gamma V(S_{t+1}) | S_t=s]$$

递归.

期望去除 $\Rightarrow E[R_{t+1} + \gamma V(S_{t+1}) | S_t=s]$

$$v(s) = E[\underbrace{R_{t+1}}_{\text{即时奖励}} + \underbrace{\gamma v(s_{t+1})}_{\text{下一状态折扣值}} | S_t = s]$$

The other form of Bellman Equation:

$$v(s) = \underbrace{R_s}_{\text{当前状态的奖励}} + \underbrace{\gamma}_{\text{折扣因子}} \underbrace{\sum_{s' \in S_{t+1}} P_{ss'} v(s')}_{\text{下一状态的价值函数}} \quad \xrightarrow{\text{转移函数}}$$

<贝尔曼方程的矩阵形式> $\downarrow$ 贝尔曼方程为线性方程组

$$\vec{v} = R + \gamma P \vec{v}$$

$$\begin{bmatrix} v(1) \\ \vdots \\ v(n) \end{bmatrix} = \begin{bmatrix} R_1 \\ \vdots \\ R_n \end{bmatrix} + \gamma \begin{bmatrix} P_{11} & \dots & P_{1n} \\ \vdots & & \vdots \\ P_{n1} & \dots & P_{nn} \end{bmatrix} \begin{bmatrix} v(1) \\ \vdots \\ v(n) \end{bmatrix}$$

$$\vec{v} = (I - \gamma P)^{-1} R \quad \text{计算复杂度为 } O(15^3)$$

直接求解适用于小型子马尔可夫奖励过程

<马尔可夫决策过程> (MDP)

MDP是具有决策性质的 MRP, 由元组 (S, A, P, R, γ) 组成:

S, P, R, γ 在 MRP 中已有提及, A 为有限动作的集合.

而且 P 和 R_s 有所改变

$$\left. \begin{aligned} P_{ss'}^a &= P[S_{t+1} = s' | S_t = s, A_t = a] \\ R_s^a &= E[R_{t+1} | S_t = s, A_t = a] \end{aligned} \right\}$$

不只是依赖于随机性决定, 还对 action 有关系.

<策略, Policy>

Policy 定义了智能体的行为.

马尔可夫决策过程中, Policy 仅取决于当前状态.

input: s output: a 动作

$$\pi(a|s) = P[A_t = a | S_t = s]$$

在策略下, 随机性由

$$\begin{bmatrix} \pi(a_1|s) \\ \vdots \\ \pi(a_n|s) \end{bmatrix} \Rightarrow \text{策略!}$$

有了 $\pi(a|s)$, 即 Policy, P 和 R 发生改变, Markov chain, 事 MRP, MDP 都发生变化: 马尔可夫决策过程

$$\begin{aligned} \vec{P}^{\pi} &= \sum_{a \in A} \pi(a|s) \vec{P}_{ss'}^a \\ \vec{R}^{\pi} &= \sum_{a \in A} \pi(a|s) \vec{R}_s^a \end{aligned}$$

<状态价值函数 (State-Value)> ← 对不同的 s 求解的

$$V_{\pi}(s) = E_{\pi}[G_t | S_t = s]$$

遵从策略 π

$V_{\pi}(s)$ 为从 s 出发遵循 π 得到的期望回报

如何求解?

由于 π 的加入, $V_{\pi}(s)$ 可以区分不同 π 和不同 s 的区别, 但是如果要评估同一个 π , 就没有办法, 由此引入:

<动作价值函数, action-Value> ← 对同一个 π 求解的

$$Q^{\pi}(s, a) = E_{\pi}[G_t | S_t = s, A_t = a]$$

遵从策略 π

$Q^{\pi}(s, a)$ 为从 s 出发遵循 π 对当前状态 s 执行 a 得到的期望回报

<贝尔曼期望方程, Bellman Expectation Equation> Bellman Equation + action

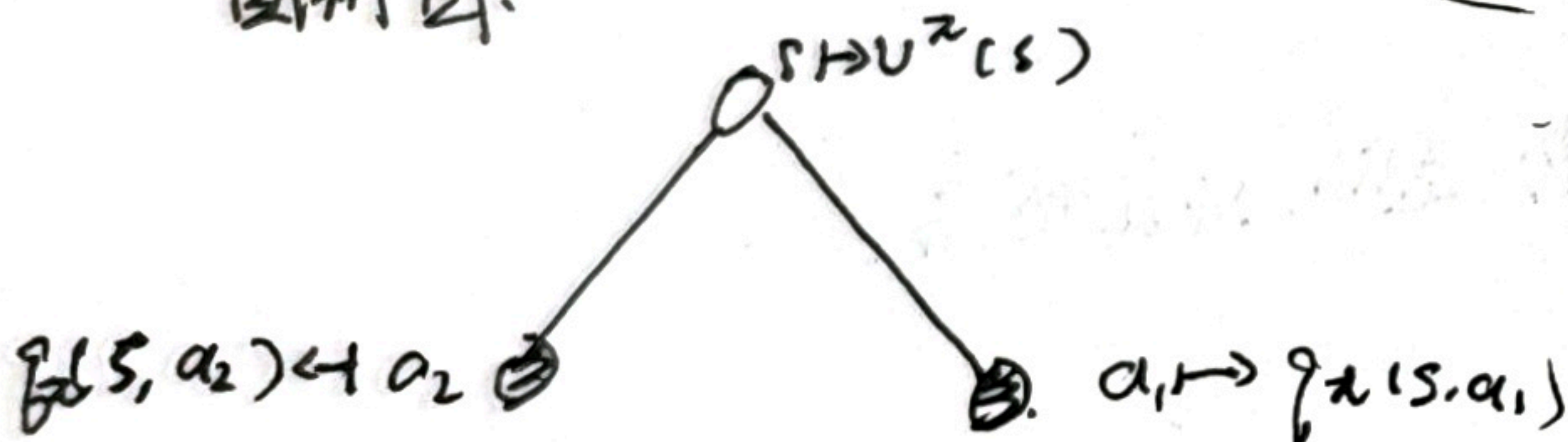
$$V_{\pi}(s) = E_{\pi}[R_{t+1} + \gamma V_{\pi}(S_{t+1}) | S_t = s] \leftarrow \begin{array}{l} a \text{ 不固定, 随机 } A \\ \text{也随机, Different} \end{array}$$

$$Q^{\pi}(s, a) = E_{\pi}[R_{t+1} + \gamma Q^{\pi}(S_{t+1}, A_{t+1}) | S_t = s, A_t = a] \leftarrow \begin{array}{l} a \text{ 固定, 随机 } A \\ \text{随机} \end{array}$$

图例:

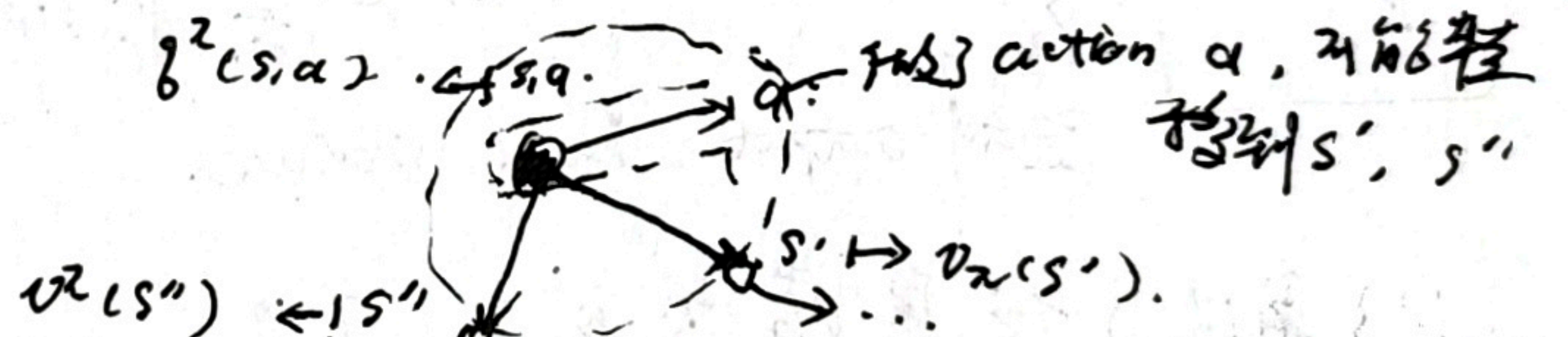
$V_{\pi}(s)$ 比 $Q^{\pi}(s, a)$ 多了 A_t 的随机性

