

Learning-augmented Mechanism Design

Jakob Barkalaia

July 3, 2025

- 1 Preliminaries
- 2 Minimizing egalitarian social cost
- 3 Minimizing utilitarian social cost

"Learning-Augmented Mechanism Design: Leveraging Predictions for Facility Location", 2022

By Priyank Agrawal, Eric Balkanski , Vasilis Gkatzelis, Tingting Oua , and Xizhi Tan (Columbia University / Drexel University)

Table of Contents

- 1 Preliminaries
- 2 Minimizing egalitarian social cost
- 3 Minimizing utilitarian social cost

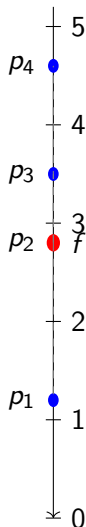
Warmup: Facility Location on the Line

- Given: n preferences $P = (p_1, p_2, \dots, p_n) \subset \mathbb{R}$
- Choose: Facility location $f \in \mathbb{R}$
- Constraints: Agent i incurs cost $|f - p_i|$
- Objective: Minimize a social cost function $C(f, P)$

Definition (Egalitarian /Utilitarian cost 1D)

Egalitarian: $C^e := \max_{p \in P} |f - p|$

Utilitarian: $C^u := \frac{1}{n} \sum_{p \in P} |f - p|$



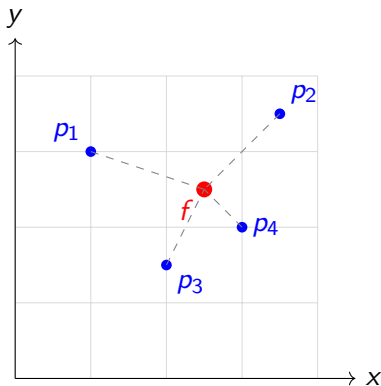
The problem

- Given: n preferences of agents
 $P = (p_1, p_2, \dots, p_n) \subset \mathbb{R}^2$
- Choose: Facility location $f \in \mathbb{R}^2$
- Constraints: Each agent suffers cost $d(f, p_i)$ (Euclidean distance)
- Objective: Minimize social cost function $C : (f, P) \rightarrow \mathbb{R}$

Definition

Egalitarian Cost $C^e := \max_{p \in P} d(f, p)$

Utilitarian Cost $C^u := \sum_{p \in P} \frac{d(f, p)}{n}$



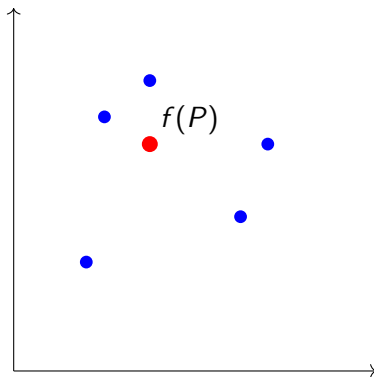
Agents may be strategic

Definition (**Strategyproofness**)

Mechanism $f : \mathbb{R}^{2n} \rightarrow \mathbb{R}^2$ is **strategyproof** iff for all instances $P \in \mathbb{R}^{2n}$, all $i \in [n]$, $p'_i \in \mathbb{R}^2$ it is the case that $d(p_i, f(P)) \leq d(p_i, f(P_{-i}, p'_i))$

In other words: It is the dominant strategy for every player to truthfully report preferences

Coordinatewise Median Mechanism (CM)

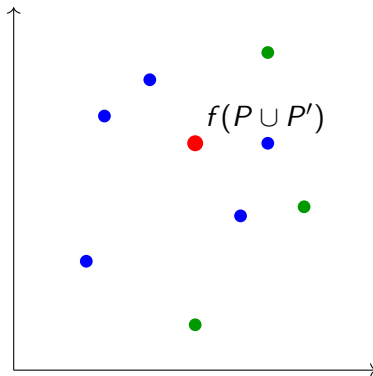


Definition (Coordinatewise Median Mechanism)

Given: Preferences $P = \{(x_1, y_1), \dots, (x_n, y_n)\} \subset \mathbb{R}^2$

Output: $\rightsquigarrow f(P) := (\text{Median}(x_1, \dots, x_n), \text{Median}(y_1, \dots, y_n))$

Generalized Coordinatewise Median (GCM)



Definition (Generalized CM)

Given: Preferences $P = \{(x_1, y_1), \dots, (x_n, y_n)\}$ and multiset $P' \subset \mathbb{R}^2$

Output: $\rightsquigarrow f(P) := \text{CM}(P \cup P')$

Median to the rescue

Theorem

*The **GCM** mechanism is deterministic, strategyproof and anonymous (= invariance under permutations of the agents).*

Learning-Augmented Mechanism Design

Now: We are given prediction $\hat{o}$ of optimal facility location $o(P)$

Definition (Consistency and Robustness)

Let C be some social cost function and f some mechanism.

f is α -**consistent**, if an α -approximation is achieved for $\hat{o} = o(P)$:

$$\max_P \left\{ \frac{C(f(P, \hat{o} = o(P)), P)}{C(o(P), P)} \right\} \leq \alpha$$

f is β -**robust**, if a β -approximation is achieved for any prediction $\hat{o}$:

$$\max_{P, \hat{o}} \left\{ \frac{C(f(P, \hat{o}), P)}{C(o(P), P)} \right\} \leq \beta$$

