



Microsoft
Research



NEURAL INFORMATION
PROCESSING SYSTEMS

Efficient and Near-Optimal Algorithm for Contextual Duelling Bandits with Offline Regression Oracles

Aadirupa Saha (UIC), Robert E Schapire (MSR)

Corresponding: aadirupa.saha@gmail.com



39th Annual Conference on Neural Information Processing Systems

Learning from Preference (Relative) Feedback



--- How much you score it out of 5?



Rankings (Relative)

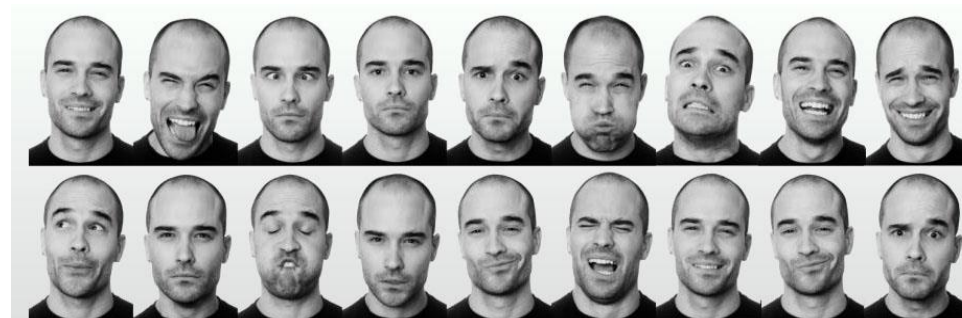
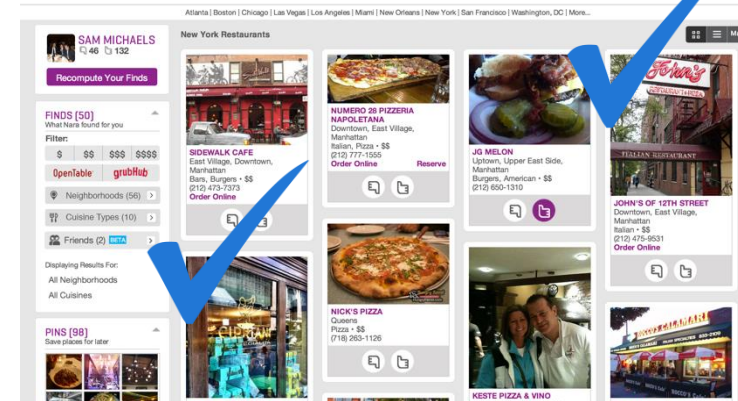
--- Do you like movie A over B?

Multitude of Applications

Learning from Preferences



Recommender Systems

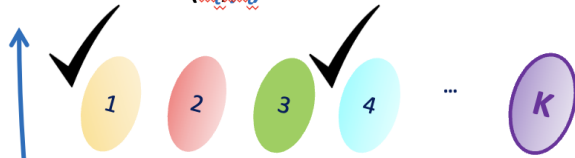


Performance Measure?

Problem setup: Dueling Bandits

At round $t = 1, 2, 3 \dots T$

Select **two** arms (x_t, y_t)



repeat

Observe (noisy) **comparison** $o_t \in \{0, 1\}$ $\sim P :=$

Condorcet winner (i^*)

Yue et al. The K-armed dueling bandits problem. Jour. Of CSS 2012.

Yue and Joachims. Beat the mean bandit. ICML 2011.

