

Learning-Augmented Facility Location Mechanisms for Envy Ratio

H. AZIZ, Y. Guo, A. Lam, and H. Zhou @NeurIPS 2025

Joseph Chuang-Chieh Lin

Department of Computer Science & Engineering,
National Taiwan Ocean University

Fall 2026



Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions
- 3 Randomized Mechanisms without Predictions
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions
- 3 Randomized Mechanisms without Predictions
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



The Paper in One Sentence

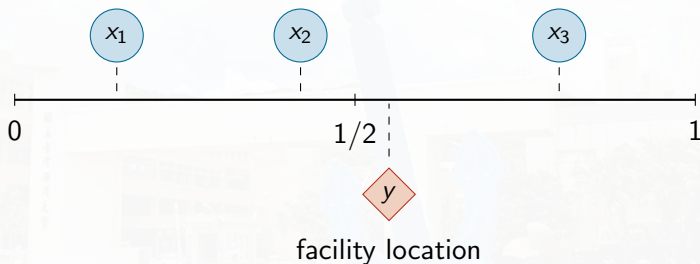
Main theme

Can machine-learned **predictions** help us design **strategyproof** facility-location mechanisms with better **fairness guarantees**?

- Setting: one facility on the interval $[0, 1]$.
- Agents want the facility close to their locations.
- Fairness objective: minimize the **envy ratio**.
 - Firstly proposed by Ding et al. [Comput. Ind. Eng. 2020].
- Prediction: an estimate $\hat{y}$ of the optimal facility location.



Classical Facility Location on a Line



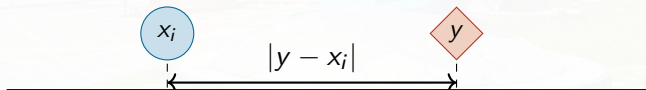
- Agent i has a private location $x_i \in [0, 1]$.
- The mechanism chooses a facility location $y \in [0, 1]$.

Utilities

Utility of agent i

$$u(y, x_i) = 1 - |y - x_i|.$$

- Utility is larger when the facility is closer.
- The interval length is normalized to 1.
- The utility is always in $[0, 1]$.



Why Utility Ratios?

A distance-ratio objective can become unbounded if the facility is exactly at an agent's location.



- Hence the paper uses **utility** $1 - |y - x_i|$.
- Envy ratio compares utility levels, not raw distances.

Envy Ratio

Definition (Envy ratio)

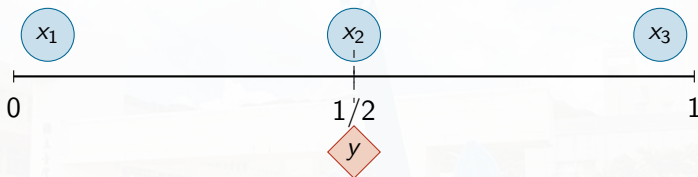
Given a location profile $\mathbf{x}$ and a facility location y ,

$$ER(y, \mathbf{x}) = \max_{i \neq j} \frac{u(y, x_i)}{u(y, x_j)}.$$

- The numerator is the utility of a best-off agent.
- The denominator is the utility of a worst-off agent.
- Lower is better; the ideal value is 1.



Example: Best-off vs. Worst-off



center agent has utility 1;
endpoint agents have utility $1/2$

$$ER(y, \mathbf{x}) = \frac{1}{1/2} = 2.$$

Optimal Envy Ratio

Let

$$lm(\mathbf{x}) = \min_i x_i, \quad rm(\mathbf{x}) = \max_i x_i.$$

Theorem (Ding et al. 2020)

The facility location that *minimizes envy ratio* is

$$mid(\mathbf{x}) = \frac{lm(\mathbf{x}) + rm(\mathbf{x})}{2}.$$



Approximation Ratio

Definition (Approximation ratio)

A mechanism f has approximation ratio ρ if

$$\rho = \sup_{\mathbf{x} \in [0,1]^n} \frac{ER(f(\mathbf{x}), \mathbf{x})}{ER(OPT(\mathbf{x}), \mathbf{x})}.$$

- Here $OPT(\mathbf{x}) = \text{mid}(\mathbf{x})$ for envy ratio.
- Since OPT is best possible, $\rho \geq 1$.



Approximation Ratio

Definition (Approximation ratio)

A mechanism f has approximation ratio ρ if

$$\rho = \sup_{\mathbf{x} \in [0,1]^n} \frac{ER(f(\mathbf{x}), \mathbf{x})}{ER(OPT(\mathbf{x}), \mathbf{x})}.$$

- Here $OPT(\mathbf{x}) = \text{mid}(\mathbf{x})$ for envy ratio.
- Since OPT is best possible, $\rho \geq 1$.
- The challenge: **require strategyproofness as well.**



Strategyproofness

Definition (Strategyproofness)

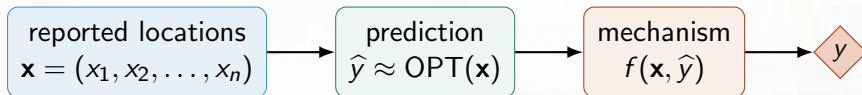
A mechanism f is strategyproof if no agent can improve by misreporting:

$$u(f(\mathbf{x}), x_i) \geq u(f(\mathbf{x}_{-i}, x'_i), x_i)$$

for every profile $\mathbf{x}$, agent i , and false report x'_i of agent i .



Learning-Augmented Mechanism Design



- Prediction may come from historical data or a machine-learning model.
- **Goal:** exploit accurate predictions, but retain guarantees under inaccurate predictions (robustness).

Consistency and Robustness

$$\rho(\mathbf{x}, \hat{y}) = \frac{\text{ER}(f(\mathbf{x}, \hat{y}), \mathbf{x})}{\text{ER}(\text{OPT}(\mathbf{x}), \mathbf{x})}$$

γ -consistency

Approximation ratio when the prediction is correct:

$$\hat{y} = \text{OPT}(\mathbf{x}).$$

That is,

$$\gamma = \sup_{\mathbf{x} \in [0,1]^n} \rho(\mathbf{x}, \text{OPT}(\mathbf{x})).$$

β -robustness

Worst-case approximation ratio under arbitrary prediction:

$$\hat{y} \in [0, 1].$$

That is,

$$\beta = \sup_{\substack{\mathbf{x} \in [0,1]^n \\ \hat{y} \in [0,1]}} \rho(\mathbf{x}, \hat{y}).$$



Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions**
- 3 Randomized Mechanisms without Predictions
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



Baseline Without Predictions [Ding et al. 2020]

For deterministic strategyproof mechanisms without predictions:

- The **constant** mechanism $f(x) = 1/2$ is strategyproof.
- The approximation ratio 2 is best possible [Ding et al. 2020].



ignore all reports

Question

Can predictions improve the consistency–robustness trade-off?



α -Bounding Interval Mechanism

(Extending Ding et al.'s result by considering predictions)

Choose a parameter $\alpha \in [1, 2]$.

α -BIM (α -Bounding Interval Mechanism)

Define the trusted interval

$$\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right].$$

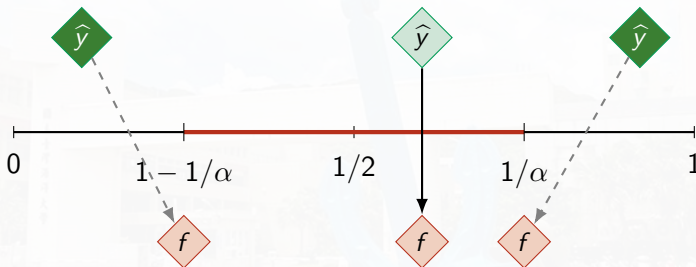
Return $\hat{y}$ if it lies in this interval; otherwise return the nearest endpoint.

$$f(\mathbf{x}, \hat{y}) = \Pi_{[1-1/\alpha, 1/\alpha]}(\hat{y}),$$

where Π_I denotes **projection onto interval** I .



Diagram of α -BIM



- Inside the interval: trust the prediction.
- Outside the interval: clip it to protect robustness.

Extreme Choices of α

Parameter	Trusted interval	Behavior
$\alpha = 1$	$[0, 1]$	always trust $\hat{y}$
$\alpha = 2$	$[1/2, 1/2]$	always output $1/2$

- Smaller α : more prediction-sensitive.
- Larger α : more conservative.



Why Is α -BIM Strategyproof?

Key observation

The output of α -BIM depends only on the prediction $\hat{y}$ and parameter α ,
not on the reported locations $\mathbf{x}$.

- If an agent changes her reported location, the output remains the same.
- Therefore no misreport can improve her utility.
- The mechanism is also **anonymous**: agent labels do not matter.



Reduction to Two Agents

Let's **simplify the space of location profiles** which need to be considered.

Lemma 3.1

For any instance $\mathbf{x}$ and any distribution P of facility locations, there always exists a **2-agent instance $\mathbf{x}' = (\text{lm}(\mathbf{x}), \text{rm}(\mathbf{x}))$** such that

$$\mathbb{E}_{y \in P} \left[\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \right] \leq \mathbb{E}_{y \in P} \left[\frac{\text{ER}(y, \mathbf{x}')}{\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}')} \right].$$

- For analyzing the approximation ratio, it suffices to consider the two extreme agents $\mathbf{x}' = (\text{lm}(\mathbf{x}), \text{rm}(\mathbf{x}))$.



Proof of Lemma 3.1: Setup

- Sort the agents:

$$x_1 \leq x_2 \leq \cdots \leq x_n.$$

- Let

$$m = \text{mid}(\mathbf{x}) = \frac{x_1 + x_n}{2}.$$

- Let u^{**} and u^* be the maximum and minimum utilities *under the midpoint mechanism*.

$$\text{ER}(m, \mathbf{x}) = \frac{u^{**}}{u^*}.$$

- y : output of the mechanism.



Proof of Lemma 3.1: Case 1

Case 1: $y \in [x_1, x_n]$

Let $\Delta = |y - \text{mid}(\mathbf{x})|$. When the facility moves from $\text{mid}(\mathbf{x})$ to y :
maximum utility $\leq u^{**} + \Delta$, minimum utility $\geq u^* - \Delta$.

