

RL 基础

集中不等式:

① Markov Inequality

$X \in \mathbb{R}^+$, 有 $\Pr(X \geq s) \leq \frac{E(X)}{s}$

$$\begin{aligned} \text{Proof: } E(X) &= \int_0^s x p(x) dx + \int_s^\infty x p(x) dx \\ &\geq \int_s^\infty x p(x) dx \\ &\geq s \int_s^\infty p(x) dx = s \Pr(X \geq s) \\ \Rightarrow \Pr(X \geq s) &\leq \frac{E(X)}{s} \end{aligned}$$

② Chebyshev's Inequality

$$\begin{aligned} \Pr(|X - E(X)| \geq \epsilon) &= \Pr((X - E(X))^2 \geq \epsilon^2) \\ &\leq \frac{E[(X - E(X))^2]}{\epsilon^2} = \frac{\text{VAR}(X)}{\epsilon^2} \end{aligned}$$

③ Chernoff Lemma $\rightarrow$ Chernoff Bound

$$\text{ma: } \Pr(X - E(X) \geq \epsilon) \leq \min_{\lambda \geq 0} \frac{E[\exp(\lambda(X - E(X)))]}{\exp(\lambda\epsilon)} \quad \text{and} \quad \Pr(X \geq \epsilon) \leq \min_{\lambda \geq 0} \frac{E[\exp(\lambda X)]}{\exp(\lambda\epsilon)}$$

$$\text{proof: } \Pr(X - E(X) \geq \epsilon) \stackrel{\lambda \geq 0}{=} \Pr(\exp(\lambda(X - E(X))) \geq \exp(\lambda\epsilon)) \leq \frac{E[\exp(\lambda(X - E(X)))]}{\exp(\lambda\epsilon)} \xrightarrow{\text{MGF}} \frac{E[\exp(\lambda X)]}{\exp(\lambda\epsilon + \lambda E(X))}$$

$$\text{Chernoff Inequality } X_i \sim p_i(\cdot) \quad \Pr\left(\sum_{i=1}^n X_i - nE[X] \geq \epsilon\right) \leq \min_{\lambda \geq 0} \left[\prod_{i=1}^n (E[\exp(\lambda(X_i - E(X_i)))]) \right] \exp^{-\lambda\epsilon}$$

$$\text{Proof: } \text{MGF}_{X_1 + X_2 + \dots + X_n}(\lambda) = \prod_{i=1}^n \text{MGF}_{X_i}(\lambda)$$

Chernoff Bound: Get the optimal λ for a particular distribution

④ Sub-Gaussian

If $X \sim \text{subG}(c^2)$, then we have $E[\exp^{\lambda X}] \leq \exp(\frac{\lambda^2 c^2}{2})$ the MGF of $N(0, c^2)$

扩展: If $X \sim N(\mu, c^2)$, then we have $E[\exp^{\lambda X}] = e^{\lambda \mu + \frac{\lambda^2 c^2}{2}}$

Proof: $M_X(t) = E[\exp^{\lambda X}] = \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}c} \exp(-\frac{(x-\mu)^2}{2c^2} + \lambda x) dx$

$$= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}c} \exp(-\frac{[x-(\mu+\lambda c^2)]^2}{2c^2} + \lambda \mu + \frac{\lambda^2 c^2}{2}) dx$$

$$= e^{\lambda \mu + \frac{\lambda^2 c^2}{2}} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}c} \exp(-\frac{[x-(\mu+\lambda c^2)]^2}{2c^2}) dx = 1$$

$$= e^{\lambda \mu + \frac{\lambda^2 c^2}{2}}$$

⑤ Hoeffding Lemma $\rightarrow$ Hoeffding Bound

Hoeffding Lemma: (Loose Version) $E[\exp^{\lambda(X-E[X])}] \leq \exp(\frac{\lambda^2(b-a)^2}{2})$, $\lambda \in \mathbb{R}$, $X \in [a, b]$

(Strict Version) $E[\exp^{\lambda(X-E[X])}] \leq \exp(-\frac{\lambda^2(b-a)^2}{8})$, $\lambda \in \mathbb{R}$, $X \in [a, b]$
Tight

