

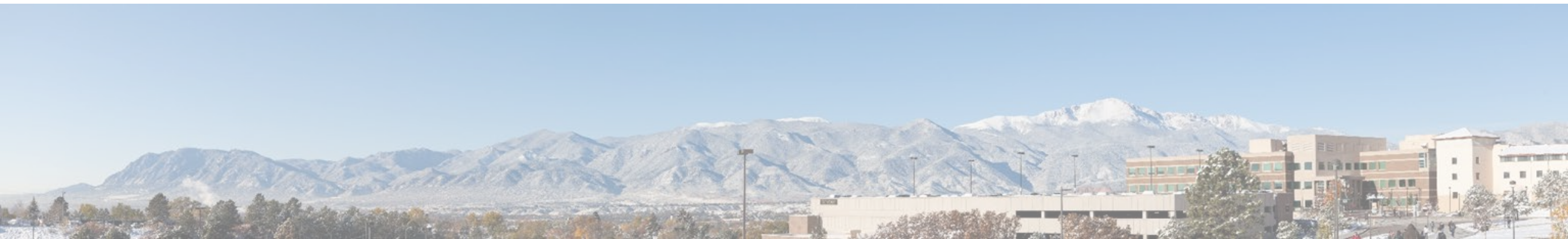
A Little Anarchy Never Hurt Anyone: Beyond Worst-Case Equilibria in Games

KTH, Stockholm: Game On!
June 9, 2026

Philip N. Brown

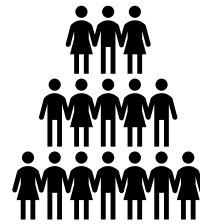
Associate Professor, Computer Science
University of Colorado, Colorado Springs

Visiting Professor, DISMA
Politecnico di Torino



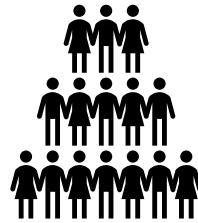
Agenda: Mathematics for Multiagent Decision-making

Social
Systems

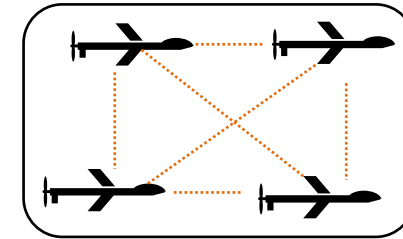


Policy

Agenda: Mathematics for Multiagent Decision-making

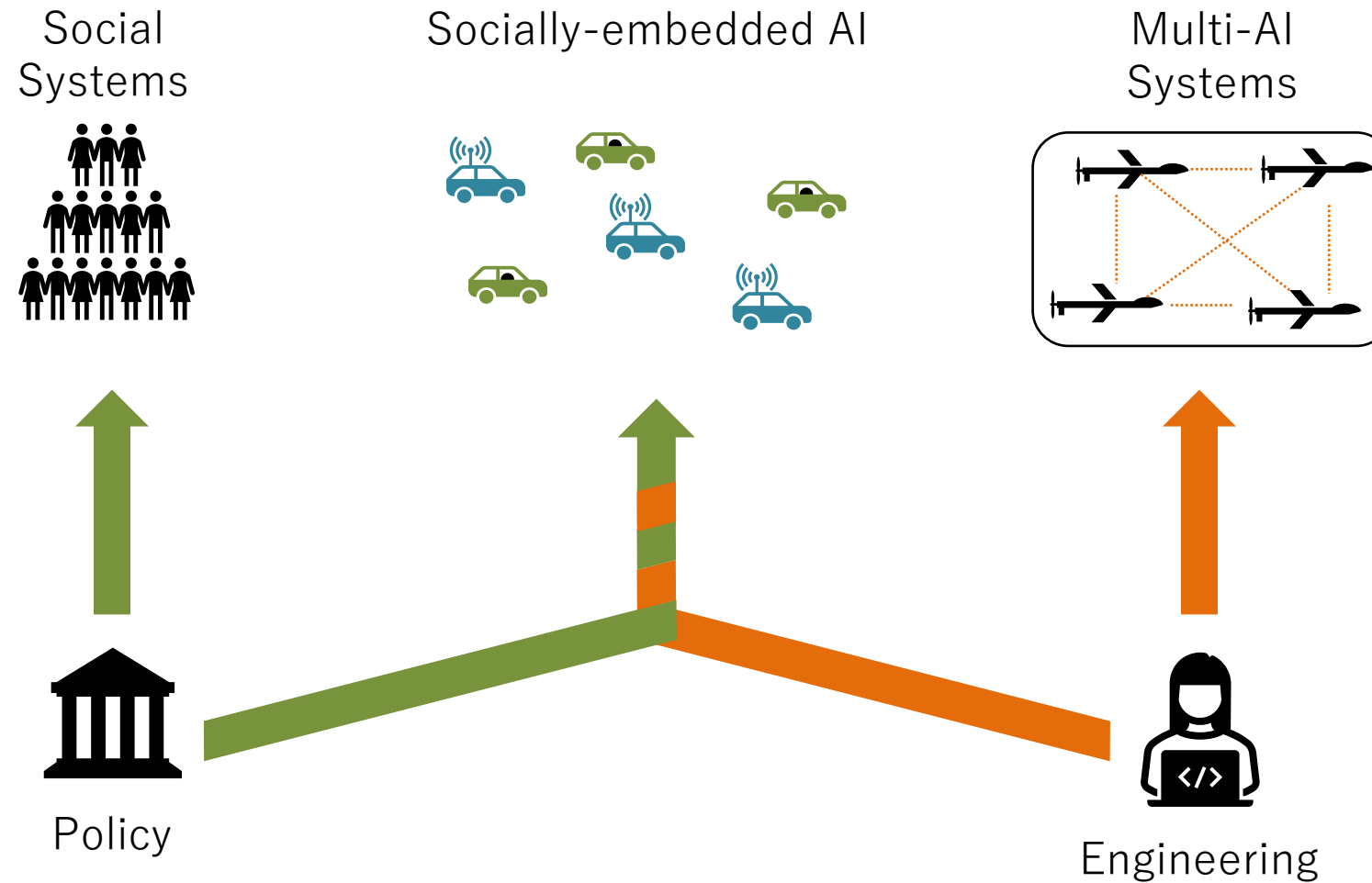
Social
Systems

Policy

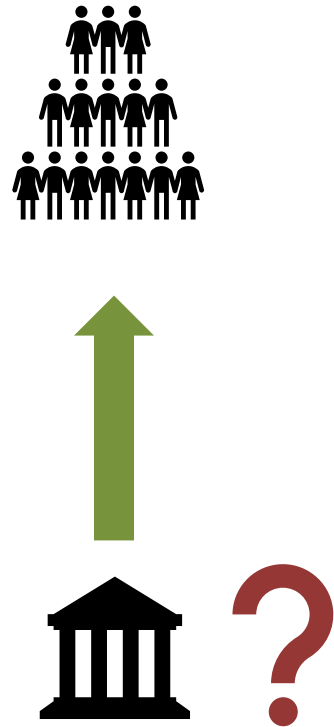
Multi-AI
Systems

Engineering

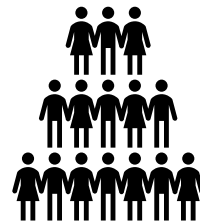
Agenda: Mathematics for Multiagent Decision-making



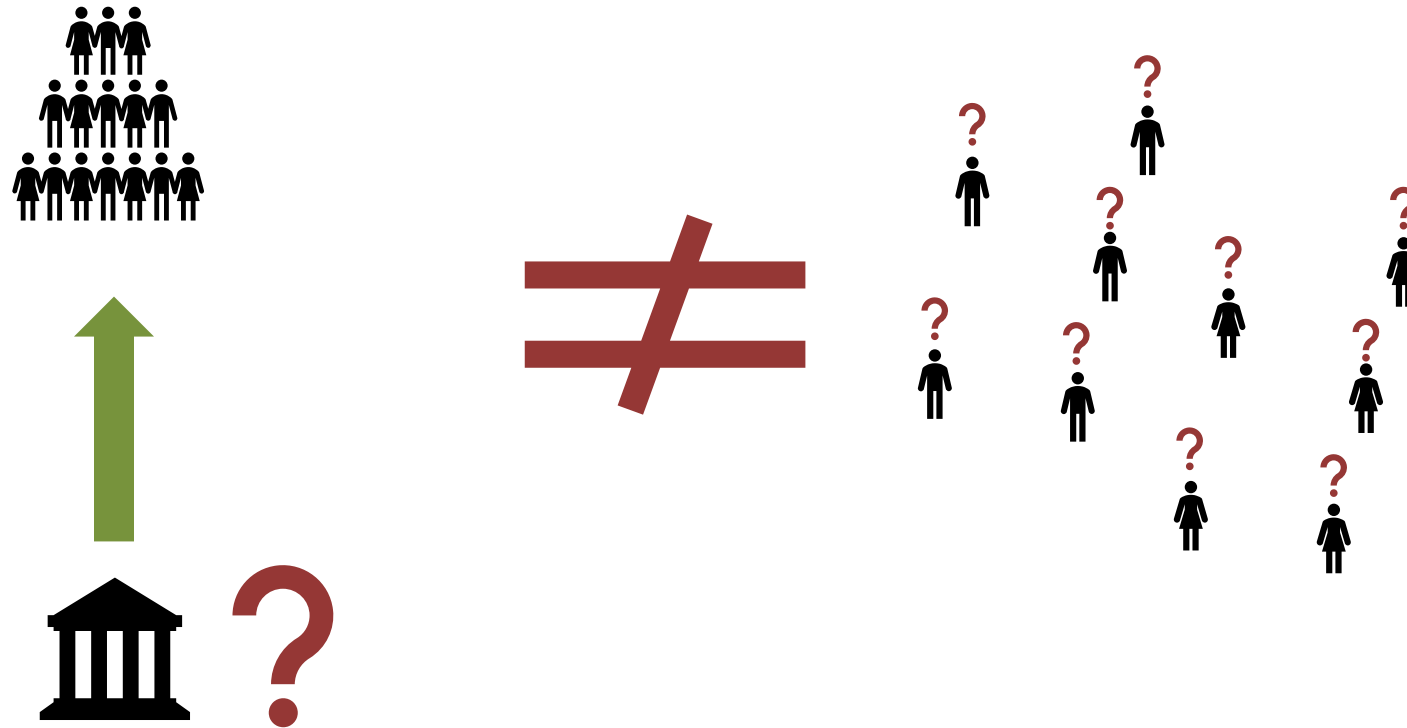
Q: How good is Multiagent Coordination?



