

(Online) Multi-Armed Bandit Learning

"online learning with partial feedback"

Model

- k actions, n rounds
- action j 's payoff in round i : $V_{j,i}$
- in round i :
 - (a) choose action j^i
 - (b) learn payoff $V_{j^i,i}$
 - (c) obtain payoff $V_{j^i,i}$
- ALG's payoff = $\sum_{i=1}^n V_{j^i,i}$

Goal: total payoff close to best in hindsight.

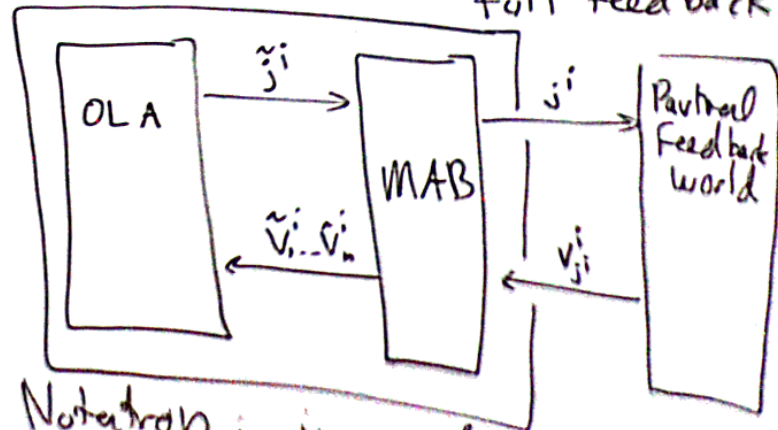
Note: identical to online learning except only learn $V_{j^i,i}$ not $V_{1,i}, \dots, V_{k,i}$

Note: if we don't take action j , don't learn if j is good or not.

Challenge: "tradeoff of exploration and exploitation"

Reducing MAB to Online Learning

"reduce partial feedback to full feedback"



Notation: in round i

- probabilities of actions: $\Pi^i = (\pi_{1,i}, \dots, \pi_{k,i})$
 $\pi_{j,i} = P(j^i = j)$
- choose action $j^i \sim \Pi^i$
- payoffs $V^i = (V_{1,i}, \dots, V_{k,i})$

• expected payoff:

$$\begin{aligned} E[V_j^i] &= \sum_{j=1}^K E[V_j^i | j^i=j] P(j^i=j) \\ &= \sum_{j=1}^K \frac{V_j^i}{\pi_j^i} \times \pi_j^i \\ &= V^i \cdot \Pi^i \quad (\text{vector dot product}) \end{aligned}$$

Challenge 1: what to report to OLA?

Idea 1: report "unbiased estimator" of true payoffs.

- if MAB uses probabilities $\Pi^i = (\pi_1^i, \dots, \pi_K^i)$
- samples $j^i \sim \Pi^i$
- real payoffs $V^i = (V_1^i, \dots, V_K^i)$
- only has $V_{j^i}^i$

• report: $\tilde{V}^i = (0, \dots, \frac{V_{j^i}^i}{\pi_{j^i}^i}, \dots, 0)$

Lemma 1: expected report is equal to actual payoff

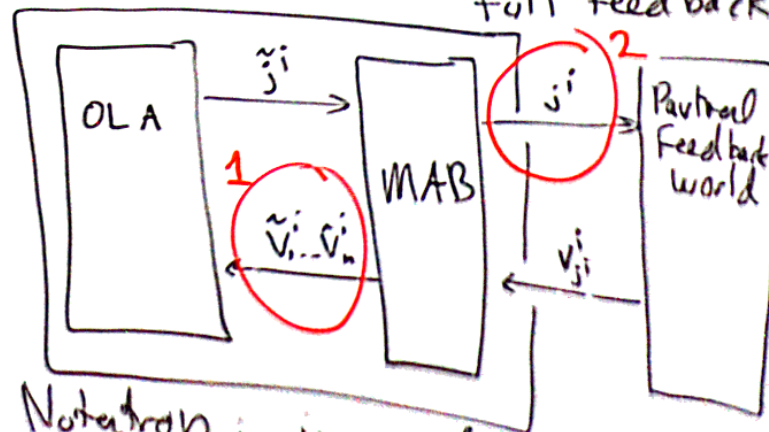
Proof $E[\tilde{V}_j^i] = E[\tilde{V}_j^i | j^i=j] P(j^i=j) + E[\tilde{V}_j^i | j^i \neq j] P(j^i \neq j)$

$$= V_j^i \cdot \frac{\pi_j^i}{\pi_j^i} \times \pi_j^i + 0 \cdot \pi_{\neq j}^i = V_j^i$$

