

Learning-Augmented Facility Location Mechanisms for the Envy Ratio Objective

Haris Aziz¹, Yuhang Guo¹, Alexander Lam², Houyu Zhou¹
¹UNSW Sydney, ²Hong Kong Polytechnic University

Introduction

- Facility location **mechanism design** on a compact interval.
- Envy ratio objective: a **fairness metric** defined as the **maximum** ratio between the utilities of any two agents [DLC⁺20].
- Open problem: no randomized mechanism exists breaking the approximation ratio bound of 2 w.r.t. the envy ratio objective.
- Learning-augmentation: incorporating **ML/DL based prediction model** from historical data with **mechanism design** [ABG⁺22].
- An ideal augmented mechanism: optimum with accurate predictor (**consistency**) and worst-case guarantee in general (**robustness**).

Results Summary

(Our results are marked with gray background)

Mechanism	Classical	Learning-augmented
Deterministic	2	α -consistent, $\frac{\alpha}{\alpha-1}$ -robust (optimal)
Randomized	UB: 1.8944 LB: 1.1257	(2 - c ²)-consistent, (2 + c)-robust, c ∈ ($\frac{1}{4}, \frac{1}{2}$) $\frac{7}{4}$ -consistent, $\frac{9}{4}$ -robust, c ∈ [0, $\frac{1}{4}$]

Model

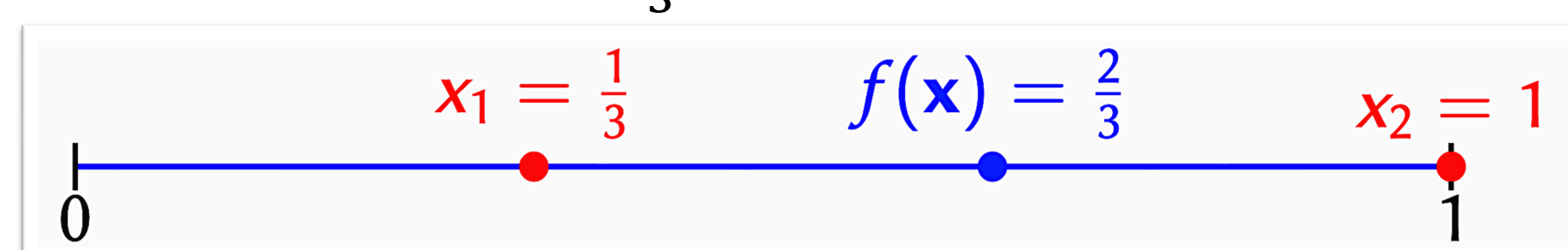
- A set of agents $N = \{1, 2, \dots, n\}$ are located on an interval with **private** locations $x = (x_1, x_2, \dots, x_n) \in [0, 1]^n$.
- A deterministic (resp. randomized) mechanism $f(x): [0, 1]^n \rightarrow [0, 1]$ (resp. $[0, 1]^n \rightarrow \Delta([0, 1])$) locates the facility at **some point** (resp. with a **distribution over some points**) on the interval.
- For any location distribution $\Delta([0, 1])$, utility of agent i :
 $u_i(\Delta([0, 1]), x_i) = 1 - \mathbb{E}_{y \sim \Delta([0, 1])} [|y - x_i|]$.

- Envy ratio objective**: $ER(f(x), x) = \max_{i \neq j} \frac{u_i(f(x), x_i)}{u_j(f(x), x_j)}$.

