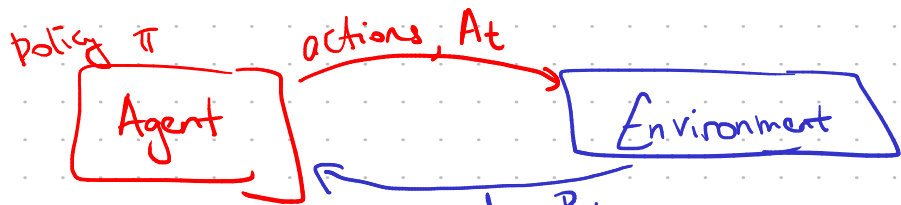


① Markov decision process (MDP)

Maximization : / episodic task



- A_t - код
- R_t - текущая награда
- S_t - состояние /

state, S_t

π

$S_0, A_0, R_1, S_1, A_1, R_2, S_2, A_2, \dots, S_t, A_t, R_{t+1}, S_{t+1}, A_{t+1}$

P

P

P

policy

prob. seq.

τ_i^* - optimal strategy

$$\pi: S \rightarrow A \quad - \text{deterministic policy}$$
$$\pi: S \rightarrow \text{Prob}(A) \text{ - stochastic policy}$$

$$\pi(a|s) = p(A_t = a | S_t = s)$$

$$p(r, s' | s, a) = p(R_{t+1} = r, S_{t+1} = s' | S_t = s, A_t = a)$$

Episodic vs continuous tasks

$$t = 1, 2, \dots, T$$

$$G_t = \sum_{s=t+1}^T R_s \xrightarrow{T} \max$$

$$\delta \approx 1$$

$\delta = 1$

$0 \rightarrow 0 \rightarrow 0 \rightarrow \dots \rightarrow 0 \rightarrow 0$

S_t T R_t

$R=0$

Problems

$$t = 1, 2, \dots, T, \dots$$

$$G_t = \sum_{s=t+1}^{\infty} \gamma^{s-t-1} \cdot R_s \xrightarrow{\pi} \max$$

 $\delta < 1$

1) Credit assignment

2) Exploration vs. exploitation

② Multicomed bandits

$$a_x = \operatorname{argmax}_a \mathbb{E}_{p(z|a)}[z]$$

$$|S|=1, \quad A = a_1 a_2 \dots a_n$$

Env - $p(z|a)$

1) A/B testing : депыс N жана θ_a , θ_b - max.

4) Greedy: — дернуть 1 рр за канту

- for $t = 1, 2, \dots$

Binary bandits
 $\mathcal{K} = \{0, 1\}$

$$A_t = \operatorname{argmax}_i \hat{\mu}_i(t)$$

$$\hat{\mu}_i(t) = \frac{1}{n_i} \sum_{t: A_t = i} R_t$$



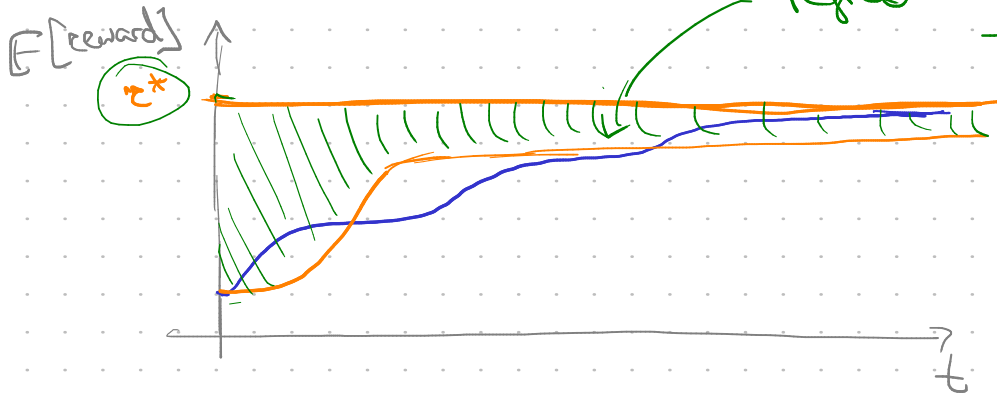
3) ϵ -Greedy

ϵ -soft strategies

$$A_t = \begin{cases} \arg \max_i \hat{\mu}_i(t), & \text{с вер. } 1 - \epsilon \\ \text{Unit}(\{1..K\}), & \text{с вер. } \epsilon. \end{cases}$$