Therefore,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{(u^{**} + \Delta)/(u^* - \Delta)}{u^{**}/u^*}.$$

Since the right-hand side is non-increasing in u^{**} ,

$$\frac{(u^{**} + \Delta)/(u^* - \Delta)}{u^{**}/u^*} \leq \frac{u^* + \Delta}{u^* - \Delta}.$$

$$\boxed{\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{u^* + \Delta}{u^* - \Delta}}$$



Proof of Lemma 3.1: Case 1

Case 1: $y \in [x_1, x_n]$

Let $\Delta = |y - \text{mid}(\mathbf{x})|$. When the facility moves from $\text{mid}(\mathbf{x})$ to y :
maximum utility $\leq u^{**} + \Delta$, minimum utility $\geq u^* - \Delta$.

Therefore,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{(u^{**} + \Delta)/(u^* - \Delta)}{u^{**}/u^*}.$$

Since the right-hand side is non-increasing in u^{**} , (so let's set $u^{**} = u^*$)

$$\frac{(u^{**} + \Delta)/(u^* - \Delta)}{u^{**}/u^*} \leq \frac{u^* + \Delta}{u^* - \Delta}.$$

$$\boxed{\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{u^* + \Delta}{u^* - \Delta}}$$



Proof of Lemma 3.1: Reduction to Two Agents

- Now consider the 2-agent instance: $\mathbf{x}' = (x_1, x_n)$.
- For this instance, the midpoint is the same: $\text{mid}(\mathbf{x}') = \text{mid}(\mathbf{x})$.



Proof of Lemma 3.1: Reduction to Two Agents

- Now consider the 2-agent instance: $\mathbf{x}' = (x_1, x_n)$.
- For this instance, the midpoint is the same: $\text{mid}(\mathbf{x}') = \text{mid}(\mathbf{x})$.
- Moreover, under the midpoint mechanism, both endpoint agents receive utility u^* , so $\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}') = 1$.



Proof of Lemma 3.1: Reduction to Two Agents

- Now consider the 2-agent instance: $\mathbf{x}' = (x_1, x_n)$.
- For this instance, the midpoint is the same: $\text{mid}(\mathbf{x}') = \text{mid}(\mathbf{x})$.
- Moreover, under the midpoint mechanism, both endpoint agents receive utility u^* , so $\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}') = 1$.
- For $y \in [x_1, x_n]$, the two endpoint utilities become $u^* + \Delta$ and $u^* - \Delta$.
- Hence,

$$\frac{\text{ER}(y, \mathbf{x}')}{\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}')} = \frac{u^* + \Delta}{u^* - \Delta}.$$

- Combining with the previous slide gives

$$\boxed{\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{\text{ER}(y, \mathbf{x}')}{\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}')}}.$$



Proof of Lemma 3.1: Case 2 and Conclusion

Case 2: $y \notin [x_1, x_n]$

By symmetry, assume $y > x_n$. Then x_1 has min. utility and x_n has max. utility.

Since $d(\text{mid}(\mathbf{x}), x_1) = d(\text{mid}(\mathbf{x}), x_n) = 1 - u^*$, we have

$$u(y, x_1) = -y + \text{mid}(\mathbf{x}) + u^* \text{ and } u(y, x_n) = 2 - y + \text{mid}(\mathbf{x}) - u^*.$$



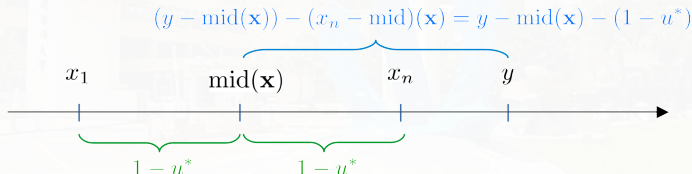
Proof of Lemma 3.1: Case 2 and Conclusion

Case 2: $y \notin [x_1, x_n]$

By symmetry, assume $y > x_n$. Then x_1 has min. utility and x_n has max. utility.

Since $d(\text{mid}(\mathbf{x}), x_1) = d(\text{mid}(\mathbf{x}), x_n) = 1 - u^*$, we have

$$u(y, x_1) = -y + \text{mid}(\mathbf{x}) + u^* \text{ and } u(y, x_n) = 2 - y + \text{mid}(\mathbf{x}) - u^*.$$



Proof of Lemma 3.1: Case 2 and Conclusion

Case 2: $y \notin [x_1, x_n]$

By symmetry, assume $y > x_n$. Then x_1 has min. utility and x_n has max. utility.

Since $d(\text{mid}(\mathbf{x}), x_1) = d(\text{mid}(\mathbf{x}), x_n) = 1 - u^*$, we have

$$u(y, x_1) = -y + \text{mid}(\mathbf{x}) + u^* \text{ and } u(y, x_n) = 2 - y + \text{mid}(\mathbf{x}) - u^*.$$

Thus,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{\frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}}{u^{**}/u^*} \leq \frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}.$$



Proof of Lemma 3.1: Case 2 and Conclusion

Case 2: $y \notin [x_1, x_n]$

By symmetry, assume $y > x_n$. Then x_1 has min. utility and x_n has max. utility.

Since $d(\text{mid}(\mathbf{x}), x_1) = d(\text{mid}(\mathbf{x}), x_n) = 1 - u^*$, we have

$$u(y, x_1) = -y + \text{mid}(\mathbf{x}) + u^* \text{ and } u(y, x_n) = 2 - y + \text{mid}(\mathbf{x}) - u^*.$$

Thus,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{\frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}}{u^{**}/u^*} \leq \frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}.$$

The same expression as that for the 2-agent instance $\mathbf{x}' = (x_1, x_n)$.



Proof of Lemma 3.1: Case 2 and Conclusion

Case 2: $y \notin [x_1, x_n]$

By symmetry, assume $y > x_n$. Then x_1 has min. utility and x_n has max. utility.

Since $d(\text{mid}(\mathbf{x}), x_1) = d(\text{mid}(\mathbf{x}), x_n) = 1 - u^*$, we have

$$u(y, x_1) = -y + \text{mid}(\mathbf{x}) + u^* \text{ and } u(y, x_n) = 2 - y + \text{mid}(\mathbf{x}) - u^*.$$

Thus,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{\frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}}{u^{**}/u^*} \leq \frac{2 - y + \text{mid}(\mathbf{x}) - u^*}{-y + \text{mid}(\mathbf{x}) + u^*}.$$

The same expression as that for the 2-agent instance $\mathbf{x}' = (x_1, x_n)$.

Therefore, for every $y \in P$,

$$\frac{\text{ER}(y, \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \frac{\text{ER}(y, \mathbf{x}')}{\text{ER}(\text{mid}(\mathbf{x}'), \mathbf{x}')}.$$

Taking expectation over $y \in P$ proves the lemma.

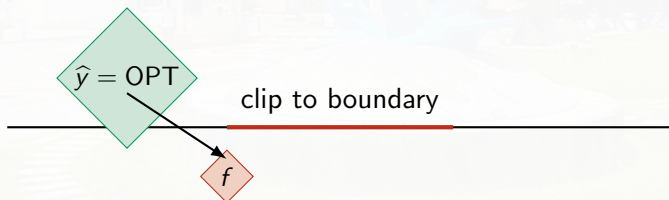


Consistency of α -BIM

Theorem

When $\hat{y} = \text{OPT}(\mathbf{x})$, α -BIM achieves α -consistency.

- If the optimal midpoint lies inside the trusted interval, the mechanism is optimal.
- If it lies outside, clipping moves the facility to the nearest endpoint.
- The worst-case utility loss is bounded by a factor α .



Consistency of α -BIM

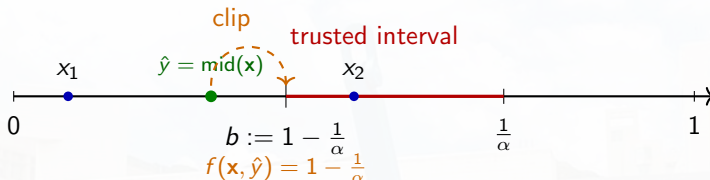
- By Lemma 3.1, it suffices to consider a two-agent instance $\mathbf{x} = (x_1, x_2)$, $x_1 < x_2$, with $\text{mid}(\mathbf{x}) = \frac{x_1 + x_2}{2}$.
- Assume that the prediction is **accurate**: $\hat{y} = \text{mid}(\mathbf{x})$.
 - If $\hat{y} \in [1 - \frac{1}{\alpha}, \frac{1}{\alpha}]$, then α -BIM returns $f(\mathbf{x}, \hat{y}) = \hat{y} = \text{mid}(\mathbf{x})$, so the mechanism is optimal on this instance.
 - Otherwise, by symmetry, it suffices to consider $\hat{y} = \text{mid}(\mathbf{x}) < 1 - \frac{1}{\alpha}$.
 - Let $b := 1 - \frac{1}{\alpha}$. In the nontrivial case,

$$x_1 < \hat{y} = \text{mid}(\mathbf{x}) < b,$$

and α -BIM clips the prediction to $f(\mathbf{x}, \hat{y}) = b$.

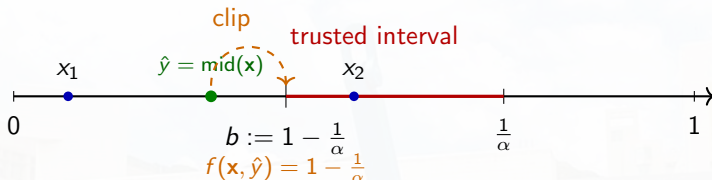


Consistency of α -BIM: The Clipped Case



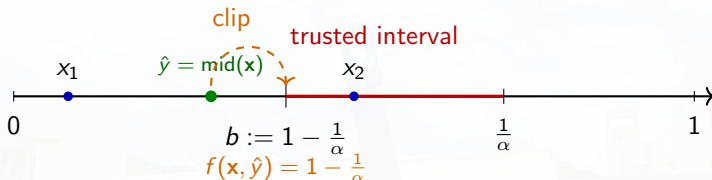
- Because $b > \text{mid}(\mathbf{x})$, we have $|b - x_2| \leq |b - x_1|$.
 $\Rightarrow u(b, x_2) \geq u(b, x_1)$, so x_2 attains the max. utility and x_1 the min. utility.