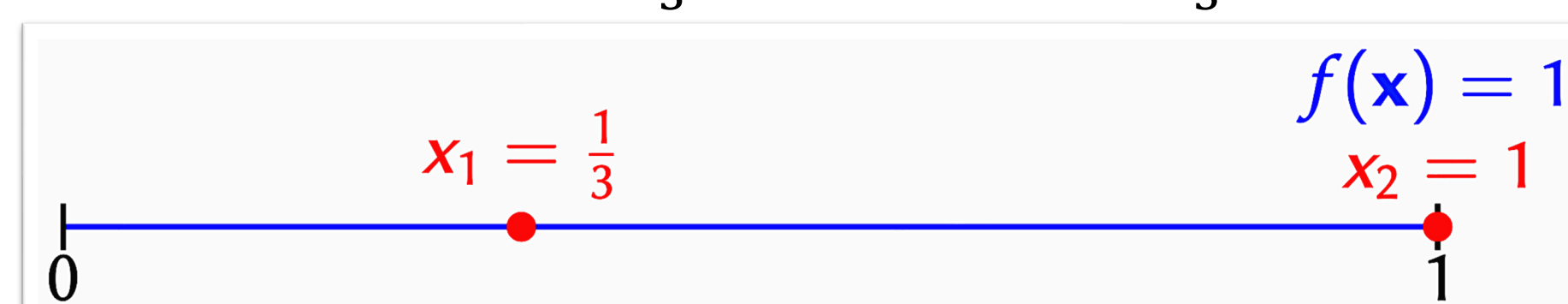
- Approximation ratio**: $\rho = \sup_{x \in [0, 1]^n} \frac{ER(f(x), x)}{ER(OPT(x), x)}$.

- Example**: consider the instance $x = (x_1 = \frac{1}{3}, x_2 = 1)$.

If $f(x) = \frac{2}{3}$, $ER(f(x), x) = \frac{1 - (\frac{2}{3} - \frac{1}{3})}{1 - (1 - \frac{2}{3})} = 1$. (OPT)



If $f(x) = 1$, $ER(f(x), x) = \frac{1 - (1 - 1)}{1 - (1 - \frac{1}{3})} = 3$. ($\rho = \frac{ER(1, x)}{ER(\frac{2}{3}, x)} = 3$)



- Strategyproofness**: no agents has incentive to misreport his/her location information.

Model

- Learning-augmented mechanism $f(x, \hat{y})$ with $\hat{y}$ as extra input (representing prediction of optimal facility location).
- Consistency**: approximation ratio when $\hat{y}$ is always accurate.

$$\gamma = \sup_{x \in [0, 1]^n} \frac{ER(f(x, \hat{y}), x)}{ER(OPT(x), x)}, \text{ with } \hat{y} = OPT(x).$$

- Robustness**: approximation ratio for any arbitrary $\hat{y}$.

$$\beta = \sup_{x \in [0, 1]^n, \hat{y} \in [0, 1]} \frac{ER(f(x, \hat{y}), x)}{ER(OPT(x), x)}, \text{ for arbitrary } \hat{y}.$$

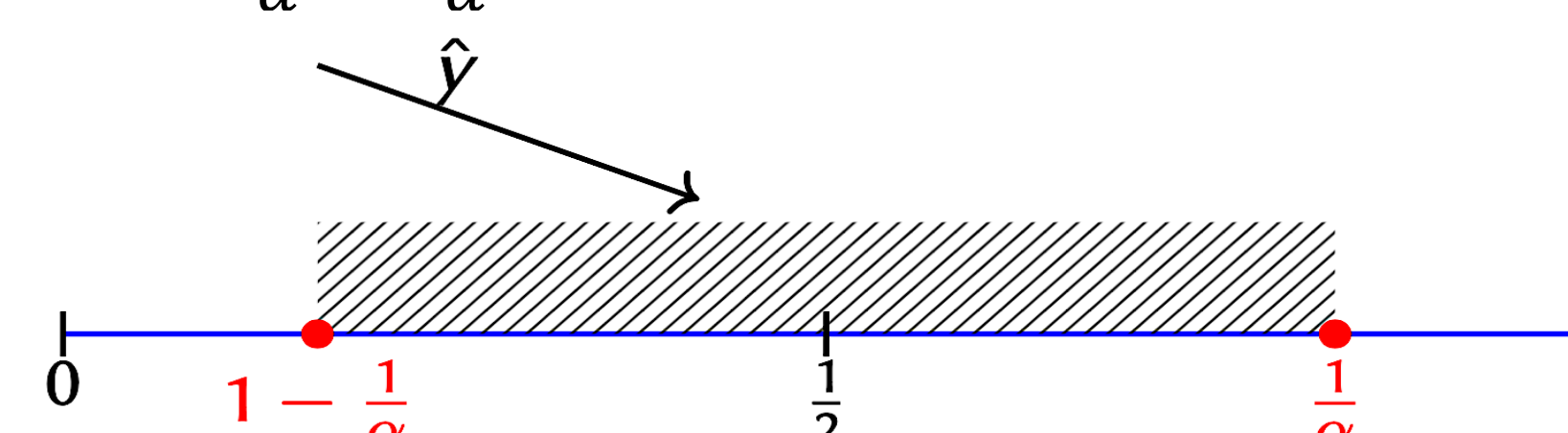
Deterministic Mechanism

In classical setting [DLC⁺20]: Placing the facility at $f(x) = \frac{1}{2}$ achieves an approximation ratio of 2 (tight).