Regret / цена ошибки

$$\text{regret} = \sum_{t=1}^{\infty} \mathbb{E} [z^* - z_t] \quad \text{learning}$$



$$\hat{\mu}(n) = \frac{1}{n} \sum_{i=1}^n z_i \quad z_1, z_2, \dots, z_n, \textcircled{z_{n+1}}$$

$$\hat{\mu}(n+1) = \frac{1}{n+1} \sum_{i=1}^{n+1} z_i = \frac{n}{n+1} \cdot \hat{\mu}(n) + \frac{1}{n+1} \cdot z_{n+1} = \hat{\mu}(n) + \frac{1}{n+1} (z_{n+1} - \hat{\mu}(n))$$

$$\boxed{\text{Новая оценка} = \text{Старая оценка} + d_n \cdot \left(\text{Мгновенное значение} - \text{Старая оценка} \right)}$$

$$d_n = \frac{1}{n} \Rightarrow \hat{\mu} = \text{Avg}(z_1 - z_n)$$

$$\sum_n d_n = \infty, \quad \sum_n d_n^2 < \infty \quad \Rightarrow \text{оценка сходятся}$$

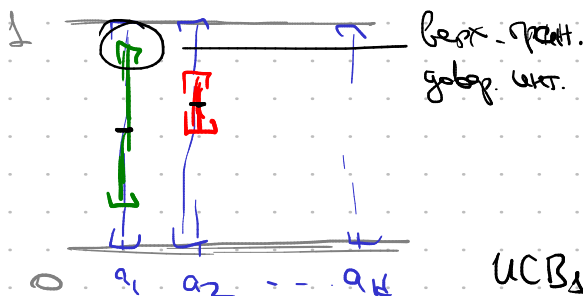
Nonstationary bandits

$d_n = \alpha = \text{const}$; EMA exp. mov. avg

$$\hat{\mu}(n+1) = \hat{\mu}(n) + d(z_{n+1} - \hat{\mu}(n)) = d \cdot z_{n+1} + (1-d)\hat{\mu}(n)$$

$$= d \cdot z_{n+1} + d \cdot (1-d)z_n + d \cdot (1-d)^2 z_{n-1} + \dots$$

3) Optimism under uncertainty



$$\text{Priority}_i(t) = \hat{\mu}_i(t) + \text{Exploration bonus}$$

UCB - upper confidence bound

$$\text{UCB}_i: \text{Priority}_i(t) = \hat{\mu}_i(t) + c \cdot \sqrt{\frac{\log t}{n_i}} \quad \text{брош. } n_i \rightarrow \infty$$

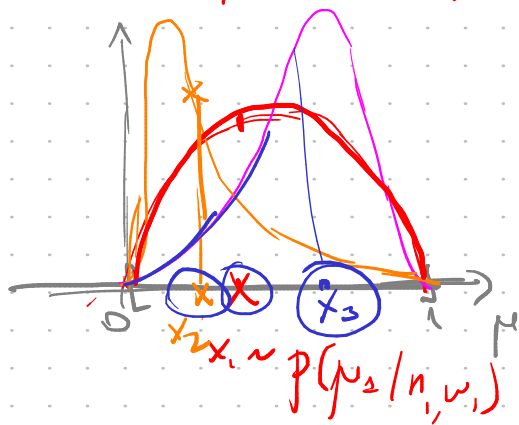
Thompson sampling

n trials, w - wins,

$$p(\mu_i) = \text{Beta}(1, 1)$$



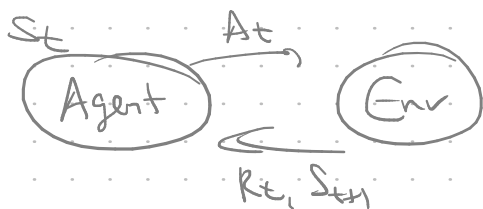
$$p(\mu_i | n, w) = \text{Beta}(w+1, n-w+1)$$



