

- Multi-armed bandits (MAB)

So far in online learning: agent receive instructive feedback

MAB: A basic model of agent making sequential decisions and receiving evaluative feedback.

A slot machines

① Setup: action set $A = \{1, \dots, A\}$. (arm set) 
"one armed bandit".

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

agent takes action a_t (arm)

agent receives reward $r_t = f^*(a_t) + \underline{\epsilon}_t$ (zero mean, 1-SG noise)

$f^*: A \rightarrow [0, 1]$,
agent's expected reward $= \mathbb{E} \left[\sum_{t=1}^T r_t \right] = \mathbb{E} \left[\sum_{t=1}^T f^*(a_t) \right]$

optimal strategy if f^* is known: take $a^* = \underset{a \in A}{\operatorname{argmax}} f^*(a)$

optimal expected reward $= T \cdot f^*(a^*)$.

Performance measure: regret:

$$\text{Reg}(T) = T f^*(a^*) - \mathbb{E} \left[\sum_{t=1}^T f^*(a_t) \right] = \mathbb{E} \left[\sum_{t=1}^T (f^*(a^*) - f^*(a_t)) \right]$$

How to design a strategy for the agent?

* need to learn the f^* fn thru noisy observations

(exploration) $\rightarrow$ Take diverse actions

* need to take actions a_t that has large f^* value

(exploitation) $\rightarrow$ Take actions that maximizes f^* .

Conflicting!

② Algorithms for MAB.

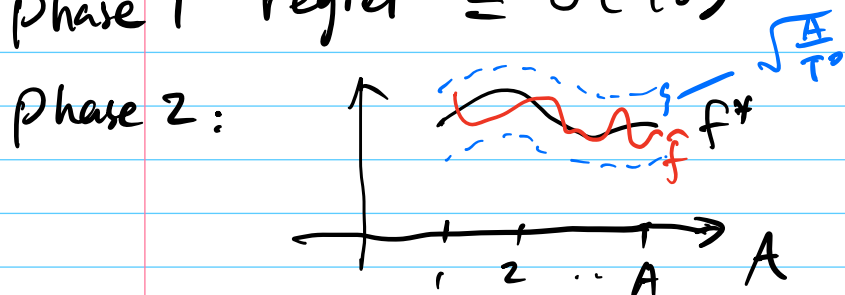
First idea: explore-then-exploit (commit)

Phase 1: use first T_0 rounds to estimate f^* by taking actions in a round robin fashion $\rightarrow \hat{f}$

Phase 2: from T_0+1 round on. take action $\hat{a} = \arg \max_{a \in A} \hat{f}(a)$

Analysis sketch:

Phase 1 regret $\leq O(T_0)$



Hoeffding
 $\downarrow$
 $\max_a |\hat{f}(a) - f^*(a)| \leq O(\sqrt{\frac{A}{T_0}})$

$$\Rightarrow f^*(a^*) - f^*(\hat{a}) \leq O(\sqrt{\frac{A}{T_0}})$$

$$\Rightarrow \text{Phase 2 regret} \leq O((T - T_0) \sqrt{\frac{A}{T_0}})$$

$$\Rightarrow \text{Reg}(T) \leq O(T_0 + T \sqrt{\frac{A}{T_0}}) = O(A^{\frac{1}{3}} T^{\frac{2}{3}})$$

choosing T_0 optimally

Second idea: intersperse explore & exploit
 ϵ -greedy.

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

Flip a coin Z w/ head prob. ϵ .

(Explore) If $Z = \text{head}$, take $a_t \sim \text{unif}(\{1 \dots A\})$

(Exploit) If $Z = \text{tail}$, take $a_t = \underset{a}{\hat{a}_t = \arg\max} \hat{f}_t(a)$

(Notation: $\hat{f}_t(a) = \text{avg reward of } a \text{ at } t$)

$$\frac{\text{total reward of } a}{\text{\# times } a \text{ taken}} = \frac{\sum_{i=1}^{t-1} \mathbb{I}(a_i = a) r_i}{N_{t-1}(a)} \rightarrow \sum_{i=1}^{t-1} \mathbb{I}(a_i = a)$$

Analysis sketch

— at step t , each a has $\approx t \frac{\epsilon}{A}$ reward samples

$$\Rightarrow \max_a |\hat{f}_t(a) - f^*(a)| \leq O\left(\sqrt{\frac{A}{t\epsilon}}\right).$$

$$\Rightarrow f^*(a^*) - f^*(\hat{a}_t) \leq O\left(\sqrt{\frac{A}{t\epsilon}}\right)$$

$\Rightarrow$ "on average", a_t is $\left(\epsilon + \sqrt{\frac{A}{t\epsilon}}\right)$ -optimal

$$\text{— Reg}(T) \leq O\left(\sum_{t=1}^T \left(\epsilon + \sqrt{\frac{A}{t\epsilon}}\right)\right)$$

$$= O\left(\epsilon T + \sqrt{\frac{AT}{\epsilon}}\right).$$

setting ϵ optimally

$$= O\left(A^{\frac{1}{3}} T^{\frac{2}{3}}\right).$$

Can we do better?