Consistency of α -BIM: The Clipped Case



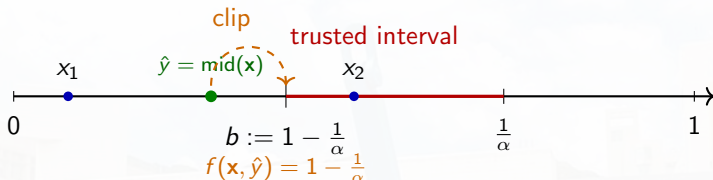
- Because $b > \text{mid}(\mathbf{x})$, we have $|b - x_2| \leq |b - x_1|$.
- $\Rightarrow u(b, x_2) \geq u(b, x_1)$, so x_2 attains the max. utility and x_1 the min. utility.
- Moreover, $u(b, x_2) \leq 1$, $u(b, x_1) = 1 - (b - x_1) \geq 1 - b = \frac{1}{\alpha}$.

Consistency of α -BIM: The Clipped Case



- Because $b > \text{mid}(\mathbf{x})$, we have $|b - x_2| \leq |b - x_1|$.
 $\Rightarrow u(b, x_2) \geq u(b, x_1)$, so x_2 attains the max. utility and x_1 the min. utility.
- Moreover, $u(b, x_2) \leq 1$, $u(b, x_1) = 1 - (b - x_1) \geq 1 - b = \frac{1}{\alpha}$.
- Therefore,
$$\text{ER}(b, \mathbf{x}) = \frac{u(b, x_2)}{u(b, x_1)} \leq \frac{1}{1/\alpha} = \alpha.$$

Consistency of α -BIM: The Clipped Case



- Because $b > \text{mid}(\mathbf{x})$, we have $|b - x_2| \leq |b - x_1|$.
- $\Rightarrow u(b, x_2) \geq u(b, x_1)$, so x_2 attains the max. utility and x_1 the min. utility.
- Moreover, $u(b, x_2) \leq 1$, $u(b, x_1) = 1 - (b - x_1) \geq 1 - b = \frac{1}{\alpha}$.
- Therefore,
$$\text{ER}(b, \mathbf{x}) = \frac{u(b, x_2)}{u(b, x_1)} \leq \frac{1}{1/\alpha} = \alpha.$$
- Since $\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x}) = 1$, we conclude that

$$\frac{\text{ER}(f(\mathbf{x}, \text{mid}(\mathbf{x})), \mathbf{x})}{\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x})} \leq \alpha.$$



Robustness of α -BIM

Theorem

Under arbitrary prediction, α -BIM achieves

$$\frac{\alpha}{\alpha - 1}\text{-robustness.}$$

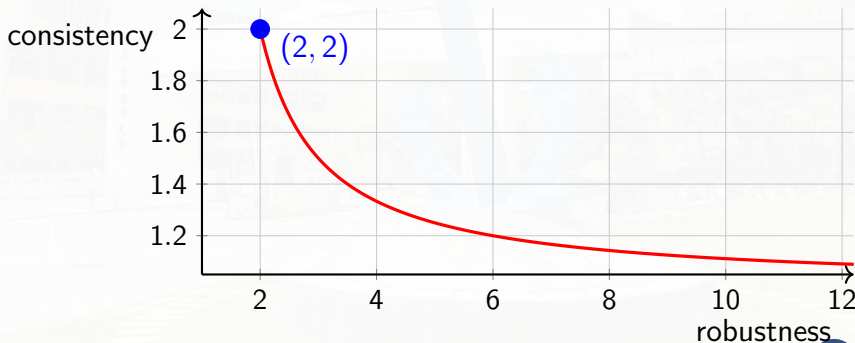
- The chosen facility is always inside $[1 - 1/\alpha, 1/\alpha]$.
- Therefore every agent has utility at least $1 - 1/\alpha$.
- The best-off utility is at most 1.

$$\text{ER}(f(x, \hat{y}), x) \leq \frac{1}{1 - 1/\alpha} = \frac{\alpha}{\alpha - 1}.$$



Consistency–Robustness Trade-off

$$\text{consistency} = \alpha, \quad \text{robustness} = \frac{\alpha}{\alpha - 1}.$$



- Does α -BIM achieve the “best” consistency and robustness guarantees among all Strategyproof and anonymous deterministic mechanisms?



Optimality Among Deterministic Mechanisms

Theorem (Optimality)

For every $\alpha \in (1, 2]$ and every $\varepsilon > 0$, no deterministic, strategyproof, anonymous mechanism can achieve both

$$(\alpha - \varepsilon)\text{-consistency} \quad \text{and} \quad \left(\frac{\alpha}{\alpha - 1} - \varepsilon \right)\text{-robustness}.$$

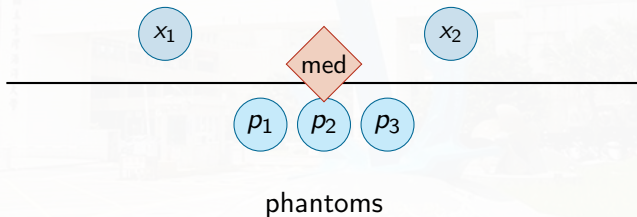
α -BIM lies on the deterministic Pareto frontier among all strategyproof and anonymous deterministic mechanisms.



Proof Intuition: Phantom Mechanisms

Proposition 2, Moulin @Public Choice 1980

A deterministic strategyproof and anonymous mechanism on a line must be a phantom mechanism with $n + 1$ constant (fixed) **phantom points**.



- Robustness forces all phantoms into the BIM interval.
- Then consistency cannot beat the BIM guarantee.

Example: A Phantom Mechanism for One Agent

Suppose the mechanism has **one** reported location x and **two fixed phantom points** $a \leq b$. The output is

$$f(x) = \text{median}\{x, a, b\} = \begin{cases} a, & x < a, \\ x, & a \leq x \leq b, \\ b, & x > b. \end{cases}$$

- This is exactly a **clipping rule**.



Optimality of α -BIM: Proof Idea

Here ‘phantom’ locations may be a function of the prediction.

Moulin’s Characterization (Proposition 2)

Every deterministic, anonymous, and strategyproof mechanism on a line is a **phantom mechanism**.

For every prediction $\hat{y}$, there exist $n + 1$ phantom points

$$p_1(\hat{y}), p_2(\hat{y}), \dots, p_{n+1}(\hat{y}),$$

such that the facility location is

$$f(\mathbf{x}, \hat{y}) = \text{median}(x_1, x_2, \dots, x_n, p_1(\hat{y}), p_2(\hat{y}), \dots, p_{n+1}(\hat{y})).$$

- The phantom locations may depend on the prediction $\hat{y}$.
- They do *not* depend on the reported locations.



Optimality of α -BIM: Step 1 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- Assume $1 < \alpha \leq 2$, so that $1 - \frac{1}{\alpha} \leq \frac{1}{\alpha}$.
- We show that every phantom point must satisfy

$$p_i(\hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right].$$



Optimality of α -BIM: Step 1 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- Assume $1 < \alpha \leq 2$, so that $1 - \frac{1}{\alpha} \leq \frac{1}{\alpha}$.
- We show that every phantom point must satisfy

$$p_i(\hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right].$$

- Suppose, for contradiction, $p_i < 1 - \frac{1}{\alpha}$. Since p_i is fixed for this prediction, consider the (adversarial) location profile

$$\underbrace{p_i, \dots, p_i}_{n-1 \text{ agents}}, 1.$$

- Because p_i is the median, $f(\mathbf{x}, \hat{y}) = p_i$.



Optimality of α -BIM: Step 2 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- The facility is placed at p_i .



Optimality of α -BIM: Step 2 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- The facility is placed at p_i .
- Hence, $ER = \frac{u_{\max}}{u_{\min}} = \frac{1}{1 - (1 - p_i)} = \frac{1}{p_i}$.



Optimality of α -BIM: Step 2 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- The facility is placed at p_i .
- Hence, $ER = \frac{u_{\max}}{u_{\min}} = \frac{1}{1 - (1 - p_i)} = \frac{1}{p_i}$.
- Since $p_i < 1 - \frac{1}{\alpha}$, we obtain the robustness will be

$$\frac{1}{p_i} > \frac{1}{1 - \frac{1}{\alpha}} = \frac{\alpha}{\alpha - 1}.$$



Optimality of α -BIM: Step 2 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- The facility is placed at p_i .
- Hence, $ER = \frac{u_{\max}}{u_{\min}} = \frac{1}{1 - (1 - p_i)} = \frac{1}{p_i}$.
- Since $p_i < 1 - \frac{1}{\alpha}$, we obtain the robustness will be

$$\frac{1}{p_i} > \frac{1}{1 - \frac{1}{\alpha}} = \frac{\alpha}{\alpha - 1}.$$

- Therefore, robustness $> \frac{\alpha}{\alpha - 1}$, contradicting the assumed robustness guarantee.



Optimality of α -BIM: Step 2 (in order to make $\beta \leq \alpha/(\alpha - 1)$)

- The facility is placed at p_i .
- Hence, $ER = \frac{u_{\max}}{u_{\min}} = \frac{1}{1 - (1 - p_i)} = \frac{1}{p_i}$.
- Since $p_i < 1 - \frac{1}{\alpha}$, we obtain the robustness will be

$$\frac{1}{p_i} > \frac{1}{1 - \frac{1}{\alpha}} = \frac{\alpha}{\alpha - 1}.$$

- Therefore, robustness $> \frac{\alpha}{\alpha - 1}$, contradicting the assumed robustness guarantee.
- Hence every phantom point must lie inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$.



Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, 1.$$

Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, 1.$$

- Since every phantom lies inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$, the median (i.e., y) must also lie inside this interval.



Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, \quad 1.$$

- Since every phantom lies inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$, the median (i.e., y) must also lie inside this interval. $\Rightarrow y := f(\mathbf{x}, \hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$.



Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, \quad 1.$$

- Since every phantom lies inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$, the median (i.e., y) must also lie inside this interval. $\Rightarrow y := f(\mathbf{x}, \hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$.
- Claim:** $\gamma \geq \alpha$. (note that $\text{ER}(\text{OPT}(\mathbf{x}), \mathbf{x}) = 1$.)



Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, \quad 1.$$

- Since every phantom lies inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$, the median (i.e., y) must also lie inside this interval. $\Rightarrow y := f(\mathbf{x}, \hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$.
- Claim:** $\gamma \geq \alpha$. (note that $\text{ER}(\text{OPT}(\mathbf{x}), \mathbf{x}) = 1$.)

- The correct prediction is $\hat{y} = \text{OPT}(\mathbf{x}) = \text{mid}(\mathbf{x}) = \frac{1 + \frac{1}{\alpha}}{2}$.