Objective: Find the **BEST** arm(s)

[[Regret minimization (or)
PAC best-arm identification]]

Regret :=

$$\sum_{t=1}^T \frac{\left[P(i^*, x_t) - \frac{1}{2} \right] + \left[P(i^*, y_t) - \frac{1}{2} \right]}{2}$$

P: Preference Matrix ($K = 5$)

	1	2	3	4	5
1	0.5	0.53	0.54	0.56	0.6
2	0.47	0.5	0.53	0.58	0.61
3	0.46		0.5	0.54	0.57
4	0.44	0.42		0.5	0.51
5	0.4	0.39	0.43	0.49	0.5

Limitations of the Existing Framework

Regret :=

$$\sum_{t=1}^T \frac{\left[P(i^*, x_t) - \frac{1}{2} \right] + \left[P(i^*, y_t) - \frac{1}{2} \right]}{2}$$

1. Fixed benchmark (i^*) over T
2. Stationary Preferences over T
3. No contextualization!!

Best-Response Regret!

A Strong Notion of Regret!



Standard Regret $:= \max_{i^* \in [K]} \sum_{t=1}^T \frac{\{[P_t(i^*, x_t) - \frac{1}{2}] + [P_t(i^*, y_t) - \frac{1}{2}]\}}{2}$

Fixed benchmark over T

Strongest opponent

Best-response Regret (BR-Reg_T) $:= \sum_{t=1}^T \max_{i_t^* \in [K]} \left[\frac{\{[P_t(i_t^*, x_t) - \frac{1}{2}] + [P_t(i_t^*, y_t) - \frac{1}{2}]\}}{2} \right]$

However...

Rock Paper Scissor

$$P := \begin{pmatrix} 0 & 1 & -1 \\ -1 & 0 & 1 \\ 1 & -1 & 0 \end{pmatrix}$$

Learner's choice (x_t, y_t)	Adversary's choice (i^*)
(1,2)	(1)
(2,3)	(2)
(3,1)	(3)

Learner is always forced to suffer a loss of $\frac{1}{2} \rightarrow \text{Reg}_T = O(T)$!
(too bad)

$$\tilde{P} = 2P - 1$$

	1	2	3	4	5
1	0.5	0.53	0.54	0.56	0.6
2	0.47	0.5	0.53	0.58	0.61
3	0.46	0.47	0.5	0.54	0.57
4	0.44	0.42	0.46	0.5	0.51
5	0.4	0.39	0.43	0.49	0.5

$$P \in [0,1]^{K \times K}$$

Preference matrix as a 2 player Zero-Sum Game

	1	2	3	4	5
1	0	0.06	0.08	0.12	0.20
2	-0.06	0	0.06	0.16	0.22
3	-0.08	-0.06	0	0.08	0.04
4	-0.12	-0.16	-0.08	0	0.02
5	-0.20	-0.22	-0.07	-0.02	0

$$\tilde{P} \in [-1,1]^{K \times K}$$

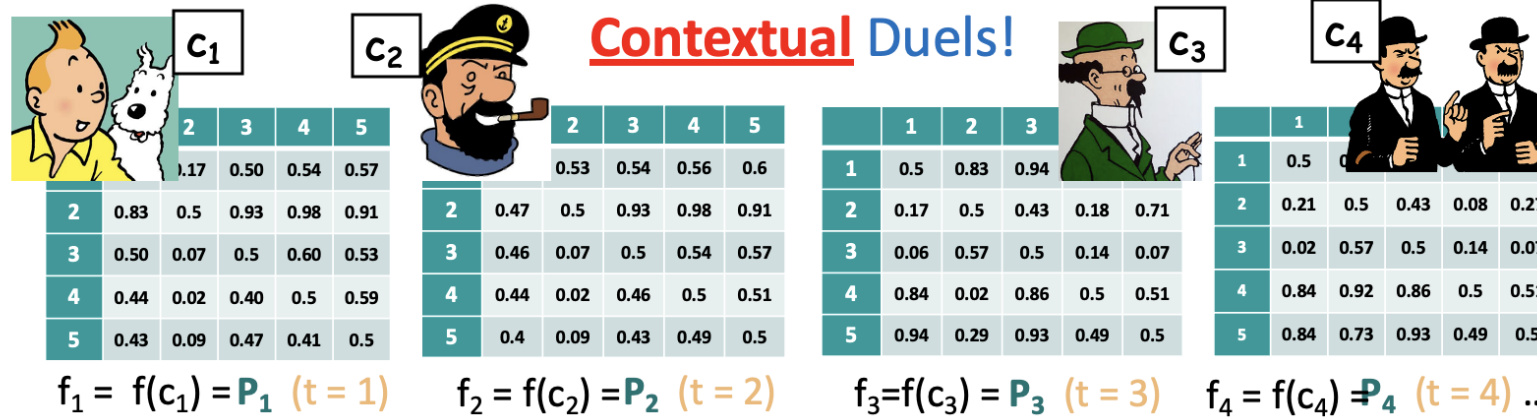
(0.06, -0.06)