A better idea: optimism in face of uncertainty.
(abbrev. optimism principle)

In words:
Act according to the best plausible world.
"optimistic world model"

Why it works:

If the optimistic world model is

correct ✓

no regret

✓

wrong

learn new things

avoid making the same mistake in the future

✓

Real world agents are using this! (perhaps w/o realizing it)
the optimism bias (Shuraf, 2011).

How to define the "best plausible world" in MAB?

For every action a , what's the highest plausible $f^*(a)$ given data?

upper confidence bound for $f^*(a)$!
(UCB)

The UCB Algorithm:

For $t = 1, \dots, T$:

- Define $UCB_t(a) = \hat{f}_t(a) + \underbrace{b_t(a)}_{\text{confidence width.} \rightarrow \text{"bonus" of action } a}$

→ define to be 0 if $N_{t-1}(a) = 0$.

$$b_t(a) = \sqrt{\frac{l}{N_{t-1}(a)+1}}, \quad l = 8 \ln(2AT).$$

- choose $a_t = \underset{a \in A}{\operatorname{argmax}} UCB_t(a)$.

(Optimism property)

Validity of the UCB's:

Lemma: there exists an event E . $P(E) \geq 1 - \frac{2}{T}$.

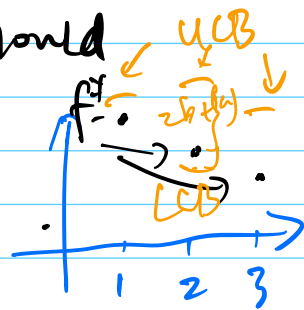
and when E happens:

$$(*) \quad \forall a, \forall t, \quad |\hat{f}_t(a) - f^*(a)| \leq b_t(a) = \sqrt{\frac{l}{N_{t-1}(a)+1}}$$

($\Rightarrow UCB_t(a) \geq \mu(a)$)

Intuition: arm a reward estimate should

have precision $\propto \frac{1}{\sqrt{\text{\#samples in arm } a}}$



pf idea:

First idea: $E_{t,a} = \{ |\hat{f}_t(a) - f^*(a)| \leq b_t(a) \}$

Hoeffding $\Rightarrow P(E_{t,a}) \geq 1 - \delta$?

Problem: $\hat{f}_t(a)$ is an average over $N_{t+1}(a)$ samples,
 $N_{t+1}(a)$ itself is random!

A fix: $E_{t,a,n} = \left\{ \left| \hat{f}_t(a) - f^*(a) \right| \leq \frac{\sqrt{\frac{1}{n+1}}}{b_{t+1}(a)}, \right\}$
 $N_{t+1}(a) = n$

Hoeffding $\Rightarrow P(E_{t,a,n}) \geq 1 - \frac{2}{AT^3}$

Define $E = \bigcap_{t=1}^T \bigcap_{a=1}^A \bigcap_{n=1}^T E_{t,a,n}$.

$$\Rightarrow P(E) \geq 1 - \frac{2}{T}.$$

can show that on E , (*) happens.

More details: Slivkins, "Introduction to Multi-armed bandits", 2019. See. 1.3.1. ✗

Regret analysis of UCB:

Thm: Assume that

UCB guarantees that $\text{Reg}(T) \leq \tilde{O}(\sqrt{AT})$

big-O ignoring
↑ log factors

Remark: — regret bound much better than $AT^{2/3}$ in ETC. & ϵ -greedy

— in general unimprovable (there is a lower)

— can also show UCB has instance dep't gr't: ^{bound}

$$\text{Reg}(T) = \tilde{O}\left(\sum_{a \neq a^*} \frac{\ln T}{\Delta a}\right) \cdot \Delta a = f^*(a^*) - f^*(a)$$
 ^{Δa larger, arm a less confus.}

Pf: Recall:
$$\text{Reg}(T) = \mathbb{E}\left[\sum_{t=1}^T \underbrace{(f^*(a^*) - f^*(a_t))}_{r(t)}\right]$$

$$= \mathbb{E}\left[\left(\sum_{t=1}^T \text{reg}(t)\right) (I(E) + I(E^c))\right]$$

Note:
$$\mathbb{E}\left[\left(\sum_{t=1}^T \text{reg}(t)\right) I(E^c)\right] \leq T \cdot \mathbb{E}[I(E^c)]$$

$$= T \cdot P(E^c)$$

$$\leq T \cdot \frac{2}{T} = 2. \quad \textcircled{1}$$

we just need to control $\sum_{t=1}^T \text{reg}(t)$ when E happens.