Proof of the Loose Version:

$$E[\exp^{\lambda(X-E[X])}] \stackrel{X \sim \text{Rad}}{=} E_X[E_{X'}[\exp^{\lambda(X-E[X])}]] \stackrel{\text{Jensen Inequality}}{\leq} E_X[E_{X'}[\exp^{\lambda(X-X')}]]$$

$$X \text{ 是 } [a, b] \text{ 上关于 } 0 \text{ 对称的 Rademacher } = \begin{cases} 1, P=\frac{1}{2} \\ -1, P=\frac{1}{2} \end{cases}, P_r(X-X') = P_r(0) = 1 \Rightarrow E_X[E_{X'}[E_{0 \sim \text{Rad}}[\exp^{\lambda(X-X')}]]]$$

$$E_{0 \sim \text{Rad}}[\exp^{\lambda X}] = \sum_{k=0}^{\infty} \frac{\lambda^k E[X^k]}{k!} \leq E_X[E_{X'}[\exp(\frac{\lambda^2(X-X')^2}{2})]]$$

$$= \sum_{k=0,2,4,\dots}^{\infty} \frac{\lambda^k E[X^k]}{k!} \quad (\because E[X^k] = 0 \text{ if } k \text{ is odd}) \leq E_X[E_{X'}[\exp(\frac{\lambda^2(a-b)^2}{2})]] \quad (\because a \leq X-X' \leq b)$$

$$= \sum_{k=0,2,4,\dots}^{\infty} \frac{\lambda^{2k}}{(2k)!} = \exp(\frac{\lambda^2(a-b)^2}{2})$$

$$\leq \sum_{k=0,1,2,\dots}^{\infty} \frac{\lambda^{2k}}{2^k k!} = \sum_{k=0,1,2,\dots}^{\infty} (\frac{\lambda^2}{2})^k \frac{1}{k!}$$

$$= \exp(\frac{\lambda^2}{2})$$

Proof of the Tight Version:

Set $Y = X - E[X]$, we have $E[Y] = 0$, $Y \in [\frac{a-E[X]}{c}, \frac{b-E[X]}{d}] = [c, d]$

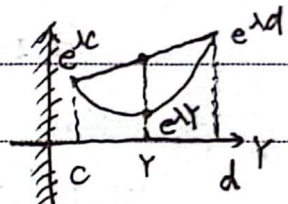
What we need to prove is: $E[\exp^{\lambda Y}] \leq \exp(\frac{\lambda^2(c-d)^2}{8})$, $\lambda \in \mathbb{R}$, $Y \in [c, d]$, $E[Y] = 0$

D Proof One (Easy to understand):

$\because e^{\lambda Y}$ is convex of Y $\therefore e^{\lambda Y} \leq \frac{Y - \frac{c-d}{2}}{\frac{d-c}{2}} e^{\lambda \frac{c+d}{2}} + \frac{Y - \frac{c+d}{2}}{\frac{d-c}{2}} e^{\lambda \frac{c+d}{2}}$

$$E[\exp^{\lambda Y}] \leq \frac{d-E[Y]}{d-c} e^{\lambda c} + \frac{E[Y]-c}{d-c} e^{\lambda d}$$

$$= \frac{d}{d-c} e^{\lambda c} + \frac{-c}{d-c} e^{\lambda d}$$

$$= \left(\frac{-c}{d-c}\right) e^{\lambda c} \left(-\frac{d}{c} + e^{\lambda(d-c)}\right)$$


$\because E[Y] = 0 \therefore c \leq 0, d \geq 0$ and c, d don't equal to 0 at the same time $\therefore c < 0, d > 0$

Set $\theta = -\frac{c}{d-c} > 0$,

$$E[\exp^{\lambda Y}] \leq \theta e^{-\lambda \theta(d-c)} \left(\frac{1}{\theta} - 1 + e^{\lambda(d-c)}\right)$$

$$= (1 - \theta + \theta e^{\lambda(d-c)}) e^{-\lambda \theta(d-c)}$$

$$\because 1 - \theta + \theta e^{\lambda(d-c)} = \theta \left(\frac{1}{\theta} - 1 + e^{\lambda(d-c)}\right) = \theta \left(-\frac{a}{b} + e^{\lambda(d-c)}\right) > 0$$

$$\therefore E[\exp^{\lambda Y}] \leq e^{\ln(1 - \theta + \theta e^{\lambda(d-c)}) - \lambda \theta(d-c)}$$

Set $u = \lambda(d-c)$: $E[e^{\lambda Y}] \leq e^{\ln(1 - \theta + \theta e^u) - \theta u}$

Define $\varphi: \mathbb{R} \rightarrow \mathbb{R}$, $\varphi(u) = \ln(1 - \theta + \theta e^u) - \theta u$, $E[e^{\lambda Y}] \leq e^{\varphi(u)}$

According to Taylor's Theorem, $\exists \xi \in [0, u]$, $\varphi(u) = \varphi(0) + u\varphi'(0) + \frac{1}{2}u^2\varphi''(\xi)$

$$\varphi(0) = 0, \varphi'(0) = \left(\frac{\theta e^u}{1 - \theta + \theta e^u} - \theta\right) \Big|_{u=0} = 0, \varphi''(\xi) = \frac{\theta e^\xi}{1 - \theta + \theta e^\xi} \left(1 - \frac{\theta e^\xi}{1 - \theta + \theta e^\xi}\right) = t(1-t) \leq \frac{1}{4}$$

$$\therefore \varphi(u) \leq 0 + 0 + \frac{1}{2}u^2 \cdot \frac{1}{4} = \frac{1}{8}\lambda^2(d-c)^2$$