Table of Contents

- 1 Preliminaries
- 2 Minimizing egalitarian social cost
- 3 Minimizing utilitarian social cost

Recap of 1-dimensional case

Definition (**MinMaxP** mechanism)

Input: $P = (p_1, \dots, p_n) \in \mathbb{R}^n$, prediction $\hat{o} \in \mathbb{R}$, let
 $p_{min} := \min_{p \in P} (p)$, $p_{max} := \max_{p \in P} (p)$

$$\rightsquigarrow f(P, \hat{o}) = \begin{cases} \hat{o} & \text{if } \hat{o} \in [p_{min}, p_{max}] \\ p_{min} & \text{if } \hat{o} < p_{min} \\ p_{max} & \text{if } \hat{o} > p_{max} \end{cases}$$

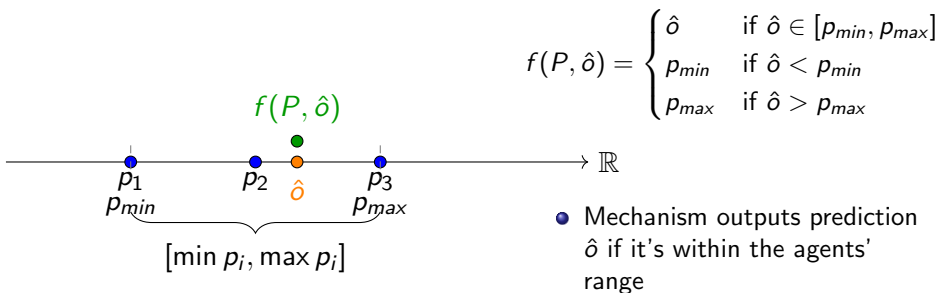
Equivalently:

$$f(P, \hat{o}) = CM(P \cup P')$$

where P' contains $n - 1$ copies of $\hat{o}$

The mechanism is strategyproof, 1-consistent and 2-robust as we have seen

Illustration of **MinMaxP** Mechanism (1D)



- Mechanism outputs prediction $\hat{o}$ if it's within the agents' range
- Otherwise, clips to nearest endpoint of the interval

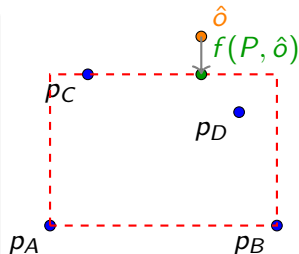
2-dimensional case: **Minimum Bounding Box** Mechanism

Definition (**Min Bounding Box** Mechanism)

Input: $P = \{(x_1, y_1), \dots, (x_n, y_n)\}$

Prediction: $\hat{o} = (\hat{x}, \hat{y})$

Output: $f(P, \hat{o}) = \text{MinMaxP}(x_1, \dots, x_n, \hat{x}),$
 $\text{MinMaxP}(y_1, \dots, y_n, \hat{y})$



2-dimensional case

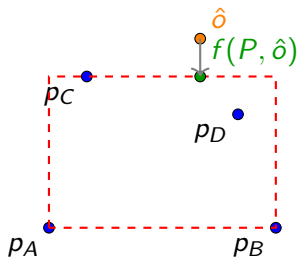
Definition (**Minimum Bounding Box** mechanism)

Input: $P = ((x_1, y_1), \dots, (x_n, y_n)) \in \mathbb{R}^{2n}$, prediction $\hat{o} = (\hat{x}, \hat{y}) \in \mathbb{R}^2$

$\rightsquigarrow f(P, \hat{o}) = (\text{MinMaxP}((x_1, \dots, x_n), \hat{x}), \text{MinMaxP}((y_1, \dots, y_n), \hat{y}))$

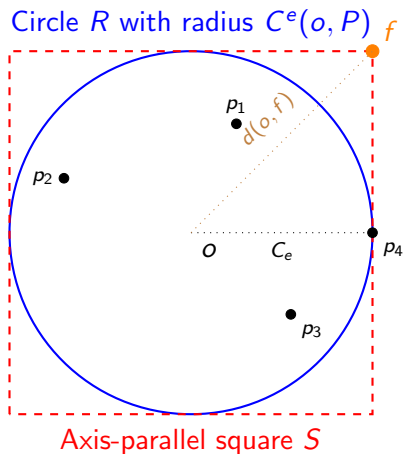
Theorem

*The **Minimum Bounding Box** mechanism is strategyproof, 1-consistent and $1 + \sqrt{2}$ robust for the egalitarian objective.*