Insight 1: can bound $\text{reg}(t)$ by $b_t(a_t)$: bonus on the action taken

$$\text{reg}(t) = f^*(a^*) - f^*(a_t)$$

validity of UCB $\leq \text{UCB}_t(a^*) - f^*(a_t)$

choice of a_t $\leq \text{UCB}_t(a_t) - f^*(a_t)$

$$= \underbrace{\hat{f}_t(a_t)} + b_t(a_t) - \underbrace{f^*(a_t)}$$

event $E \leq 2b_t(a_t)$.

Insight 2: $\sum_{t=1}^T b_t(a_t)$ is controlled.

$$b_t(a) = \sqrt{\frac{1}{N_{t-1}(a)+1}}$$

(every time a is chosen as a_t ,
its bonus $b_t(a)$ will shrink a bit)

$$\sum_{t=1}^T \text{reg}(t)$$

$$\leq 2 \sum_{t=1}^T b_t(a_t)$$

$$= 2 \sum_{a=1}^A \sum_{t: a_t=a} b_t(a)$$

(group according
to a_t)

a decreasing sequence

say, $a=1$.

$$= 2 \sum_{a=1}^A \sum_{t: a_t=a} \sqrt{\frac{1}{N_{t-1}(a)+1}}$$

$$\begin{array}{ccc} \times & \times & \times \\ \downarrow & \downarrow & \downarrow \\ 1 & 5 & 9 \end{array}$$

$$= 2 \sum_{a=1}^A \sum_{n=1}^{N_T(a)} \sqrt{\frac{1}{n}}$$

$$\sqrt{\frac{1}{0+1}} \quad \sqrt{\frac{1}{1+1}} \quad \sqrt{\frac{1}{2+1}}$$

$$= 4\sqrt{1} \sum_{a=1}^A \sqrt{N_T(a)}$$

$$\left(\sum_{n=1}^N \sqrt{\frac{1}{n}} \leq 2\sqrt{N} \right)$$

$$\leq 4\sqrt{1} \sqrt{AT}$$

$$\left(\frac{1}{A} \sum_{a=1}^A \sqrt{N_T(a)} \right)$$

$$\leq \sqrt{\frac{1}{A} \sum_{a=1}^A N_T(a)}$$

$$= \sqrt{\frac{T}{A}}$$

by Jensen & concavity
of $x \rightarrow \sqrt{x}$

Therefore,

$$\mathbb{E} \left[\left(\sum_{t=1}^T \text{reg}(t) \right) I(E) \right] \leq 4\sqrt{LAT} \quad (2)$$

Summing up (1) (2)

$$\Rightarrow \text{Reg}(T) \leq 4\sqrt{LAT} + 2 = \tilde{O}(\sqrt{AT}).$$

Remarks on designing optimism-based algs:

① Define bonus so UCB's are valid wh.p. (so $\text{reg}_t \leq b_t(a_t)$)

②. subject to ①. design bonus to be as tight as possible. (so $\sum_t b_t(a_t)$ is small)

Stochastic linear (contextual) bandits

① Setup.

Application: recommendation with personalization

MAB: recommend a single product so that it has maximum overall satisfaction over population

Contextual Bandits: for different users. recommend different products that fit their respective needs.

Protocol:

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

observes context $x_t \in \mathcal{X}$

takes action $a_t \in \mathcal{A}$

receive reward $r_t = \underbrace{f^*(x_t, a_t)}_{\text{unknown}} + \underline{\epsilon}_t$

(generalizes MAB by allowing contexts)

Goal: maximize $\mathbb{E} \left[\sum_{t=1}^T r_t \right] = \mathbb{E} \left[\sum_{t=1}^T f^*(x_t, a_t) \right]$

Assumption (Linear Realizability)

This lecture: $f^* \in \mathcal{F} = \left\{ f(x, a) = \underbrace{\langle \theta, \phi(x, a) \rangle}_{\substack{\mathbb{R}^d \quad \mathbb{R}^d}} : \|\theta\|_2 \leq 1 \right\}$ with ϕ known

— assumption: $\|\phi(x, a)\|_2 \leq 1$ for all x, a .

— Notation:

$f^* = \langle \theta^*, \phi(x, a) \rangle$; θ^* needs to be learned.

What's the best strategy had f^* been known?

$\forall x$, take action $a = \pi_{f^*}(x) = \arg \max_{a \in A} f^*(x, a)$.

Regret notion: $P \text{ Reg}(T)$: ^{optimal policy} pseudo regret

$$\text{Reg}(T) = \mathbb{E} \left[\underbrace{\sum_{t=1}^T \max_a f^*(x_t, a)}_{\text{reward by following optimal policy}} - \underbrace{\sum_{t=1}^T f^*(x_t, a_t)}_{\text{reward of the agent}} \right].$$

Practical relevance: "A contextual bandit approach to personalized News Article Recommendation".
(Li, Chu, Langford, Schapire, 2010, WWW)
Test of Time Award.