Then we get $E[\exp^{\lambda Y}] \leq \exp\left(\frac{\lambda^2(d-c)^2}{8}\right)$

② Proof Turo (Hard to understand): $M_Y(\lambda) = E[e^{\lambda Y}]$

Set $\Psi_Y(\lambda) = \ln M_Y(\lambda) = \ln E[e^{\lambda Y}]$

$$\Psi_Y(0) = \ln E[e^{0 \cdot Y}] = 0, \quad \Psi_Y'(0) = \frac{M_Y'(0)}{M_Y(0)} = \frac{E[Y e^{0 \cdot Y}]}{E[e^{0 \cdot Y}]} = E[Y] = 0$$

For any $\xi \in [0, \lambda]$,

$$\Psi_Y''(\xi) = \left(\frac{M_Y'(\xi)}{M_Y(\xi)} \right)' = \frac{M_Y''(\xi)}{M_Y(\xi)} - \frac{M_Y'(\xi)^2}{M_Y(\xi)^2} = E\left[\frac{Y^2 e^{\xi Y}}{M_Y(\xi)}\right] - E\left[\frac{Y e^{\xi Y}}{M_Y(\xi)}\right]^2$$

Define a Z with CDF:

$$F_Z(z) = \int_c^z \frac{e^{\xi y}}{M_Y(\xi)} dF_Y(y) = \int_c^z \frac{e^{\xi y}}{E[e^{\xi Y}]} dF_Y(y)$$

F_Z satisfies $F_Z(d) = \int_c^d \frac{e^{\xi y}}{E[e^{\xi Y}]} dF_Y(y) = 1$, so F_Z is well-defined, and $c \leq Z \leq d$

$$E\left[\frac{Y e^{\xi Y}}{M_Y(\xi)}\right] = \int_c^d y \cdot \frac{e^{\xi y}}{M_Y(\xi)} dF_Y(y) = E[Z]$$

$$E\left[\frac{Y^2 e^{\xi Y}}{M_Y(\xi)}\right] = \int_c^d y^2 \cdot \frac{e^{\xi y}}{M_Y(\xi)} dF_Y(y) = E[Z^2]$$

$$\Psi_Y''(\xi) = E[Z^2] - E[Z]^2 = \text{Var}[Z]$$

$$\because c \leq Z \leq d \quad \therefore \left| Z - \frac{c+d}{2} \right| \leq \frac{d-c}{2}$$

$$\therefore \text{Var}[Z] = \text{Var}\left[Z - \frac{c+d}{2}\right] \leq E\left[\left(Z - \frac{c+d}{2}\right)^2\right] \leq \frac{(d-c)^2}{4}$$

$$\therefore \exists \xi \in [0, \lambda], \quad \Psi_Y(\lambda) = \frac{\lambda^2}{2} \text{Var}[Z] \leq \frac{\lambda^2 (d-c)^2}{8}$$

Hoeffding Inequality:

For $X_1, X_2, \dots, X_n$, if X_i can be bounded, $P(X_i \in [a_i, b_i]) = 1$, Then for
(Independent, r.v.) $\bar{X} = \frac{X_1 + \dots + X_n}{n}$

it satisfies

$$P(\bar{X} - E[\bar{X}] \geq t) \leq \exp\left(-\frac{2t^2 n^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

$$P(|\bar{X} - E[\bar{X}]| \geq t) \leq 2 \exp\left(-\frac{2t^2 n^2}{\sum_{i=1}^n (b_i - a_i)^2}\right) \quad (t \geq 0)$$

Proof: Using Chernoff Bound, we have

$$P(\bar{X} - E[\bar{X}] \geq t) \leq \min_{\lambda \geq 0} \frac{E[\exp(\lambda(\bar{X} - E[\bar{X}]))]}{\exp(\lambda t)}$$

According to Hoeffding Lemma, for $\bar{X} \in [\frac{1}{n} \sum_{i=1}^n a_i, \frac{1}{n} \sum_{i=1}^n b_i]$, we have

$$E[\exp(\lambda(\bar{X} - E[\bar{X}]))] \leq \exp\left(\frac{\lambda^2}{8} \left(\frac{1}{n} \sum_{i=1}^n (b_i - a_i)^2\right)\right)$$

So we get

$$P(\bar{X} - E[\bar{X}] \geq t) \leq \min_{\lambda \geq 0} \exp\left(\underbrace{\frac{\lambda^2}{8n^2} \left(\sum_{i=1}^n b_i - \sum_{i=1}^n a_i\right)^2}_{g(\lambda)} - \lambda t\right) \stackrel{g'(\lambda)=0}{=} \exp\left(-\frac{2t^2 n^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