... hence the Learner must randomize!

Best-Response Regret

$$\sum_{t=1}^T \max_{i_t^* \in [K]} E_{(x_t, y_t) \sim p_t} \left[\frac{\{[\tilde{P}_t(i_t^*, x_t)] + [\tilde{P}_t(i_t^*, y_t)]\}}{2} \right]$$

Personalization: Contextual DB



Assume: f : Context $\mapsto$ Preference matrix; $f \in \mathcal{M}$ (f : Unknown, $\mathcal{M}$: Known)
(Given $\mathcal{C}$: Context Class)

[Dudik et al. \(2015\)](#)

π : Context $\mapsto$ Arm

Policy Regret

"Contextual" Regret:

$$\max_{\{\pi^* \in \Pi\}} \sum_{t=1}^T E_{(x_t, y_t) \sim p_t} \left[\frac{[f_t(\pi^*(c_t), x_t)] + [f_t(\pi^*(c_t), y_t)]}{2} \right]$$

- Suboptimal
- (or) Runtime in-efficient
- But not BOTH

Contextual Best-Response DB

Best Response Regret

$$\sum_{t=1}^T \max_{i_t^* \in [K]} E_{(x_t, y_t) \sim p_t} \left[\frac{[f_t(i_t^*, x_t)] + [f_t(i_t^*, y_t)]}{2} \right]$$

where $p_t \in \Delta_{K \times K}$

Strongest opponent



Offline vs Online Regression Oracle

Offline Regression Oracle

Given a function class $\mathcal{F} \subseteq \{f \mid f : \mathcal{Z} \rightarrow \mathcal{Y}\}$ and i.i.d. data

$$D_n = \{(z_i, y_i)\}_{i=1}^n, (z_i, y_i) \sim \mathcal{D},$$

the offline oracle outputs

$$\hat{f}_n \leftarrow \text{OffReg}_{\mathcal{F}}(D_n).$$

For any $\delta \in (0, 1]$, with probability at least $1 - \delta$:

$$\mathbb{E}_{(z,y) \sim \mathcal{D}} \left[(\hat{f}_n(z) - y)^2 - \inf_{f \in \mathcal{F}} (f(z) - y)^2 \right] \leq \mathcal{E}_{\text{off}, \mathcal{F}}(\delta, n).$$

Online Regression Oracle

For the same class

$$\mathcal{F} \subseteq \{f \mid f : \mathcal{Z} \rightarrow \mathcal{Y}\},$$

the online oracle receives an arbitrary sequence

$$\{(z_t, y_t)\}_{t=1}^n$$

and guarantees:

$$\frac{1}{n} \left[\sum_{t=1}^n (\hat{y}_t(z_t) - y_t)^2 - \inf_{f \in \mathcal{F}} \sum_{t=1}^n (f(z_t) - y_t)^2 \right] \leq \mathcal{E}_{\text{on}, \mathcal{F}}(n),$$

where $\mathcal{E}_{\text{on}, \mathcal{F}}(n) = o(1)$.