MAB as linear bandits:

Every ^{K-armed} MAB problem can be viewed as a K-dimensional Linear bandit problem by defining:

— $x_t = z_0 \quad \forall t$ (dummy context)

— $\phi(z_0, a) = e_a = \begin{pmatrix} 0 \\ \vdots \\ i \\ \vdots \\ 0 \end{pmatrix} \leftarrow a\text{-th location} \in \mathbb{R}^k$

— $\Theta^* = (f^*(1) \dots f^*(K))$, thus

$$f^*(a) = \langle \Theta^*, \phi(z_0, a) \rangle \quad \forall a.$$

② Designing algorithms for linear bandits:

Idea: optimism principle, again.

↓
Act according to the best plausible world.

* "world model" = reward model = Θ .

* "plausible world" = set of Θ 's that can plausibly be Θ^* .
At t : i.e. can explain observations
 $(x_s, a_s, r_s)_{s=1}^{t-1}$

Similar to constructing confidence intervals in MAB.
here we construct confidence set for Θ^* , say.

$$\mathbb{P}(\forall t, \Theta^* \in \mathbb{C}_t) \geq 1 - \frac{1}{T}.$$

Algorithm (OFUL: optimism in the face of uncertainty for Linear bandits)

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

- construct confidence set $\mathbb{C}_t$ for Θ^* .
- observe x_t .
- Take action $a_t = \arg\max_{a \in A} \max_{\theta \in \mathbb{C}_t} \langle \theta, \phi(x_t, a) \rangle$
The largest plausible reward action a can get

Q1: How to construct confidence set $\mathbb{C}_t$?

Q2: How to analyze the OFUL alg?

Q1: Recall: in MAB, confidence interval for $f^*(a)$

$$= \left[\underbrace{\hat{f}_t(a)}_{\text{center: some best guess of } f^*(a)} \pm \underbrace{b_t(a)}_{\substack{\text{width deduced} \\ \text{by concentration} \\ \text{ineq.} \\ \text{(e.g. Hoeffding)}}} \right]$$

What's our "best guess" of θ^* based on $(x_s, a_s, r_s)_{s=1}^t$?

Recall: $r_s = \langle \theta^*, \phi(x_s, a_s) \rangle + \varepsilon_s$

Estimating θ^* from data is a linear regression problem

Standard solution: least squares + regularization

$$\hat{\theta}_t(\lambda) = \underset{\theta}{\operatorname{argmin}} \sum_{s=1}^t \left(\underbrace{\langle \theta, \phi(x_s, a_s) \rangle}_{\phi_s} - r_s \right)^2 + \lambda \|\theta\|_2^2$$

How close is $\hat{\theta}_t(\lambda)$ to θ^* ?

Note: in general, cannot claim, say $\|\hat{\theta}_t(\lambda) - \theta^*\|_2 \leq \frac{1}{\sqrt{t}}$

Problem: data may be highly "degenerate."

Example

$d=2$.

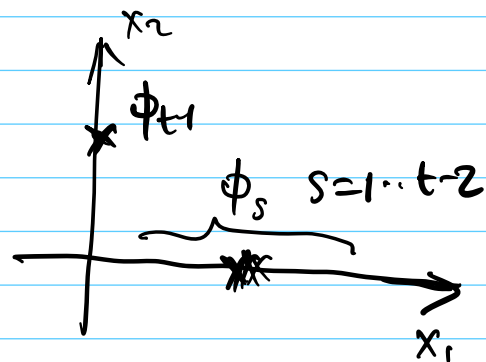
$$\phi_1 = \phi_2 = \dots = \phi_{t-2} = e_1, \quad \phi_{t-1} = e_2$$

$$r_1 = \theta_1^* + \varepsilon_1$$

$$r_2 = \theta_1^* + \varepsilon_2$$

...

$$r_{t-2} = \theta_1^* + \varepsilon_{t-2}$$



$$r_{t-1} = \theta_2^* + \varepsilon_t$$

(only one data pt is related).

The data reveals little information about θ_2^* , so cannot hope to show $|\hat{\theta}_2^t - \theta_2^*| = O(\frac{1}{\sqrt{t}})$.

In general, we will show the following estimation guarantee:

Lemma: $\exists$ event E , $P(E) \geq 1 - \frac{1}{t}$, s.t. for all t ,

$$\theta^* \in \mathcal{H}_t(\lambda) \Rightarrow \|\hat{\theta}^t(\lambda) - \theta^*\|_{V_t(\lambda)} \leq \beta_t(\lambda) = \tilde{O}(\sqrt{\lambda} + d)$$

Specifically, $\lambda = 1 \Rightarrow$

$$\theta^* \in \mathcal{H}_t(1) \Rightarrow \|\hat{\theta}^t(1) - \theta^*\|_{V_t(1)} \leq \beta_t(1) = \tilde{O}(d)$$

This, although not a pointwise guarantee, turns out to be still useful.