Inversely, $P(E[\bar{X}] - \bar{X} \geq t) \leq \exp\left(-\frac{2t^2 n^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$

Finally, we get

$$P(|\bar{X} - E[\bar{X}]| \geq t) \leq 2 \exp\left(-\frac{2t^2 n^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

Notice: Although Hoeffding Lemma is to prove the Sub-Gaussian property of bounded random variables, we don't need the assumption of bounded Hoeffding Inequality. What we need is just the Sub-Gaussian property, and

$$E[\exp^{\lambda X}] \leq f(\lambda) \quad X \sim \text{subG}\left(\frac{(b-a)^2}{4}\right)$$

3) Bernstein's Inequality (A Sharper bound using the variance of X_i)

Let $\{X_i\}$ be independent random variables, where with probability 1, $|X_i - E[X_i]| \leq R$.

Let $\text{Var}(X_i) \leq \sigma_i^2$. Then, for all $t \geq 0$,

$$P\left(\sum_{i=1}^n (X_i - E[X_i]) \geq t\right) \leq \exp\left(\frac{-t^2}{2 \sum_{i=1}^n \sigma_i^2 + \frac{2}{3} R t}\right)$$

Since this inequality is less used, we omit the proof.

① Azuma - Hoeffding Inequality (Give a bound of martingale difference)

Azuma - Hoeffding aims at solving the dependent variables

Martingale Difference:

In RL, the independent r.v. is very hard to achieve, we usually get the dependent r.v. For dependent r.v., we usually have the martingale difference

Example: Consider two random variables $X_1 \in \{0, 1\}$, $X_2 \in \{-1, 0, 1\}$ with the following distribution:

$X_1=0$	value	probability
	-1	0.50

X_2	0	0
	1	0.50

$X_1=1$	value	probability
	-1	0

X_2	0	1
	1	0

We have $E[X_2 | X_1=0] = E[X_2 | X_1=1] = E[X_2]$, then we call $\{X_i - E[X_i]\}_{i=1}^n$ is a Martingale difference. Simply explaining, the variables and randomness of the past may change the distribution of the variable I am going to take but can't change the expectation.

Azuma - Hoeffding's Inequality: Suppose $F_1, F_2, \dots, F_n$ are filtrations s.t. $F_{i-1} \subset F_i$, and $X_1, X_2, \dots, X_n$ are random variables s.t. $X_i \in F_i$. If the following conditions hold,

- $E[X_i - E[X_i] | F_{i-1}] = 0$ (Martingale Difference)

- $|X_i - E[X_i]| \leq R_i$

then we have

$$P\left(\sum_{i=1}^n (X_i - E[X_i]) \geq t\right) \leq \exp\left(-\frac{t^2}{2 \sum_{i=1}^n R_i^2}\right), \quad \forall t \geq 0$$

Proof: What we need is focusing on $E[\exp(\lambda \sum_{i=1}^n (X_i - EX_i))] \quad (\text{Need})$

Proof: Set $S_n = \sum_{i=1}^n X_i$, $Z_n = S_n - ES_n$, $W_n = e^{\lambda Z_n}$, $\mathcal{F}_1$

(Using the proof of Hoeffding lemma)

$$E[W_n | \mathcal{F}_{n-1}] = W_{n-1} E[e^{\lambda(Z_n - Z_{n-1})} | \mathcal{F}_{n-1}]$$

$$\leq W_{n-1} \frac{b_n e^{\lambda a_n} - a_n e^{\lambda b_n}}{b_n - a_n}$$

$$\therefore E[E[W_n | \mathcal{F}_{n-1}]] \leq E[W_{n-1}] \cdot \frac{b_n e^{\lambda a_n} - a_n e^{\lambda b_n}}{b_n - a_n}$$

$$\leq \prod_{i=1}^n \frac{b_i e^{\lambda a_i} - a_i e^{\lambda b_i}}{b_i - a_i}$$

$$\forall t > 0, \quad P(S_n - ES_n \geq t) = P(W_n \geq e^{\lambda t}) \leq \frac{E[W_n]}{e^{\lambda t}}, \quad \forall t > 0$$

$$\leq e^{-\lambda t} \prod_{i=1}^n \frac{b_i e^{\lambda a_i} - a_i e^{\lambda b_i}}{b_i - a_i}$$

$$\leq e^{-\lambda t} \prod_{i=1}^n e^{\lambda^2 (b_i - a_i)^2 / 8}$$

$$\leq \exp\left(-\frac{t^2}{2 \sum_{i=1}^n R_i^2}\right)$$

③ Uniform Concentration

1. (Union Bound) $P(A \cup B) \leq P(A) + P(B)$

2. Suppose $\{X_i\}_{i=1}^n$ are random variables chosen i.i.d. and $\text{supp}(X_i) = \{s \mid P(X_i = s) > 0\} \subseteq S$ where S is a discrete set with $|S| = S$.