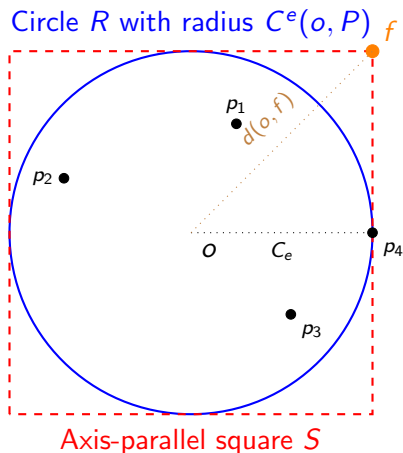
Proof $(1 + \sqrt{2})$ -Robustness

- R contains all points from P
 $\implies S$ contains all points from P



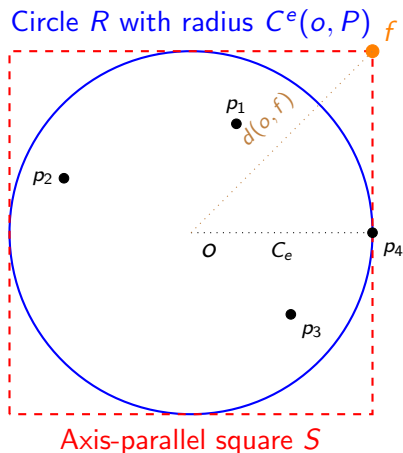
Proof $(1 + \sqrt{2})$ -Robustness

- R contains all points from P
 $\implies S$ contains all points from P
- $\implies$ All bounds for S also hold for the minimum axis-parallel bounding box



Proof $(1 + \sqrt{2})$ -Robustness

- R contains all points from P
 $\implies S$ contains all points from P
- $\implies$ All bounds for S also hold for the minimum axis-parallel bounding box
- $d(o, f) \leq \sqrt{2}C^e(o, p)$



Theorem

There is no deterministic, strategyproof, and anonymous mechanism that is $(2 - \varepsilon)$ -consistent and $(1 + \sqrt{2} - \varepsilon)$ robust with respect to the egalitarian objective for any $\varepsilon > 0$.

Theorem

There is no deterministic, strategyproof, and anonymous mechanism that is $(2 - \varepsilon)$ -consistent and $(1 + \sqrt{2} - \varepsilon)$ robust with respect to the egalitarian objective for any $\varepsilon > 0$.

Lemma (Peters et al. (2013))

*Any deterministic, strategyproof, anonymous, unanimous mechanism solving the problem is equivalent to **GCM** with $n - 1$ phantom points.*

Theorem

There is no deterministic, strategyproof, and anonymous mechanism that is $(2 - \varepsilon)$ -consistent and $(1 + \sqrt{2} - \varepsilon)$ robust with respect to the egalitarian objective for any $\varepsilon > 0$.

Lemma (Peters et al. (2013))

*Any deterministic, strategyproof, anonymous, unanimous mechanism solving the problem is equivalent to **GCM** with $n - 1$ phantom points.*

- 1. For consistency better than 2, we have to place all $n - 1$ phantom points on $\hat{o}$
- 2. If we place all $n - 1$ phantom points on $\hat{o}$, robustness is at least $1 + \sqrt{2}$

We have to place all phantoms on the prediction

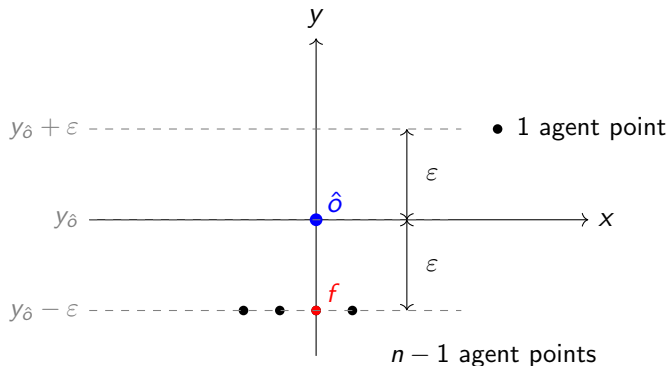
- Assume atleast one of the $n - 1$ phantoms q' is not on $\hat{o}$. Wlog $y_{q'} < y_{\hat{o}}$

We have to place all phantoms on the prediction

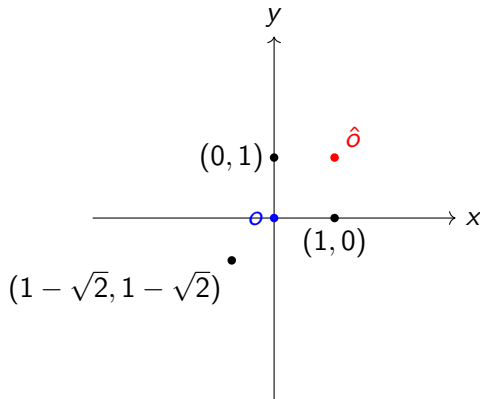
- Assume atleast one of the $n - 1$ phantoms q' is not on $\hat{o}$. Wlog $y_{q'} < y_{\hat{o}}$
- Choose $\bar{y} = \max_{q \in P: y_q < y_{\hat{o}}} y_q$ and $\varepsilon = y_{q'} - \bar{y}$

We have to place all phantoms on the prediction

- Assume atleast one of the $n - 1$ phantoms q' is not on $\hat{o}$. Wlog $y_{q'} < y_{\hat{o}}$
- Choose $\bar{y} = \max_{q \in P: y_q < y_{\hat{o}}} y_q$ and $\varepsilon = y_{q'} - \bar{y}$

