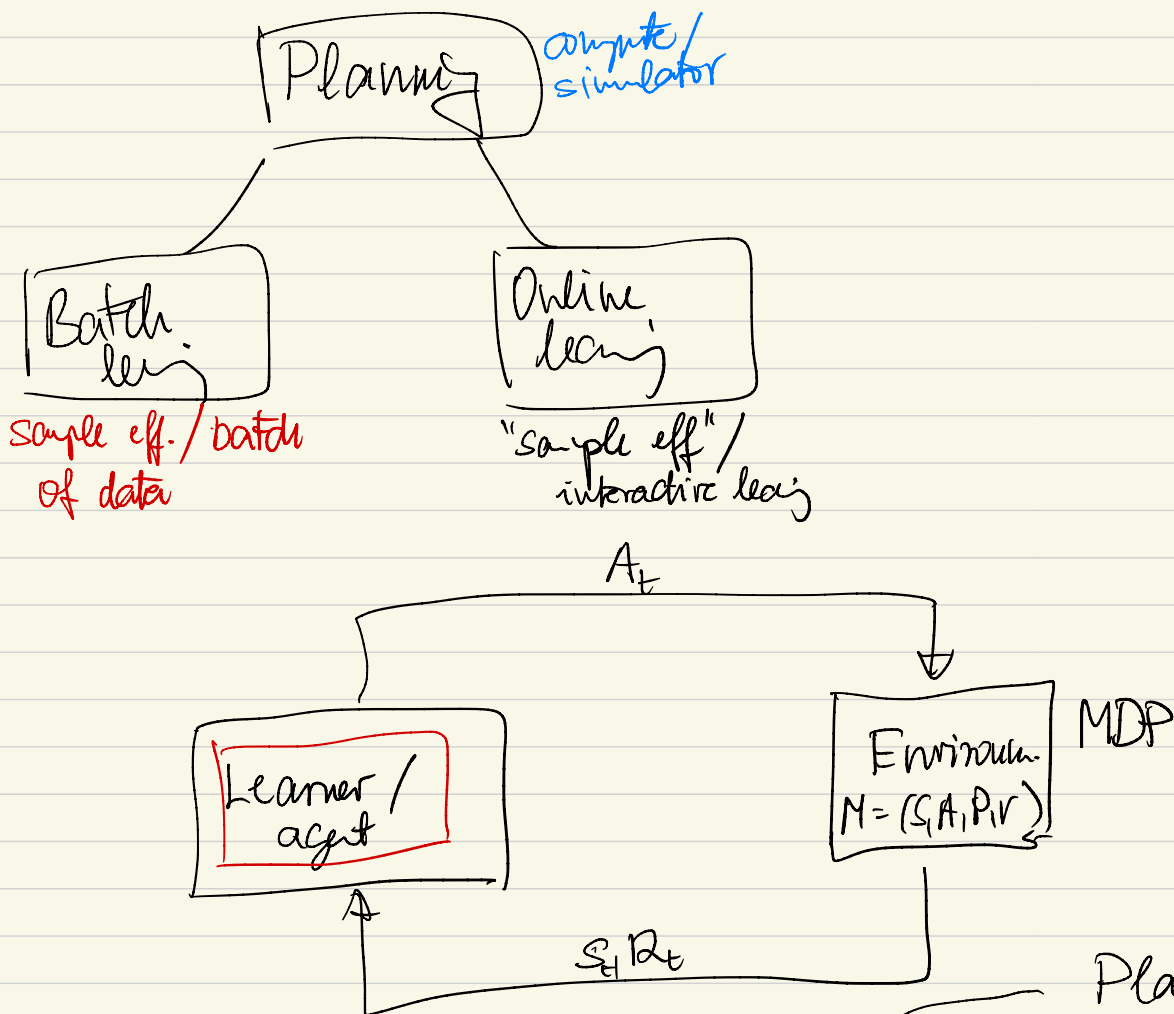


April 1

Online learning

/ Online RL



How good is a learner?