consider two cases shown below.

case 1: f fixed

$f: S \rightarrow [0, 1]$ where $f(x_i)$ are random variables chosen i.i.d. By Hoeffding Inequality

$$P\left(\left|\sum_{i=1}^n f(x_i) - E f(x_i)\right| \geq t\right) \leq 2 \exp\left(-\frac{2t^2}{n}\right)$$

Case 2: f not fixed

For example, $\hat{f} = \operatorname{argmin}_{f \in F} L(f, \{x_i\}_{i=1}^n)$, $\hat{f}$ depends on the whole dataset

$$\begin{aligned} P\left(\left|\sum_{i=1}^n (\hat{f}(x_i) - E \hat{f}(x_i))\right| \geq t\right) &\leq P(\exists f \in F, \left|\sum_{i=1}^n (f(x_i) - E f(x_i))\right| \geq t) \\ &= P\left(\bigcup_{f \in F} \left|\sum_{i=1}^n (f(x_i) - E f(x_i))\right| \geq t\right) \end{aligned}$$

Case 2a: Finite F

If F is finite, we can use Lemma 1 to get the following:

$$\begin{aligned} P\left(\bigcup_{f \in F} \left|\sum_{i=1}^n (f(x_i) - E f(x_i))\right| \geq t\right) &\leq \sum_{f \in F} P\left(\left|\sum_{i=1}^n (f(x_i) - E f(x_i))\right| \geq t\right) \\ &\leq 2|F| \exp\left(-\frac{2t^2}{n}\right) \quad \swarrow^{(*)} \end{aligned}$$

(*) is only true if the following holds with probability $1-\delta$,

$$\sup_{f \in F} \left| \frac{1}{n} \sum_{i=1}^n (f(x_i) - E f(x_i)) \right| \leq \sqrt{\frac{1}{2n} \log \frac{2|F|}{\delta}}$$

$$= O\left(\sqrt{\frac{1}{n} \log \frac{|F|}{\delta}}\right)$$

Case 2b: Infinite F

We can discretize it using ϵ -grid G_ϵ s.t. $G_\epsilon = \{0, \epsilon, 2\epsilon, 3\epsilon, \dots, \lfloor \frac{1}{\epsilon} \rfloor \epsilon\}$

$= G_\epsilon^S = \{f | f: S \rightarrow G_\epsilon\}$. We have $\forall f \in F, \exists f_\epsilon \in G_\epsilon$ s.t. $\sup_{x \in S} |f(x) - f_\epsilon(x)| \leq \epsilon$

$$\text{w.p. } 1-\delta, \forall f \in F, \left| \frac{1}{n} \sum_{i=1}^n (f(x_i) - E f(x_i)) \right| \leq \underbrace{\left| \frac{1}{n} \sum_{i=1}^n (f_\epsilon(x_i) - E f_\epsilon(x_i)) \right|}_{\text{POETRY}} + \underbrace{(f - f_\epsilon)(x_i) - E[(f - f_\epsilon)(x_i)]}_{\text{PRODUCE}}$$

⑨ Bellman Equation

Recall MDP (S, A, r, P, H) where:

environment: P, r

policy: (history independent) $\pi_h: S \rightarrow A / \Delta A$

$\pi_{h'}: H \rightarrow A / \Delta A$

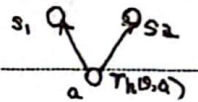
$H: (s_1, a_1, r_1, s_2, a_2, r_2, \dots)$

State-Value Function: $V_h^{\pi_h}(s) = E_{\pi_h} \left[\sum_{k=h}^H r_k(s_k, a_k) \mid s_h = s \right]$

Action-Value Function: $Q_h^{\pi_h}(s, a) = E_{\pi_h} \left[\sum_{k=h}^H r_k(s_k, a_k) \mid s_h = s, a_h = a \right]$

Bellman Equations:

$$\begin{cases} V_h^{\pi}(s) = \sum_{a \in A} Q_h^{\pi}(s, a) \pi_h(a|s) \\ Q_h^{\pi}(s, a) = r_h(s, a) + E_{s' \sim P(\cdot|s, a)} V_{h+1}^{\pi}(s') \end{cases}$$



Optimal Policy

Optimality Equation:

$$\begin{cases} V_h^*(s) := \max_{\pi_h} V_h^{\pi_h}(s) \\ Q_h^*(s, a) := \max_{\pi_h} Q_h^{\pi_h}(s, a) \end{cases} \quad \forall (s, h)$$

Optimal policy π_h^* :

$$\begin{cases} V_h^*(s) = V_h^{\pi_h^*}(s) \\ Q_h^*(s, a) = Q_h^{\pi_h^*}(s, a) \end{cases} \quad \forall (s, a, h)$$

Theorem 1 (Optimality Existence): There exists a history independent and deterministic policy π_h^* that is optimal.