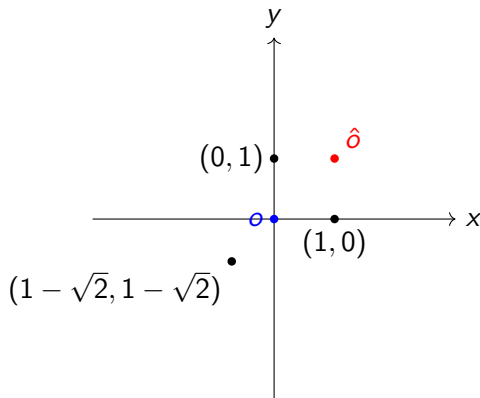


With $n-1$ phantom points on prediction , robustness at least $1 + \sqrt{2}$



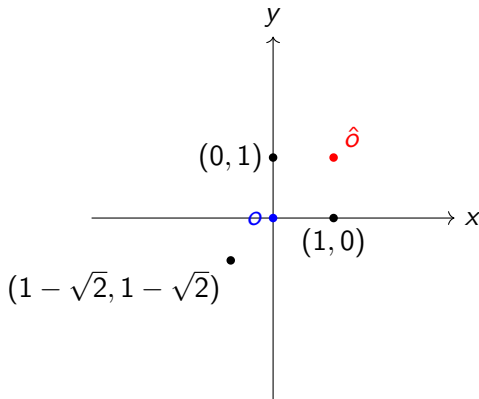
- optimum at o with cost 1

With $n-1$ phantom points on prediction , robustness at least $1 + \sqrt{2}$



- optimum at o with cost 1
- We have 2 phantom points at $\hat{o}$, hence 3 points with $x = 1$, 3 point with $y = 1$

With $n-1$ phantom points on prediction, robustness at least $1 + \sqrt{2}$



- optimum at o with cost 1
- We have 2 phantom points at $\hat{o}$, hence 3 points with $x = 1$, 3 point with $y = 1$
- f is placed at $(1, 1)$ and cost is $1 + \sqrt{2}$

2-dimensional case

Definition (Prediction Error)

Let P be an instance, $\hat{o}$ a prediction and $o(P)$ the optimal facility location. Define the prediction error

$$\eta(\hat{o}, P) := \frac{d(\hat{o}, o(P))}{C(o(P), P)}$$

as the distance to the optimal location $o(P)$ normalized by the optimal social cost

Smoothness

Lemma

Let P be an instance of the problem. For any two predictions $\hat{o}$ and $\tilde{o}$, the returned facility locations $f(P, \hat{o})$ and $f(P, \tilde{o})$ by the **Minimum Bounding Box** mechanism satisfy

$$d(f(P, \hat{o}), f(P, \tilde{o})) \leq d(\hat{o}, \tilde{o})$$

Smoothness

Lemma

Let P be an instance of the problem. For any two predictions $\hat{o}$ and $\tilde{o}$, the returned facility locations $f(P, \hat{o})$ and $f(P, \tilde{o})$ by the **Minimum Bounding Box** mechanism satisfy

$$d(f(P, \hat{o}), f(P, \tilde{o})) \leq d(\hat{o}, \tilde{o})$$

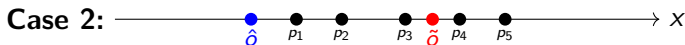
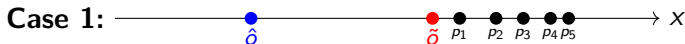


Smoothness

Lemma

Let P be an instance of the problem. For any two predictions $\hat{o}$ and $\tilde{o}$, the returned facility locations $f(P, \hat{o})$ and $f(P, \tilde{o})$ by the **Minimum Bounding Box** mechanism satisfy

$$d(f(P, \hat{o}), f(P, \tilde{o})) \leq d(\hat{o}, \tilde{o})$$



Theorem

The **Minimum Bounding Box** mechanism achieves a $\min\{1 + \eta, 1 + \sqrt{2}\}$ approximation

Proof.

- By definition $d(\hat{o}, o) = \eta C^e(o, P)$
- By previous lemma $d(f(P, \hat{o}), f(P, o)) \leq d(\hat{o}, o), \eta C^e(o, P)$



Theorem

The **Minimum Bounding Box** mechanism achieves a $\min\{1 + \eta, 1 + \sqrt{2}\}$ approximation

Proof.

- By definition $d(\hat{o}, o) = \eta C^e(o, P)$
- By previous lemma $d(f(P, \hat{o}), f(P, o)) \leq d(\hat{o}, o), \eta C^e(o, P)$

$$C^e(f(P, \hat{o}), P) = \max_{i \in [n]} d(p_i, f(P, \hat{o}))$$



Theorem

The **Minimum Bounding Box** mechanism achieves a $\min\{1 + \eta, 1 + \sqrt{2}\}$ approximation

Proof.

- By definition $d(\hat{o}, o) = \eta C^e(o, P)$
- By previous lemma $d(f(P, \hat{o}), f(P, o)) \leq d(\hat{o}, o), \eta C^e(o, P)$

$$\begin{aligned} C^e(f(P, \hat{o}), P) &= \max_{i \in [n]} d(p_i, f(P, \hat{o})) \\ &\leq \max_{i \in [n]} [d(p_i, f(P, o)) + d(f(P, o), f(P, \hat{o}))] \end{aligned}$$



Theorem

The **Minimum Bounding Box** mechanism achieves a $\min\{1 + \eta, 1 + \sqrt{2}\}$ approximation

Proof.