← 因为 a 动作的随机性



$$V_{\pi}(s) = \sum_{a \in A} \pi(a|s) Q^{\pi}(s, a)$$

消除特定动作 a 的随机性

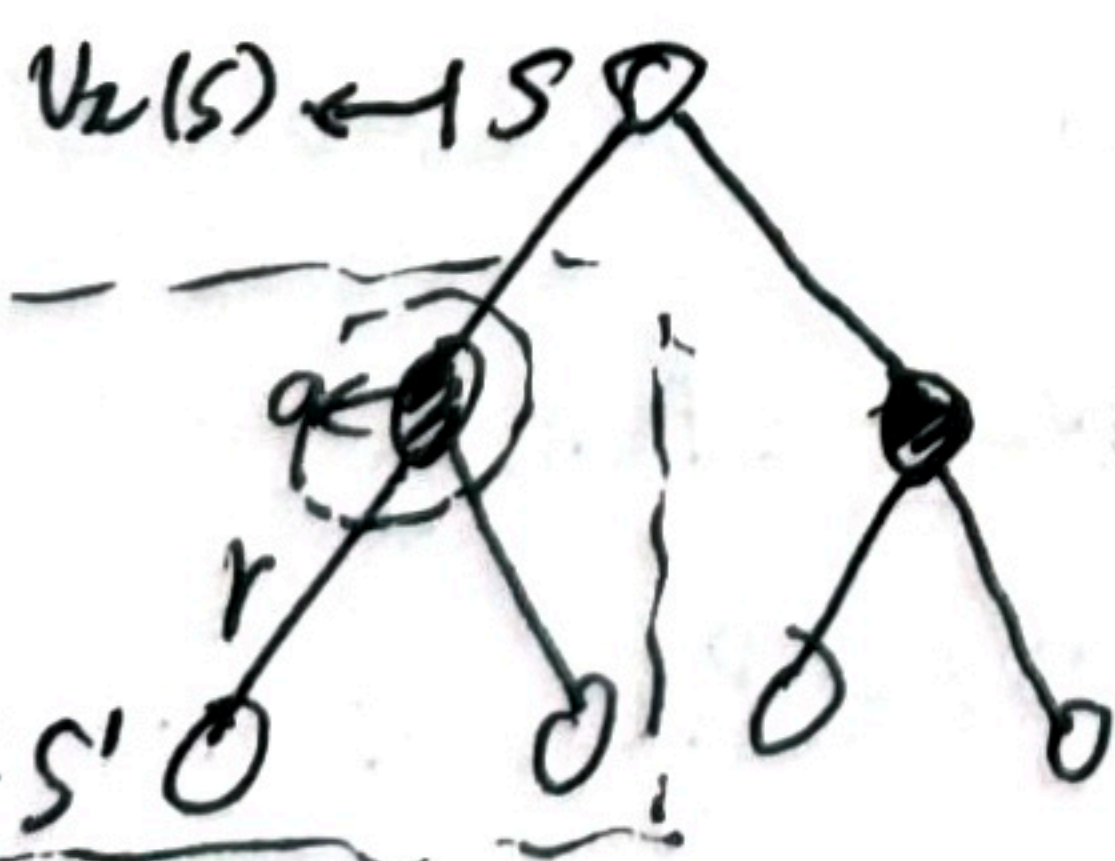


$$Q^{\pi}(s, a) = E_{\pi}[R_{t+1} + \gamma \sum_{s' \in S} P_{ss'}^a V_{\pi}(s')]$$

$$Q^{\pi}(s, a) = R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a V_{\pi}(s')$$