- Thus

$$\text{ER}(y := f(\mathbf{x}, \hat{y}), \mathbf{x}) = \frac{u(y, 1/\alpha)}{u(y, 1)} = \frac{1 - 1/\alpha + y}{y} = 1 + \frac{1 - 1/\alpha}{y} \geq \alpha.$$

Optimality of α -BIM: Conclusion (1/2)

- Fix the same prediction $\hat{y}$. Consider the location profile

$$\underbrace{\frac{1}{\alpha}, \frac{1}{\alpha}, \dots, \frac{1}{\alpha}}_{n-1 \text{ agents}}, 1.$$

- Since every phantom lies inside $\left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$, the median (i.e., y) must also lie inside this interval. $\Rightarrow y := f(\mathbf{x}, \hat{y}) \in \left[1 - \frac{1}{\alpha}, \frac{1}{\alpha}\right]$. ($y \leq 1/\alpha$)
- Claim:** $\gamma \geq \alpha$. (note that $\text{ER}(\text{OPT}(\mathbf{x}), \mathbf{x}) = 1$.)

- The correct prediction is $\hat{y} = \text{OPT}(\mathbf{x}) = \text{mid}(\mathbf{x}) = \frac{1 + \frac{1}{\alpha}}{2}$.

- Thus

$$\text{ER}(y := f(\mathbf{x}, \hat{y}), \mathbf{x}) = \frac{u(y, 1/\alpha)}{u(y, 1)} = \frac{1 - 1/\alpha + y}{y} = 1 + \frac{1 - 1/\alpha}{y} \geq \alpha.$$

Optimality of α -BIM: Conclusion (2/2)

- Hence

If robustness $\leq \frac{\alpha}{\alpha - 1}$, then consistency cannot be smaller than α .

- Thus, α -BIM simultaneously achieves the optimal consistency and robustness guarantees.



Approximation Ratio Parameterized by Prediction Error

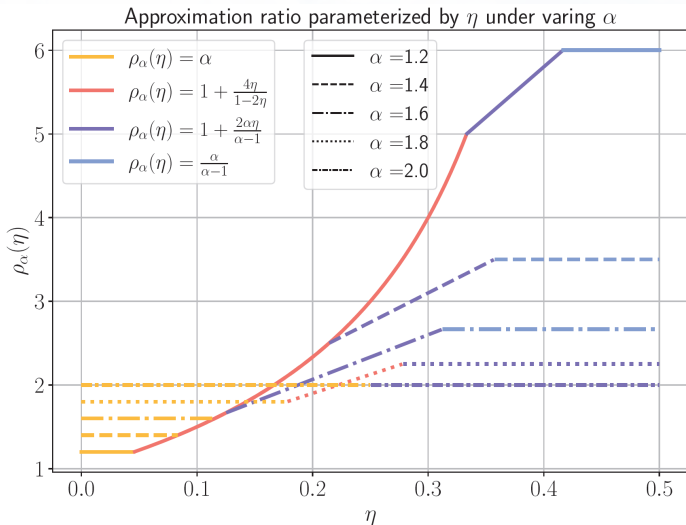
Let

$$y^* = \text{OPT}(\mathbf{x}), \quad \hat{y} : \text{prediction}, \quad \eta = |\hat{y} - y^*|.$$

- Denote $\rho_\alpha(\eta)$: the approximation ratio as a function of prediction error η .
- The ratio increases continuously from α to $\alpha/(\alpha - 1)$.



Error-Parameterized Behavior (proof skipped)



Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions
- 3 Randomized Mechanisms without Predictions**
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



Why Randomization?

- Improving worst-case guarantees while preserving **strategyproofness in expectation**.
- Ding et al. gave a lower bound near 1.0314.
- A ratio of 2 was achievable by always outputting $1/2$.
- **Open question:** Can a randomized strategyproof mechanism beat the approximation ratio of 2?



Why Randomization?

- Improving worst-case guarantees while preserving **strategyproofness in expectation**.
- Ding et al. gave a lower bound near 1.0314.
- A ratio of 2 was achievable by always outputting $1/2$.
- **Open question:** Can a randomized strategyproof mechanism beat the approximation ratio of 2?

This work's answer

Yes: a new randomized mechanism achieves ratio ≈ 1.8944 .



(α, p) -LRM Constant Mechanism

Mechanism: (α, p) -LRM Constant Mechanism

Choose parameters $\alpha \in [0, 1/2]$ and $p \in [0, 1/2]$.

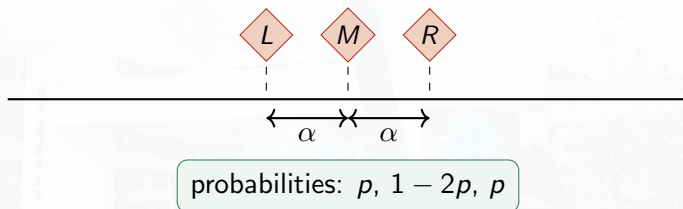
$$f(\mathbf{x}) = \begin{cases} 1/2 - \alpha, & \text{with probability } p, \\ 1/2, & \text{with probability } 1 - 2p, \\ 1/2 + \alpha, & \text{with probability } p. \end{cases}$$

The mechanism is called LRM: **left**, **right**, **midpoint**.

► Go to BAM



Diagram of LRM



- Output is **independent of reported locations**.
- Therefore it is strategyproof and anonymous.

Choosing the Parameters

The paper proves that if $\alpha > 1/4$, the approximation ratio is already at least 2.

- Hence focus on $\alpha \leq 1/4$.
- The analysis reduces to balancing two worst-case profiles:

$$x = (0, 1/2), \quad x' = (0, 1/2 + \alpha).$$



Optimal LRM Parameters

Theorem

The optimal mechanism in the (α, p) -LRM family uses

$$\alpha^* = \frac{\sqrt{5}}{2} - 1, \quad p^* = \frac{2}{5}.$$

It achieves approximation ratio

$$1 + \frac{2}{\sqrt{5}} \approx 1.8944.$$

Significance

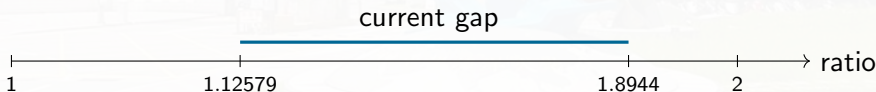
This resolves the open question of beating ratio 2 with a randomized strategyproof mechanism **without predictions**.

Lower Bound for Randomized Mechanisms

Theorem

Every randomized strategyproof mechanism without predictions has approximation ratio at least 1.12579.

- Improves the previous lower bound of about 1.0314.
- Still far from the upper bound 1.8944.



The main lemmas

Lemma 4.1

When $\alpha > \frac{1}{4}$, every (α, p) -LRM constant mechanism has an approximation ratio of at least 2.

Lemma 4.2

Let $\mathbf{x} = (x_1 = 0, x_2 = \frac{1}{2})$ and $\mathbf{x}' = (x'_1 = 0, x'_2 = \frac{1}{2} + \alpha)$. When $\alpha \leq \frac{1}{4}$, the (α^*, p^*) -LRM constant mechanism optimizes the approximation ratio for the envy-ratio objective, where

$$(\alpha^*, p^*) = \arg \min_{(\alpha, p)} \left\{ \max \left\{ \rho(\mathbf{x}), \rho(\mathbf{x}') \right\} \right\}.$$

- Lemma 3.1 $\Rightarrow$ it suffices to consider the two extreme agents.
- **Note:** Why the two profiles $\mathbf{x}, \mathbf{x}'$? Wasn't clearly proved in the paper.



The main theorem on Mechanisms without Predictions

Theorem 4.3

The

$$\left(\frac{\sqrt{5}}{2} - 1, \frac{2}{5} \right)\text{-LRM}$$

constant mechanism is anonymous and strategyproof, and achieves approximation ratio

$$1 + \frac{2}{\sqrt{5}} \approx 1.8944.$$

Moreover, this ratio is optimal among all (α, p) -LRM constant mechanisms.



Let $\alpha^* = \frac{\sqrt{5}}{2} - 1$, $p^* = \frac{2}{5}$. The mechanism selects the facility location according to

facility location	probability
$\frac{1}{2} - \alpha^* = \frac{3 - \sqrt{5}}{2}$	$\frac{2}{5}$
$\frac{1}{2}$	$1 - 2p^* = \frac{1}{5}$
$\frac{1}{2} + \alpha^* = \frac{\sqrt{5} - 1}{2}$	$\frac{2}{5}$

Optimality

For the two critical profiles $\mathbf{x} = (0, \frac{1}{2})$, $\mathbf{x}' = (0, \frac{1}{2} + \alpha^*)$, the approximation ratios are equal:

$$\rho(\mathbf{x}) = \rho(\mathbf{x}') = 1 + \frac{2}{\sqrt{5}}.$$

No other choice of (α, p) yields a smaller worst-case ratio.



Proof of Lemma 4.1 (1/2)

Let the (α, p) -LRM mechanism output

$$L = \frac{1}{2} - \alpha, \quad C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha$$

with probabilities $p, 1 - 2p, p$, respectively.



Proof of Lemma 4.1 (1/2)

Let the (α, p) -LRM mechanism output

$$L = \frac{1}{2} - \alpha, \quad C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha$$

with probabilities $p, 1 - 2p, p$, respectively. Consider the two-agent instance

$$\mathbf{x} = (x_1, x_2) = \left(0, \frac{1}{2}\right).$$

The optimal midpoint (m) gives envy ratio 1 ($\because u(m, x_1) = u(m, x_2)$)).