With learning-augmentation:

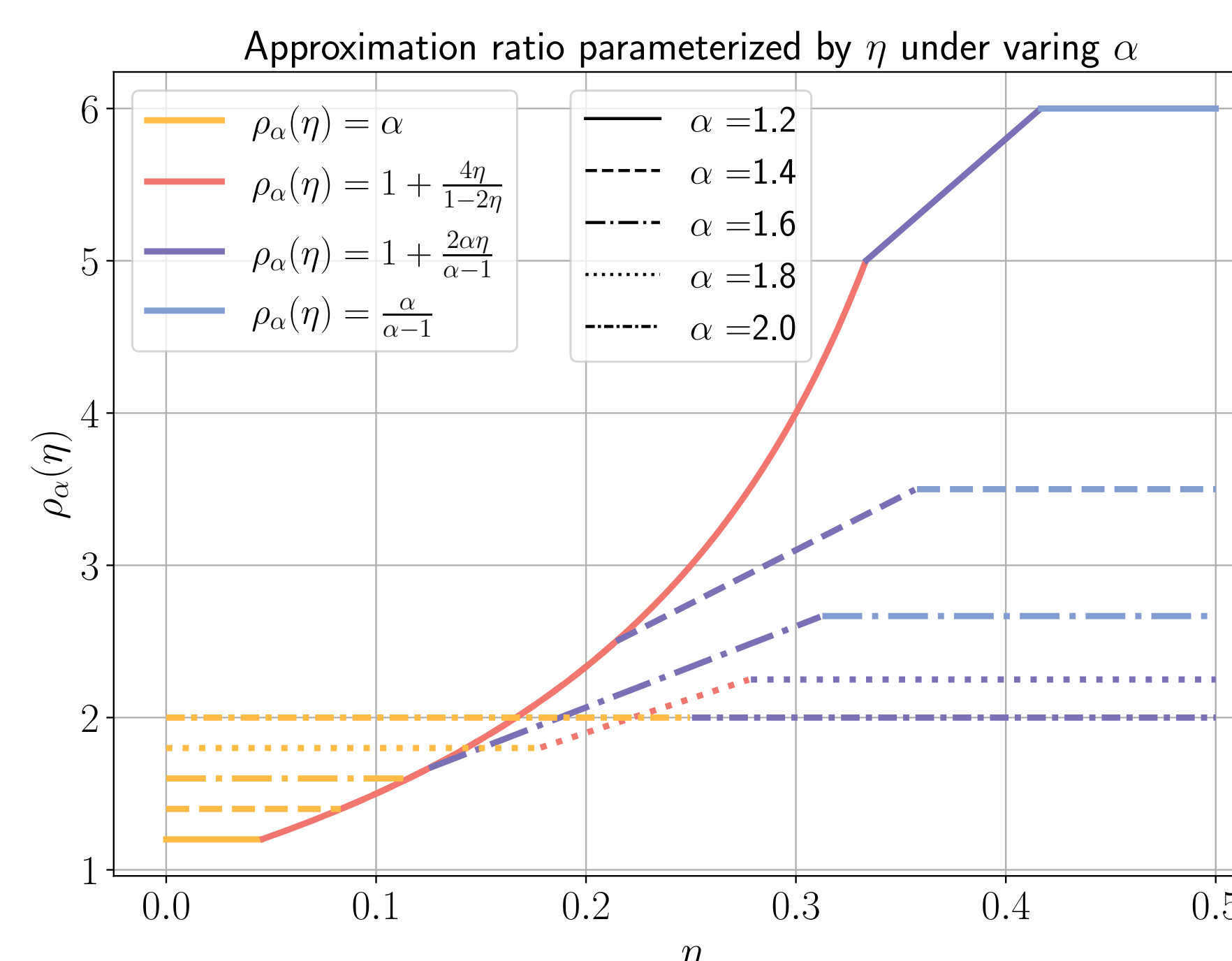
α -Bounding Interval Mechanism (α -BIM)

Return $f(x, \hat{y}) = \hat{y}$ if $\hat{y} \in [1 - \frac{1}{\alpha}, \frac{1}{\alpha}]$, otherwise return $f(x, \hat{y})$ with the nearest endpoint $1 - \frac{1}{\alpha}$ or $\frac{1}{\alpha}$ to $\hat{y}$.