- By definition $d(\hat{o}, o) = \eta C^e(o, P)$
- By previous lemma $d(f(P, \hat{o}), f(P, o)) \leq d(\hat{o}, o), \eta C^e(o, P)$

$$\begin{aligned} C^e(f(P, \hat{o}), P) &= \max_{i \in [n]} d(p_i, f(P, \hat{o})) \\ &\leq \max_{i \in [n]} [d(p_i, f(P, o)) + d(f(P, o), f(P, \hat{o}))] \\ &\leq \max_{i \in [n]} [d(p_i, o) + \eta \cdot C^e(f(P, o), P)] \end{aligned}$$



Theorem

The **Minimum Bounding Box** mechanism achieves a $\min\{1 + \eta, 1 + \sqrt{2}\}$ approximation

Proof.

- By definition $d(\hat{o}, o) = \eta C^e(o, P)$
- By previous lemma $d(f(P, \hat{o}), f(P, o)) \leq d(\hat{o}, o), \eta C^e(o, P)$

$$\begin{aligned} C^e(f(P, \hat{o}), P) &= \max_{i \in [n]} d(p_i, f(P, \hat{o})) \\ &\leq \max_{i \in [n]} [d(p_i, f(P, o)) + d(f(P, o), f(P, \hat{o}))] \\ &\leq \max_{i \in [n]} [d(p_i, o) + \eta \cdot C^e(f(P, o), P)] \\ &\leq (1 + \eta) \cdot C^e(f(P, o), P) \end{aligned}$$



Table of Contents

- 1 Preliminaries
- 2 Minimizing egalitarian social cost
- 3 Minimizing utilitarian social cost

The plan

- 1-Dimensional case: Median is optimal and strategyproof
- 2-Dimensional case: CM mechanism is a $\sqrt{2}$ -approximation and no deterministic, strategyproof and anonymous algorithm can provide a better approximation [Meir 2019]

The plan

- 1-Dimensional case: Median is optimal and strategyproof
- 2-Dimensional case: CM mechanism is a $\sqrt{2}$ -approximation and no deterministic, strategyproof and anonymous algorithm can provide a better approximation [Meir 2019]

What can we achieve if we add predictions into the mix?

The plan

- 1-Dimensional case: Median is optimal and strategyproof
- 2-Dimensional case: CM mechanism is a $\sqrt{2}$ -approximation and no deterministic, strategyproof and anonymous algorithm can provide a better approximation [Meir 2019]

What can we achieve if we add predictions into the mix?

$\implies \sqrt{2c^2 + 2}/(1 + c)$ -consistency and $\sqrt{2c^2 + 2}/(1 - c)$ -robustness

where $c \in [0, 1)$

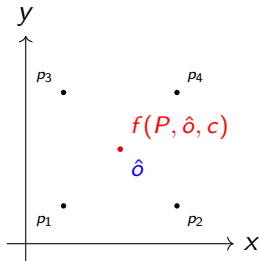
The Mechanism

Definition (Coordinate Median with Predictions (**CMP**) mechanism)

Input: Locations $P \in \mathbb{R}^{2n}$, prediction $\hat{o} \in \mathbb{R}^2$, confidence value $c \in [0, 1)$

$$\rightsquigarrow f(P, \hat{o}, c) = CM(P \cup P')$$

where P' contains cn copies of $\hat{o}$



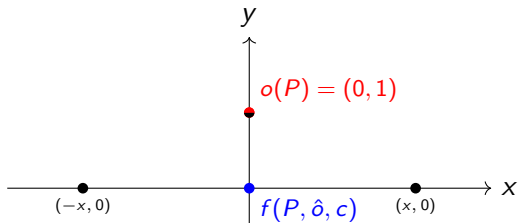
The Mechanism

Definition (Clusters-and-Opt-on-Axes Instances)

Let $c \in [0, 1)$. Define $\mathcal{P}$ to be the class of all instances $(P, \hat{o})$ such that

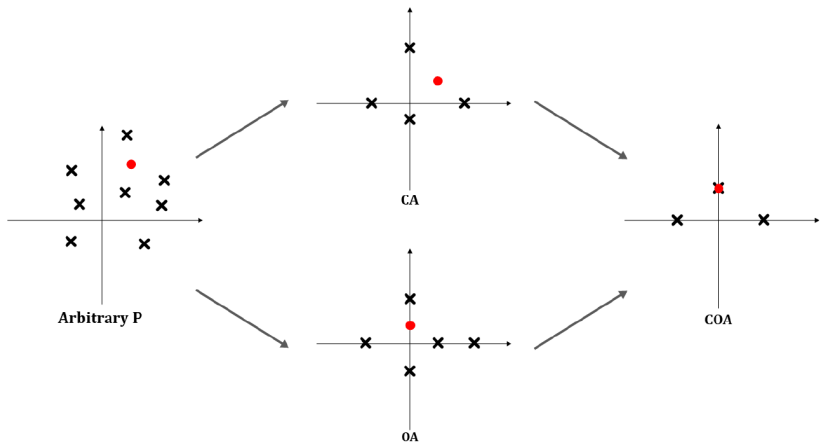
- $f(P, \hat{o}, c) = (0, 0)$
- $o(P) = (0, 1)$
- $\forall p \in P : p \in \{(0, 1), (x, 0), (-x, 0)\}$ for some $x \in \mathbb{R}_{\geq 0}$

Let $\mathcal{P}_{coa}^C(c) \subset \mathcal{P}$ where $\hat{o} = o(P)$ and $\mathcal{P}_{coa}^R \subset \mathcal{P}$ where $\hat{o} = (0, 0)$