Here: $\|x\|_M := \sqrt{x^T M x}$ For $M \succ 0$ This is a norm: $\|x\|_M = |a| \|x\|_M$
 $= \|M^{\frac{1}{2}} x\|_2$ psd $M = U(\lambda_1 \dots \lambda_d) U^T$ $\|x+y\|_M \leq \|x\|_M + \|y\|_M$
 $M^{\frac{1}{2}} = U M^{\frac{1}{2}} U^T$ $\langle x, y \rangle \leq \|x\|_M \|y\|_M$

$$V_t(\lambda) = \sum_{s=1}^{t-1} \phi_s \phi_s^T + \lambda I$$

generalized Cauchy Schwarz.

Intuition check on previous example w/ $\lambda = 0$:

$$\hat{\theta}^t = \underset{\theta}{\operatorname{argmin}} \sum_{s=1}^{t-1} (\theta_1 - (\theta_1^* + \varepsilon_s))^2 + (\theta_2 - (\theta_2^* + \varepsilon_{t-1}))^2$$

$$\downarrow$$

$$\theta_1 = \frac{1}{t-2} \sum_{s=1}^{t-2} (\theta_1^* + \varepsilon_s) \quad \theta_2 = \theta_2^* + \varepsilon_{t-1}$$

$$\text{let } z = \hat{\theta}^t - \theta^*$$

$$\Rightarrow |z_1| \leq O(\sqrt{\frac{1}{t}}), |z_2| < O(1).$$

$$* \text{ Note: } V_{t-1} = (t-1) e_1 e_1^T + 1 \cdot e_2 e_2^T = \begin{pmatrix} t-1 & 0 \\ 0 & 1 \end{pmatrix}$$

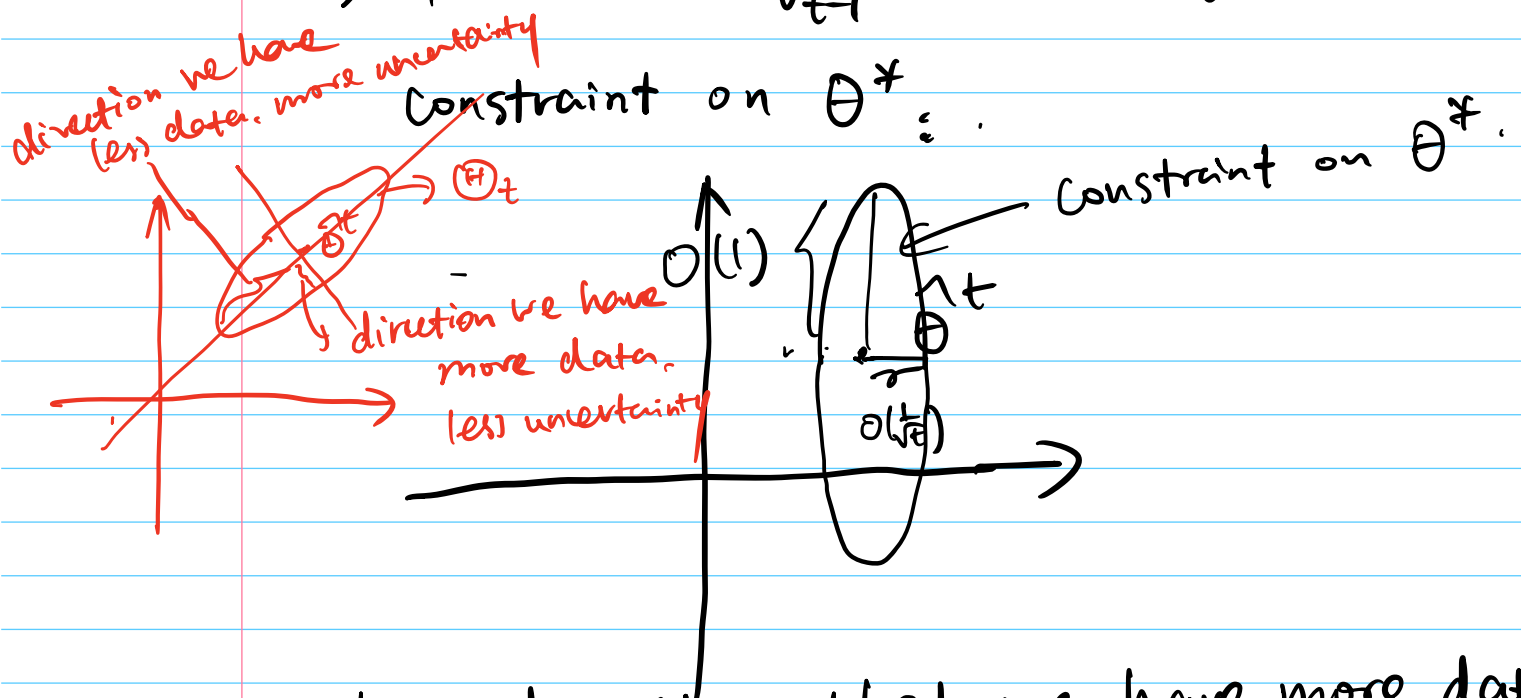
$$\Rightarrow \|z\|_{V_t}^2 = (z_1, z_2) \begin{pmatrix} t-1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}$$