- α -BIM is **strategyproof** and satisfies **α -consistency**, and $\frac{\alpha}{\alpha-1}$ -**robustness**. No deterministic strategyproof mechanism can be $(\alpha - \varepsilon)$ -consistent and $(\frac{\alpha}{\alpha-1} - \varepsilon)$ -robust for any $\varepsilon > 0$.

- (**Fine-grained analysis**) Approximation ratio parameterized by prediction error η where $\eta = \sup_{x \in [0, 1]^n} |OPT(x) - \hat{y}|$.



Reference

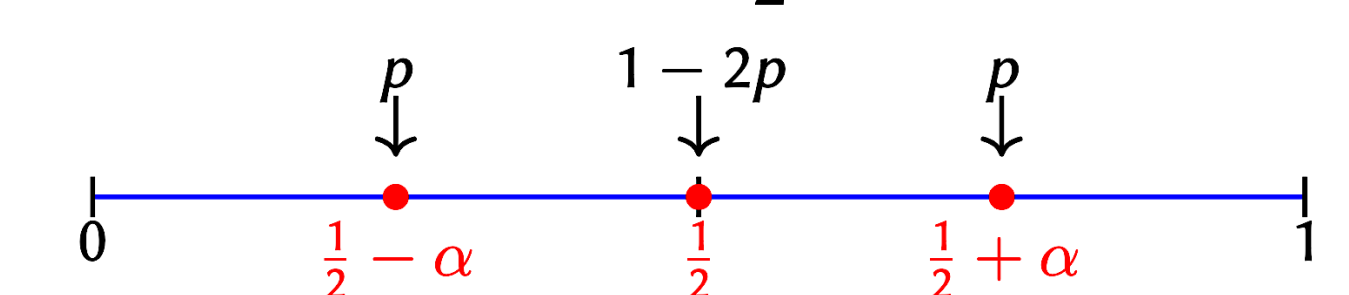
[DLC⁺20] Y. Ding, W. Liu, X. Chen, Q. Fang, and Q. Nong. Facility location game with envy ratio. Computers & Industrial Engineering, 148:106710, 2020.
 [ABG⁺22] P. Agrawal, E. Balkanski, V. Gkatzelis, T. Ou, and X. Tan. Learning-augmented mechanism design: Leveraging predictions for facility location. In Proceedings of the 23rd ACM Conference on Economics and Computation (EC), page 497–528, 2022.

Randomized Mechanism

New mechanism in classical setting:

(α, p) -LRM Constant Mechanism

Consider parameter $p \in [0, \frac{1}{2}]$, $\alpha \in [0, \frac{1}{2}]$; With Prob. p return $f(x) = \frac{1}{2} - \alpha$; With Prob. $1 - 2p$ return $f(x) = \frac{1}{2}$; With Prob. p return $f(x) = \frac{1}{2} + \alpha$.



- $(\frac{\sqrt{5}}{2} - 1, \frac{2}{5})$ -LRM constant mechanism is **strategyproof** and achieves an approximation ratio of $1 + \frac{2}{\sqrt{5}} \approx 1.8944$, which is **optimal** among all the (α, p) -LRM constant mechanisms.