Proof of Lemma 4.1 (1/2)

Let the (α, p) -LRM mechanism output

$$L = \frac{1}{2} - \alpha, \quad C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha$$

with probabilities $p, 1 - 2p, p$, respectively. Consider the two-agent instance

$$\mathbf{x} = (x_1, x_2) = \left(0, \frac{1}{2}\right).$$

The optimal midpoint (m) gives envy ratio 1 ($\because u(m, x_1) = u(m, x_2)$). For $\alpha > \frac{1}{4}$, the realized envy ratios are

$$\text{ER}(L, \mathbf{x}) = \frac{\frac{1}{2} + \alpha}{1 - \alpha}, \quad \text{ER}(C, \mathbf{x}) = 2, \quad \text{ER}(R, \mathbf{x}) = \frac{1 - \alpha}{\frac{1}{2} - \alpha}.$$



Proof of Lemma 4.1 (2/2)

Hence

$$\rho(\mathbf{x}) = p \frac{\frac{1}{2} + \alpha}{1 - \alpha} + (1 - 2p)2 + p \frac{1 - \alpha}{\frac{1}{2} - \alpha}.$$

Both outer terms are increasing in α on $(\frac{1}{4}, \frac{1}{2})$. At $\alpha = \frac{1}{4}$,

$$\frac{\frac{1}{2} + \alpha}{1 - \alpha} = 1, \quad \frac{1 - \alpha}{\frac{1}{2} - \alpha} = 3.$$

Therefore

$$\rho(\mathbf{x}) \geq p \cdot 1 + (1 - 2p)2 + p \cdot 3 = 2.$$

Thus every (α, p) -LRM constant mechanism with $\alpha > \frac{1}{4}$ has approximation ratio at least 2.



Proof of Lemma 4.2: Setup

Let

$$L = \frac{1}{2} - \alpha, \quad C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha.$$

By Lemma 3.1, it suffices to consider two-agent instances

$$\mathbf{x} = (x_1, x_2), \quad 0 \leq x_1 \leq x_2 \leq 1.$$

By symmetry, assume $\text{mid}(\mathbf{x}) \leq \frac{1}{2}$. Write $m = \text{mid}(\mathbf{x})$, $\delta = \frac{1}{2} - m$.

Let u be the utility of both agents when the facility is placed at m . Since $x_1 \geq 0$,

$$u \geq \frac{1}{2} + \delta.$$

Goal

Show that, for the optimal choice of (α, p) , it is enough to compare the two profiles

$$\mathbf{x} = \left(0, \frac{1}{2}\right), \quad \mathbf{x}' = \left(0, \frac{1}{2} + \alpha\right).$$

Proof of Lemma 4.2: Setup

Let

$$L = \frac{1}{2} - \alpha, \quad C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha.$$

By Lemma 3.1, it suffices to consider two-agent instances

$$\mathbf{x} = (x_1, x_2), \quad 0 \leq x_1 \leq x_2 \leq 1.$$

By symmetry, assume $\text{mid}(\mathbf{x}) \leq \frac{1}{2}$. Write $m = \text{mid}(\mathbf{x})$, $\delta = \frac{1}{2} - m$.

Let u be the utility of both agents when the facility is placed at m . Since $x_1 \geq 0$,

$$u \geq \frac{1}{2} + \delta. \quad (\text{e.g., } 1 - (m - x_1) = (1/2) + \delta + x_1 \geq 1/2 + \delta)$$

Goal

Show that, for the optimal choice of (α, p) , it is enough to compare the two profiles

$$\mathbf{x} = \left(0, \frac{1}{2}\right), \quad \mathbf{x}' = \left(0, \frac{1}{2} + \alpha\right).$$

Why Consider Only the Two Profiles in Lemma 4.2?

By Lemma 3.1, it suffices to analyze two-agent instances. By symmetry and Claim B.1, we may further reduce to

$$\mathbf{x} = (0, t), \quad 0 \leq t \leq 1.$$

Thus the problem becomes one-dimensional:

$$\rho(t) := \rho((0, t)).$$

Goal

Show that the worst case over all t is captured by only

$$t = C = \frac{1}{2} \quad \text{or} \quad t = R = \frac{1}{2} + \alpha.$$



Claim B.1: Reduction to $x_1 = 0$

Claim B.1

Suppose

$$\text{mid}(\mathbf{x}) \in \left[0, \frac{1}{2} - \alpha\right).$$

For any two-agent instance $\mathbf{x} = (x_1, x_2)$, let $\mathbf{x}' = (0, x_2)$. Then the approximation ratio of the (α, p) -LRM mechanism satisfies $\rho(\mathbf{x}) \leq \rho(\mathbf{x}')$.

- Thus, in this region, it suffices to analyze only instances of the form $(0, x_2)$.
- This reduces the worst-case search from two variables (x_1, x_2) to a single variable x_2 .



Proof of Claim B.1

Since $\text{mid}(\mathbf{x}) = \frac{x_1 + x_2}{2} \leq \frac{1}{2} - \alpha$, and every possible facility location satisfies

$$y \in \left\{ \frac{1}{2} - \alpha, \frac{1}{2}, \frac{1}{2} + \alpha \right\},$$

we have

$$y \geq \frac{1}{2} - \alpha \geq \text{mid}(\mathbf{x}).$$

Therefore,

$$y - x_1 \geq |y - x_2|.$$

Hence agent 1 is always the worse-off agent: $u_1(y) \leq u_2(y)$, for every y . The expected approximation ratio is therefore

$$\rho(\mathbf{x}) = p \frac{u_2(L)}{u_1(L)} + (1 - 2p) \frac{u_2(C)}{u_1(C)} + p \frac{u_2(R)}{u_1(R)}.$$



Proof of Claim B.1 (continued)

- Move the left agent from x_1 to 0, obtaining $\mathbf{x}' = (0, x_2)$.
- Observe that $u_2(y)$ does not change, while $u_1(y) = 1 - (y - x_1)$ becomes

$$u'_1(y) = 1 - y \leq 1 - (y - x_1) = u_1(y).$$

- Hence

$$\frac{u_2(y)}{u_1(y)} \leq \frac{u_2(y)}{u'_1(y)} \quad \forall y \in \{L, C, R\}.$$

- Taking the weighted average, we have

$$\rho(\mathbf{x}) \leq \rho(\mathbf{x}').$$

- Therefore every worst-case instance in this region may be assumed to satisfy $x_1 = 0$.



Piecewise Monotonicity

- For fixed (α, p) , the mechanism randomizes over L, C, R .
- The expression for $\rho(t)$ changes only when t crosses one of $L, C, R, 1 - 2\alpha$.
- When $0 \leq \alpha < \frac{1}{6}$, their order is $L < C < R < 1 - 2\alpha$.
- Appendix B.2 shows (proof omitted here), by differentiating $\rho(t)$, that the only endpoint candidates are

$$t = L, \quad C, \quad R, \quad 1 - 2\alpha.$$

So the continuous worst-case search reduces to finitely many candidates.



Eliminating the Extra Candidates

- The appendix then compares the candidate values:

$$\rho(L) \leq \rho(C), \quad \rho(1 - 2\alpha) \leq \rho(R).$$

- Therefore the candidates L and $1 - 2\alpha$ are dominated.
- Hence

$$\sup_{t \in [0,1]} \rho(t) = \max\{\rho(C), \rho(R)\}.$$

Since

$$C = \frac{1}{2}, \quad R = \frac{1}{2} + \alpha,$$

the two surviving profiles are

$$\mathbf{x} = \left(0, \frac{1}{2}\right), \quad \mathbf{x}' = \left(0, \frac{1}{2} + \alpha\right).$$



Meaning of Lemma 4.2

Lemma 4.2 says that optimizing over all instances is equivalent to optimizing over the two surviving worst cases:

$$\min_{\alpha, p} \sup_{\mathbf{z}} \rho(\mathbf{z}) = \min_{\alpha, p} \max \left\{ \rho\left(0, \frac{1}{2}\right), \rho\left(0, \frac{1}{2} + \alpha\right) \right\}.$$

- $(0, \frac{1}{2})$ is one extremal constraint.
- $(0, \frac{1}{2} + \alpha)$ is the other extremal constraint.
- Every other profile has approximation ratio no larger than one of these two.

So the optimal (α, p) is found by balancing these two cases.



Proof of Lemma 4.2: Case $m \in [L, C]$

- Assume

$$m = \frac{1}{2} - \delta \in \left[\frac{1}{2} - \alpha, \frac{1}{2} \right], \quad 0 \leq \delta \leq \alpha.$$

The distances from m to the three possible facility locations are $\alpha - \delta$, δ , $\alpha + \delta$.

Thus

$$\rho(\mathbf{x}) \leq p \frac{u + \alpha - \delta}{u - (\alpha - \delta)} + (1 - 2p) \frac{u + \delta}{u - \delta} + p \frac{u + (\alpha + \delta)}{u - (\alpha + \delta)}.$$



Proof of Lemma 4.2: Case $m \in [L, C]$

- Assume

$$m = \frac{1}{2} - \delta \in \left[\frac{1}{2} - \alpha, \frac{1}{2} \right], \quad 0 \leq \delta \leq \alpha.$$

The distances from m to the three possible facility locations are $\alpha - \delta$, δ , $\alpha + \delta$.

Thus

$$\rho(\mathbf{x}) \leq p \frac{u + \alpha - \delta}{u - (\alpha - \delta)} + (1 - 2p) \frac{u + \delta}{u - \delta} + p \frac{u + (\alpha + \delta)}{u - (\alpha + \delta)}.$$

- Since the right-hand side is non-increasing in u , and $u \geq \frac{1}{2} + \delta$, we get

$$\begin{aligned} \rho(\mathbf{x}) &\leq p \frac{1 + 2\alpha}{1 + 4\delta - 2\alpha} + (1 - 2p)(1 + 4\delta) + p \frac{1 + 4\delta + 2\alpha}{1 - 2\alpha} \\ &\leq 2p \frac{1 + 2\alpha}{1 - 2\alpha} + (1 - 2p)(1 + 4\delta). \end{aligned}$$



Proof of Lemma 4.2: Case $m \in [L, C]$

- Assume

$$m = \frac{1}{2} - \delta \in \left[\frac{1}{2} - \alpha, \frac{1}{2} \right], \quad 0 \leq \delta \leq \alpha.$$

The distances from m to the three possible facility locations are $\alpha - \delta$, δ , $\alpha + \delta$.