$$= (t-1) z_1^2 + z_2^2$$

$$\leq O\left(t \cdot \frac{1}{t} + 1 \cdot \frac{1}{t}\right) = O(1)$$

* Why is $\|z\|_{V_{t-1}}^2 \leq O(1)$ guarantee useful?

$$\Rightarrow \|\theta^* - \hat{\theta}^t\|_{V_{t-1}}^2 \leq O(1) \text{ gives a}$$



For those directions that we have more data, our estimate on θ^* is more accurate.

Proof of the Lemma: ($V_{t-1}(\lambda)$ abbrev. V_{t-1} below)

$$\text{First, } \hat{\theta}_t(\lambda) = V_{t-1}^{-1} \left(\sum_{s=1}^{t-1} \phi_s r_s \right)$$

$$= V_{t-1}^{-1} \left(\underbrace{\sum_{s=1}^{t-1} \phi_s \phi_s^T}_{V_{t-1}} \theta^* + \underbrace{\sum_{s=1}^{t-1} \phi_s \varepsilon_s}_{m_{t-1}} \right)$$

$$= \theta^* - V_{t-1}^{-1} \lambda \cdot \theta^* + V_{t-1}^{-1} m_{t-1}$$

$$\Rightarrow \hat{\theta}_t(\lambda) - \theta^* = -V_{t-1}^{-1} \lambda \theta^* + V_{t-1}^{-1} m_{t-1}$$

$$\Rightarrow \|\hat{\theta}_t(\lambda) - \theta^*\|_{V_t} \leq \|V_{t-1}^{-1} \lambda \theta^*\|_{V_t} + \|V_{t-1}^{-1} m_{t-1}\|_{V_t}$$

defn of Mahalanobis norm

$$= \|\lambda \theta^*\|_{V_{t-1}^{-1}} + \|m_{t-1}\|_{V_{t-1}^{-1}}$$

w.p. $1 - \frac{1}{T^2}$

$$\leq \sqrt{\lambda} + \sqrt{2 \ln(T^2) + d \ln(1 + \frac{t}{d\lambda})}$$

union bound over all $t \Rightarrow$ Lemma.

In the last step, to deal with $\| \underbrace{m_{t-1}}_{\text{martingale}} \|_{V_{t-1}^{-1}}$, we use

self-normalized Tail Inequality for vector-valued martingales:

~~Lemma~~ (Abbasi-Yadkori, Pál, Szepesvári, 2011)

$\forall t,$

$$\mathbb{P} \left(\left\| \sum_{s=1}^t \phi_s \varepsilon_s \right\|_{V_t^{-1}(\lambda)} \leq \sqrt{2 \ln \frac{1}{\delta} + \ln \underbrace{\frac{\det(V_t(\lambda))}{\det(V_0(\lambda))}}_{\leq d \ln(1 + \frac{t}{d\lambda})}} \right) \geq 1 - \delta$$

Intuition: when all $\varepsilon_s \sim N(0, 1)$ and ϕ_s 's are chosen ahead of time

$$m_t = \sum_{s=1}^t \phi_s \varepsilon_s \stackrel{d}{\approx} N(0, V_t(\theta))$$

$$V_t^{-\frac{1}{2}}(\lambda) \cdot m_t \stackrel{d}{\approx} N(0, I_d)$$

$$\| V_t^{-\frac{1}{2}}(\lambda) m_t \|_2 \stackrel{\text{whp.}}{\leq} O(d + \ln \frac{1}{\delta})$$

problem with this intuition: ϕ_s 's are not fixed:

ϕ_s depends on $\phi_1, \epsilon_1, \dots, \phi_{s-1}, \epsilon_{s-1}$.