"Regret"

Cumulative regret
over time

undiscounted

fixed-horizon

infinite horizon

$$R_n = v_n^*(S_0) - \sum_{t=0}^{n-1} R_t$$
$$\rightarrow R_n = g^* n - \sum_{t=0}^{n-1} R_t$$

$\uparrow$ opt. arg. max.

Planning

$$\hat{v} - v^*$$

"simple" regret

Episodes of length $H > 0$, $h = 0, 1, \dots, H-1$

$1 \leq k \leq K$: # episodes.

$$S_0^{(k)} \sim \mu, A_0^{(k)}, S_1^{(k)}, A_1^{(k)}, \dots, S_{H-1}^{(k)}, A_{H-1}^{(k)}, S_H^{(k)}$$
$$S_{h+1}^{(k)} \sim P_{A_h^{(k)}}(S_h^{(k)})$$

$$M = (S, A, P, r)$$

For simplicity: r is known

$$R_k = \sum_{h=0}^K \left(v_0^*(S_0^{(k)}) - \sum_{h=0}^{H-1} r_{A_h^{(k)}}(S_h^{(k)}) \right)$$

↑
regret

$$\frac{R_K}{K}$$

PAC $\iff$ Regret

MDP

$$H = 1$$

$\iff$ contextual
bandits

$$|S| = 1$$

finite-armed bandits

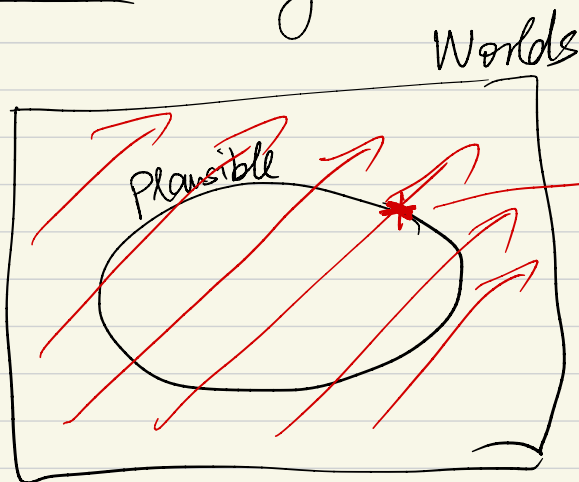
regret is
unknown

$$[\epsilon\text{-greedy}] R_n = \Theta(n^{2/3}), \underline{R_n = O(\sqrt{n})}$$

Principle: Optimism in the Face of Uncertainty

Necessary: Trying sg gives info about how good ' sg ' is

suff. for opt. to be "opt": try x does not give "much" info about y



follow policy that gives best value in the best world

Lai & Robbins '85

P.R. Kumar '83

$H > 0$ fixed horizon, episodic setting, $0 \leq \gamma < 1$

$$M = (S, A, \underline{P}, r)$$

$$k = 1, \dots, K$$

↑ knows
does not know

$$a \vee b = \max(a, b)$$

$$P_a^{(k)}(s, s') = \frac{N_k(s, a, s')}{1 + N_k(s, a)} = 0$$

$$C_k = \{ P = (P_a(s)) \mid \forall s, a : \|P_a^{(k)}(s) - P_a(s)\|_1 \leq \beta(N_k(s, a)) \}$$

$$\beta: \mathbb{N} \rightarrow (0, \infty)$$

$$\beta(0) = 1 \rightarrow$$

Goal of choosing β :

- ① $P^* \in C_k \quad \forall k$
wp. $1-\delta$
- ② C_k "not too large"

$$\pi_k = \operatorname{argmax}_{\pi} v_{\tilde{P}_k}^{\pi}(S_0^{(k)})$$

$$\tilde{P}_k = \operatorname{argmax}_{P \in C_k} v_P^*(S_0^{(k)})$$

Follow π_k up to the end of the episode.

Step 1: $C_k = ?$

Step 2: $R_k \leq ?$