

1)

Machine Learning

Supervised Learning

$$D = \{(x, y)\}$$

$$f: x \mapsto y$$

Unsupervised Learning

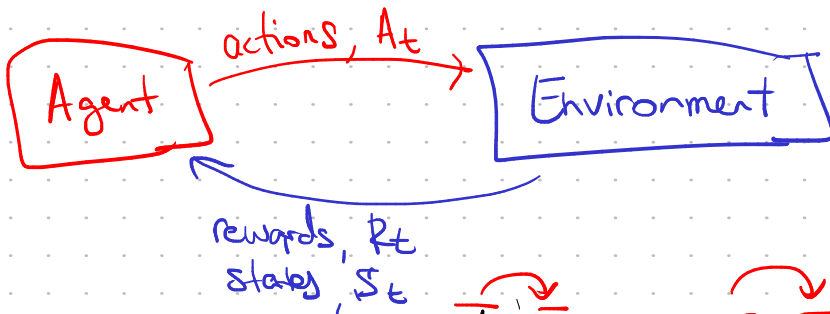
$$D = \{x\}$$

$$p(x)$$

Reinforcement Learning

2)

Markov decision process (MDP)



$$\pi: S \rightarrow \text{Prob}(A(S))$$

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

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

episodic tasks

Maxwell:

action — 0/1/2/3
reward — 0/1/2/3
state — позиция + черта

Episodic vs continuous

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

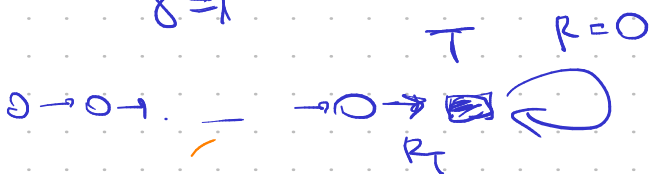
$$G_t = \sum_{k=1}^T \gamma^k R_{t+k} \rightarrow \max$$

$$\gamma = 1$$

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

$$G_t = \sum_{k=1}^{\infty} \gamma^k R_{t+k} \rightarrow \max$$

$$\gamma < 1$$



1) Credit assignment

2) Exploration vs. exploitation

3 Multiarmed bandits

$$|S|=1, A = a_1, a_2, \dots, a_K$$

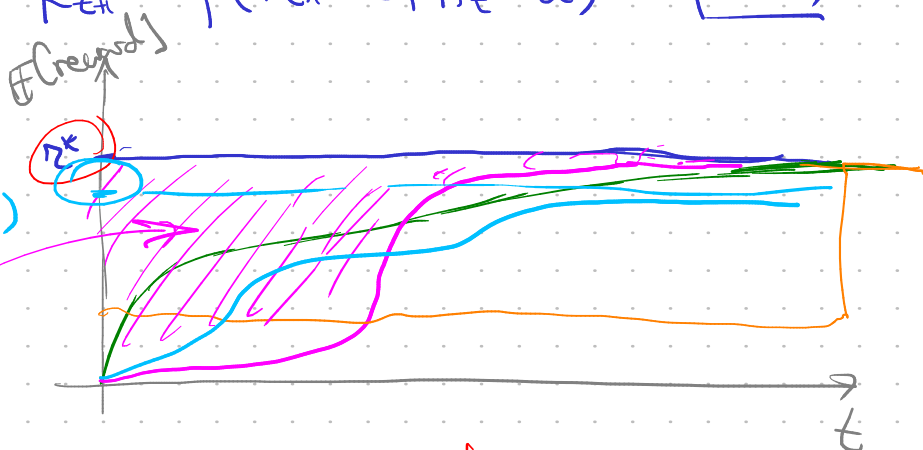
$$R_{t+1} \sim p(R_{t+1}=z | A_t=a) \quad p(z|a)$$

$$a_* = \arg \max_a \mathbb{E} p(z|a) [z]$$

$$R_1 - R_t$$

$$\hat{\mu}_i = \frac{1}{t} \sum R_j$$

Regret



Greedy:

- зберігає 1 та 3 в пам'яті
- for $t=1, \dots$

Binary bandits
 $z \in \{0, 1\}$

Thompson sampling

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

$$A_t = \arg \max_i \hat{\mu}_i(t)$$

ϵ -Greedy

- for $t=1:$

ϵ -soft strategies

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

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

$$z_{n+1}$$

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

новий оцінка = старий оцінка + α_n (новий зразок - старий оцінка)

$$\alpha_n \sim \frac{1}{n}$$

$$\sum_{n=1}^{\infty} \alpha_n = \infty$$

$$\sum_{n=1}^{\infty} \alpha_n^2 < \infty$$

$$\alpha_n = \alpha = \text{const}$$

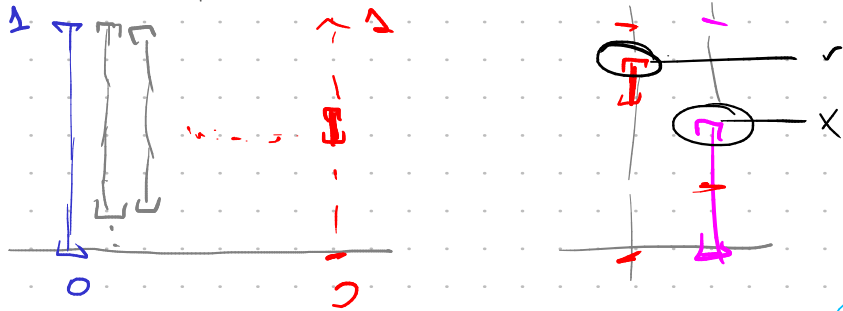
$$\begin{aligned} \hat{\mu}(n+1) &= \hat{\mu}_n + \alpha (z_{n+1} - \hat{\mu}_n) = \alpha \cdot z_{n+1} + (1-\alpha) \hat{\mu}_n = \\ &= \alpha z_{n+1} + (1-\alpha) (\alpha z_n + (1-\alpha) \hat{\mu}_{n-1}) = \dots \end{aligned}$$

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

exponential moving average

nonstationary bandits

④ Optimism under uncertainty



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

UCB - Upper Confidence Bound

$$\text{UCBs: } \text{Priority}_i(t) = \hat{\mu}_i(t) + c \cdot \sqrt{\frac{\log t}{n_i}}$$

← zero mean
← mean of i-th population

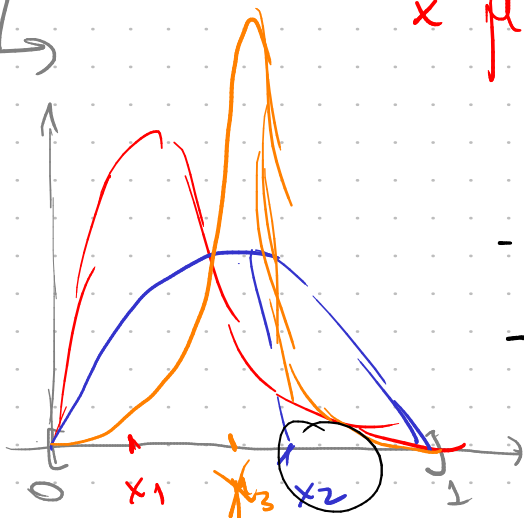
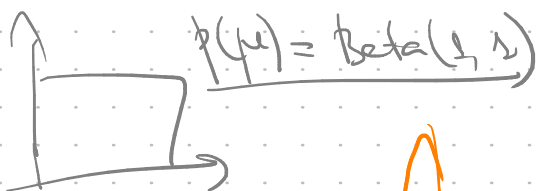
Thm. $c = \sqrt{2} \Rightarrow$