Challenge: "tradeoff of exploration and exploitation"

Reducing MAB to Online Learning

"reduce partial feedback to full feedback"



Notation: in round i

- probabilities of actions: $\Pi^i = (\pi_1^i, \dots, \pi_K^i)$
- $\pi_j^i = P(j^i=j)$
- choose action $j^i \sim \Pi^i$
- payoffs $V^i = (V_1^i, \dots, V_K^i)$

Note: for actual payoffs $v_j \in [0, h]$

- reported payoffs $\tilde{v}_j \in \{0, v_j/\pi_j\}$

- upper bound on reported payoffs is

$$\hat{h} = \max_{i,j} v_j/\pi_j$$

- if π_j is small, $\hat{h}$ is big.

overall OLA regret $= O(h\sqrt{\frac{\log k}{n}})$

Challenge 2: keep $\hat{h}$ small.

Idea 2: pick random action with some minimal probability ϵ .

- if OLA recommends $\hat{\pi}^i$

- MAB uses π^i with $\pi_j^i = (1-\epsilon)\hat{\pi}_j^i + \epsilon/k$

- then $\pi_j^i \geq \epsilon/k$

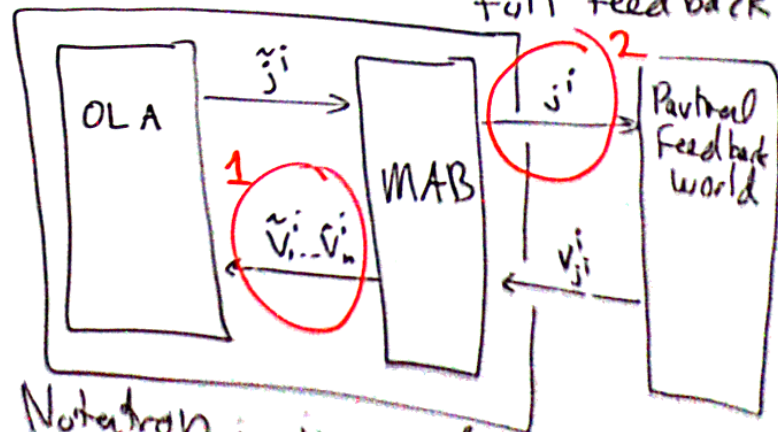
Lemma 2: if $\pi_j^i \geq \epsilon/k$ then $\tilde{v}_j \leq \hat{h} = \frac{kh}{\epsilon}$

Proof: $\tilde{v}_j \leq v_j/\pi_j \leq v_j k \leq h k \frac{\epsilon}{\epsilon} = \hat{h}$

Challenge: "tradeoff of exploration and exploitation"

Reducing MAB to Online Learning

"reduce partial feedback to full feedback"



Notation: in round i

- probabilities of actions: $\pi^i = (\pi_1^i, \dots, \pi_k^i)$

$$\pi_j^i = P(j^i = j)$$

- choose action $j^i \sim \pi^i$

- payoffs $v^i = (v_1^i, \dots, v_k^i)$

ALG: MAB reduction to OLA

in round i :

1) $\tilde{\pi}^i \leftarrow \text{OLA}$

2) draw $j^i \sim \tilde{\pi}^i$

$$\pi_j^i = (1-\epsilon)\pi_j^{i-1} + \epsilon/k$$

3) take action j^i ; learn $v_{j^i}^i$

4) report to OLA

$$\tilde{v}_j^i = \begin{cases} v_{j^i}^i / \pi_j^i & \text{if } j = j^i \\ 0 & \text{otherwise} \end{cases}$$