$$x_i \sim p(\mu_i | D)$$

$$A_t = \underset{i}{\operatorname{argmax}} x_i$$

④ Value functions



Reward: $R_t, R_{t+1}, R_{t+2}, \dots$

Return: $G_t = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots$

$$G_t = R_{t+1} + \gamma \cdot G_{t+1}$$

(State) value function:

$$V_\pi(s) = \mathbb{E}_\pi[G_t | S_t = s] = \mathbb{E}_\pi \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} | S_t = s \right]$$

$$V_*(s) = \max_{\pi} V_\pi(s) = V_{\pi_*}(s) \quad \mathbb{E}_{R_{t+1}, S_{t+1}, A_{t+1}, R_{t+2}, \dots}$$

State-action value function:

$$Q_\pi(s, a) = \mathbb{E}_\pi[G_t | S_t = s, A_t = a] = \mathbb{E}_\pi \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} | S_t = s, A_t = a \right]$$

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

$$\pi_*(s) = \underset{a}{\operatorname{argmax}} Q_*(s, a)$$

$$Q_\pi(s, a) = \mathbb{E}_\pi[R_{t+1} + \gamma \cdot G_{t+1} | S_t = s, A_t = a] =$$

$$= \sum_{s', a'} p(s', a' | s, a) \cdot \mathbb{E}_\pi[r + \gamma \cdot G_{t+1} | S_{t+1} = s']$$

$$Q_\pi(s, a) = \sum_{s'} p(s' | s, a) (r + \gamma \cdot V_\pi(s'))$$

$$V_\pi(s) = \mathbb{E}_\pi[R_{t+1} + \gamma G_{t+1} | S_t = s] = \mathbb{E}_{A_t} \left[\mathbb{E}_\pi[\dots] \right] =$$

$$= \sum_a \pi(a|s) \cdot \mathbb{E}_\pi [R_{t+1} + \gamma G_{t+1} | S_t = s, A_t = a]$$

$$V_\pi(s) = \sum_a \pi(a|s) Q(s, a)$$

$$V_\pi(s) = \sum_a \pi(a|s) \sum_{s'} p(s'|s, a) (r + \gamma \cdot V_\pi(s'))$$

Bellman equations

$$\bar{x} = f(x)$$

$$Q_\pi(s, a) = \sum_{s'} p(s'|s, a) (r + \gamma \cdot \sum_{a'} \pi(a'|s') Q(s', a'))$$

Алгоритм:

- для каждого π :

- решить уравнения, найти $Q_\pi(s, a) \forall s, a$

$$\pi_*(s) = \operatorname{argmax}_a \left[\max_{\pi} Q_\pi(s, a) \right]$$

Проблемы:

1) π слишком много

2) уравнений слишком много $|S|, |S| \times |A|$

3) Часто мы не знаем $p(s', a)$

$$\begin{aligned} 1) V_*(s) &= \max_{\pi} V_\pi(s) = \max_{\pi} \sum_a \underbrace{\pi(a|s)}_{\substack{|S| \\ |S| \times |A|}} \sum_{s'} p(s'|s, a) (r + \gamma V_\pi(s')) = \\ &= \max_a \sum_{s'} p(s'|s, a) (r + \gamma \max_{\pi} V_\pi(s')) \end{aligned}$$

$$V_*(s) = \max_a \sum_{s'} p(s'|s, a) (r + \gamma \cdot V_*(s'))$$

$$\bar{x} = f(x)$$

$$Q_*(s, a) = \sum_{s'} p(s'|s, a) (r + \gamma \cdot \max_{a'} Q_*(s', a'))$$

Bellman equations

$$\bar{\theta} = f(\theta)$$

2) Tabular RL

s	$V(s)$
s_1	$V(s_1)$
s_2	$V(s_2)$



Approximate RL

$$s - \boxed{\mu} - V(s)$$

$$s - \boxed{\mu} - Q(s, a)$$

5 Policy improvement theorem

Thm. Ecan ^{designe} π, π'

$$\forall s \quad Q_{\pi}(s, \pi'(s)) \geq V_{\pi}(s) = Q_{\pi}(s, \pi(s))$$

$$\text{to } \forall s \quad V_{\pi'}(s) \geq V_{\pi}(s).$$

Proof:

$$V_{\pi}(s) \leq Q_{\pi}(s, \pi'(s)) = \mathbb{E}_{\pi}[R_{t+1} + \gamma G_{t+1} | S_t = s, A_t = \pi'(s)]$$