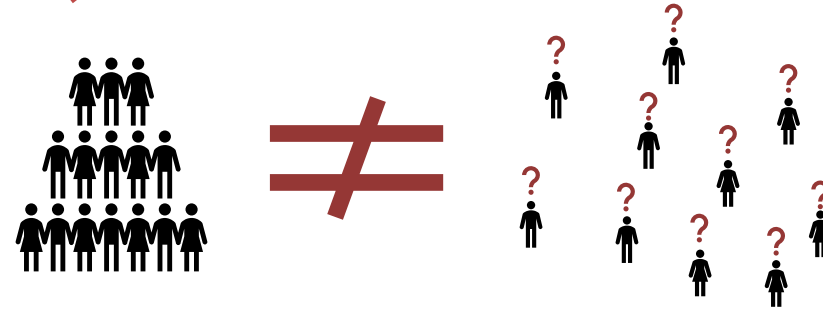
Q: How good is Multiagent Coordination?



Q: How good is Multiagent Coordination?



~~bad~~
Q: How ~~good~~ is Multiagent Coordination?



Prisoner's Dilemma

Tragedy of the Commons

Free-Rider Problems

Inefficiency of Equilibria

Price of Anarchy

Price of Anarchy

Agenda [Papadimitriou, STOC'01]:

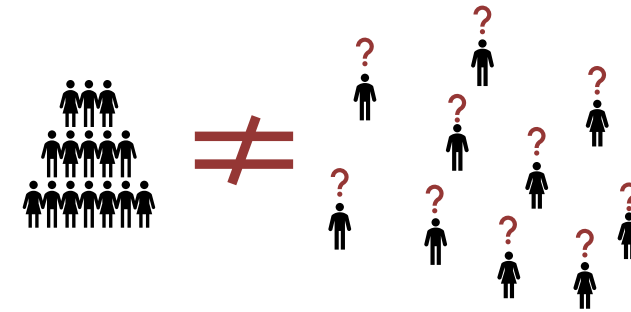
- Class of games $\mathcal{G}$
- Emergent behavior $E(G)$, $G \in \mathcal{G}$
- Performance metric $P(e)$, $e \in E(G)$

$$\text{PoA}(\mathcal{G}) = \max_{G \in \mathcal{G}} \max_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \geq 1$$

(if P is a **cost**)

$$\text{PoA}(\mathcal{G}) = \min_{G \in \mathcal{G}} \min_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \leq 1$$

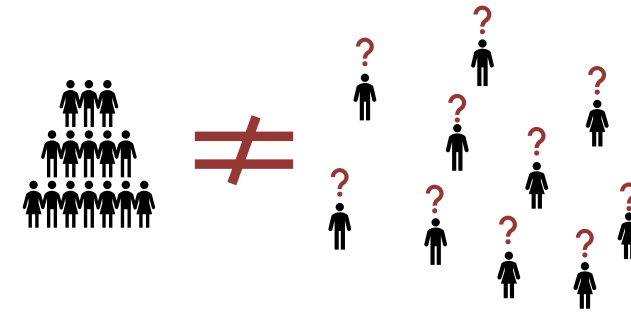
(if P is a **benefit**)



Price of Anarchy

Agenda [Papadimitriou, STOC'01]:

- Class of games $\mathcal{G}$
- Emergent behavior $E(G)$, $G \in \mathcal{G}$
- Performance metric $P(e)$, $e \in E(G)$



$$\text{PoA}(\mathcal{G}) = \max_{G \in \mathcal{G}} \max_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \geq 1$$

(if P is a **cost**)

$$\text{PoA}(\mathcal{G}) = \min_{G \in \mathcal{G}} \min_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \leq 1$$

(if P is a **benefit**)

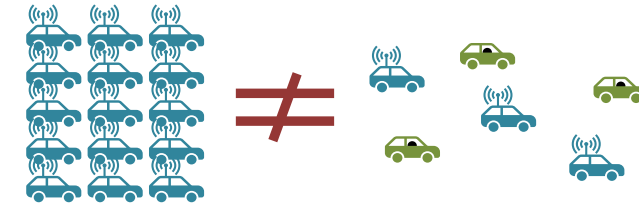
Sample of PoA Bounds

1 st -price auction (revenue) [Hartline, Hoy, Taggart EC'14]	3.16	}	Cost-Minimization Games
Finite-player routing [Christodoulou, Koutsoupias STOC'05]	2.5		
Generalized 2 nd -Price Auction [Leme, Tardos FOCS'10]	1.618		
Continuous-player routing [Roughgarden, Tardos STOC'00]	1.333		
1 st -price auction (welfare) [Roughgarden, Syrgkanis, Tardos JAIR'17]	0.63	}	Utility-Maximization Games
Submodular resource allocation [Vetta STOC'01]	0.5		

Price of Anarchy

Agenda [Papadimitriou, STOC'01]:

- Class of games $\mathcal{G}$
- Emergent behavior $E(G)$, $G \in \mathcal{G}$
- Performance metric $P(e)$, $e \in E(G)$



$$\text{PoA}(\mathcal{G}) = \max_{G \in \mathcal{G}} \max_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \geq 1$$

(if P is a **cost**)

$$\text{PoA}(\mathcal{G}) = \min_{G \in \mathcal{G}} \min_{e \in E(G)} \frac{P(e)}{P(\text{OPT})} \leq 1$$

(if P is a **benefit**)

Sample of PoA Bounds

1 st -price auction (revenue) [Hartline, Hoy, Taggart EC'14]	3.16	}	Cost-Minimization Games
Finite-player routing [Christodoulou, Koutsoupias STOC'05]	2.5		
Generalized 2 nd -Price Auction [Leme, Tardos FOCS'10]	1.618		
Continuous-player routing [Roughgarden, Tardos STOC'00]	1.333		
1 st -price auction (welfare) [Roughgarden, Syrgkanis, Tardos JAIR'17]	0.63	}	Utility-Maximization Games
Submodular resource allocation [Vetta STOC'01]	0.5		

Optimizing Price of Anarchy

Idea: Machine agents can be designed!



Optimizing Price of Anarchy



Idea: Machine agents can be designed!

Agenda [Marden, Shamma, Brown, Chandan, Paccagnan, Ferguson, ...]:

- *Parameterized* Class of games $\mathcal{G}(\theta)$
- Emergent behavior $E(G(\theta))$, $G \in \mathcal{G}(\theta)$
- Performance metric $P(e; \theta)$, $e \in E(G)$

$$\min_{\theta} \text{PoA}(\mathcal{G}(\theta))$$

(if P is a **cost**)

$$\max_{\theta} \text{PoA}(\mathcal{G}(\theta))$$

(if P is a **benefit**)

Select θ to push $\text{PoA}(\mathcal{G}(\theta))$ towards 1

Optimizing Price of Anarchy



Idea: Machine agents can be designed!

Agenda [Marden, Shamma, Brown, Chandan, Paccagnan, Ferguson, ...]:

- *Parameterized* Class of games $\mathcal{G}(\theta)$
- Emergent behavior $E(G(\theta))$, $G \in \mathcal{G}(\theta)$
- Performance metric $P(e; \theta)$, $e \in E(G)$

$$\min_{\theta} \text{PoA}(\mathcal{G}(\theta))$$

(if P is a **cost**)

$$\max_{\theta} \text{PoA}(\mathcal{G}(\theta))$$

(if P is a **benefit**)

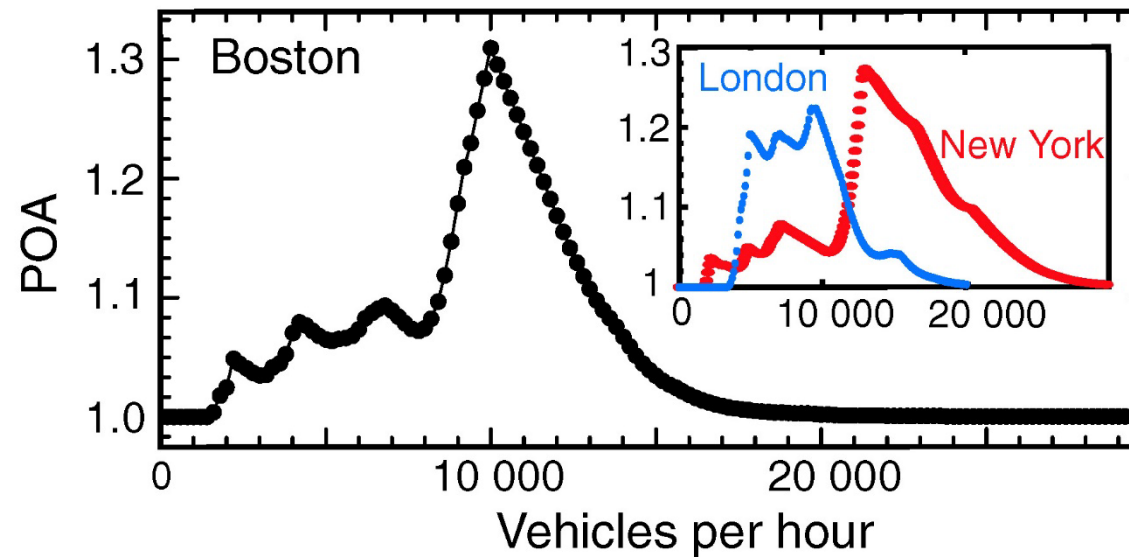
Select θ to push $\text{PoA}(\mathcal{G}(\theta))$ towards 1

2.5 $\rightarrow$ ~ 2
atomic congestion

0.5 $\rightarrow$ ~ 0.63
weighted set cover

But: is **bad PoA** predictive of **bad performance**?