Thus

$$\rho(\mathbf{x}) \leq p \frac{u + \alpha - \delta}{u - (\alpha - \delta)} + (1 - 2p) \frac{u + \delta}{u - \delta} + p \frac{u + (\alpha + \delta)}{u - (\alpha + \delta)}.$$

- Since the right-hand side is non-increasing in u , and $u \geq \frac{1}{2} + \delta$, we get

$$\begin{aligned} \rho(\mathbf{x}) &\leq p \frac{1 + 2\alpha}{1 + 4\delta - 2\alpha} + (1 - 2p)(1 + 4\delta) + p \frac{1 + 4\delta + 2\alpha}{1 - 2\alpha} \\ &\leq 2p \frac{1 + 2\alpha}{1 - 2\alpha} + (1 - 2p)(1 + 4\delta). \end{aligned}$$

This is **increasing in δ** . Since $\frac{1}{2} - \alpha = m = \frac{x_1 + x_2}{2} \Rightarrow x_1 + x_2 = 1 - 2\alpha$ and $x_1 = 0$, so the worst case in this region occurs at

$$\delta = \alpha, \quad \mathbf{x} = (0, 1 - 2\alpha).$$



Proof of Lemma 4.2: Case $m < L$

- Now consider the case $m < L = \frac{1}{2} - \alpha$.
- For every possible facility location $y \in \{L, C, R\}$, we have $y \geq m$. Hence $y - x_1 \geq |y - x_2|$, for any $y \in \{\frac{1}{2} - \alpha, \frac{1}{2}, \frac{1}{2} + \alpha\}$.
- Therefore agent x_1 is weakly worse off than agent x_2 for every realized facility location y .



Proof of Lemma 4.2: Case $m < L$

- Now consider the case $m < L = \frac{1}{2} - \alpha$.
- For every possible facility location $y \in \{L, C, R\}$, we have $y \geq m$. Hence $y - x_1 \geq |y - x_2|$, for any $y \in \{\frac{1}{2} - \alpha, \frac{1}{2}, \frac{1}{2} + \alpha\}$.
- Therefore agent x_1 is weakly worse off than agent x_2 for every realized facility location y .
- Recall that we only need to consider $\mathbf{x} = (0, x_2)$ and let $t := x_2$. Then every remaining instance can be written as $\mathbf{x} = (0, t)$, $0 \leq t \leq 1$.
- For convenience, define $\rho(t) := \rho((0, t))$.



Proof of Lemma 4.2: The Range $\alpha < 1/6$

- Assume $0 \leq \alpha < \frac{1}{6}$. Then $L < C < R < 1 - 2\alpha$.
- For $\mathbf{x} = (0, t)$, direct differentiation gives the following monotonicity table.

range of t	sign / behavior of $\rho'(t)$	candidate worst case
$[0, L]$	$\frac{2p}{1+2\alpha} + 2(1-2p) + \frac{2p}{1-2\alpha} \geq 0$	$t = L$
$(L, C]$	$\frac{8p\alpha}{1-4\alpha^2} + 2(1-2p) \geq 0$	$t = C$
$(C, R]$	$D = \frac{8p\alpha}{1-4\alpha^2} + 4p - 2$ constant	$t = C$ or $t = R$
$(R, 1-2\alpha]$	$-\frac{4p}{1-4\alpha^2} - 2(1-2p) \leq 0$	$t = R$

Thus the only candidates are

$$t = L, \quad C, \quad R, \quad 1 - 2\alpha.$$

Direct substitution gives

$$\rho(L) \leq \rho(C), \quad \rho(1 - 2\alpha) \leq \rho(R).$$

Hence, when $\alpha < 1/6$, the worst cases reduce to

$$(0, C) = \left(0, \frac{1}{2}\right) \quad \text{and} \quad (0, R) = \left(0, \frac{1}{2} + \alpha\right).$$



Why is $\frac{1}{6}$ the breakpoint?

- If we have assumed $L < C < R < 1 - 2\alpha$, then $R = \frac{1}{2} + \alpha < 1 - 2\alpha \iff \alpha < \frac{1}{6}$.
- For $\alpha > \frac{1}{6}$, $R > 1 - 2\alpha$.



Proof of Lemma 4.2: Excluding $\alpha \geq 1/6$

- For the two candidate profiles, define $F_\alpha(p) := \rho(0, \frac{1}{2})$, $G_\alpha(p) := \rho(0, \frac{1}{2} + \alpha)$.
- When $\alpha \in [1/6, 1/4]$, direct calculation gives $F_\alpha(p) = 2(1 - 2p) + p \frac{4 - 4\alpha}{1 - 4\alpha^2}$, and $G_\alpha(p) = p \frac{1 + 2\alpha}{2 - 4\alpha} + (1 - 2p)(2 - 2\alpha) + p \frac{2}{1 - 2\alpha}$. (Note: $F'_\alpha(p) \leq 0$, $G_\alpha(p) > 0$.)
- For fixed α , the best p equalizes the two affine functions: $F_\alpha(p) = G_\alpha(p)$.

Solving this gives $p = \frac{16\alpha^3 - 4\alpha}{32\alpha^3 - 4\alpha^2 - 28\alpha + 3}$.

- Then the substitution yields

$$F_\alpha(p) = G_\alpha(p) = \frac{8\alpha^2 - 56\alpha + 6}{32\alpha^3 - 4\alpha^2 - 28\alpha + 3}.$$

- This expression is increasing on $\alpha \in [\frac{1}{6}, \frac{1}{4}]$. Therefore its minimum is attained at $\alpha = \frac{1}{6}$, $p = \frac{4}{11}$, with value $\frac{21}{11} \approx 1.909$.



Proof of Lemma 4.2: Conclusion (1/2)

- From the previous slide, every mechanism with $\alpha \in [\frac{1}{6}, \frac{1}{4}]$ has approximation ratio at least

$$\frac{21}{11} \approx 1.909$$

on one of the two profiles $(0, \frac{1}{2})$, $(0, \frac{1}{2} + \alpha)$.

- But for

$$\alpha^* = \frac{\sqrt{5}}{2} - 1 < \frac{1}{6}, \quad p^* = \frac{2}{5},$$

the two-profile maximum is

$$1 + \frac{2}{\sqrt{5}} \approx 1.8944.$$

- Hence no globally optimal (α, p) lies in $[\frac{1}{6}, \frac{1}{4}]$.



Proof of Lemma 4.2: Conclusion (2/2)

- Therefore the global optimum lies in the range $\alpha < \frac{1}{6}$, where the previous reduction shows that the only worst-case profiles are $(0, \frac{1}{2})$ and $(0, \frac{1}{2} + \alpha)$.
- Thus the optimal parameters can be obtained by solving

$$(\alpha^*, p^*) = \arg \min_{\alpha, p} \left\{ \max \left\{ \rho \left(0, \frac{1}{2} \right), \rho \left(0, \frac{1}{2} + \alpha \right) \right\} \right\}.$$



Proof of Theorem 4.3: Setup

Theorem 4.3

The $\left(\frac{\sqrt{5}}{2} - 1, \frac{2}{5}\right)$ -LRM constant mechanism is anonymous, strategyproof, and achieves approximation ratio

$$1 + \frac{2}{\sqrt{5}} \approx 1.8944.$$

- The mechanism returns

$$\frac{1}{2} - \alpha, \quad \frac{1}{2}, \quad \frac{1}{2} + \alpha$$

with probabilities p , $1 - 2p$, p , resp.

- It is anonymous and strategyproof because the output distribution is independent of the agents' reports.
- By Lemma 4.2, it suffices to solve $\min_{\alpha, p} \max \{\rho(\mathbf{x}), \rho(\mathbf{x}')\}$, where $\mathbf{x} = (0, \frac{1}{2})$, $\mathbf{x}' = (0, \frac{1}{2} + \alpha)$.



Proof of Theorem 4.3: Ratio on $\mathbf{x}$

- For $\mathbf{x} = (0, \frac{1}{2})$, the approximation ratio is

$$\rho(\mathbf{x}) = p \cdot \frac{1 - \alpha}{1 - (\frac{1}{2} - \alpha)} + (1 - 2p) \cdot \frac{1}{1 - \frac{1}{2}} + p \cdot \frac{1 - \alpha}{1 - (\frac{1}{2} + \alpha)}.$$

- Hence

$$\rho(\mathbf{x}) = p \cdot \frac{1 - \alpha}{\frac{1}{2} + \alpha} + 2(1 - 2p) + p \cdot \frac{1 - \alpha}{\frac{1}{2} - \alpha}.$$

- Equivalently,

$$\rho(\mathbf{x}) = 2(1 - 2p) + p \cdot \frac{4 - 4\alpha}{1 - 4\alpha^2}.$$



Proof of Theorem 4.3: Ratio on $\mathbf{x}'$

For

$$\mathbf{x}' = \left(0, \frac{1}{2} + \alpha\right),$$

there are two parameter regimes.

Case 1: $\alpha \in [0, \frac{1}{6})$

$$\rho(\mathbf{x}') = p \cdot \frac{2 - 4\alpha}{1 + 2\alpha} + (1 - 2p)(2 - 2\alpha) + p \cdot \frac{2}{1 - 2\alpha}.$$

Case 2: $\alpha \in [\frac{1}{6}, \frac{1}{4}]$

$$\rho(\mathbf{x}') = p \cdot \frac{1 + 2\alpha}{2 - 4\alpha} + (1 - 2p)(2 - 2\alpha) + p \cdot \frac{2}{1 - 2\alpha}.$$

Note: $u_1(L) = 1 - ((\frac{1}{2} - \alpha) - 0) = \frac{1}{2} + \alpha$, $u_2(L) = 1 - ((\frac{1}{2} + \alpha) - (\frac{1}{2} - \alpha)) = 1 - 2\alpha$.

Thus we solve the minimax problem separately in the two regimes.



