

Online Bidding

"learning to bid in repeated auctions"
Model

- repeated first-price auction
 - highest bid wins
 - winner pays bid.

• static value v

• in round i :

- bid b^i
- competing bid $\hat{b}^i$
- win if $b^i > \hat{b}^i$
 - payoff: $(v - b^i)$
- lose o.w.
 - payoff: 0

• feedback model

(cf. online learning) full: learn $\hat{b}^i$

(cf. MAB learning) partial: learn "win" or "lose"

Discretization

"discretize action space, run learning alg"

Examples: for bids in $[1, h]$

- linear discretization

$$b_j = 1 + \epsilon j \quad \text{for } j \in \{0, \dots, k\}$$

$k = \frac{h-1}{\epsilon}$

- geometric discretization

$$b_j = (1 + \epsilon)^j \quad \text{for } j \in \{0, \dots, k\}$$

$$k = \log_{1+\epsilon} h \approx \frac{1}{\epsilon} \ln h$$

Recall Thm: n rounds, k actions, payoff $\in [0, 1]$

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

Cor: for linear discretization

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

Cor: for geometric discretization

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

Compare $\boxed{\text{OPT (of all bids in } [1, h])}$
 $\boxed{\text{OPT (of discretized bids)}}$

e.g. linear discretization

$$\text{OPT}(\text{discrete}) \geq \text{OPT}(\text{all}) - \epsilon n$$

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

$$(1 - \epsilon) \text{OPT}(\text{discrete}) \geq (1 - \epsilon)^2 \text{OPT}(\text{all})$$

$$(1 - 2\epsilon + \epsilon^2) \text{OPT}(\text{all})$$

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

to go from $\boxed{\text{discrete OPT}}$
to $\boxed{\text{continuous OPT}}$

Q: can these bounds be improved?
A: yes

Idea: non uniform bounds on payoffs of actions.

- if we take action b_j :
- payoff is 0 or $v - b_j = b_j$

Discretization

"discretize action space, run learning alg"

Examples: for bids in $[1, h]$

- linear discretization

$$b_j = v - \epsilon j \text{ for } j \in \{0, \dots, k\}$$

$$k = \frac{h}{\epsilon}$$

- geometric discretization

$$b_j = v(1 + \epsilon)^j \text{ for } j \in \{0, \dots, k\}$$

winning payoffs like " ϵj " or " $(1 + \epsilon)^j$ " $k = \log_{1+\epsilon} h \approx \frac{1}{\epsilon} \ln h$

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

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

Cor: for linear discretization

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

Cor: for geometric discretization

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

Full Feedback

"Faster learning with full feedback"

Recall: FTPL

- learning rate ϵ
- hallucinate $v_j^0 = h \times$ "# tails of ϵ -bias coin in a row"
- let $V_j^i = \sum_{r=1}^i v_j^r$

• in round i choose $j^i = \arg \max_j v_j^0 + V_j^{i-1}$

Recall Analysis:

- stability lemma: FTPL & BTPL do the same thing up to $(1-\epsilon)$

• small perturbation:
 $BTPL \geq OPT - \epsilon \underbrace{(\text{MAXPIRB})}_{\text{is small}}$

Note: for geometric discretization action b_j payoff is in $[0, h_j = (1+\epsilon)^j]$

Idea: hallucinate $v_j^0 = h_j \times$ "# ϵ -biasing in a row"

Discretization

"discretize action space, run learning alg"

Examples: for bids in $[1, h]$

- linear discretization

$$b_j = v - \epsilon j \quad \text{for } j \in \{0, \dots, k\}$$

$k = \frac{h}{\epsilon}$

- geometric discretization

$$b_j = v(1+\epsilon)^j \quad \text{for } j \in \{0, \dots, k\}$$

winning payoffs like " ϵj " or " $(1+\epsilon)^j$ " $k = \log_{1+\epsilon} h \approx \frac{1}{\epsilon} \ln h$

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

$$EW \geq (1-\epsilon) \boxed{OPT} - \frac{h}{\epsilon} \ln k$$

Cor: for linear discretization

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

Cor: for geometric discretization:

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

Full Feedback

"Faster learning with Full Feedback"

Recall: FTPL

- learning rate ϵ
- hallucinate $v_j^0 = h \times$ "# tails of ϵ -bias coin in a row"
- let $V_j^i = \sum_{r=1}^i v_j^r$

• in round i choose $j^i = \arg \max_j v_j^0 + V_j^{i-1}$

Recall Analysis:

(1) stability lemma: FTPL & BTPL do the same thing w.p. $(1-\epsilon)$

(2) small perturbation:
 $BTPL \geq OPT - \underbrace{E[\text{MAXPTRB}]}_{\text{is small}}$

Note: for geometric discretization action b_j payoff is in $[0, h_j = (1+\epsilon)^j]$

Idea: hallucinate $v_j^0 = h_j \times$ "# ϵ -biasing in a row"

Lemma: Stability (same argument)

FTPL & BTPL do the same thing with probability $(1-\epsilon)$.

Lemma: small perturbation

$$E[\text{MAXPTRB}] \leq h/\epsilon^2$$

[c.f. $h/\epsilon \log k$]

Proof:

- $E[\text{geometric w. rate } \epsilon] = \frac{1-\epsilon}{\epsilon}$
 - $E[\max_j v_j^0] \leq E[\sum_j v_j^0] = \sum_j \frac{1-\epsilon}{\epsilon} h_j = \frac{1-\epsilon}{\epsilon} \sum (1+\epsilon)^j$
 - $H = \sum_{j=0}^k (1+\epsilon)^j = (1+\epsilon)^k \sum_{j=0}^k (1+\epsilon)^{j-k}$
 $\leq (1+\epsilon)^k \sum_{j=0}^{\infty} (1+\epsilon)^{j-k}$ " $j' = k-j$ "
 $= (1+\epsilon)^k \frac{(1+\epsilon)}{\epsilon} = \frac{1+\epsilon}{\epsilon} h$
 - $E[\max_j v_j^0] \leq \frac{1+\epsilon}{\epsilon} h$
- $\leq (1+\epsilon)h/\epsilon^2 \leq h/\epsilon^2$

Thm: FTPL with geometric distribution

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

[c.f. $\frac{h}{\epsilon^2} \ln \frac{1}{\epsilon} \ln h$ min]

Lemma: Stability (same argument)

FTPL & BTPL do the same thing with probability $(1-\epsilon)$.

Lemma: small perturbation

$$E[\text{MAXPTRB}] \leq h/\epsilon^2$$

[c.f. $h/\epsilon \log k$]

Proof:

- $E[\text{geometric w. rate } \epsilon] = \frac{1-\epsilon}{\epsilon}$.
- $E[\max_j v_j^0] \leq E[\sum_j v_j^0]$
 $= \sum_j \frac{1-\epsilon}{\epsilon} h_j = \frac{1-\epsilon}{\epsilon} \sum (1+\epsilon)^j$
- $H = \sum_{j=0}^k (1+\epsilon)^j = (1+\epsilon)^k \sum_{j=0}^k (1+\epsilon)^{j-k}$
 $\leq (1+\epsilon)^k \sum_{j'=0}^{\infty} (1+\epsilon)^{j'} \quad "j' = k-j"$
 $= (1+\epsilon)^k \frac{(1+\epsilon)}{\epsilon} = \frac{1+\epsilon}{\epsilon} h.$
- $E(\max_j v_j^0) \leq \frac{(1+\epsilon)h}{(1-\epsilon)\epsilon} \leq h/\epsilon^2$