无随机性

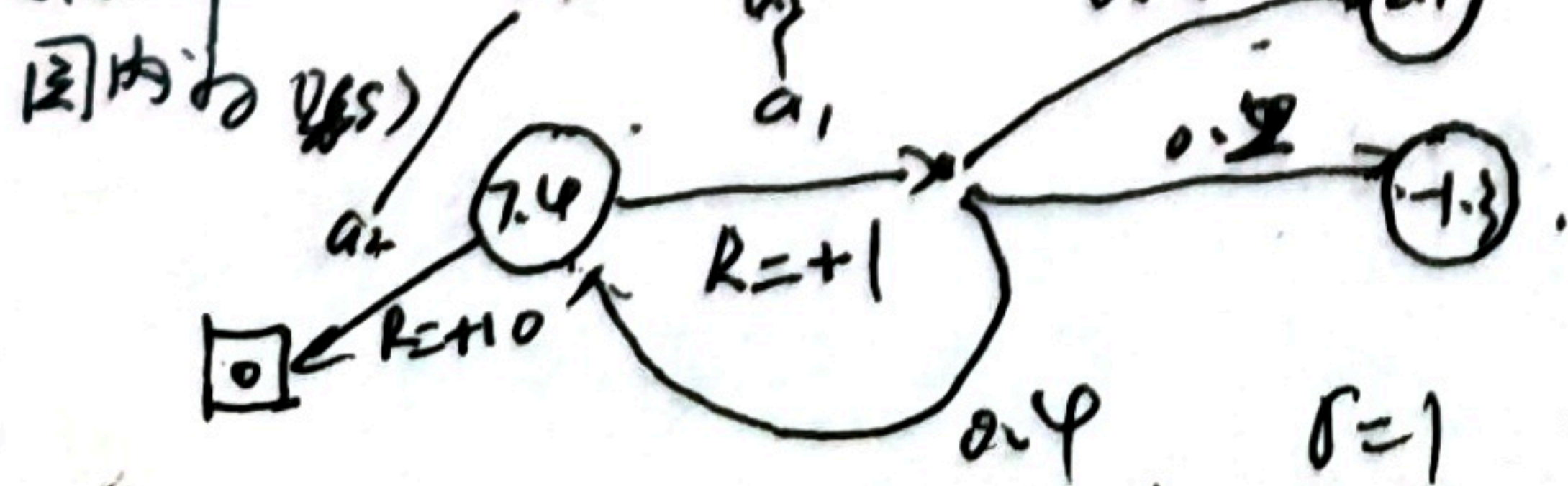
加上特定动作 a 的随机性, 指定 a



$$V_{\pi}(s) = \sum_{a \in A} \pi(a|s) \left(R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a V_{\pi}(s') \right)$$

$$Q^{\pi}(s, a) = R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a \sum_{a' \in A} \pi(a'|s) Q^{\pi}(s', a')$$

example



$$V_{\pi}(s) = 0.5 \times (1 + 0) + 0.5 \times (1 + 0.4 \times 2.7 + 0.2 \times 1.3 + 0.4 \times 7.4)$$

2.47488

< 最优价值函数, Optimal Value Function >

最优值函数

最优状态价值函数

$$V_{\pi}^*(s) = \max_{\pi} V_{\pi}(s)$$

最优动作价值函数

$$Q_{\pi}^*(s) = \max_{\pi} Q_{\pi}(s, a)$$

最优价值函数已知, MDP 求解

都是对策略取最优

Bellman equation

$$V_{\pi}^*(s) =$$

$$Q_{\pi}^*(s, a) = E[R_{t+1} + \gamma \max_{a'} Q^*(s', a') | s, a]$$

< 最优策略, Optimal Policy >

< 偏序 > 对于任意状态 s 都有 $V_{\pi}(s) \geq V_{\pi'}(s)$, $\pi \geq \pi'$

在 MDP 中, 至少存在一个策略不低于其他策略, 即 $\pi^* \geq \pi$, $\forall \pi$, $\pi^* \rightarrow$ optimal policy