Proof of Theorem 4.3: Case $\alpha \in [0, \frac{1}{6})$

To minimize

$$\max\{\rho(\mathbf{x}), \rho(\mathbf{x}')\},$$

the optimum equalizes the two expressions:

$$\rho(\mathbf{x}) = \rho(\mathbf{x}').$$

Using the formulas from the previous slides gives

$$2(1 - 2p) + p \cdot \frac{4 - 4\alpha}{1 - 4\alpha^2} = p \cdot \frac{2 - 4\alpha}{1 + 2\alpha} + (1 - 2p)(2 - 2\alpha) + p \cdot \frac{2}{1 - 2\alpha}.$$

Solving for p , we get

$$p = \frac{4\alpha^2 - 1}{2(4\alpha^2 - 2\alpha - 1)}.$$

Substituting this into $\rho(\mathbf{x})$ gives

$$\rho(\mathbf{x}) = \rho(\mathbf{x}') = \frac{2\alpha + 2}{-4\alpha^2 + 2\alpha + 1}.$$



Proof of Theorem 4.3: Optimizing Case 1

We now minimize

$$R(\alpha) = \frac{2\alpha + 2}{-4\alpha^2 + 2\alpha + 1}$$

over

$$\alpha \in \left[0, \frac{1}{6}\right).$$

Its derivative is

$$R'(\alpha) = \frac{8\alpha^2 + 16\alpha - 2}{(-4\alpha^2 + 2\alpha + 1)^2}.$$

Thus the unique nonnegative critical point is

$$\alpha^* = \frac{\sqrt{5}}{2} - 1.$$

Substituting into the formula for p ,

$$p^* = \frac{4(\alpha^*)^2 - 1}{2(4(\alpha^*)^2 - 2\alpha^* - 1)} = \frac{2}{5}.$$

Therefore

$$R(\alpha^*) = 1 + \frac{2}{\sqrt{5}} \approx 1.8944.$$



Proof of Theorem 4.3: Case $\alpha \in [\frac{1}{6}, \frac{1}{4}]$

Now use the second formula for $\rho(\mathbf{x}')$:

$$\rho(\mathbf{x}') = p \cdot \frac{1+2\alpha}{2-4\alpha} + (1-2p)(2-2\alpha) + p \cdot \frac{2}{1-2\alpha}.$$

Again equalize

$$\rho(\mathbf{x}) = \rho(\mathbf{x}').$$

Solving gives

$$p = \frac{16\alpha^3 - 4\alpha}{32\alpha^3 - 4\alpha^2 - 28\alpha + 3}.$$

Substituting back gives

$$\rho(\mathbf{x}) = \rho(\mathbf{x}') = \frac{8\alpha^2 - 56\alpha + 6}{32\alpha^3 - 4\alpha^2 - 28\alpha + 3}.$$

This expression is increasing on

$$\left[\frac{1}{6}, \frac{1}{4}\right].$$

Thus the best value in this regime occurs at

$$\alpha = \frac{1}{6}, \quad p = \frac{4}{11}.$$



Proof of Theorem 4.3: Comparing the Two Regimes

- For $\alpha \in [0, \frac{1}{6})$, the optimal parameters are

$$\alpha^* = \frac{\sqrt{5}}{2} - 1, \quad p^* = \frac{2}{5},$$

with approximation ratio $1 + \frac{2}{\sqrt{5}} \approx 1.8944$.

- For $\alpha \in [\frac{1}{6}, \frac{1}{4}]$, the best parameters are

$$\alpha = \frac{1}{6}, \quad p = \frac{4}{11},$$

with approximation ratio $\frac{21}{11} \approx 1.909$.

- Since $1 + \frac{2}{\sqrt{5}} < \frac{21}{11}$, the global optimum among all (α, p) -LRM constant mechanisms is

$$(\alpha^*, p^*) = \left(\frac{\sqrt{5}}{2} - 1, \frac{2}{5} \right).$$



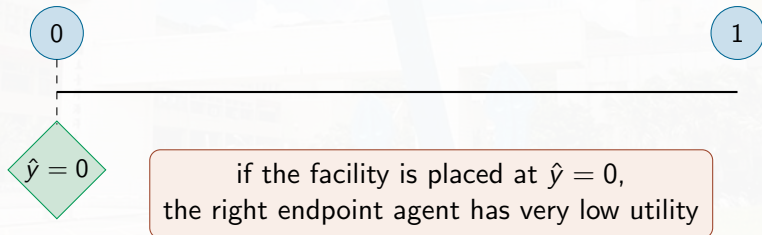
Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions
- 3 Randomized Mechanisms without Predictions
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



The Main Difficulty

- ★ A prediction-sensitive randomized mechanism might put positive probability on $\hat{y}$.



- More trust in $\hat{y}$ helps consistency.
- Too much trust can destroy robustness.

Bias-Aware Mechanism: Idea

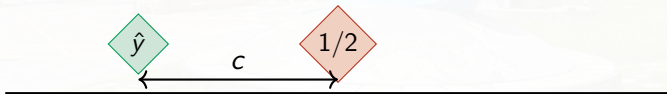
- Let

$$c = \left| \hat{y} - \frac{1}{2} \right|.$$

This is the **bias** of the prediction away from the center.

Design principle

Trust the prediction less when it is closer to an endpoint.



Bias-Aware Mechanism (BAM)

Mechanism: Bias-Aware Mechanism (BAM)

Given prediction $\hat{y}$, compute

$$c = \left| \hat{y} - \frac{1}{2} \right|, \quad p = \frac{1}{2} - c.$$

Then:

$$f(\mathbf{x}, \hat{y}) = \begin{cases} \hat{y}, & \text{w.p. } p, \\ 1/2, & \text{w.p. } 1 - p. \end{cases}$$



Bias-Aware Mechanism (BAM)

Mechanism: Bias-Aware Mechanism (BAM)

Given prediction $\hat{y}$, compute

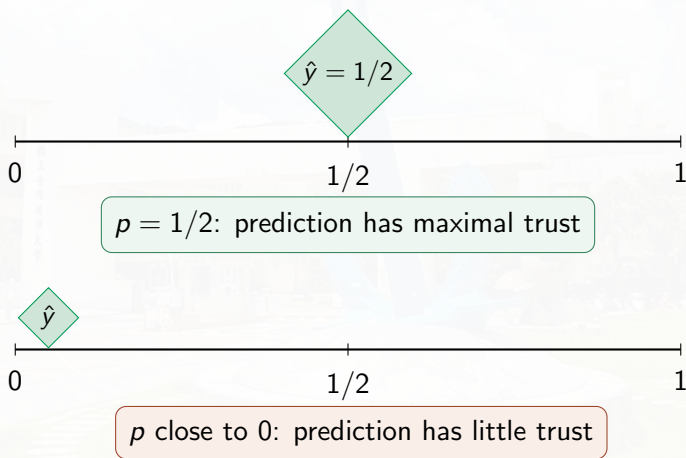
$$c = \left| \hat{y} - \frac{1}{2} \right|, \quad p = \frac{1}{2} - c.$$

Then:

$$f(\mathbf{x}, \hat{y}) = \begin{cases} \hat{y}, & \text{w.p. } p, \\ 1/2, & \text{w.p. } 1 - p. \end{cases}$$

Again, the output distribution is independent of reported locations, so BAM is strategyproof.

BAM Probability Schedule



BAM Guarantees

Theorem 4.5

BAM is anonymous and strategyproof. Its guarantees are:

bias c	consistency	robustness
$0 \leq c < 1/4$	$7/4$	$9/4$
$1/4 \leq c \leq 1/2$	$-4c^2 + 2$	$c + 2$

- If the prediction is central, BAM safely uses it with some probability.
- If the prediction is extreme, BAM mostly falls back to $1/2$.



Robustness Intuition for BAM

Assume $\hat{y} \leq 1/2$.



Robustness Intuition for BAM

Assume $\hat{y} \leq 1/2$. $\Rightarrow c = \frac{1}{2} - \hat{y} \Rightarrow \hat{y} = \frac{1}{2} - c = p$.



Robustness Intuition for BAM

Assume $\hat{y} \leq 1/2$. $\Rightarrow c = \frac{1}{2} - \hat{y} \Rightarrow \hat{y} = \frac{1}{2} - c = p$.

- Probability of using prediction: $p = \hat{y}$.
- Probability of using center: $1 - p = 1 - \hat{y}$.
- Any agent's utility from $\hat{y}$ is at least $\hat{y}$.
- Any agent's utility from $1/2$ is at least $1/2$.

Worst case

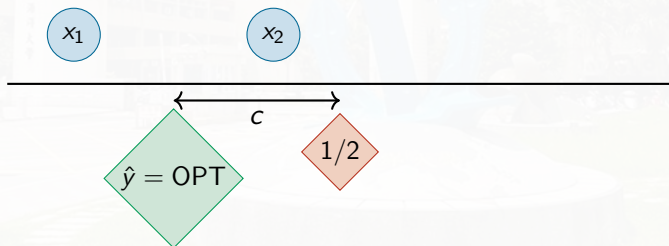
A bad case places one agent near the chosen safe center and another near the far endpoint.



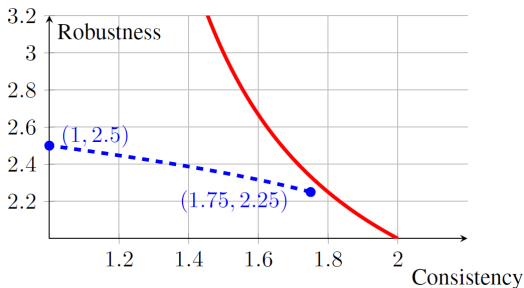
Consistency Intuition for BAM

If $\hat{y} = \text{OPT}(\mathbf{x})$, then placing the facility at $\hat{y}$ is optimal.

- With probability p , BAM uses the optimal location.
- With probability $1 - p$, it uses the center $1/2$.
- The **loss depends on how far the optimal location is from $1/2$.**



BAM vs. α -BIM



BAM gives a better consistency–robustness profile than the deterministic frontier α -BIM.

Why Not Combine BIM with LRM?

A tempting idea:

- use $\hat{y}$ inside the BIM interval;
- use the LRM mechanism outside the interval.

Paper's conclusion

This hybrid is worse than the original α -BIM in its guarantees.

- Randomization helps in the no-prediction setting.
- But blindly inserting it into a prediction framework can harm consistency and robustness.



Proof of Theorem 4.5: Reduction

- By Lemma 3.1, it suffices to consider two-agent instances $\mathbf{x} = (x_1, x_2)$, $x_1 < x_2$.
- By symmetry, assume $\hat{y} \leq \frac{1}{2}$. Then $c = \frac{1}{2} - \hat{y}$, $p = \frac{1}{2} - c = \hat{y}$.
- So BAM becomes:

$$\begin{cases} \text{output } \hat{y}, & \text{with probability } \hat{y}, \\ \text{output } \frac{1}{2}, & \text{with probability } 1 - \hat{y}. \end{cases}$$

- For a two-agent instance, the optimal envy ratio is

$$\text{ER}(\text{mid}(\mathbf{x}), \mathbf{x}) = 1.$$

- Thus, it suffices to upper-bound the envy ratio achieved by BAM.



Robustness of BAM (1/2)

- Consider an arbitrary prediction $\hat{y} \leq \frac{1}{2}$.
 - For the output $\hat{y}$, every agent has utility at least $\hat{y}$.
 - For the output $\frac{1}{2}$, every agent has utility at least $\frac{1}{2}$.