Q1 summary:

$\forall t$. define $\Theta_t = \{ \theta : \|\theta - \hat{\theta}^t(\cdot)\|_{V_{t+1}(\cdot)} \leq \beta_t(\cdot) \}$

$$P(\forall t, \theta^* \in \Theta_t) \geq 1 - \frac{1}{T}.$$

Algorithm induced by this Θ_t construction:

$$a_t = \operatorname{argmax}_{a \in A} \left[\max_{\theta \in \Theta_t} \langle \theta, \phi(x_t, a) \rangle \right]$$

Interpretation:
 $=: \text{UCB}_t(x_t, a)$

$$\max_{x: \|x\|_{V_{t+1}(\cdot)} \leq \beta_t(\cdot)} \langle \hat{\theta}^t(\cdot) + x, \phi(x_t, a) \rangle$$

Exercise

$$\left[\begin{array}{l} \theta^* \in \Theta_t \Rightarrow \text{this } \square \\ \geq \langle \theta^*, \phi(x_t, a) \rangle \\ = f^*(x_t, a) \end{array} \right]$$

$$\langle \hat{\theta}^t(\cdot), \phi(x_t, a) \rangle + \beta_t(\cdot) \|\phi(x_t, a)\|_{V_{t+1}(\cdot)}^{-1}$$

$$\phi(z_0, a) = e_a$$

$$V_{t+1}(\cdot) = \sum_{s=1}^t \phi(z_s, a_s) \phi(z_s, a_s)^T + I$$

↑
Estimated reward

$$= (N_{t+1}(\cdot) + I)$$

$\beta_t(\cdot)$

uncertainty / exploration bonus

For MAB: $\sqrt{K}$

$$\sqrt{\frac{1}{N_{t+1}(a) + 1}}$$

Q2 Regret analysis

Thm: on E ($P(E) \geq 1 - \frac{1}{T}$), $P \text{Reg}(T) \leq \tilde{O}(d\sqrt{T})$.

($\Rightarrow \text{Reg}(T) \leq \tilde{O}(d\sqrt{T})$.)

(For MAB, what regret bound do we get?)

Pf: Recall $P \text{Reg}(T) = \sum_{t=1}^T \text{reg}(t)$, where

$$\text{reg}(t) = \max_a f^*(x_t, a) - f^*(x_t, a_t)$$

How to bound? Similar to MAB analysis

Step 1: $\text{reg}(t) \leq 2 b_t(a_t)$

why: $\text{reg}(t) = \max_a f^*(x_t, a) - f^*(x_t, a_t)$

$$\leq \max_a \text{UCB}_t(a) - f^*(x_t, a_t)$$

$$= \text{UCB}_t(a_t) - f^*(x_t, a_t)$$

$$= \langle \hat{\theta}^{t(1)}, -\theta^*, \phi(x_t, a_t) \rangle + b_t(a_t)$$

$$\begin{aligned} &\leq \underbrace{\langle \hat{\theta}^{t(1)} - \theta^*, \phi(x_t, a_t) \rangle}_{\leq b_t(a_t)} + b_t(a_t) \\ &\leq \underbrace{\|\hat{\theta}^{t(1)} - \theta^*\|_{V_{t+1}^{-1}}}_{\leq \beta_{t+1}} \|\phi(x_t, a_t)\|_{V_{t+1}^{-1}} + b_t(a_t) \end{aligned}$$

$$= 2 b_t(a_t)$$

Step 2: bounding $\sum_t b_t(a_t)$

how:
$$\sum_t b_t(a_t) = \sum_t \beta_t(l) \|\phi(x_t, a_t)\|_{V_{t+1}(l)}^{-1}$$

$$\leq \underbrace{\left(\max_t \beta_t(l) \right)}_{\tilde{O}(\sqrt{d})} \underbrace{\left(\sum_{t=1}^T \|\phi_t\|_{V_{t+1}(l)}^{-1} \right)}_{(*)}$$

$$(*) = \sum_{t=1}^T \|\phi_t\|_{V_{t+1}(l)}^{-1}$$

$$\leq \sqrt{\sum_{t=1}^T \|\phi_t\|_{V_{t+1}(l)}^2} \cdot \sqrt{T}$$

$$= \sqrt{\sum_{t=1}^T \phi_t^T V_{t+1}(l)^{-1} \phi_t} \cdot \sqrt{T}$$

Note: $V_t(l) \leq 2 V_{t+1}(l) \Rightarrow V_{t+1}(l)^{-1} \leq 2 V_t(l)^{-1}$

special case: $d=1$.