Simulated traffic on real urban networks:



Most flow conditions: $\text{PoA} < 1.1$ (10% worse than optimal)

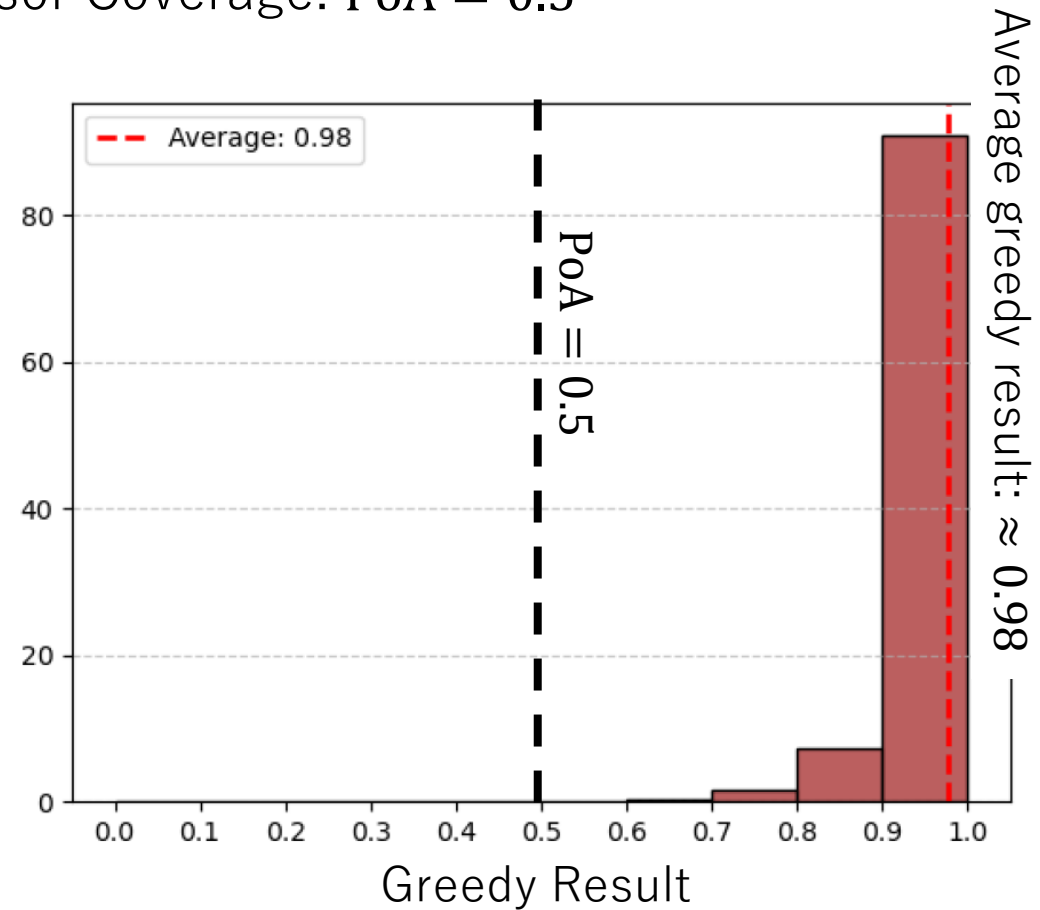
Theoretical upper bound: $\text{PoA} \approx 3.5$ (**250%** worse than optimal)

Youn et. al, "Price of Anarchy in Transportation Networks," Phys. Rev. Lett. 101, 128701, 2008

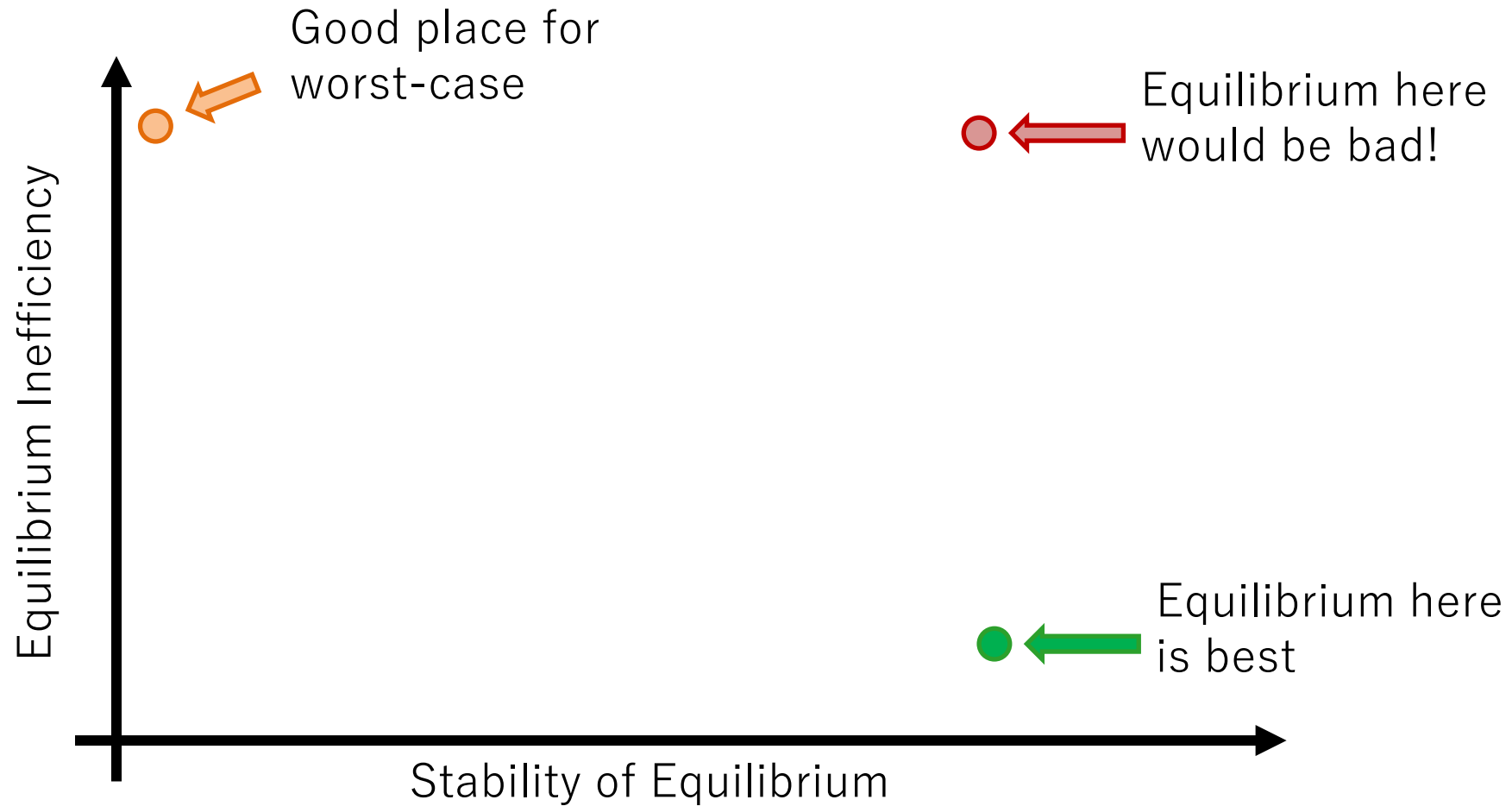
Multiagent Autonomous Sensor Coverage: $\text{PoA} = 0.5$

Agenda:

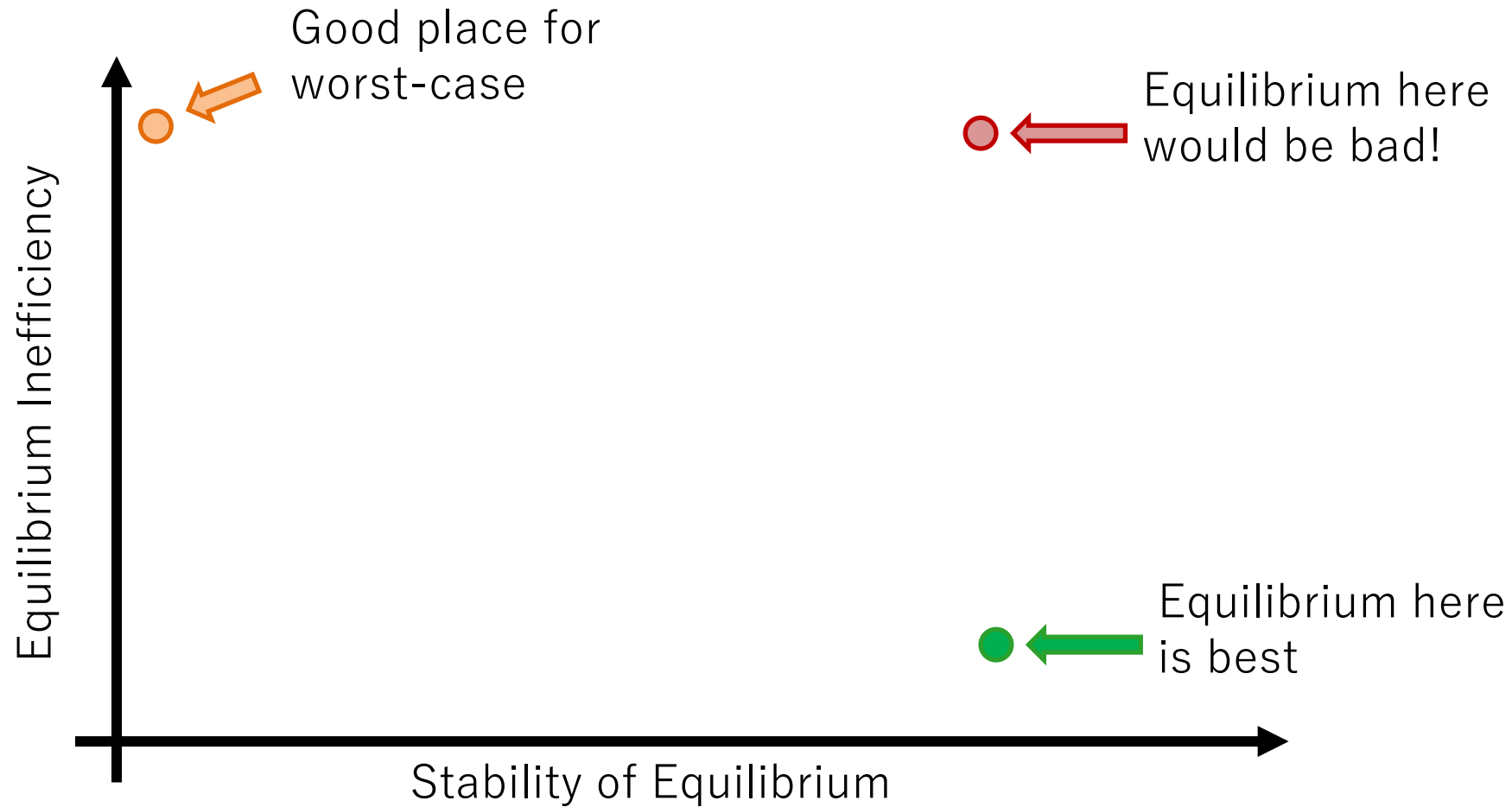
- Random game instances
- Execute greedy algorithm
(terminates at PNE)
- Histogram of results:



Guanchu He, Colton Hill, Joshua Seaton, **PNB**, “Randomized Algorithms and Neural Networks for Communication-Free Multiagent Singleton Set Cover,” *Games* 2026

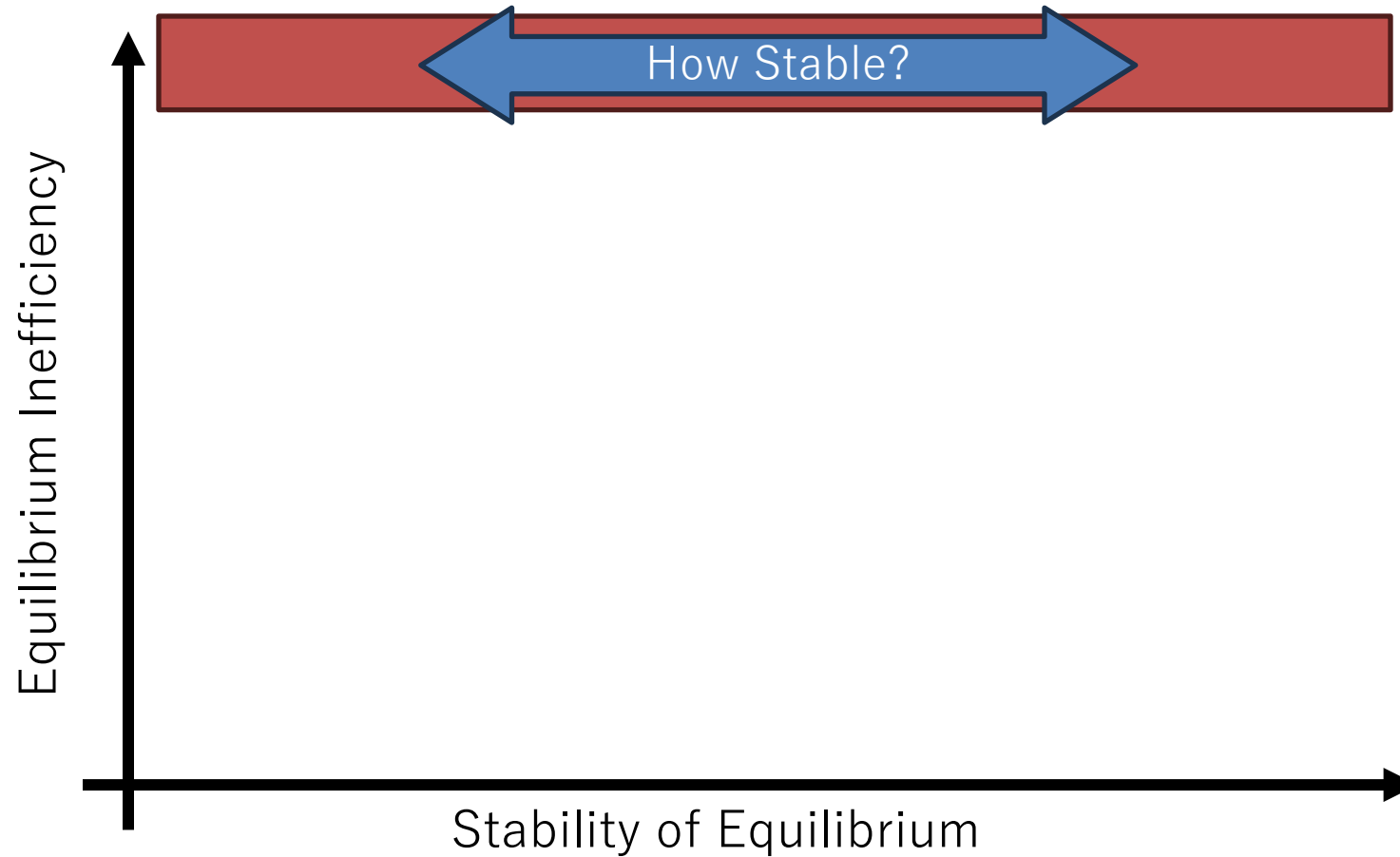


“Stability” $\approx$ agent satisfaction with equilibrium



Main Question: how “scary” are those *really bad* equilibria?

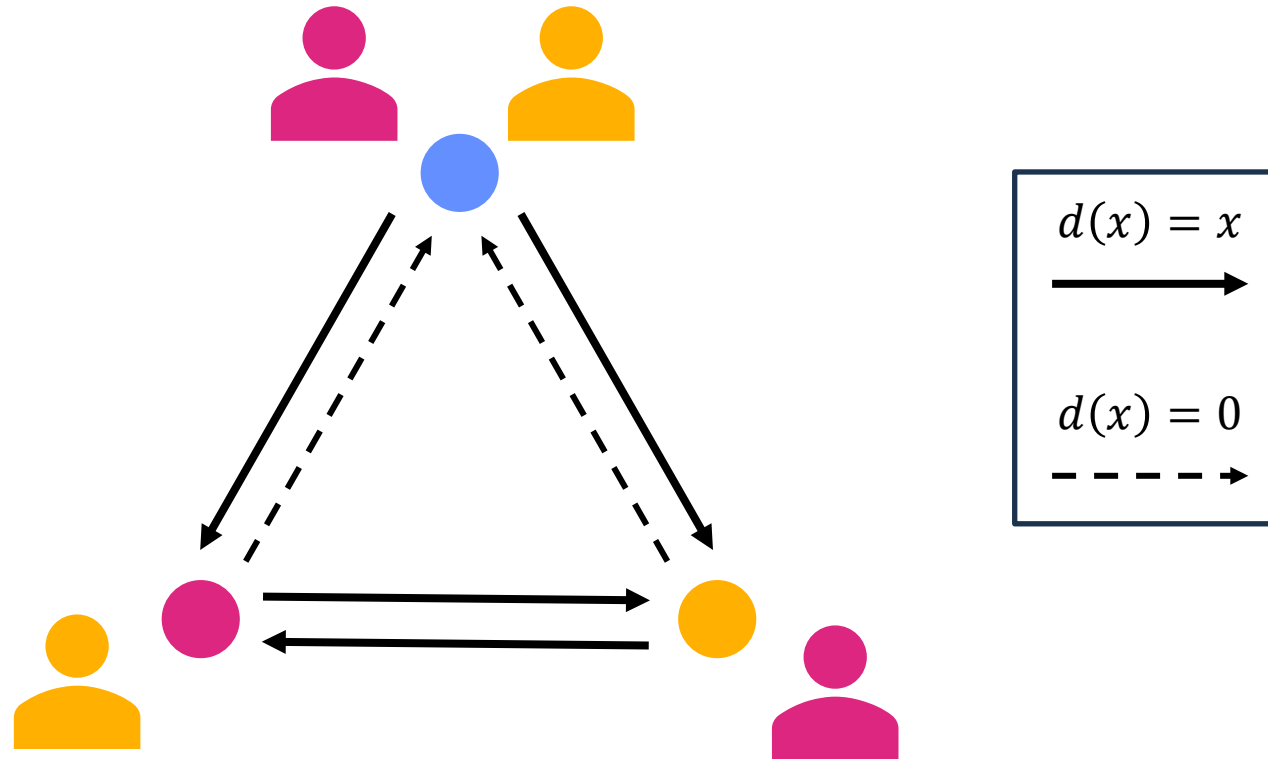
“Stability” $\approx$ agent satisfaction with equilibrium



Main Question: how “scary” are those *really bad* equilibria?

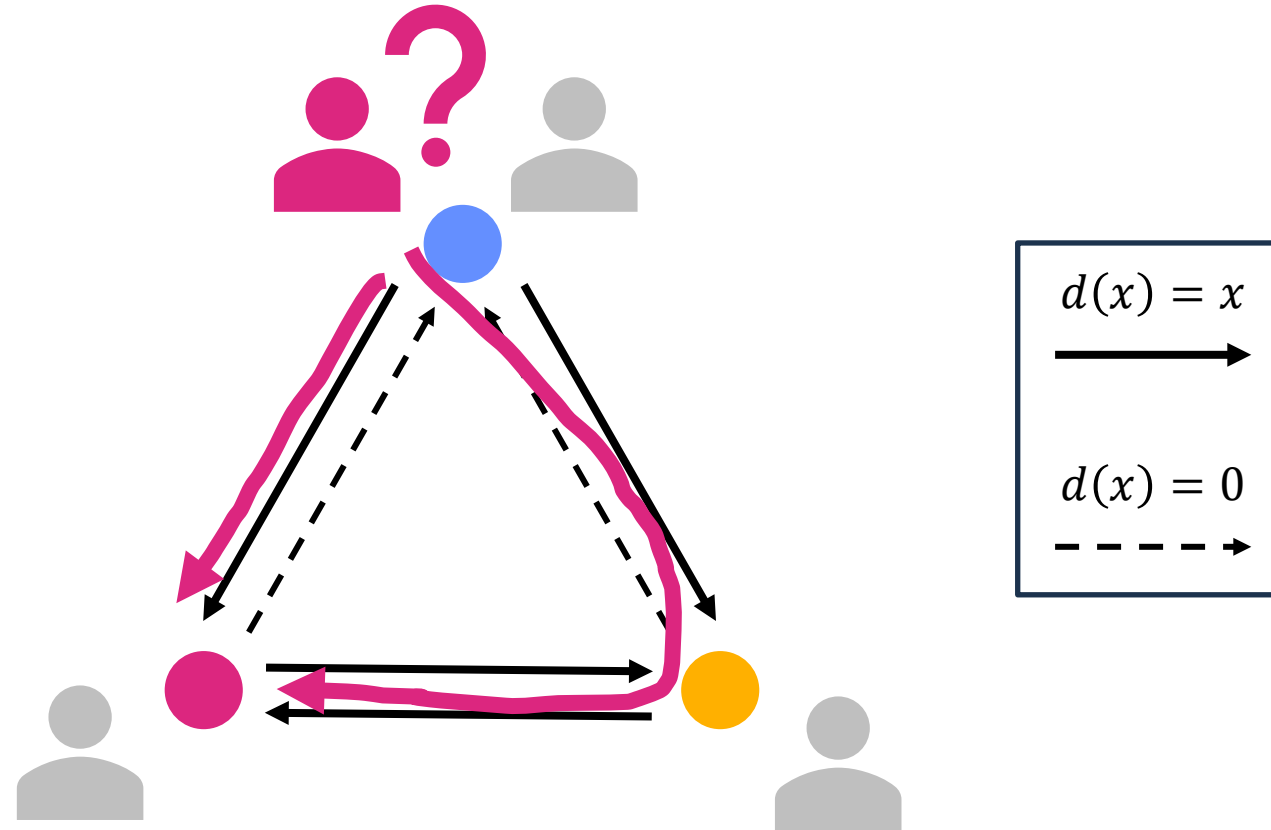
Opening agenda: look @ worst-case instances