Robustness of BAM (1/2)

- Consider an arbitrary prediction $\hat{y} \leq \frac{1}{2}$.
 - For the output $\hat{y}$, every agent has utility at least $\hat{y}$.
 - For the output $\frac{1}{2}$, every agent has utility at least $\frac{1}{2}$.
- The worst case occurs when the right agent is at $x_2 = 1$.



Robustness of BAM (1/2)

- Consider an arbitrary prediction $\hat{y} \leq \frac{1}{2}$.
 - For the output $\hat{y}$, every agent has utility at least $\hat{y}$.
 - For the output $\frac{1}{2}$, every agent has utility at least $\frac{1}{2}$.
- The worst case occurs when the right agent is at $x_2 = 1$.
- Moreover, it suffices to consider $x_1 \in [\hat{y}, \frac{1}{2}]$.



Robustness of BAM (1/2)

- Consider an arbitrary prediction $\hat{y} \leq \frac{1}{2}$.
 - For the output $\hat{y}$, every agent has utility at least $\hat{y}$.
 - For the output $\frac{1}{2}$, every agent has utility at least $\frac{1}{2}$.
- The worst case occurs when the right agent is at $x_2 = 1$.
- Moreover, it suffices to consider $x_1 \in [\hat{y}, \frac{1}{2}]$.
 - After setting $x_2 = 1$, the two realized envy ratios are

$$\frac{1 - |x_1 - \hat{y}|}{\hat{y}} \quad \text{and} \quad \frac{1 - |x_1 - \frac{1}{2}|}{\frac{1}{2}}.$$

- Moving $x_1 < \hat{y}$ to $\hat{y}$, or moving $x_1 > \frac{1}{2}$ to $\frac{1}{2}$, weakly increases both ratios. Hence, for the worst-case analysis, we may assume $x_1 \in [\hat{y}, \frac{1}{2}]$.



Robustness of BAM (2/2)

- Then the envy ratio is bounded by

$$\text{ER}(f(\mathbf{x}, \hat{y}), \mathbf{x}) \leq \hat{y} \cdot \frac{1 - (x_1 - \hat{y})}{\hat{y}} + (1 - \hat{y}) \cdot \frac{1 - (\frac{1}{2} - x_1)}{\frac{1}{2}}.$$

- Simplifying, $\text{ER}(f(\mathbf{x}, \hat{y}), \mathbf{x}) \leq 2 + x_1(1 - 2\hat{y})$.
- Since $x_1 \leq \frac{1}{2}$, we have $\text{ER}(f(\mathbf{x}, \hat{y}), \mathbf{x}) \leq 2 + \frac{1}{2} - \hat{y} = 2 + c$.
- Thus BAM is $(2 + c)$ -robust.



Robustness Bound by Ranges

- We have shown $\beta(c) \leq 2 + c$.
- Therefore, when $c \in [\frac{1}{4}, \frac{1}{2}]$, BAM is $(2 + c)$ -robust.



Robustness Bound by Ranges

- We have shown $\beta(c) \leq 2 + c$.
- Therefore, when $c \in [\frac{1}{4}, \frac{1}{2}]$, BAM is $(2 + c)$ -robust.
- When $c \in [0, \frac{1}{4})$, we have

$$2 + c < 2 + \frac{1}{4} = \frac{9}{4}.$$

- Hence BAM is $\frac{9}{4}$ -robust.



Consistency of BAM: Setup

- Now assume the prediction is accurate: $\hat{y} = \text{mid}(\mathbf{x})$.
- Let $\delta = \frac{x_2 - x_1}{2}$. Then $x_1 = \hat{y} - \delta$, $x_2 = \hat{y} + \delta$.
- Since $x_1 \geq 0$, we have $\delta \leq \hat{y}$.
- BAM outputs $\hat{y}$ w.p. $\hat{y}$, and $\frac{1}{2}$ w.p. $1 - \hat{y}$.
- If the facility is placed at $\hat{y} = \text{mid}(\mathbf{x})$, the envy ratio is 1.
- Thus only the outcome $\frac{1}{2}$ may create envy.



Consistency: Case 1, $\hat{y} \leq \frac{1}{4}$

- Assume $\hat{y} \leq \frac{1}{4}$. Then $x_2 = \hat{y} + \delta \leq 2\hat{y} \leq \frac{1}{2}$.
- When the facility is placed at $\frac{1}{2}$, agent x_2 is closer to the facility than agent x_1 . Hence

$$\text{ER}\left(\frac{1}{2}, \mathbf{x}\right) = \frac{1 - (\frac{1}{2} - \hat{y} - \delta)}{1 - (\frac{1}{2} - \hat{y} + \delta)}.$$

- This expression is increasing in δ , and $\delta \leq \hat{y}$, so

$$\text{ER}\left(\frac{1}{2}, \mathbf{x}\right) \leq \frac{\frac{1}{2} + 2\hat{y}}{\frac{1}{2}} = 1 + 4\hat{y}.$$

- Therefore, $\gamma \leq \hat{y} \cdot 1 + (1 - \hat{y})(1 + 4\hat{y}) = 1 + 4\hat{y} - 4\hat{y}^2$.
- Since $c = \frac{1}{2} - \hat{y}$, we have $1 + 4\hat{y} - 4\hat{y}^2 = 2 - 4c^2$.
- Thus, for $c \in [\frac{1}{4}, \frac{1}{2}]$, we have $\gamma \leq 2 - 4c^2$.



Consistency: Case 2, $\hat{y} > \frac{1}{4}$ (1/2)

- Assume $\hat{y} > \frac{1}{4}$. Note that $\hat{y} = \text{mid}(\mathbf{x})$ and $\hat{y} - \delta = x_1$, $\hat{y} + \delta = x_2$, $c < \frac{1}{4}$.
- When the facility is placed at $\frac{1}{2}$, the worst case occurs at $\delta = \hat{y}$ ($= \text{mid}(\mathbf{x})$).



Consistency: Case 2, $\hat{y} > \frac{1}{4}$ (1/2)

- Assume $\hat{y} > \frac{1}{4}$. Note that $\hat{y} = \text{mid}(\mathbf{x})$ and $\hat{y} - \delta = x_1$, $\hat{y} + \delta = x_2$, $c < \frac{1}{4}$.
- When the facility is placed at $\frac{1}{2}$, the worst case occurs at $\delta = \hat{y}$ ($= \text{mid}(\mathbf{x})$).
- Let $c = \frac{1}{2} - \hat{y}$. At facility location $\frac{1}{2}$,

$$\text{ER}\left(\frac{1}{2}, \mathbf{x}\right) = \frac{1 - |\frac{1}{2} - x_2|}{1 - |\frac{1}{2} - x_1|} = \frac{1 - |c - \delta|}{1 - (\frac{1}{2} - \hat{y} + \hat{y} - x_1)} = \frac{1 - |c - \delta|}{1 - c - \delta}.$$

This is nondecreasing in $\delta \in [0, \hat{y}]$. Since $\delta \leq \hat{y}$, its maximum occurs at

$$\delta = \hat{y}, \quad (\text{check } \text{ER}'\left(\frac{1}{2}, \mathbf{x}\right) \text{ of two cases: } |c - \delta| = c - \delta \text{ or } \delta - c)$$

which corresponds to $\mathbf{x} = (0, 2\hat{y})$.

- Then, $\text{ER}\left(\frac{1}{2}, \mathbf{x}\right) \leq \frac{1 + c - \hat{y}}{1 - c - \hat{y}} \leq \frac{1 + \frac{1}{2} - 2\hat{y}}{\frac{1}{2}} = 3 - 4\hat{y}$.



Consistency: Case 2, $\hat{y} > \frac{1}{4}$ (2/2)

- Therefore, $\gamma \leq \hat{y} \cdot 1 + (1 - \hat{y})(3 - 4\hat{y}) = 3 - 6\hat{y} + 4\hat{y}^2$.
- Using $c = \frac{1}{2} - \hat{y}$, we obtain that

$$3 - 6\hat{y} + 4\hat{y}^2 = 1 + 2c + 4c^2.$$

- For $c \in [0, \frac{1}{4})$, this is maximized at $c = \frac{1}{4}$, so

$$1 + 2c + 4c^2 \leq 1 + \frac{1}{2} + \frac{1}{4} = \frac{7}{4}.$$

- Hence, for $c \in [0, \frac{1}{4})$, we have $\gamma \leq \frac{7}{4}$.



Proof of Theorem 4.5: Conclusion

- Combining the two parts:

$$\beta \leq 2 + c \text{ and } \gamma \leq \begin{cases} 2 - 4c^2, & c \in [\frac{1}{4}, \frac{1}{2}], \\ \frac{7}{4}, & c \in [0, \frac{1}{4}). \end{cases}$$

- Therefore BAM satisfies

$$\begin{cases} (2 - 4c^2)\text{-consistency and } (2 + c)\text{-robustness,} & c \in [\frac{1}{4}, \frac{1}{2}], \\ \frac{7}{4}\text{-consistency and } \frac{9}{4}\text{-robustness,} & c \in [0, \frac{1}{4}). \end{cases}$$

- Thus Theorem 4.5 follows.



Outline

- 1 Motivation and Model
- 2 Deterministic Mechanisms with Predictions
- 3 Randomized Mechanisms without Predictions
- 4 Randomized Mechanisms with Predictions
- 5 Concluding & Summary



Summary of Mechanisms

Mechanism	Prediction?	Main guarantee
Constant 1/2	no	deterministic ratio 2
α -BIM	yes	α -consistent, $\alpha/(\alpha - 1)$ -robust; deterministic optimal
LRM constant	no	randomized ratio $1 + 2/\sqrt{5} \approx 1.8944$
BAM	yes	randomized; improves over α -BIM trade-off

Open Directions

- Tight bounds for randomized mechanisms with predictions.
- Closed-form error-parameterized approximation for BAM.
- Extensions to other fairness objectives in facility location.
- Extensions beyond the one-dimensional interval.



Main Referenced Works

- [1] Haris Aziz, Yuhang Guo, Alexander Lam, and Houyu Zhou.
Learning-Augmented Facility Location Mechanisms for Envy Ratio.
NeurIPS 2025.
- [2] Y. Ding, W. Liu, X. Chen, Q. Fang, and Q. Nong.
Facility location game with envy ratio.
Computers & Industrial Engineering, 2020.
- [3] H. Moulin.
On strategy-proofness and single peakedness.
Public Choice, 1980.



Discussions