$$\pi^*(a|s) = \begin{cases} 1 & \text{if } a = \arg \max_{a \in A} Q_{\pi^*}(s, a) \\ 0 & \end{cases}$$

arg $\Rightarrow$ argument, 自变量

$\Rightarrow$ 只留最优的动作, 舍弃其他

< 贝尔曼最优方程 > (Bellman Optimality Equation) 非线性, 不通过矩阵求解

$$V_{\pi}(s) = \sum_{a \in A} \pi(a|s) Q_{\pi}(s, a) = \max_a Q_{\pi}(s, a) \rightarrow V_{\pi}^*(s)$$

$$Q_{\pi}(s, a) = R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a V_{\pi}(s')$$

other form

$$V_{\pi}^*(s) = \max_a (R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a V_{\pi}^*(s'))$$

$$Q_{\pi}^*(s, a) = R_s^a + \gamma \sum_{s' \in S} P_{ss'}^a \max_{a'} Q_{\pi}^*(s', a')$$



三. From Q-learning to Actor Critic.

$Q(s, a)$ 的求解: $\rightarrow$ Sarsa 算法更新 $Q(s, a)$, 策略评估, 用贪婪优化.

< Sarsa > on-policy $S, A \rightarrow R, S' \rightarrow A \rightarrow Sarsa$. ϵ -Greedy

$$Q(s, a) = E[G_t | S_t = s, A_t = a]$$

$$= E[R_{t+1} + \gamma Q(S_{t+1}, A_{t+1}) | S_t = s, A_t = a]$$

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

G_t 用 ϵ -greedy 策略

在策略中的 Sarsa 依赖于 Bellman 期望方程.

- 1: 随机初始化 $Q(s, a)$, $\forall s \in S, a \in A(s)$. 令 $Q(\text{terminal-state}, \cdot) = 0$.
- 2: REPEAT (对于每个 episode):
- 3: 初始化 S .
- 4: 使用从 Q 得到的策略 (eg. ϵ -greedy) 在状态 S 时选择动作 A .
- 5: REPEAT (对于 episode 中的每一步):
- 6: 执行动作 A , 观察得到 R, S'
- 7: 根据从 Q 中得到的策略 (eg. ϵ -greedy) 在状态 S' 时选择动作 A' .
- 8: $Q(s, a) \leftarrow Q(s, a) + \alpha [R + \gamma Q(s', a') - Q(s, a)]$
- 9: $S \leftarrow S'; A \leftarrow A'$
- 10: 直到 S 为终止状态.

< On-policy, off-policy, 在线, 离线 >

< Off-policy >

agent 自身

其他玩家, 举例子.

目标策略: 用来学习的策略. 行为策略: 生成训练样本的策略.

评估目标策略 $v(a|a)$ 以计算 $V(s)$ 或 $Q(s, a)$.

同时.

同时遵循行为策略 $\mu(a|s)$ $\{S, A, R, \dots, R_T\} \sim \mu$

< Q-Learning > off-policy 依赖于 Bellman 最优 equation G_t

Q 更新的核心. 偷别人的东西.

$$Q(S_t, A_t) \leftarrow Q(S_t, A_t) + \alpha (R_{t+1} + \gamma Q(S_{t+1}, A') - Q(S_t, A_t))$$

行为策略和评估策略都随 Q 的不断更新而不断优化. ϵ -greedy 策略. 其他策略, 例子中的

目标策略是贪婪的, 行为策略是 ϵ -greedy 的.