Proof: Let $\pi_h^*(s) = \operatorname{argmax}_a Q_h^*(s, a)$ be history independent and deterministic.

We can now prove by induction:

1 Base case: $H: Q_H^*(s, a) = r_H(s, a) = Q_H^{\pi_h^*}(s, a)$

2 $Q_h^* = Q_h^{\pi_h^*} \Rightarrow V_h^* = V_h^{\pi_h^*}$

$$V_h^*(s,a) = \max_{\pi_h} \max_{\pi_{h+1} \dots \pi_H} \sum_{a \in A} Q_h^{\pi}(s,a) \cdot \pi_h(a|s)$$

$$= \max_{\pi_h} \sum_{a \in A} Q_h^*(s,a) \pi_h(a|s)$$

$$= \max_a Q_h^*(s,a) = Q_h^*(s, \pi_h^*(s)) = V_h^*(s)$$

$$\textcircled{3} V_{h+1}^* = V_{h+1}^{\pi^*} \Rightarrow Q_h^* = Q_h^{\pi^*}$$

$$Q_h^*(s,a) = r_h(s,a) + \max_{\pi} (P V_{h+1}^{\pi})(s,a) = r_h(s,a) + (P V_{h+1}^*)(s,a) = r_h(s,a) + (P V_{h+1}^{\pi^*})(s,a) = \textcircled{1}$$

Bellman Optimality Equation:
$$\begin{cases} V_h^*(s) = \max_a Q_h^*(s,a) \\ Q_h^*(s,a) = r_h(s,a) + (P V_{h+1}^*)(s,a) \end{cases}$$

Goal of RL: find the $\pi_h^*(s) = \operatorname{argmax}_a Q_h^*(s,a)$

Bellman Equation:
$$\begin{cases} V^{\pi}(s) = \sum_{a \in A} Q^{\pi}(s,a) \pi(a|s) \\ Q^{\pi}(s,a) = r(s,a) + \gamma (P V^{\pi})(s,a) \end{cases}$$

Discounting Factor γ MDP(S, A, r, P, γ)

2) Planning in MDP

Algorithm 1 Policy Evaluation (DP) (Infinite Horizon)

for $h = H, H-1, \dots, 1$ do

$$Q_h^{\pi}(s,a) = r_h^{\pi}(s,a) + (P_h V_{h+1}^{\pi})(s,a)$$

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

until $\max_s |V_{h+1}^{\pi}(s) - V_h^{\pi}(s)| < \theta \times$

Algorithm 2 Policy Iteration

for $t = 1, 2, \dots, H$ do

Run Policy Evaluation to evaluate $Q_h^{\pi^{t-1}}(s,a) \forall s,a,h$

$$\pi_h^t(s) \leftarrow \operatorname{argmax}_{a \in A} Q_h^{\pi^{t-1}}(s,a)$$

Algorithm 1 Policy Evaluation (Fixed Point)

$$Q_0 \in [0, \beta]^{SA}$$

$$* (\gamma + \gamma^2 + \gamma^3 + \dots < \frac{1}{1-\gamma} = \beta)$$

for $t = 0, 1, \dots$ do

$$Q_{t+1} = T^{\pi} Q_t = r + \gamma P^{\pi} Q_t$$

$$\text{lemma: } \|Q_t - Q^{\pi}\|_{\infty} \leq \gamma^t \|Q_0 - Q^{\pi}\|_{\infty} \leq \gamma^t \beta$$

Algorithm 2 Value Iteration

$$Q_0 \in [0, \beta]^{SA}$$

for $t = 0, 1, \dots$ do

$$Q_{t+1} = T^* Q_t$$

(Finite Horizon)

Finite Algorithm 3 Value Iteration

for $h = H, H-1, \dots, 1$ do

$$Q_h^*(s, a) = r_h(s, a) + (P_h V_{h+1}^*(s, a))$$

$$V_h^*(s) = \max_{a \in A} Q_h^*(s, a)$$

$$T_h^*(s) = \operatorname{argmax}_{a \in A} Q_h^*(s, a)$$

(★) Theorem: T_h^* is optimal $\forall h \geq 1$

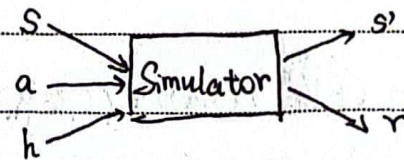
Proof: Using induction:

Base case: $t=0$ $Q_H^0(s, a) = r_H(s, a) = G$

If t is true, then at $t+1$, $\forall h \geq H-t$,

$$T_h^{t+1}(s) = \operatorname{argmax}_a Q_h^{t+1}(s, a) = \operatorname{argmax}_a Q_h^*(s, a) = T_h^*(s)$$

II) Generative Models

MDP(S, A, r, P, H) with P, r unknown, but we have access to a simulator

Goal is to find some policy $\hat{T}_h$ so that $V_h^*(s_i) - V_h^{\hat{T}_h}(s_i) \leq \epsilon$ for all i , and the sample complexity is how many queries needed to ask the simulator to find ϵ -optimal $\hat{T}_h$.