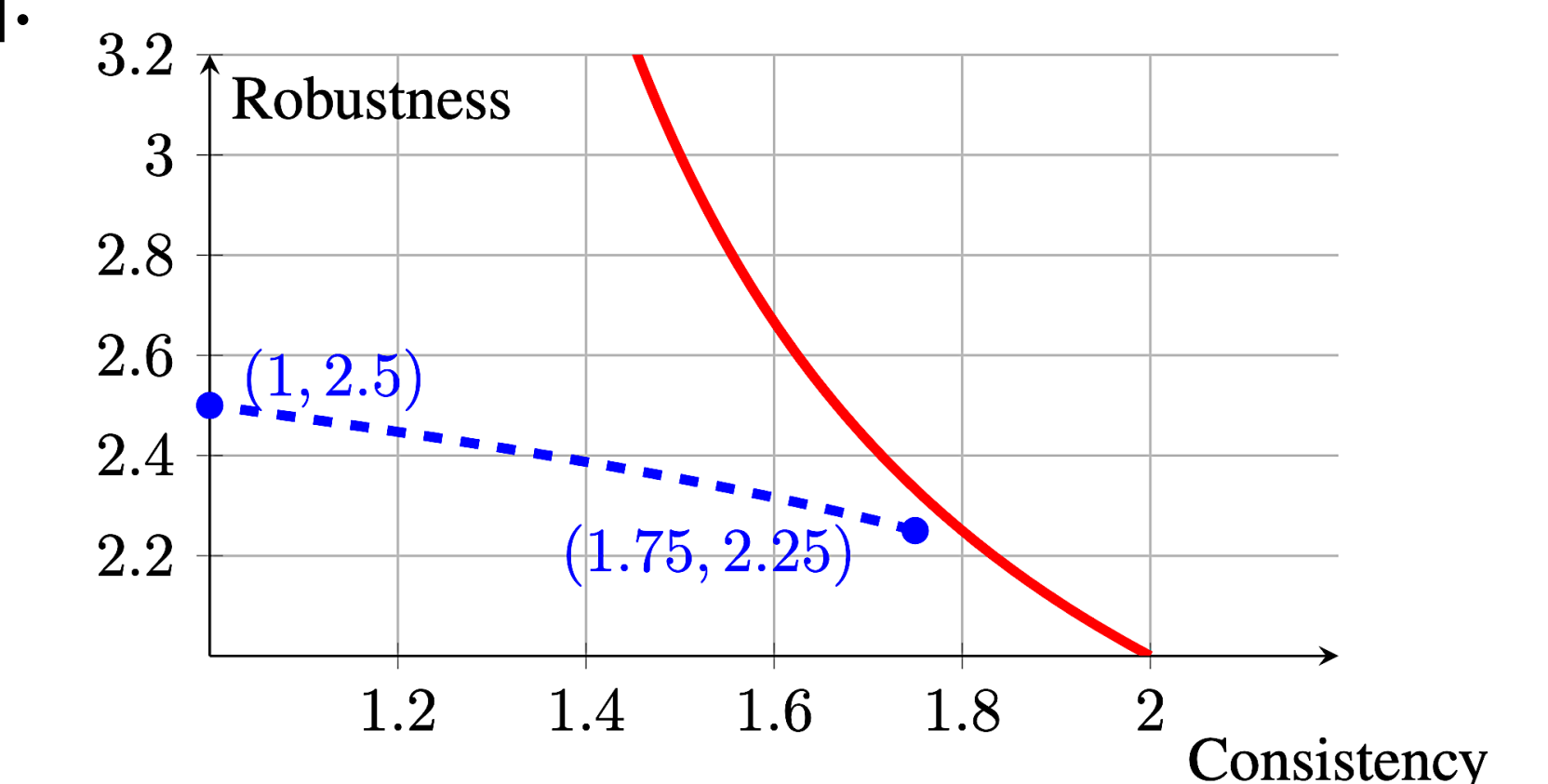
With learning-augmentation:

Bias-Aware Mechanism (BAM)