Offline Oracle ($\text{OffReg}_{\mathcal{F}}$)	Online Oracle ($\text{OnReg}_{\mathcal{F}}$)
✓ Batch processing	✓ Real-time & adaptive
✓ Higher accuracy potential	✓ Stronger guarantee (adversarial-robust)
✓ Works only when test $\approx$ train	✓ Distribution-free (no i.i.d.)
✓ Easier to implement	✓ Can convert to offline (not vice-versa)
	✗ Harder to construct



Result (Finite Actions)

Algorithm 1 Double-Monster: Efficient Contextual Dueling Bandit with Offline Regressors

- 1: **input** Epoch schedule $0 = \tau_0 < \tau_1 < \tau_2 < \dots$. Confidence parameter δ . Tuning parameter $\gamma_1, \gamma_2, \dots$.
- 2: **Arm set:** $[K]$. An instance of `OffReg` for function class $\mathcal{M}$
- 3: **for** epoch $m = 1, 2, \dots$ **do**
- 4: Let $\gamma_m = \left(\frac{K}{\sqrt{\mathcal{E}_{\text{off}, \mathcal{M}}(\delta/(2m^2), \tau_{m-1} - \tau_{m-2})}} \right)$ (for epoch 1, $\gamma_1 = 1$).
- 5: Compute $\widehat{\mathbf{M}}_m \leftarrow \text{OffReg}\left(\{(c_\tau, x_\tau, y_\tau), o_\tau\}_{\tau=\tau_{m-2}+1}^{\tau_{m-1}}\right)$, i.e.
 $\widehat{\mathbf{M}}_m = \arg \min_{f \in \mathcal{M}} \sum_{\tau=\tau_{m-2}+1}^{\tau_{m-1}} (f(c_\tau)[x_\tau, y_\tau] - \tilde{o}_\tau)^2$ via the **offline least squares oracle**,
 where $\tilde{o}_\tau = 2o_\tau - 1$, for any $\tau \in [\tau_{m-2} + 1, \tau_{m-1}]$.
- 6: **for** round $t = \tau_{m-1} + 1, \dots, \tau_m$ **do**
- 7: Observe context $c_t \in \mathcal{C}$.
- 8: Compute $\mathbf{p}_t \sim \text{Cvx-Constraint-Solver}(c_t, \widehat{\mathbf{M}}_m, \gamma_m)$ s.t. $\mathbf{p}_t \in \Delta_{K \times K}$ satisfies:

$$\forall i, j \in [K] : \mathbf{E}_{x \sim \mathbf{p}_t^i} [\widehat{\mathbf{M}}_m(c_t)[i, x]] + \mathbf{E}_{y \sim \mathbf{p}_t^j} [\widehat{\mathbf{M}}_m(c_t)[j, y]] + \frac{2}{\gamma_m} \frac{1}{p_t(i, j)} \leq \frac{5K^2}{\gamma_m}.$$
- 9: Sample $(x_t, y_t) \sim \mathbf{p}_t$ and observe preference feedback $o_t \sim \text{Ber}\left(\frac{f^*(c_t)[x_t, y_t] + 1}{2}\right)$.
- 10: **end for**
- 11: **end for**

Theorem (Regret Analysis of Algorithm 1): With an epoch schedule $\tau_1, \dots, \tau_m$ such that $\tau_m \geq 2^m$ for $m \leq \log T$, with probability at least $(1 - \delta)$ for any $\delta \in (0, 1)$ the best-response regret (BR-Reg $_T$) of **Double-Monster** after T rounds is bounded by:

$$O\left(K \sum_{m=1}^{m(T)} \sqrt{\mathcal{E}_{\text{off}, \mathcal{M}}(\delta/(2m^2), \tau_{m-1} - \tau_{m-2})} (\tau_m - \tau_{m-1})\right)$$

$$\approx O(K \sqrt{T \log(|\mathcal{M}| T \log T)})$$

for special classes!