Atomic congestion game with linear cost functions:



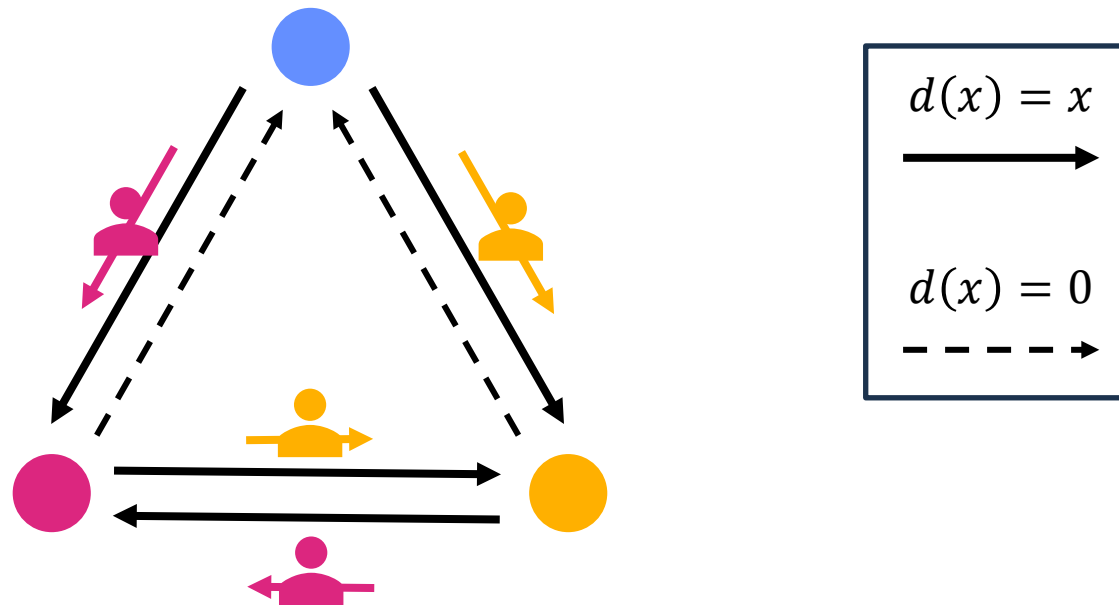
Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game with linear cost functions:



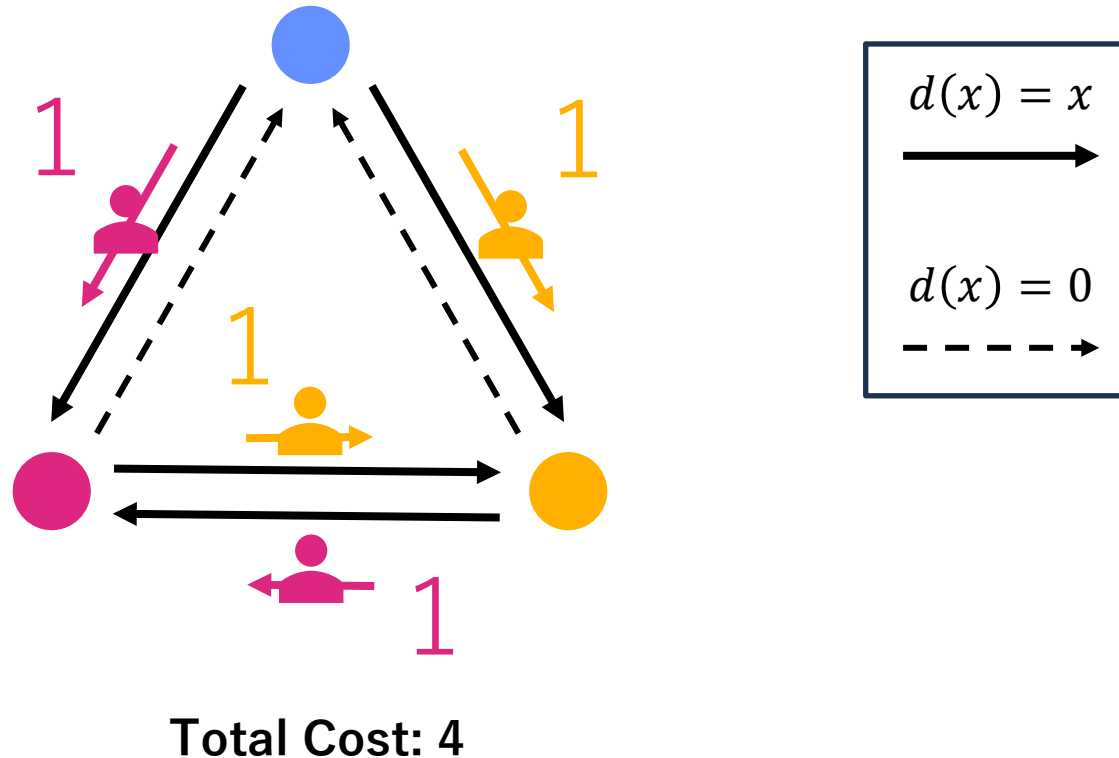
Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game with linear cost functions:



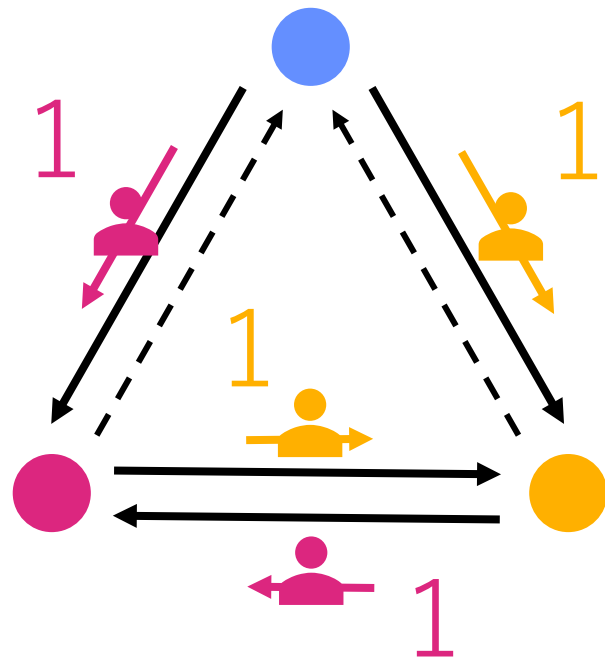
Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game with linear cost functions:

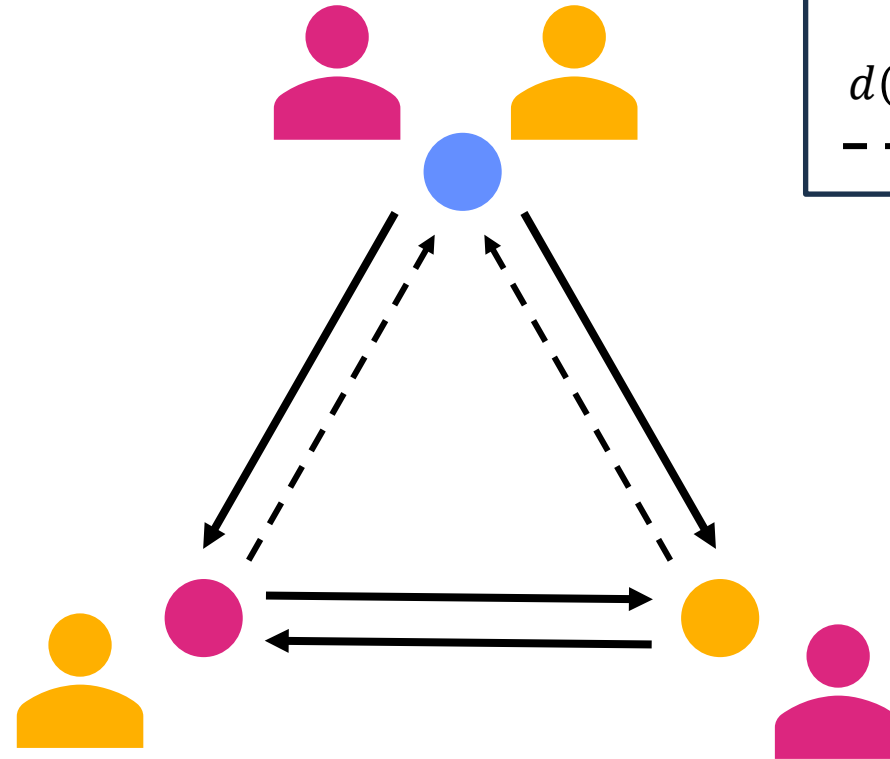


Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game:



Total Cost: 4

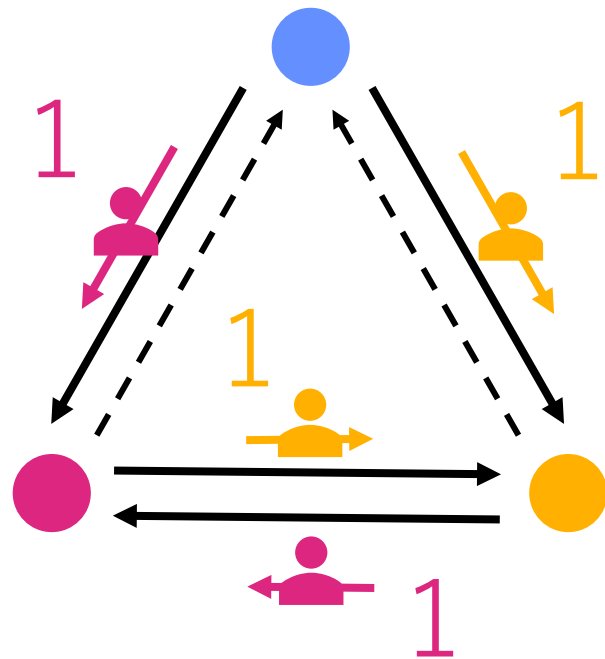


$$d(x) = x$$

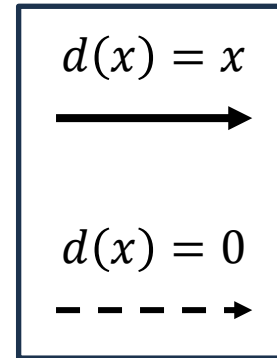
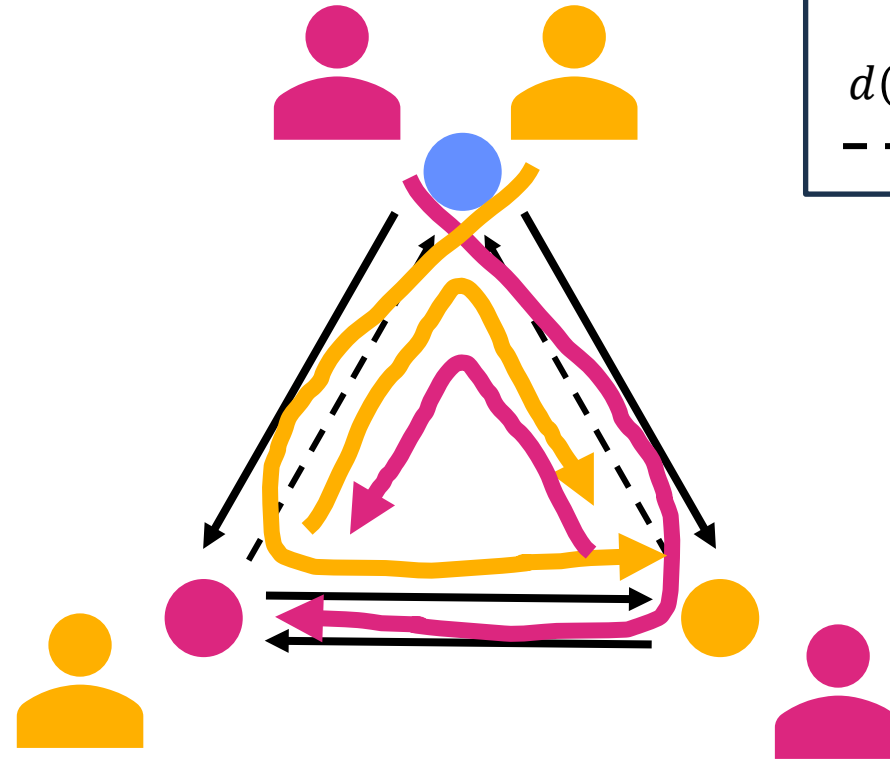
$$d(x) = 0$$

Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game:

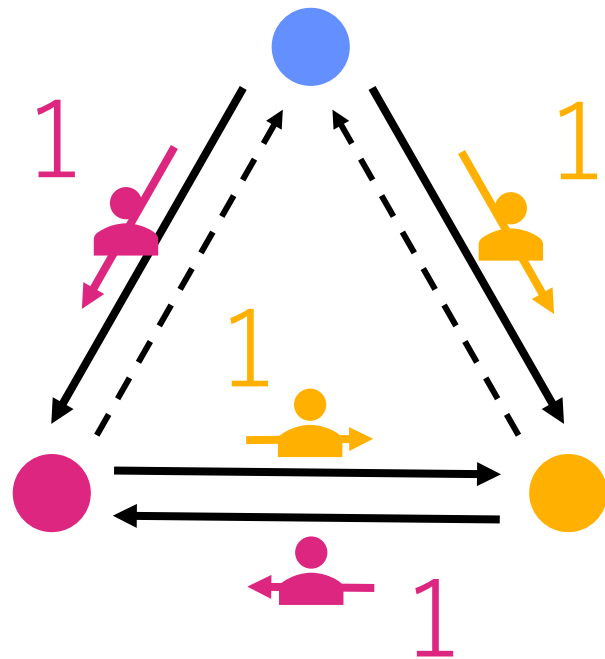


Total Cost: 4

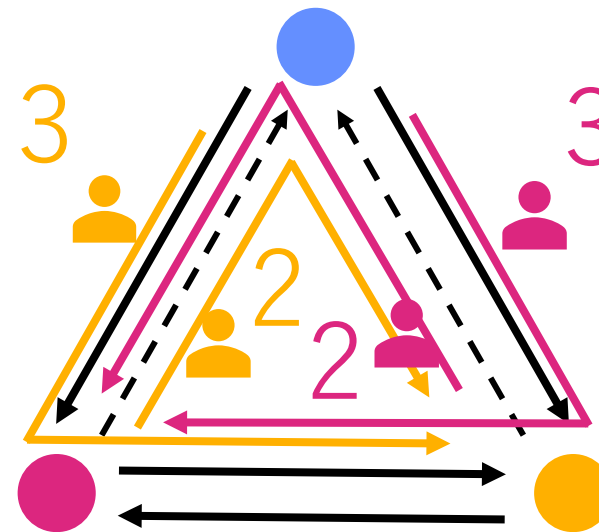


Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

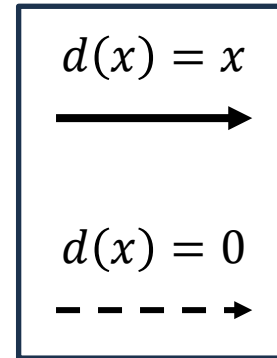
Atomic congestion game:



Total Cost: 4

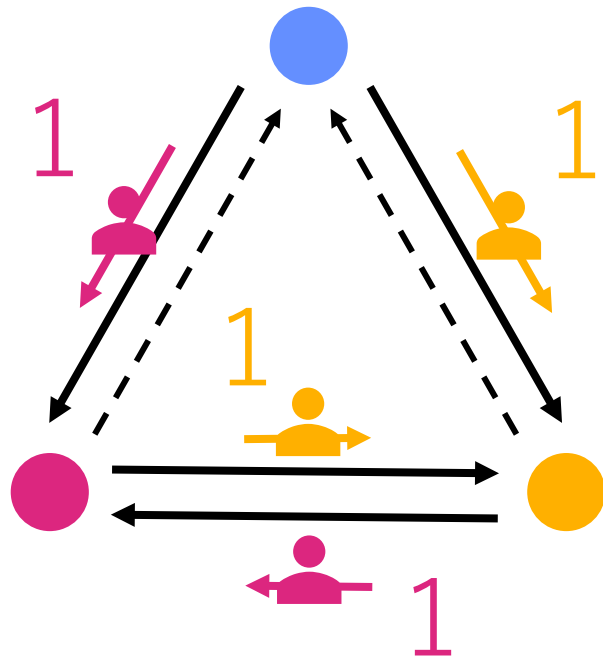


Total Cost: 10

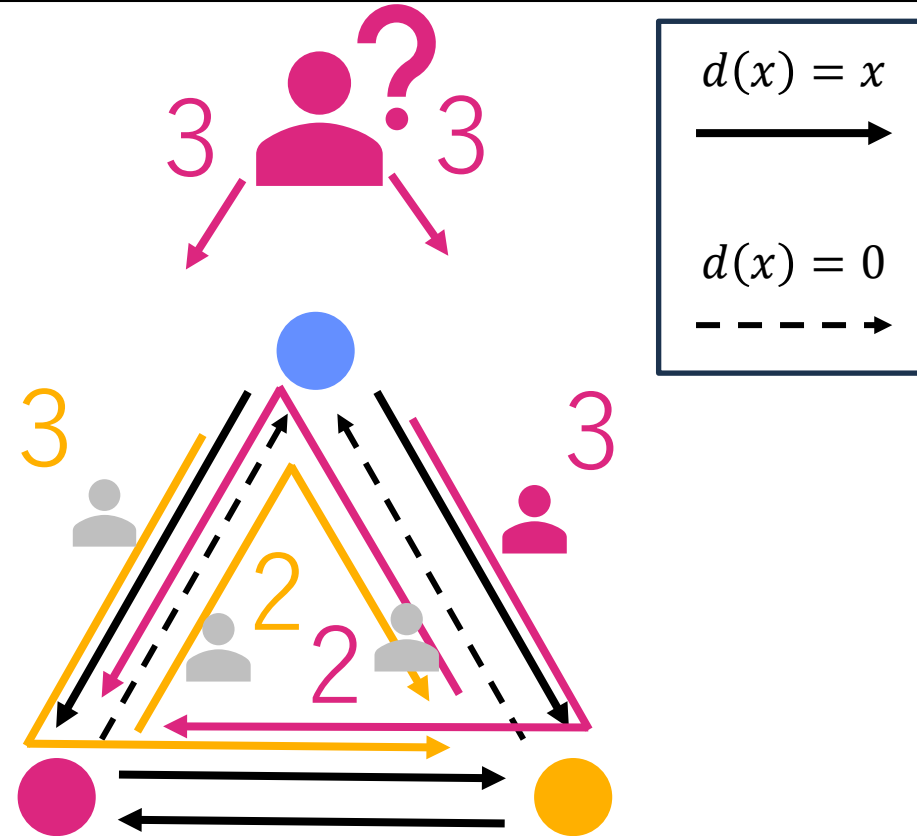


Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

Atomic congestion game:

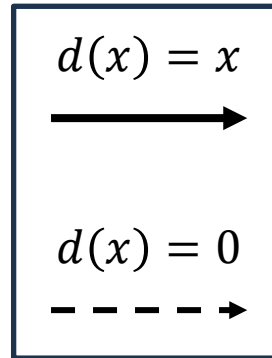


Total Cost: 4



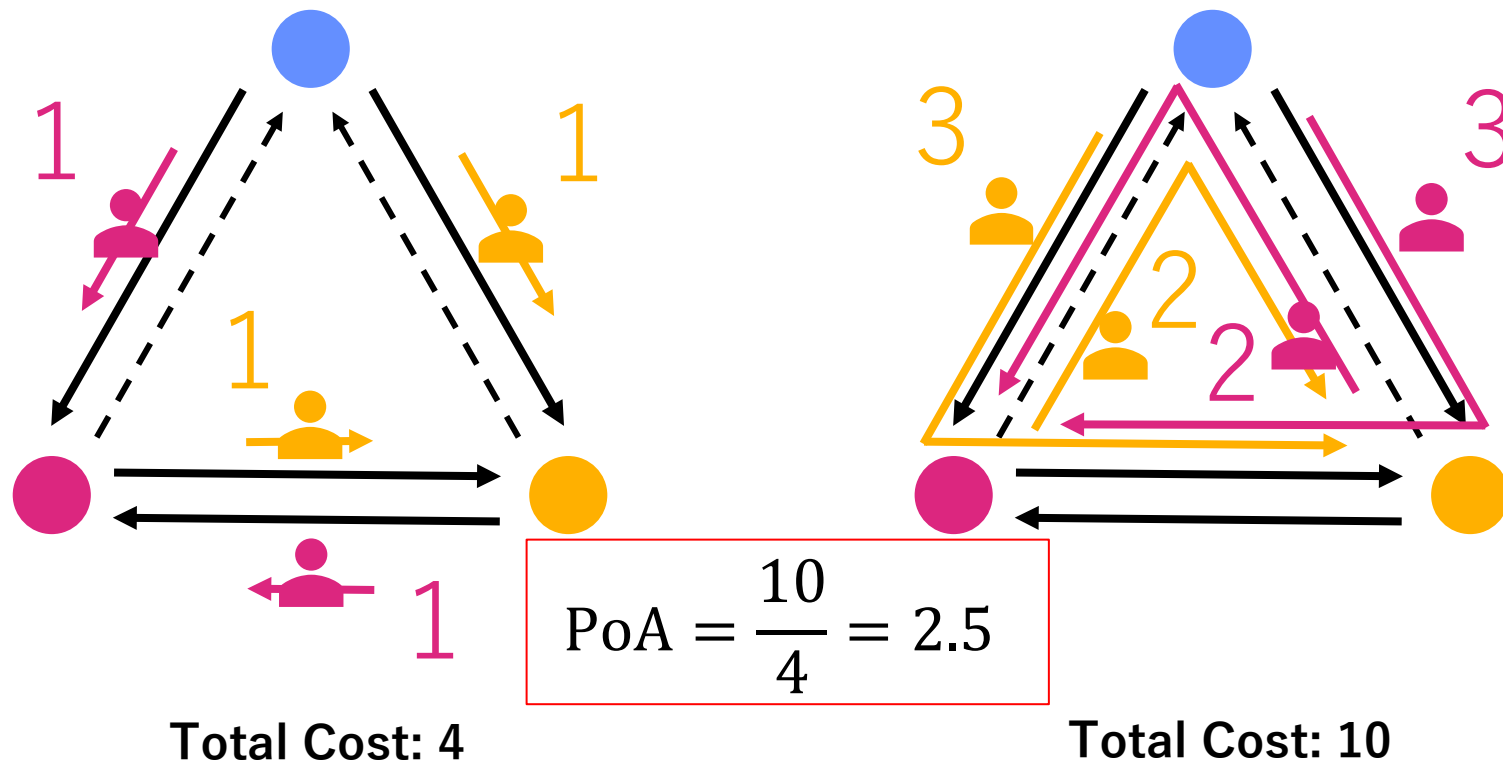
Fact: this is a Nash equilibrium!

Total Cost: 10

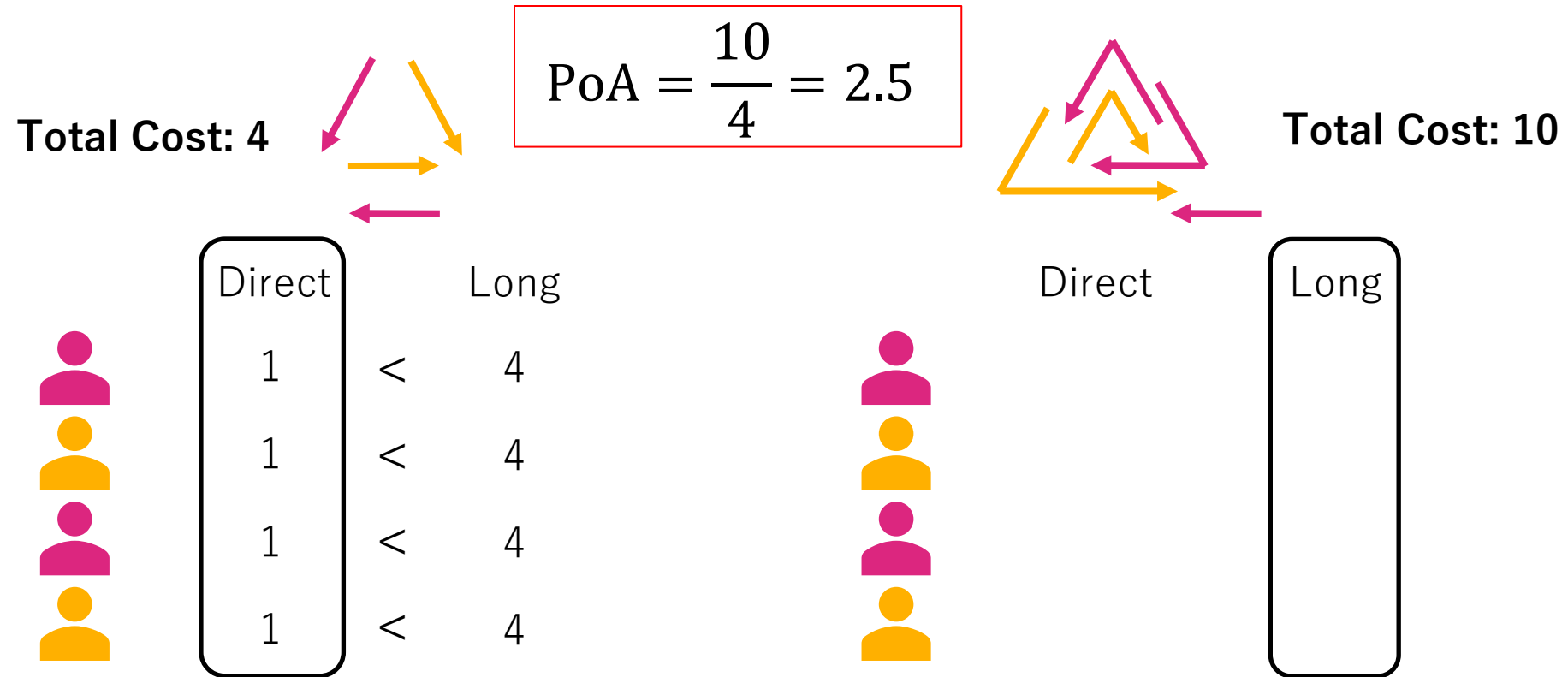


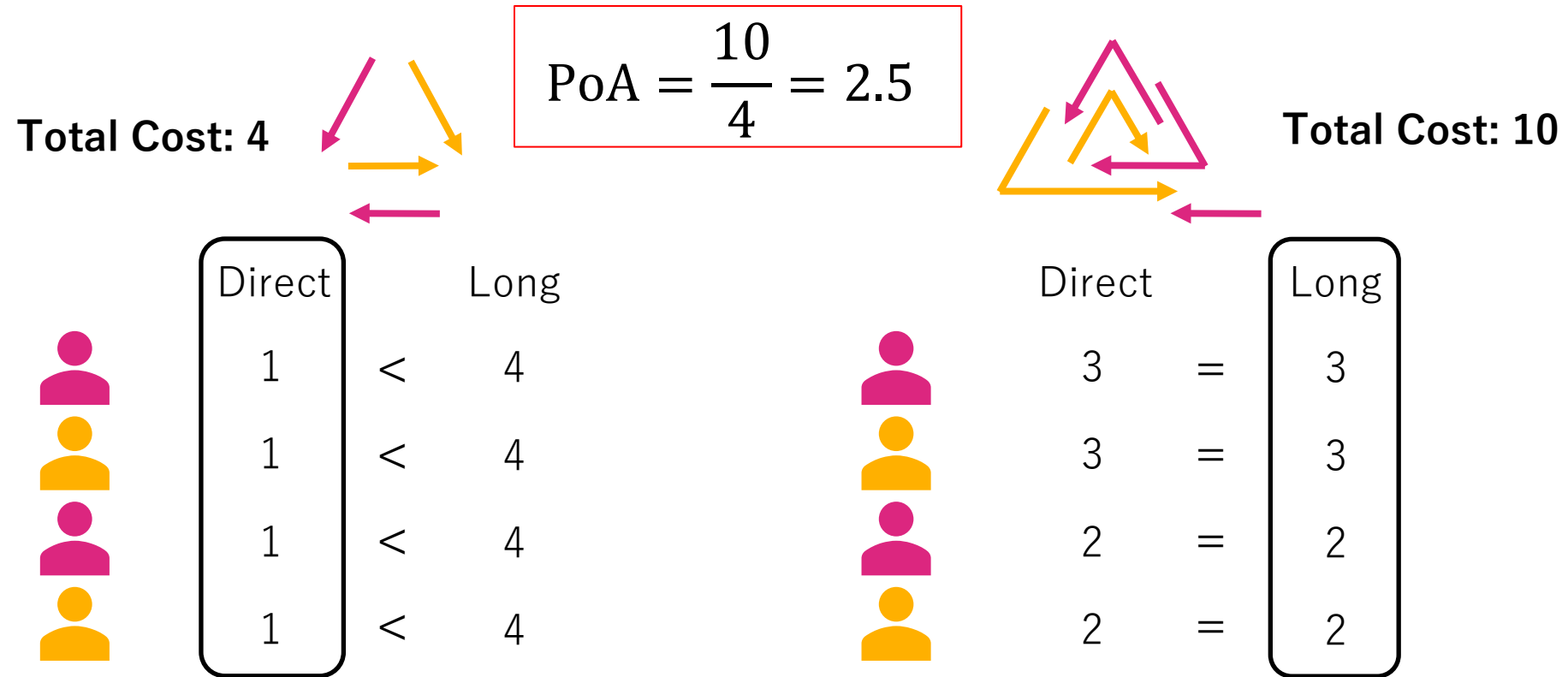
Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05