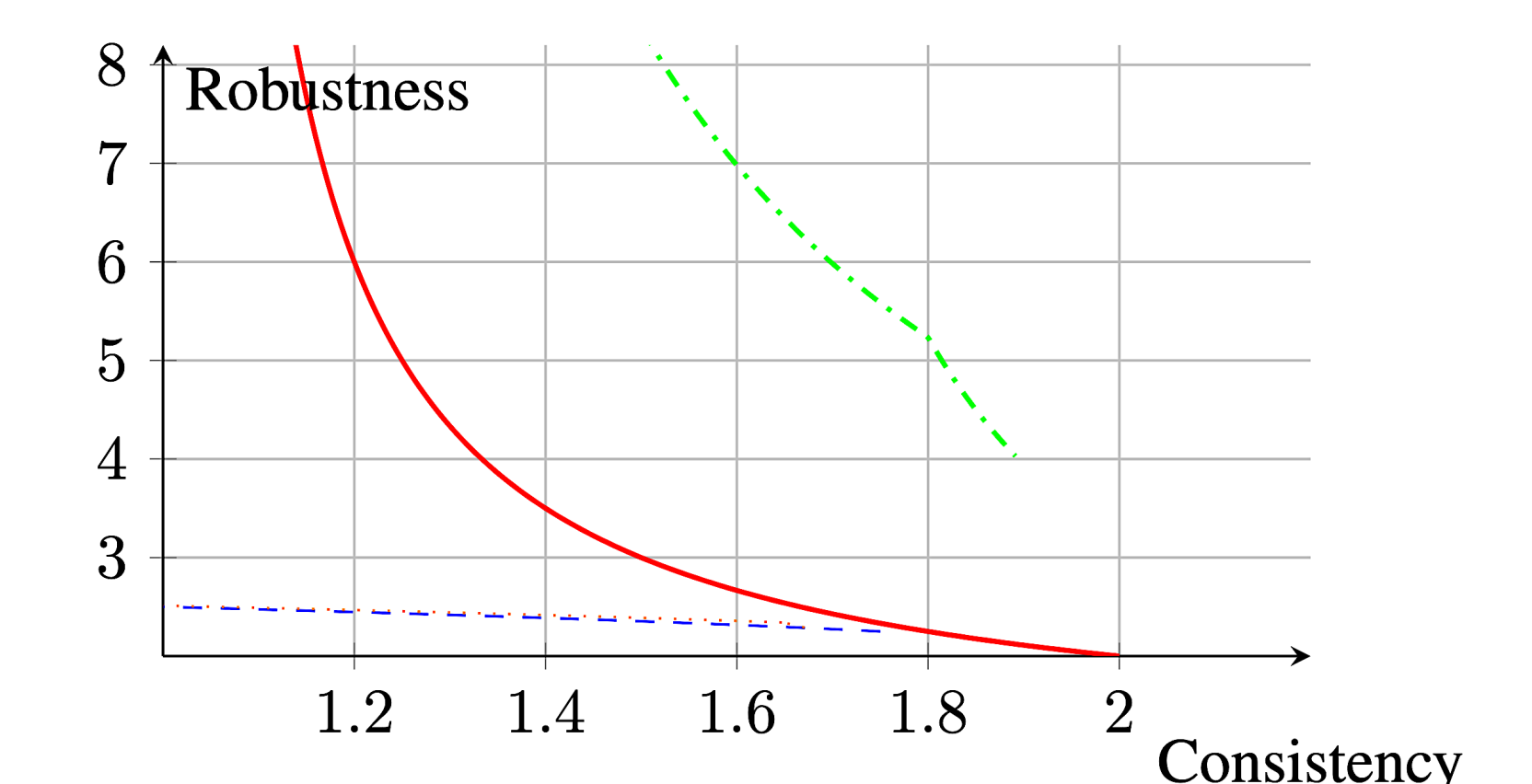
Compute bias $c \leftarrow \left| \hat{y} - \frac{1}{2} \right|$ and probability $p \leftarrow \frac{1}{2} - c$; With Prob. p return $f(x, \hat{y}) = \hat{y}$. With Prob. $1 - p$ return $f(x, \hat{y}) = \frac{1}{2}$.

Key idea of BAM: For any learning-augmented mechanism with bounded robustness, when $\hat{y} \rightarrow 0$ or $\hat{y} \rightarrow 1$, $\text{Prob}[f(x, \hat{y}) = \hat{y}] \rightarrow 0$.

- BAM is **strategyproof** and satisfies **$(2 - c^2)$ -consistency**, **$(2 + c)$ -robustness**, when $c \in (\frac{1}{4}, \frac{1}{2})$ and $\frac{7}{4}$ -**consistency**, $\frac{9}{4}$ -**robustness**, when $c \in [0, \frac{1}{4}]$.



Comparison of α -BIM (red solid) and BAM (blue dashed)



Comparison of α -BIM (red solid), BAM (blue dashed), α -Bounding Interval Randomized Mechanism (green dashdotted), and Biased-aware LRM Mechanism (orange dotted)