Experiments

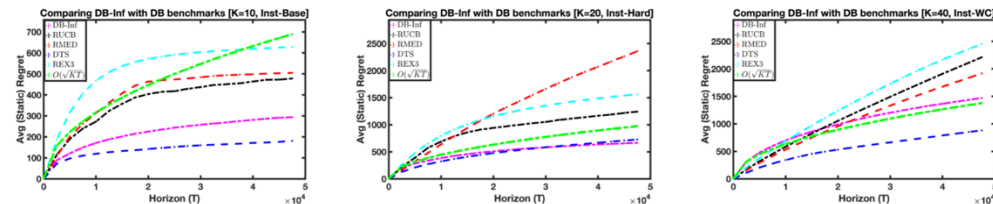


Figure 2: Comparing with Finite-armed Non-Contextual Baselines

Continuous Action Spaces

- **Setup:** Continuous decision (action) space $\mathcal{K} \subset \mathbb{R}^d$
- **Context Embedding:** $\phi^* : \mathcal{C} \rightarrow \mathbb{R}^d$, $\phi^* \in \Phi$ [Given Φ : **Embedding Class**]
- **(Expected) Action Scores:** $\mathbb{E}[s_t \mid c_t = c, \mathbf{x}_t = \mathbf{x}] = \langle \mathbf{x}, \phi^*(c) \rangle$, $\mathbf{x} \in \mathcal{K}, c \in \mathcal{C}$
- **Preference Model:** $P_t(\mathbf{x}, \mathbf{y}) = \sigma(\langle \phi^*(c_t), \mathbf{x} - \mathbf{y} \rangle)$, $\sigma(z) = \frac{1}{1+e^{-z}}$, $z \in \mathbb{R}$
Learner observes feedback $o_t \sim \text{Ber}(P_t(\mathbf{x}_t, \mathbf{y}_t))$.
- **Best-Response Regret:**

$$\text{BR-Reg}_T^{\text{cont}} = \sum_{t=1}^T \mathbb{E}_{x_t \sim \mathcal{D}} \left[\max_{\mathbf{a}^* \in \mathcal{K}} \langle \mathbf{x}^*, \phi^*(x_t) \rangle - \frac{1}{2} \langle \mathbf{x}_t + \mathbf{y}_t, \phi^*(c_t) \rangle \right]$$

Main Algorithm

Algorithm 2 Double-Monster-Inf (for Continuous Action Spaces)

- 1: **input** Epoch schedule $0 = \tau_0 < \tau_1 < \tau_2 < \dots$. Confidence parameter δ . Tuning parameter $\gamma_1, \gamma_2, \dots$
- 2: Arm set: $\mathcal{K} \subset \mathbb{R}^d$. An instance of `OffReg` for function class Φ
- 3: **for** epoch $m = 1, 2, \dots$ **do**
- 4: Set $\gamma_m = \left(\frac{\sqrt{d}}{2\sqrt{2\mathcal{E}_{\text{off}, \Phi}(\delta/(2m^2), \tau_{m-1} - \tau_{m-2})}} \right)$ (for epoch 1, $\gamma_1 = 1$).
- 5: Compute $\phi_m \leftarrow \text{OffReg}\left(\{(c_\tau, \mathbf{x}_\tau, \mathbf{y}_\tau), o_\tau\}_{\tau=\tau_{m-2}+1}^{\tau_{m-1}-1}\right)$, i.e.
 $\phi_m = \arg \min_{\phi \in \Phi} \sum_{\tau=\tau_{m-2}+1}^{\tau_{m-1}-1} (\sigma(\langle \phi(c_\tau), \mathbf{x}_\tau - \mathbf{y}_\tau \rangle) - o_\tau)^2$ via the **offline least squares oracle**
- 6: **for** round $t = \tau_{m-1} + 1, \dots, \tau_m$ **do**
- 7: Observe context $x_t \in \mathcal{X}$.
- 8: Compute $\mathbf{p}_t \sim \text{Cont-CvxConstraint-Solver}(x_t, \phi_m, \gamma_m)$ i.e.:

$$\mathbf{p}_t \leftarrow \arg \max_{\mathbf{q} \in \Delta_{\mathcal{K} \times \mathcal{K}}} \left[\langle \mathbf{a}\mathbf{b}_{\mathbf{q}}, \phi_m(c_t) \rangle + \frac{1}{\gamma_m} \log \det(\mathbf{H}_{\mathbf{q}}) \right],$$

where $\mathbf{a}\mathbf{b}_{\mathbf{q}} = \mathbf{E}_{(\mathbf{a}, \mathbf{b}) \sim \mathbf{q}}[(\mathbf{a} + \mathbf{b})/2]$, $\mathbf{H}_{\mathbf{q}} = \mathbf{E}_{(\mathbf{a}, \mathbf{b}) \sim \mathbf{q}}[(\mathbf{a} - \mathbf{b})(\mathbf{a} - \mathbf{b})^\top] + c\mathbf{I}_d$.
- 9: Sample $(\mathbf{x}_t, \mathbf{y}_t) \sim \mathbf{p}_t$ and observe preference feedback $o_t \sim \text{Ber}\left(\sigma(\langle \phi^*(c_t), \mathbf{x}_t - \mathbf{y}_t \rangle)\right)$.
- 10: **end for**
- 11: **end for**



Results

Theorem (Regret Analysis of Algorithm 2): Assume $\forall c \in \mathcal{C}, \|\phi^*(x)\| \leq D$, for some $D > 0$. Then with an epoch schedule $\tau_1, \dots, \tau_m$ such that $\tau_m \geq 2^m$ for $m \leq \log T$, with probability at least $(1 - \delta)$ for any $\delta \in (0, 1)$ the best-response regret in T rounds is bounded by:

$$O\left(\sum_{m=2}^{m(T)} \sqrt{d \mathcal{E}_{\text{off}, \Phi}(\delta/(2m^2), \tau_{m-1} - \tau_{m-2})(\tau_m - \tau_{m-1})}\right) \cdot \approx \tilde{O}(\sqrt{dT})$$

for special classes!

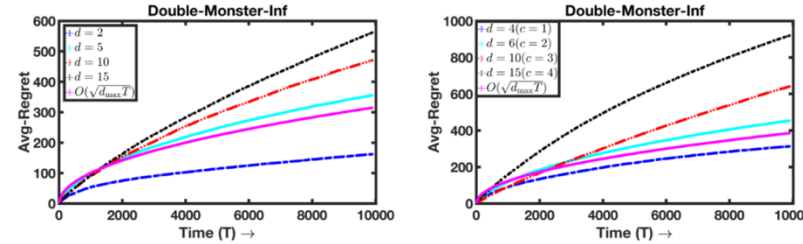


Figure 1: (Left) Inst-1 (linear transformation) and (Right) Inst-2 (quadratic transformation)

Special Classes

- **Finite class $\mathcal{M}$:** $\mathcal{E}_{\text{off}, \mathcal{M}}(\delta, n) = O\left(\frac{\log(|\mathcal{M}|n)}{n\delta}\right)$
- **Nonparametric class (with empirical entropy $O(\varepsilon^{-p})$):**

$$\mathcal{E}_{\text{off}, \mathcal{M}}(\delta, n) = O\left(n^{-2/(2+p)} \log(1/\delta)\right)$$

- **Linear predictors (d -dimensional):** $\mathcal{E}_{\text{off}, \mathcal{M}}(\delta, n) = O\left(\frac{d \log(n/d)}{n\delta}\right)$
- **Deep MLP / ReLU nets (Sobolev smoothness β):**

$$\mathcal{E}_{\text{off}, \mathcal{M}}(\delta, n) = \tilde{O}\left(n^{-\beta/(\beta+d)}\right)$$

Paper



Poster + More



Thanks!

Follow-up/ questions?

aadirupa.saha@gmail.com

<https://aadirupa.github.io/>

Looking forward :)

