

CS 396: Online Markets

Lecture 10: Learning to Bid

Last Time:

- auction theory
- second-price auction
- first-price auction
- complete information analysis (Nash equilibrium)
- incomplete information analysis (Bayes-Nash equilibrium)

Today:

- learning to bid
 - full feedback
 - partial feedback
 - equilibrium of no-regret learning (coarse correlated equilibrium)
-

Exercise: Optimal Bid

Recall:

- cumulative distribution function: $F_X(z) = \Pr[X \leq z]$
- uniform distribution on $[0, 1]$: $F_X(z) = z$
- first-price auction: highest bidder wins, winner pays bid.

Setup:

- you are bidding in a first-price auction
- other bidders with i.i.d. uniform bids on $[0, 1]$

Questions: your value is $v = 1/2$

- What should you bid against one other bidder?
 - What should you bid against two other bidders?
-

Online Bidding

“bidding in a repeated auction”

Model:

- repeated first-price auction
 - highest bidder wins (random tie-breaking)
 - winner pays bid
- static value $v \in [1, h]$
- in round i :
 - bid b^i
 - competing bid $\hat{b}^i \in [1, h]$
 - win if $b^i > \hat{b}^i$
 - * payoff: $v - b^i$
 - lose otherwise
 - * payoff: 0
- feedback model:
 - full: learn $\hat{b}_i^i$
 - partial: learn “win” or “lose”

Discretization

“discretize bid space, run learning algorithm”

Examples: for bidder with value v

- linear discretization: $b_j = v - \epsilon j$ for $k = h/\epsilon$
- geometric discretization: $b_j = v - (1 + \epsilon)^j$ for $k = \log_{1+\epsilon} h \approx \frac{1}{\epsilon} \ln h$

Note: e.g.,

- with $b_j = v - (1 + \epsilon)^j$,
- on winning utility is $v - b_j = (1 + \epsilon)^j$.

Recall Thm: n rounds, k actions, payoffs in $[0, h]$:

$$EW \geq (1 - \epsilon) \text{OPT} - h/\epsilon \ln k$$

Cor: EW and linear discretization is:

$$EW \geq (1 - 2\epsilon) \text{OPT} - \frac{h}{\epsilon^2} \ln h/\epsilon$$

Cor: EW and geometric discretization is:

$$EW \geq (1 - 2\epsilon) \text{OPT} - \frac{h}{\epsilon^2} \ln \frac{1}{\epsilon} \ln h$$

Can solve for ϵ to optimize per-round regret

Q: can these bounds be improved?

A: yes

Idea: non-uniform bounds on payoffs

- action j is bid b_j , payoff is either 0 or b_j
- upper bound of h is lose.

Full Feedback

“faster learning for full feedback”

Recall: Follow the Perturbed Leader (FTPL)

- learning rate ϵ
- hallucinate: $v_j^0 = h \times$ “num tails of ϵ -bias coin flipped in a row”
- let $V_j^i = \sum_{r=1}^i v_j^r$
- in round i choose: $j^i = \operatorname{argmax}_j v_j^0 + V_j^{i-1}$

Recall Analysis:

- stability lemma: FTPL and BTPL do the same thing w.p. $(1 - \epsilon)$
- small perturbation:
BTPL $\geq$ OPT $-\mathbf{E}[\text{MAXPTRB}]$

Idea: non-uniform hallucination:

- payoffs of action j are in $[0, h_j]$
- hallucinate: $v_j^0 = h_j \times$ “num tails of ϵ -bias coin flipped in a row”

Lemma: hallucination for geometric discretizing
 $\mathbf{E}[\text{MAXPTRB}] \leq h/\epsilon^2$

Proof:

- $\mathbf{E}[\text{geometric with rate } \epsilon] = 1 - \epsilon/\epsilon$
- $\mathbf{E}[\max_j v_j^0] \leq \mathbf{E}[\sum_j v_j^0] = \sum_j h_j (1 - \epsilon)/\epsilon$
- geometric $h_j = (1 + \epsilon)^j$
- $\sum_{j=0}^k (1 + \epsilon)^j \leq h \sum_{j=0}^{\infty} (1 + \epsilon)^{-j} \leq h \frac{1 + \epsilon}{\epsilon}$
- $\mathbf{E}[\text{MAXPTRB}] \leq h/\epsilon^2$

Thm: FTPL and geometric discretization is
 $\text{ALG} \geq (1 - 2\epsilon) \text{OPT} - \frac{h}{\epsilon^2}$

Partial Feedback

“faster learning for partial feedback”

Recall: MAB Reduction to OLA

In round i :

1. $\pi \leftarrow \text{OLA}$
2. draw $j^i \sim \tilde{\pi}$ with

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

3. take action j^i
4. report $\tilde{v}$ to OLA with

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

Recall Thm:

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

Recall Analysis:

Challenge 2: keep $\tilde{h}$ small

Idea 2: pick random action with some minimal probability ϵ/k

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

Q: improve this with $h_j = (1 + \epsilon)^j$?

A: geometric exploration

- explore bid $b_j = (1 + \epsilon)^j$ with prob. $\propto (1 + \epsilon)^{-j}$
- $H = \sum_{j=0}^k (1 + \epsilon)^j \approx h/\epsilon$
- $\tilde{h}_j = (1 + \epsilon)^j H / \epsilon (1 + \epsilon)^j \approx h/\epsilon^2$

Thm: MAB with geometric exploration $\mathbf{E}[\text{MAB}] \geq (1 - 3\epsilon) \text{OPT} - h/\epsilon^3 \ln h/\epsilon^2$

Equilibrium of No-regret Learning

“outcomes of games under learning”

Recall: mixed Nash: players indepently randomize

Q: would learning in repeated game converge to independent randomization?

A: not generally.

Defn: coarse correlated equilibrium (CCE)

- mediator offers joint distributions of actions (a_R, a_C)
- players either
 - follow mediator
 - pick any fixed outcome
- CCE if best response is to follow mediator

Example: rock-paper-sissors

	Rock	Paper	Sissors
Rock	-6,-6	-1,1	1,-1
Paper	1,-1	-6,-6	-1,1
Sissors	-1,1	1,-1	-6,-6

Q: Nash?

A: uniform mixing

Q: other CCE?

A: uniform mixing over $\{(R, P), (R, S), (P, R), (P, S), (S, R), (S, P)\}$

- payoff from following mediator: 0
- payoff from any fixed action: $0 \times 2/3 - 6 \times 1/3 = -2$

Thm: play is no-regret iff distribution of play is CCE.

- suppose $((a_R^0, a_C^0), \dots, (a_R^n, a_C^n))$ is no-regret
- no-regret for player R, for all a_R^* :

$$\sum_i R_{a_R^i, a_C^i} \geq \sum_i R_{a_R^*, a_C^i}$$
- consider mediator:
 - pick i uniformly from round $\{1, \dots, n\}$
 - recommend a_R^i to R and a_C^i to C .
- no-regret $\Leftrightarrow$ CCE.