$A \leq B \Leftrightarrow A - B \geq 0$. positive semi definite order of matrices

property: $A \leq B \Rightarrow \forall x. x^T A x \leq x^T B x$

$A \leq B \Rightarrow B^{-1} \leq A^{-1}$.

$$\leq \sqrt{\sum_{t=1}^T \phi_t^T V_t(l)^{-1} \phi_t} \cdot \sqrt{2T}$$

The elliptic potential lemma (EPL)

suppose $u_1, \dots, u_T \in \mathbb{R}^d$. $A_t = \mu I + \sum_{s=1}^t u_s u_s^T$, then

$(A_0 = \mu I)$

$$\sum_{t=1}^T \|u_t\|_{A_t^{-1}}^2 \leq \ln \frac{\det(A_T)}{\det(A_0)}$$

if all $\|u_t\| \leq D$

$$\leq d \ln \left(1 + \frac{D^2 T}{d\mu} \right) = \tilde{O}(d)$$

$$\leq \tilde{O}(\sqrt{dT})$$

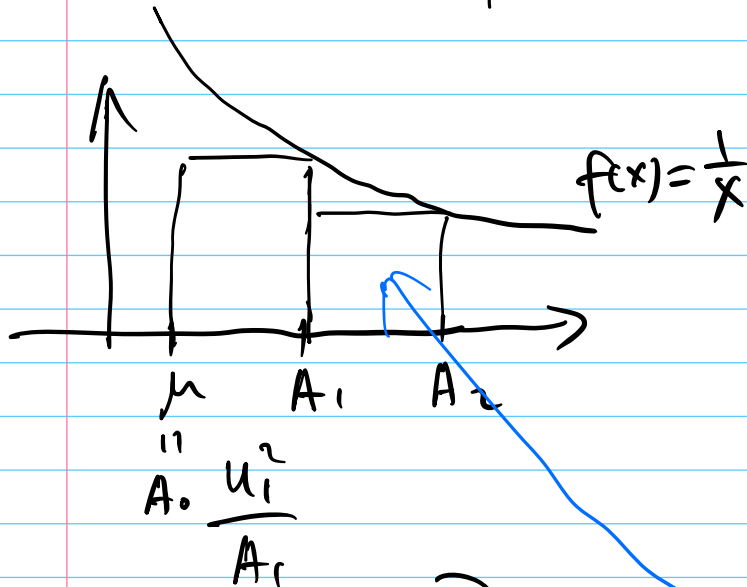
$$\Rightarrow \sum_{t=1}^T b_t(a_t) \leq \tilde{O}(\sqrt{d}) \cdot \tilde{O}(\sqrt{dT}) = \tilde{O}(d\sqrt{T})$$

Intuition of EPL

when $d=1$.

$$A_t = \mu + \sum_{s=1}^t u_s^2$$

$$\sum_{t=1}^T \frac{u_t^2}{A_t} \leq \int_{\mu}^{A_T} \frac{1}{x} dx = \ln \frac{A_T}{\mu}$$



pf idea of EPL: (1) show that $\|u_t\|_{A_t}^2 \leq \ln \frac{\det(A_t)}{\det(A_{t-1})}$

Note: $\ln \frac{\det(A_t)}{\det(A_{t-1})}$

$$(1-d): \frac{u_t^2}{A_t} \leq \ln \frac{A_t}{A_{t-1}}$$

$$= \ln (\det(A_{t-1}^{-\frac{1}{2}} A_t A_{t-1}^{-\frac{1}{2}}))$$

$$= \ln (\det(I + A_{t-1}^{-\frac{1}{2}} u_t u_t^T A_{t-1}^{-\frac{1}{2}}))$$

Exercise (Hint: $\det(M) = \prod_{i=1}^d \lambda_i(M)$)

$$= \ln(1 + \|u_t\|_{A_{t-1}^{-1}}^2)$$

$$\geq \frac{\|u_t\|_{A_{t-1}^{-1}}^2}{1 + \|u_t\|_{A_{t-1}^{-1}}^2}$$

$$(\forall x \geq 0, \ln(1+x) \geq \frac{x}{1+x})$$

$$= \|u_t\|_{A_t^{-1}}^2$$

$$(A_t^{-1} = (A_t + u_t u_t^T)^{-1})$$

Sherman-Morrison

$$= A_{t-1}^{-1} - \frac{A_{t-1}^{-1} u_t u_t^T A_{t-1}^{-1}}{1 + u_t^T A_{t-1}^{-1} u_t}$$

(2) show $\ln \frac{\det(A_T)}{\det(A_0)} = \tilde{O}(d)$

Fact: for p.s.d matrices $A = U \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_d \end{pmatrix} U^T$, $\det(A) = \prod_{i=1}^d \lambda_i$

$$\text{For } A_0 = \begin{pmatrix} \mu & & \\ & \ddots & \\ & & \mu \end{pmatrix} \Rightarrow \det(A_0) = \mu^d$$

$$\text{For } A_T = \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_d \end{pmatrix} \Rightarrow \det(A_T) = \prod_{i=1}^d \lambda_i \stackrel{\text{AM-GM}}{\leq} \left(\frac{1}{d} \sum_{i=1}^d \lambda_i \right)^d$$

$$= \left(\frac{1}{d} \text{tr}(A_T) \right)^d$$

$$\leq \left(\frac{1}{d} (d\mu + TD^2) \right)^d$$

~~*~~