Thm: for payoffs in $[0, \tilde{h}]$, if OLA satisfies

$$E[\text{OLA}] \geq (1-\epsilon)\text{OPT}(\tilde{v}) - \frac{\tilde{h}}{\epsilon} \ln k$$

then for payoffs in $[0, \tilde{h}]$, MAB satisfies

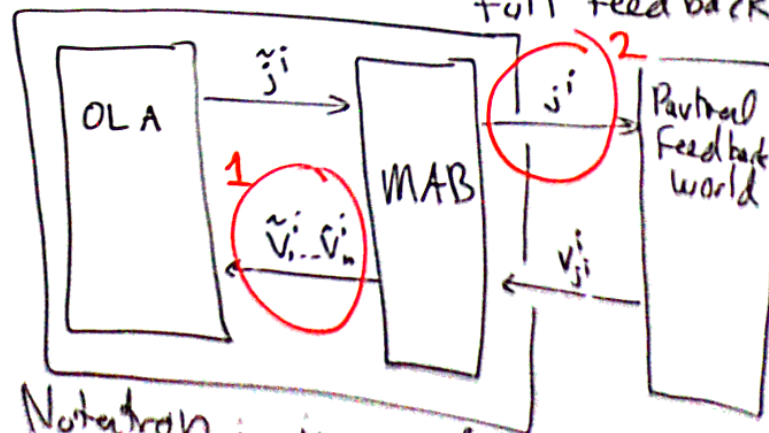
$$E[\text{MAB}] \geq (1-2\epsilon)\text{OPT}(\tilde{v}) - \frac{\tilde{h}k}{\epsilon^2} \ln k$$

Cor: MAB with Exp Weights has vanishing regret.

Challenge: "tradeoff of exploration and exploitation"

Reducing MAB to Online Learning

"reduce partial feedback to full feedback"



Notation: in round i

- probabilities of actions: $\pi^i = (\pi_1^i, \dots, \pi_k^i)$
 $\pi_j^i = P(j^i = j)$
- choose action $j^i \sim \pi^i$
- payoffs $v^i = (v_1^i, \dots, v_k^i)$

3. combine

$$E(\text{MAB}) \geq \frac{(1-\epsilon)^2 \sum v_i^j}{(1-2\epsilon+\epsilon^2) \text{OPT}} - R$$

$$\geq (1-2\epsilon) \text{OPT} - R$$

$$= (1-2\epsilon) \text{OPT} - \frac{h}{\epsilon} \ln k$$

$$= (1-2\epsilon) \text{OPT} - \frac{hk}{\epsilon^2} \ln k$$

Thm: for profits in $[0, \tilde{h}]$, if OLA satisfies

$$E[\text{OLA}] \geq (1-\epsilon) \text{OPT}(\tilde{v}_1) - \frac{\tilde{h}}{\epsilon} \ln k$$

then for profits in $[0, h]$, MAB satisfies

$$E[\text{MAB}] \geq (1-2\epsilon) \text{OPT}(v) - \frac{hk}{\epsilon^2} \ln k$$

(or: MAB with Exp weights has vanishing regret).

Analysis

"online learning works on unbiased estimators"

$$0. R = \frac{\tilde{h}}{\epsilon} \ln k$$

$$j^x = \arg\max_j \sum_i v_j^i$$

1. OLA's guarantee on $\tilde{v}_1, \dots, \tilde{v}_n$

$$\text{OLA} = \sum_i \hat{\pi}^i \cdot \tilde{v}_i \geq (1-\epsilon) \sum_i \tilde{v}_i^{j^x} - R$$

$$E[\text{OLA}] = \sum_i E[\hat{\pi}^i \cdot \tilde{v}_i] \geq (1-\epsilon) \sum_i E[\tilde{v}_i^{j^x}] - R$$

$$= \sum_i E[\hat{\pi}^i \cdot v_i] \geq (1-\epsilon) \sum_i v_i^{j^x} - R$$

2. MAB's performance on $v_1, \dots, v_n$

$$\text{MAB} = \sum_i \pi^i \cdot v_i$$

$$= (1-\epsilon) \sum_i \hat{\pi}^i \cdot v_i + \frac{\epsilon}{k} \sum_i v_i^{j^x}$$

$$\geq (1-\epsilon) \sum_i \hat{\pi}^i \cdot v_i$$