Lemma (COA instances always contain a worst-case instance)

Let $r(P, \hat{o}, c)$ be the achieved approx-ratio. For any $c \in [0, 1]$, the CMP-mechanism is α -consistent and β -robust where
 $\alpha = \max_{P \in \mathcal{P}_{coa}^C} r(P, \hat{o} = o(P), c)$ and $\beta = \max_{P \in \mathcal{P}_{coa}^R} r(P, \hat{o} = (0, 0), c)$

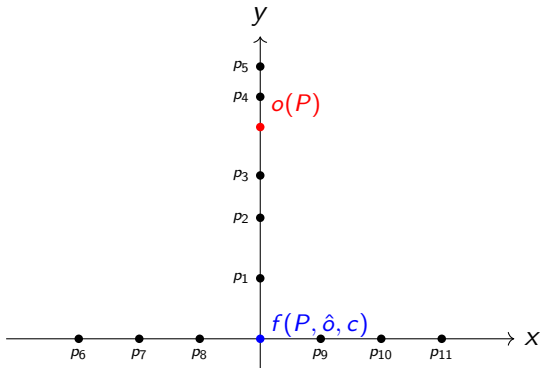


Definition (Optimal-on-Axes family)

For some c , $\hat{o}$, define $\mathcal{P}_{oa}$ be the family of multisets P such that

- $f(P, \hat{o}, c) = (0, 0)$
- $x_o(P) = 0, y_o(P) > 0$
- $\forall p \in P : p \in A_x \cup A_y$

Let $\mathcal{P}_{oa}^C(c)$ be the family when $\hat{o} = o(P)$ and $\mathcal{P}_{oa}^R(c)$ when $\hat{o} = (0, 0)$

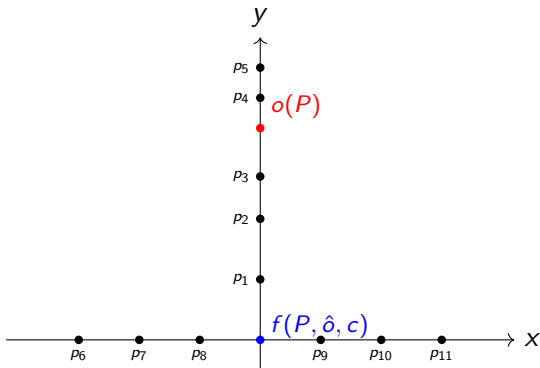


Lemma (Convert OA instance to COA or make it strictly worse)

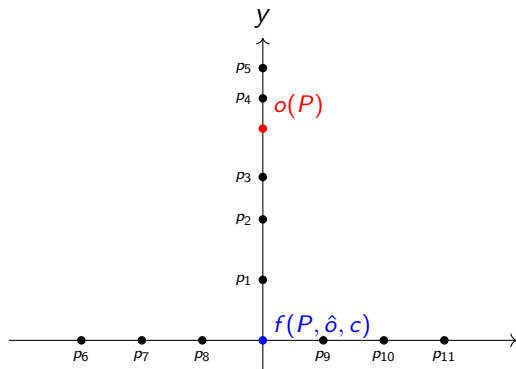
For any $c \in [0, 1)$, $P \in \mathcal{P}_{oa}^C$

- There is some Q such that $r(Q, o(Q), c) > r(P, o(P), c)$
- OR there is some $Q \in \mathcal{P}_{coa}^C$ such that $r(Q, o(Q), c) \geq r(P, o(P), c)$

(This holds analogously for the robustness case)

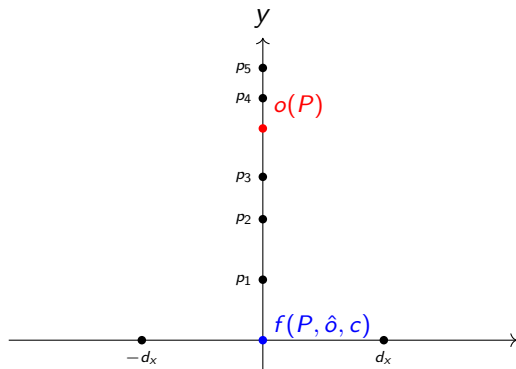


OA to COA



OA to COA

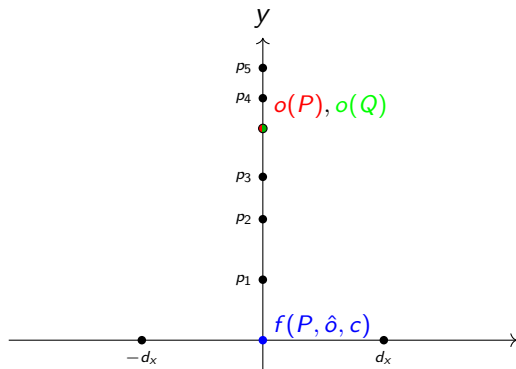
Step 1: Move x-Points to $d_x, -d_x$
where $d_x := \text{average distance to } f$.



OA to COA

Step 1: Move x-Points to $d_x, -d_x$
where $d_x := \text{average distance to } f$.

Step 2: In the new instance Q we
have $o(P) = o(Q)$.