Generative Value Iteration:

Input: n For all $(s, a, h) \in S \times A \times H$ doQuery (s, a, h) and collect n_h samples $\{s'_1, s'_2, \dots, s'_{n_h}\}$

Estimate transition probability by

$$\hat{P}_h(s'|s, a) = \frac{1}{n_h} \sum \mathbb{I}[s'_i = s']$$

for all $s' \in S$.For $h = H$ to 1 do

$$\text{Compute } \hat{Q}_h^*(s, a) = r_h(s, a) + (\hat{P}_h V_{h+1}^*)(s, a)$$

$$\text{Compute } V_h^*(s) = \max_{a \in A} \hat{Q}_h^*(s, a)$$

$$\text{Output: } \hat{T}_h(s) = \operatorname{argmax}_a \hat{Q}_h^*(s, a)$$

(Coarse Analysis)

Analysis 1:

Theorem: If we chose $n \geq C \frac{H^4 S}{\epsilon^2} l$, where $l = \log \frac{HSA}{p}$, then with probability 1 the policy resulting from value iteration will be ϵ -optimal.

Lemma 1. For any policy π and any $(s, h) \in S \times [H]$,

$$\underbrace{V_h^\pi(s)}_{\text{true}} - \underbrace{\hat{V}_h^\pi(s)}_{\text{estimated}} = E_{\pi, \pi} \left[\sum_{i=h}^H (P_i - \hat{P}_i) V_{i-1}^\pi(s_i, a_i) \mid S_h = s \right]$$

Proof. Note that

$$\begin{aligned} Q_h^\pi(s, a) - \hat{Q}_h^\pi(s, a) &= (P V_{h+1}^\pi)(s, a) - (\hat{P} \hat{V}_{h+1}^\pi)(s, a) = [(P - \hat{P}) V_{h+1}^\pi] + (\hat{P} (V_{h+1}^\pi - \hat{V}_{h+1}^\pi)) \\ \Rightarrow V_h^\pi(s) - \hat{V}_h^\pi(s) &= E_{\pi} [(P - \hat{P}) V_{h+1}^\pi(s_{h+1}, a_{h+1}) \mid S_h = s] + E_{\pi, \pi} [V_{h+1}^\pi(s_{h+1}) - \hat{V}_{h+1}^\pi(s_{h+1}) \mid S_h = s] \\ &= E_{\pi, \pi} \left[\sum_{i=h}^H (P_i - \hat{P}_i) V_{i-1}^\pi(s_i, a_i) \mid S_h = s \right] \end{aligned}$$

Lemma 2. With probability $\geq 1-p$, we have $V(s, a, h) \in S \times A \times [H]$, $\forall V \in [0, H]^S$,

$$|(P - \hat{P}) V(s, a)| \leq C \cdot H \cdot \frac{S l}{n}, \text{ where } l = \log \frac{HSA}{p}$$

Using Lemma 1, 2, Theorem can be get.

(Refined Analysis)

Analysis 2:

Theorem: If we chose $n \geq C \cdot \frac{H^4 l}{\epsilon^2}$ where $l = \log \frac{HSA}{p}$, then ...

⑫ Q-learning

$$Q_h^t(s, a) = (1 - \alpha_t) Q_h^{t-1}(s, a) + \alpha_t (r_h(s, a) + V_{h+1}^t(s'))$$

↙ learning rate

Algorithm Q-learning

Input: learning rate $\{\alpha_t\}_{t=1}^{\infty}$

Initialize: $Q_h^0(s, a) = 0, V(s, a, h) \in S \times A \times [H]$

for $t=1, \dots$ do

for $h=1, \dots, H$ do

for $(s, a) \in S \times A$ do

sample next state s' from $P_h(\cdot | s, a)$

Update

$$Q_h^t(s, a) = (1 - \alpha_t) Q_h^{t-1}(s, a) + \alpha_t (r_h(s, a) + V_{h+1}^t(s'))$$

for $s \in S$ do

$$V_h^t(s) = \max_{a \in A} Q_h^t(s, a)$$

Remark:

① 1. Model-free learning

2. Incremental update using one noisy sample. Space complexity: $O(HSA)$

Using the update rule we get:

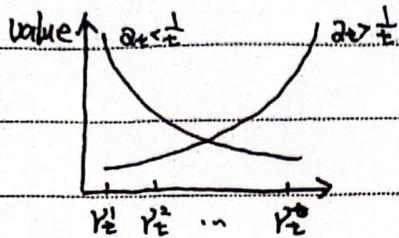
$$Q_h^t(s, a) = \sum_{i=1}^t \gamma_i^i (r_h(s, a) + V_{h+1}^{i+1}(s')) = r_h(s, a) + \sum_{i=1}^t \gamma_i^i V_{h+1}^{i+1}(s')$$

$$\gamma_i^i = \alpha_i \prod_{j=i+1}^t (1 - \alpha_j) \text{ and } \sum_{i=1}^t \gamma_i^i = 1$$