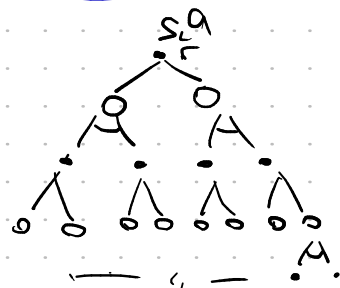
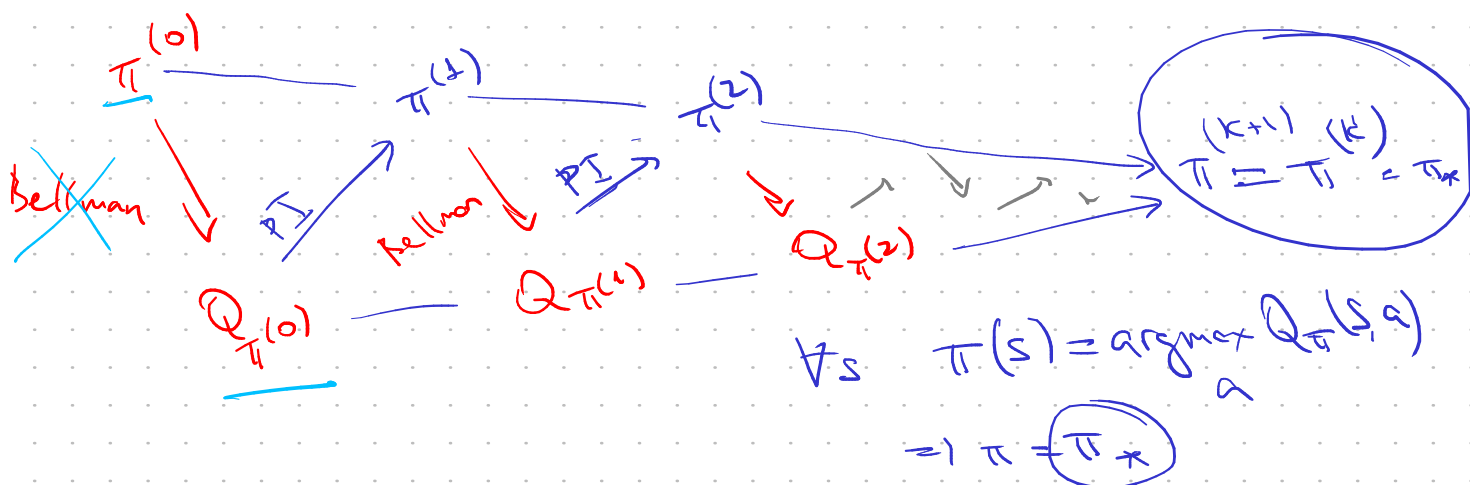
$$= \mathbb{E}_{R_{t+1}, S_{t+1}} [R_{t+1} + \gamma \cdot \underbrace{\mathbb{E}_{\pi}[G_{t+1} | S_{t+1}]}_{V_{\pi}(S_{t+1})} | S_t = s, A_t = \pi'(s)] \leq$$

$$\leq \mathbb{E}_{R_{t+1}, S_{t+1}} [R_{t+1} + \gamma \cdot \mathbb{E}_{\pi}[R_{t+2} + \gamma \cdot \mathbb{E}_{\pi}[G_{t+2} | S_{t+2}] | S_{t+1}, A_{t+1} = \pi'(S_{t+1})] | S_t = s, A_t = \pi'(s)] \leq \dots \leq V_{\pi'}(s) \quad \text{QED}$$

$\pi \rightsquigarrow \pi'$ - why?

Policy iteration:

$$\pi \rightarrow Q_{\pi} \rightarrow \pi'(s) = \arg\max_a Q_{\pi}(s, a)$$



Indirect RL

π, V, Q

Direct RL

Monte-Carlo
TD-learning
policy gradient

Bellman equations

planning

Experience

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

MDP model

$p(r, s' | s, a)$

Model learning