目标策略简化:

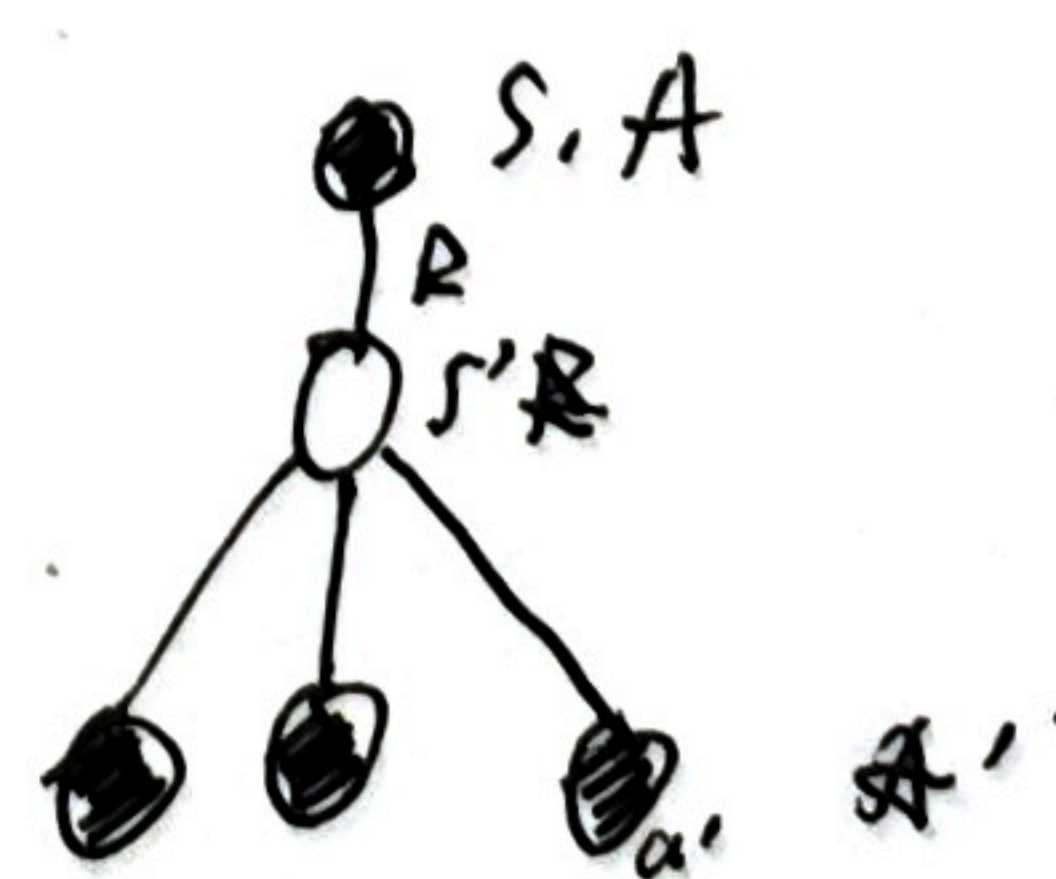
$$R_{t+1} + \gamma Q(S_{t+1}, A')$$

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

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

Q-Learning

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



Sarsamax (Q-learning)

- 1: 随机初始化 $Q(s, a)$, $\forall s \in S, a \in A(s)$, 将 $Q(\text{terminal-state}, \cdot) = 0$.
- 2: REPEAT (对于每一个 episode):
- 3: 初始化 s .
- 4: REPEAT (对于 episode 中每一步):
- 5: 根据从 Q 学到的策略 (eg. ϵ -greedy) 在 S 中选择 A .
- 6: 执行 A , 观察到 R, s'
- 7: $Q(s, A) \leftarrow Q(s, A) + \alpha [R + \gamma \max_a Q(s', a) - Q(s, A)]$
- 8: $s \leftarrow s'$
- 9: 直到 s 为终止状态.

我更新了
Q 表
操作者共同
优化 Q.

不需要自己选择动作.

未建模, 因为操作者别人.

Sarsa 走远路, 采取点收益最高的路径.

Q-Learning 走近路, 点收益不如 Sarsa, 所谓高费险中求!