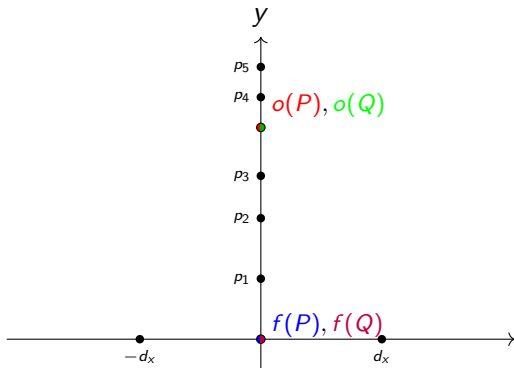


OA to COA

Step 1: Move x-Points to $d_x, -d_x$
where $d_x := \text{average distance to } f$.

Step 2: In the new instance Q we
have $o(P) = o(Q)$.

Step 3: In the new instance Q we
have $f(P) = f(Q)$ hence social cost
doesn't change, optimal cost weakly
improves $\implies r(Q) \geq r(P)$.



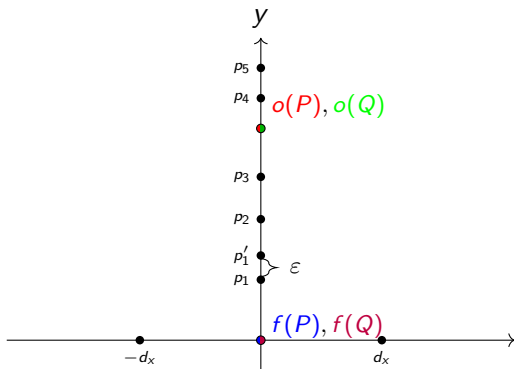
OA to COA

Step 1: Move x-Points to $d_x, -d_x$
where $d_x := \text{average distance to } f$.

Step 2: In the new instance Q we
have $o(P) = o(Q)$.

Step 3: In the new instance Q we
have $f(P) = f(Q)$ hence social cost
doesn't change, optimal cost weakly
improves $\implies r(Q) \geq r(P)$.

Step 4: If not all y-points are on
 $o(Q)$, move one towards it to obtain
a strictly worse instance



Lemma

Let

$$\alpha := \max_{P \in \mathcal{P}_{oa}^C \cup \mathcal{P}_{ca}^C} r(P, o(P), c)$$

and

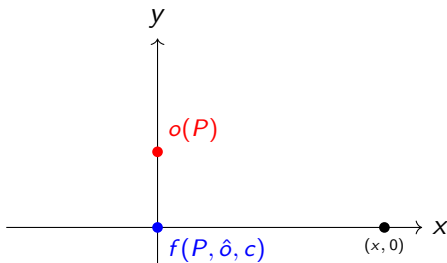
$$\beta := \max_{P \in \mathcal{P}_{oa}^R \cup \mathcal{P}_{ca}^R} r(P, (0, 0), c)$$

then the **CMP** mechanism is α -consistent and β -robust

Consistency and Robustness

Lemma

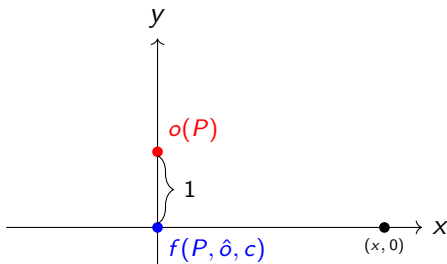
The **CMP** mechanism is with parameter $c \in [0, 1)$ is $\frac{\sqrt{2c^2+2}}{1+c}$ -consistent and $\frac{\sqrt{2c^2+2}}{1-c}$ -robust



Consistency and Robustness

Lemma

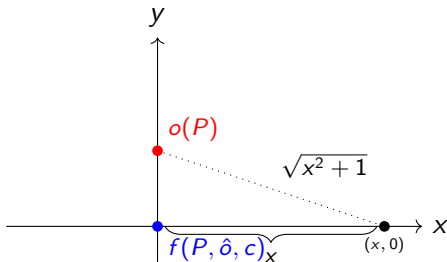
The **CMP** mechanism is with parameter $c \in [0, 1)$ is $\frac{\sqrt{2c^2+2}}{1+c}$ -consistent and $\frac{\sqrt{2c^2+2}}{1-c}$ -robust



Consistency and Robustness

Lemma

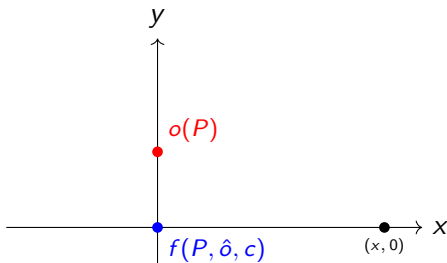
The **CMP** mechanism is with parameter $c \in [0, 1)$ is $\frac{\sqrt{2c^2+2}}{1+c}$ -consistent and $\frac{\sqrt{2c^2+2}}{1-c}$ -robust



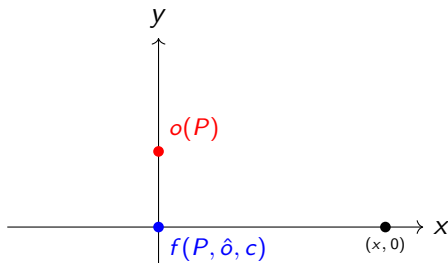
Consistency and Robustness

Lemma

The **CMP** mechanism is with parameter $c \in [0, 1)$ is $\frac{\sqrt{2c^2+2}}{1+c}$ -consistent and $\frac{\sqrt{2c^2+2}}{1-c}$ -robust