$$\text{Regret} = O(\sqrt{k \cdot T \cdot \log T}), n$$

sa cyson-pyry definen $O(\log T)$ per

$$\Delta_i = |\mu_* - \mu_i|$$

Thompson sampling

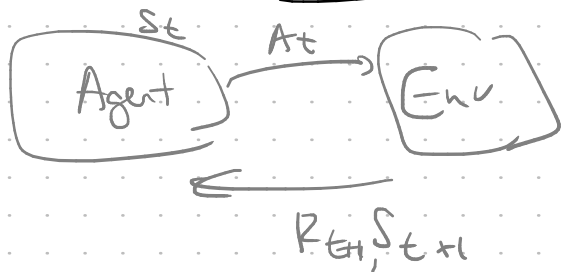


$$x \sim \mu^w (1-\mu)^{n-w} \propto p(\mu|D) = \text{Beta}(w+1, n-w+1)$$

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

$$A_t = \arg\max x_i$$

5) Value functions & Bellman equations



$\pi(a|s)$ - policy

π^* ?

$p(z, s' | s, a)$ - MDP dynamics

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

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

(State) value function:

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

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

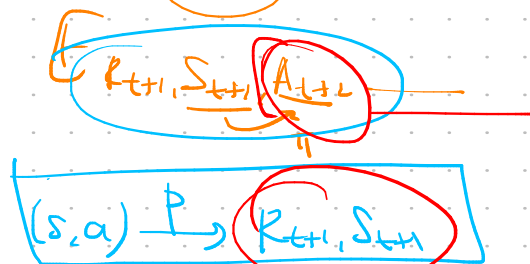


State-action value function:

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

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

$$\pi_*(s) = \arg \max_a Q_*(s, a)$$



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

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

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

$a \in \mathcal{A}(s)$

$a \in \mathcal{A}$

$$V_{\pi}(s) = E_{\pi} [R_{t+1} + \gamma \cdot G_{t+1} | s_t = s] = E_{A_t} [E_{\pi}[\cdot]]$$

$$= \sum_a \pi(a|s) \cdot Q_{\pi}(s, a)$$

$$V_{\pi}(s) = \sum_a \pi(a|s) \cdot Q_{\pi}(s, a)$$

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

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

Bellman equations

Αλγόριθμοι :

- για καθέναν $\pi: \mathcal{S} \rightarrow \mathcal{A}$
- περνάει γι' $\mathcal{S}$, νόημα $Q_{\pi}(s, a)$
- $\mathcal{S}$ & $\mathcal{A}$ περσ $\pi_*(s) = \arg \max_a \max_{\pi} Q_{\pi}(s, a)$

Προσέγγιση :

- 1) π continuous μνοο
- 2) $|\mathcal{S}|$ continuous μνοο, $|\mathcal{S}| \cdot |\mathcal{A}|$
- 3) γ ατο μτ ηε γνοο $p(z, s' | s, a)$

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

$$= \max_a \sum_{z, s'} p(z, s' | s, a) (z + \gamma \cdot \max_{\pi} V_{\pi}(s'))$$

$|\mathcal{S}| \cdot |\mathcal{A}|$

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

Bellman equations

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

Tabular RL

$$\pi_*(s) = \arg \max_a Q_*(s, a)$$

Approximate RL
 $\frac{s}{a} = \text{ML} - Q$

s	V
s_1	$V(s_1)$
s_2	
$\vdots$	
s_N	$V(s_N)$

$$V(s_0) = \dots$$

$$V(s_1) = \max_{a_1, a_2, a_3} \dots V(s_i) \dots$$

$$V(s_2) = \max \{ \dots \}$$

⑥ Policy improvement theorem

Given two policies π, π'

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

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

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

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

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

$$\leq \dots \mathbb{E}[R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} \dots | S_t = s, A_t = \pi'(s), A_{t+1} = \pi'(S_{t+1}), \dots]$$

$$= V_{\pi'}(s)$$

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