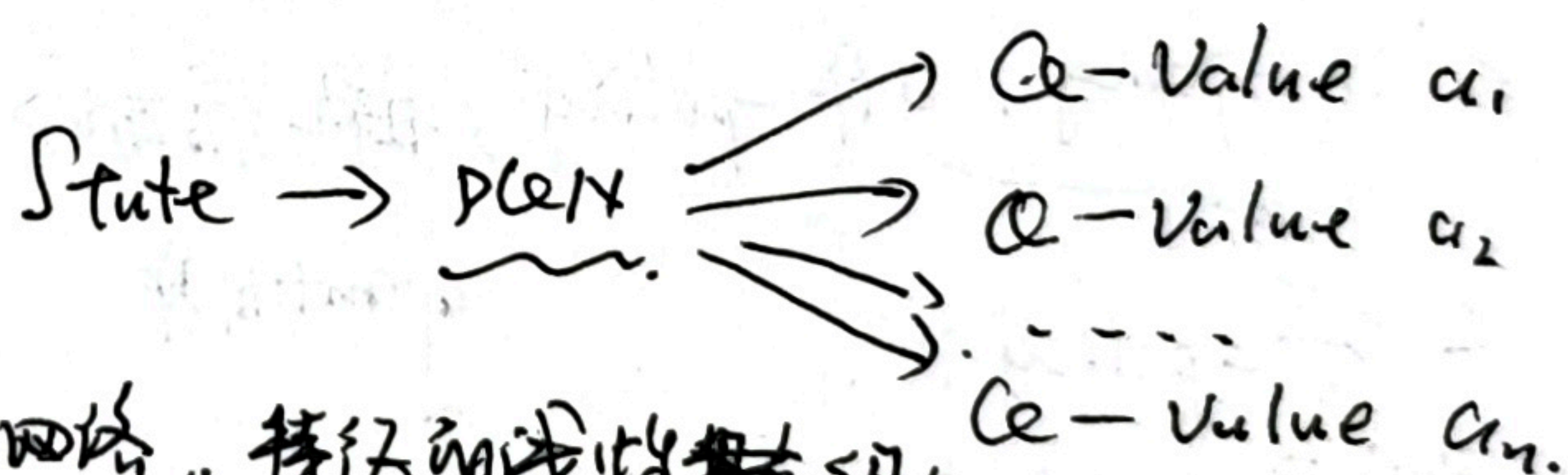
< DQN >

如果 S 有限, 上述方法可用表格或其他优化方法, 但是有无限状态呢?

以前的书会建立一个 Q -table, 学习的过程就是填表的过程, 但是太多的数据, 无法存储
解决方法: (用函数逼近价值函数, $\text{input} \rightarrow s, \text{output } Q$)

$$\hat{V}(s; w) \approx V_w(s) \quad \text{or} \quad \hat{Q}(s, a; w) \approx Q_w(s, a).$$

input
↓
system
↓
output



考虑用微分的函数逼近器。——神经网络, 特征和线性组合。
此外, 用于非平稳非独立分布的数据训练方法。Q-L 中的数据特征

< 线性最小二乘 >

线性最小二乘 = 矩阵求逆

非线性最小二乘 = < Gradient Descent >

While True.

$$\nabla_w L(w) = \nabla_w \frac{1}{N} \sum_{i=1}^N L_i(f(x_i, w), y_i)$$

$$w \leftarrow w - \alpha \nabla_w L(w) \quad \Delta w = \alpha (y_i - f(x_i, w)) \frac{\partial f}{\partial w}$$

$$GD, SGD, MBGD, GD \text{ with Momentum} = -\frac{1}{2} \alpha \nabla_w L(w)$$

折扣中.

$$w \leftarrow w + \Delta w$$

$$\Delta w = -\frac{1}{2} \alpha \nabla_w J(w) = \alpha E_w [(V_w(s) - \hat{V}(s, w)) \nabla_w \hat{V}(s, w)]$$

$$\pm SGD \text{ 折扣 } \Delta w = \alpha [(V_w(s) - \hat{V}(s, w)) \nabla_w \hat{V}(s, w)]$$

状态 s 可以用特征向量表示.