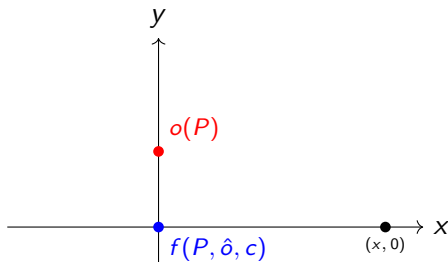


Consistency proof



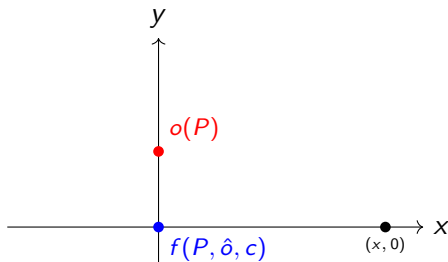
- $z \leq \frac{1-c}{2}n$ where z is number of agent points with $y = 1$

Consistency proof



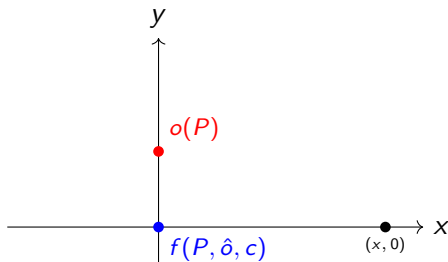
- $z \leq \frac{1-c}{2}n$ where z is number of agent points with $y = 1$
- $\alpha = \frac{C^u(f(P, \hat{o}=o(P), c), P)}{C^u(o(P), P)} = \frac{\frac{1+c}{2}nx + \frac{1-c}{2}n}{\frac{1+c}{2}n\sqrt{x^2+1}}$

Consistency proof



- $z \leq \frac{1-c}{2}n$ where z is number of agent points with $y = 1$
- $\alpha = \frac{C^u(f(P, \hat{o}=o(P), c), P)}{C^u(o(P), P)} = \frac{\frac{1+c}{2}nx + \frac{1-c}{2}n}{\frac{1+c}{2}n\sqrt{x^2+1}}$
- $\alpha' = \frac{1+c-(1-c)x}{(1+c)(1+x^2)\sqrt{1+x^2}} = 0 \implies x = \frac{1+c}{1-c}$

Consistency proof



- $z \leq \frac{1-c}{2}n$ where z is number of agent points with $y = 1$
- $\alpha = \frac{C^u(f(P, \hat{o}=o(P), c), P)}{C^u(o(P), P)} = \frac{\frac{1+c}{2}nx + \frac{1-c}{2}n}{\frac{1+c}{2}n\sqrt{x^2+1}}$
- $\alpha' = \frac{1+c-(1-c)x}{(1+c)(1+x^2)\sqrt{1+x^2}} = 0 \implies x = \frac{1+c}{1-c}$
- $\implies \alpha = \frac{\sqrt{2c^2+2}}{1+c}$

Optimality

Lemma

Any deterministic, strategyproof and anonymous mechanism with $\frac{\sqrt{2c^2+2}}{1+c}$ -consistency has a robustness no better than $\frac{\sqrt{2c^2+2}}{1+c}$

Approximation in relation to the prediction error

Lemma

Let P be an instance of the problem. For any two prediction $\hat{o}$ and $\tilde{o}$ the returned facility locations $f(P, \hat{o}, c)$ and $f(P, \tilde{o}, c)$ by the **CMP** mechanism satisfy

$$d(f(P, \hat{o}, c), f(P, \tilde{o}, c)) \leq d(\hat{o}, \tilde{o})$$

Approximation in relation to the prediction error

Lemma

Let P be an instance of the problem. For any two prediction $\hat{o}$ and $\tilde{o}$ the returned facility locations $f(P, \hat{o}, c)$ and $f(P, \tilde{o}, c)$ by the **CMP** mechanism satisfy

$$d(f(P, \hat{o}, c), f(P, \tilde{o}, c)) \leq d(\hat{o}, \tilde{o})$$

Theorem

The **CMP** mechanism achieves a $\min\left\{\frac{\sqrt{2c^2+2}}{1+c} + \eta, \frac{\sqrt{2c^2+2}}{1-c}\right\}$ approximation



Agrawal, Priyank et al. (2022). “Learning-Augmented Mechanism Design: Leveraging Predictions for Facility Location”. In: *Proceedings of the 23rd ACM Conference on Economics and Computation*. EC '22. Boulder, CO, USA: Association for Computing Machinery, pp. 497–528. ISBN: 9781450391504. DOI: 10.1145/3490486.3538306. URL: <https://doi.org/10.1145/3490486.3538306>.



Meir, Reshef (2019). “Strategyproof Facility Location for Three Agents on a Circle”. In: *Algorithmic Game Theory: 12th International Symposium, SAGT 2019, Athens, Greece, September 30 – October 3, 2019, Proceedings*. Athens, Greece: Springer-Verlag, pp. 18–33. ISBN: 978-3-030-30472-0.