(Generative)

Comparing to Value Iteration: $Q_h(s,a) = r_h(s,a) + \sum_{i=1}^n \frac{1}{n} V_{h+1}(s_i')$

1. When $\alpha_t = \frac{1}{n}$ and $r_t^i = \frac{1}{n}$, Equal
2. When $\alpha_t > \frac{1}{n}$, Q-learning favors later samples.
3. When $\alpha_t < \frac{1}{n}$, Q-learning favors earlier samples.



sually, we choose $\alpha_t > \frac{1}{n}$.

⑭ Multi-arm bandit (MAB)

Exploration vs. Exploitation

$$\text{regret}(T) = T \cdot r^* - \sum_{t=1}^T r_{i_t}$$

$$0 \leq \text{regret}(T) \leq T$$

Conversion:

Lemma 1. If algorithm A has regret $c \cdot T^{1-\alpha}$, $\alpha \in (0,1]$. Then A finds an ϵ -optimal distribution of arm in $(\frac{c}{\epsilon})^{\frac{1}{\alpha}}$ samples. Conversely, $c \cdot \epsilon^{-\frac{1}{\alpha}}$ samples $\Rightarrow 2 \cdot c^{\frac{1}{1+\alpha}} \cdot T^{\frac{\alpha}{1+\alpha}}$

$$\sqrt{T} \text{ Regret} \rightarrow \frac{1}{\epsilon^2} \text{ samples}$$

$$\frac{1}{\epsilon^2} \text{ samples} \rightarrow T^{\frac{2}{3}}$$

Proof: a. $\text{regret}(T) = T \cdot r^* - \sum_{t=1}^T r_{i_t} \leq c \cdot T^{1-\alpha}$

$$\frac{1}{T} \sum_{t=1}^T r_{i_t} \geq r^* - \frac{c}{T^\alpha}$$

$$\frac{c}{T^\alpha} = \epsilon \Rightarrow T = \left(\frac{c}{\epsilon}\right)^{\frac{1}{\alpha}}$$

b. $\text{regret}(T) \leq c \cdot \epsilon^{-\frac{1}{\alpha}} \times 1 + (T - c \cdot \epsilon^{-\frac{1}{\alpha}}) \times \epsilon$

$$\leq c \cdot \epsilon^{-\frac{1}{\alpha}} + T \epsilon$$

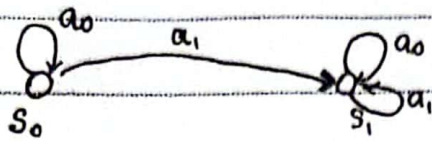
$$= 2 \cdot c^{\frac{1}{1+\alpha}} T^{\frac{\alpha}{1+\alpha}}$$

⑮ Exploration in MDPs

A contextual MAB can be viewed as an MDP without transition

$$\text{exploration} \left\{ \begin{array}{l} \text{MABs: } \mathcal{O}(A|\epsilon^2) \\ \text{MDP: } \Omega(2^H) \text{ (worst)} \end{array} \right.$$

Combinatorial lock:



$$r_H(s_0) = 1, \quad r_h(s) = 0$$

$\underbrace{\{a_0, a_0, \dots, a_0\}}_H$ is the optimal, however, we need $\Omega(2^H)$ to get the po

MPP is more hard for exploration than MAB.

1) UCB-VI (Upper Confidence Bound Value Iteration)

Algorithm UCB-VI

Initialize: dataset $\mathcal{D} := \emptyset$; $Q_h(s, a) := H$ for $h \in [H]$; $V_{H+1}(s) = 0$.

for episode $k = 1$ to K do

(Part 1: Value iteration with bonus)

for $h = H$ to 1 do

for $(s, a) \in \mathcal{S} \times \mathcal{A}$ do

if $(s, a) \in \mathcal{D}$ then

Compute $\hat{P}_h(s'/s, a) := \frac{N_h(s, a, s')}{N_h(s, a)}$ for all $s' \in \mathcal{S}$, where $N_h(s, a, s')$ and $N_h(s, a)$

$Q_h(s, a) := \min \{ H, r_h(s, a) + (\hat{P}_h V_{h+1}(s, a) + b(N_h(s, a))) \};$

for $s \in \mathcal{S}$ do

$V_h(s) := \max_{a \in \mathcal{A}} Q_h(s, a);$

for $h = 1$ to H do

Take action $a_n := \operatorname{argmax}_{a \in \mathcal{A}} Q_h(s_n, a)$