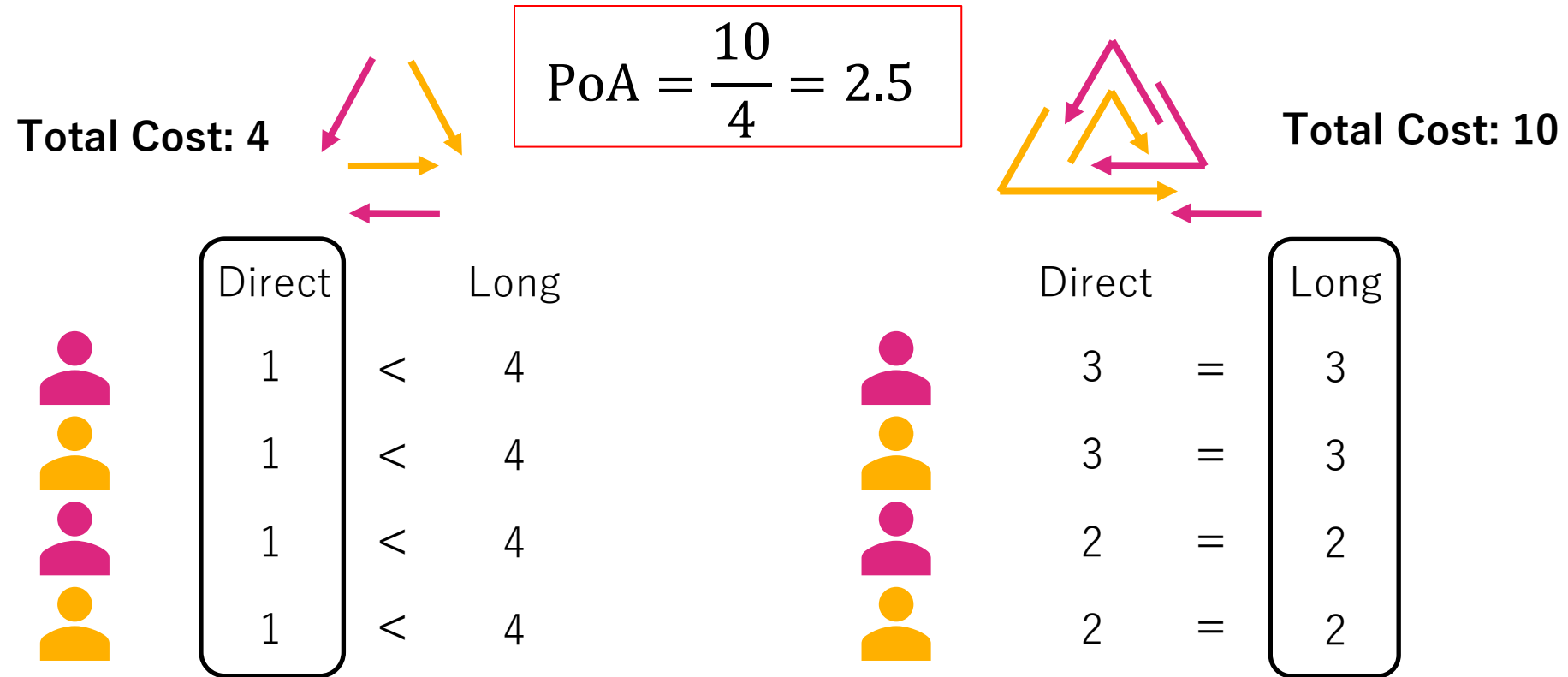
Atomic congestion game:



Christodoulou, Koutsoupias, "The price of anarchy in finite congestion games," STOC'05





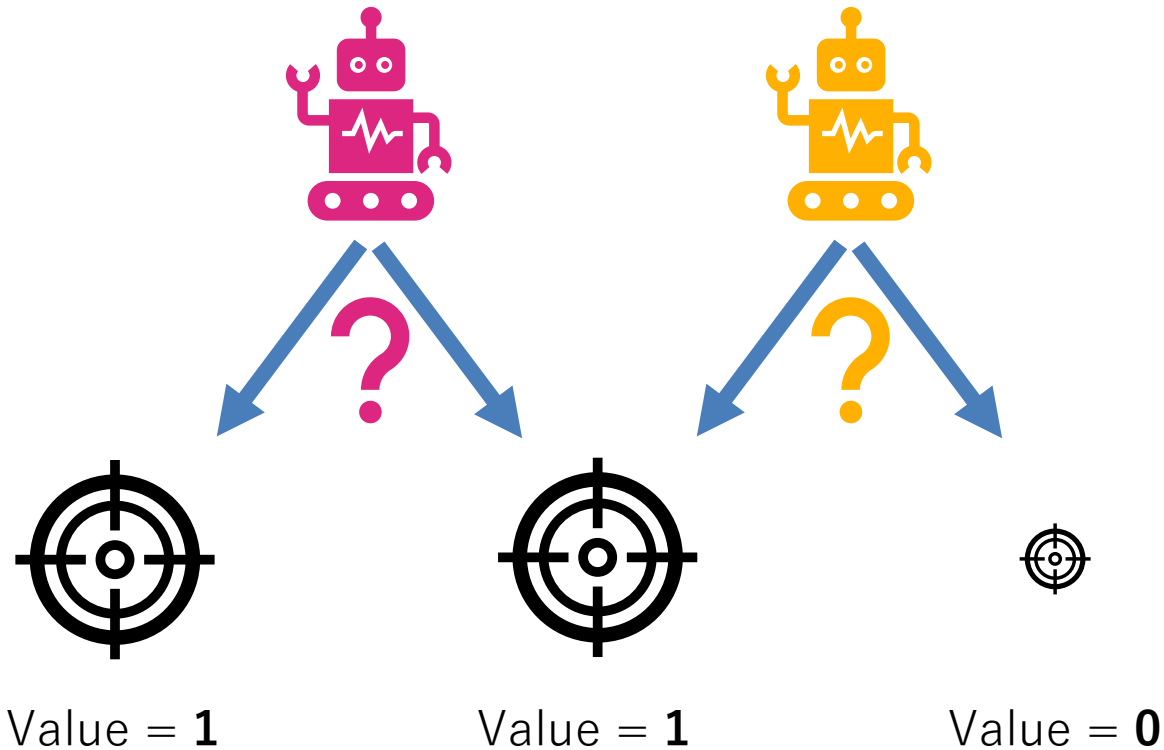


Our Observation:

“High Stability”
(agents prefer NOT to change)

“Zero Stability”
(all agents are indifferent)
deviation → behavior cascade to optimal

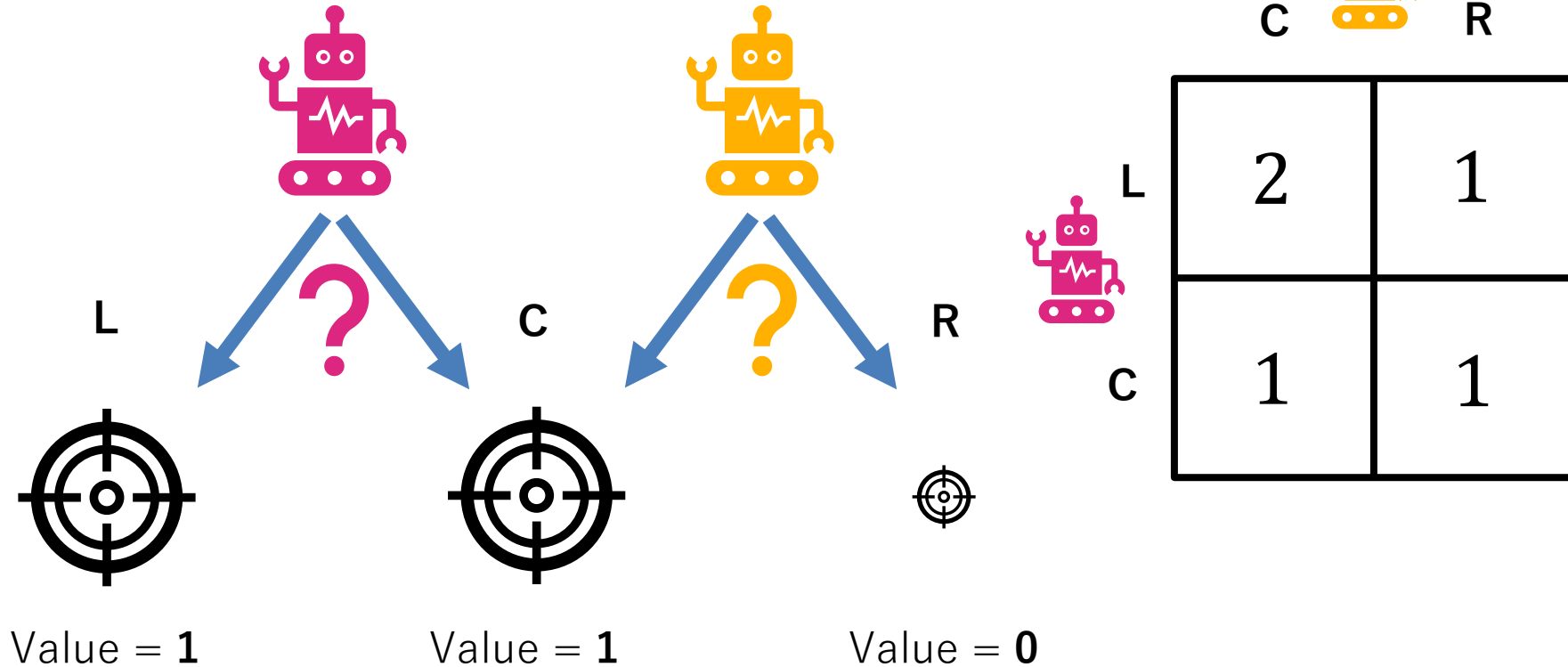
Max-Weight Set Cover:



System-level objective: cover the **maximum total value** of targets

Agent Decision: maximize system-level objective

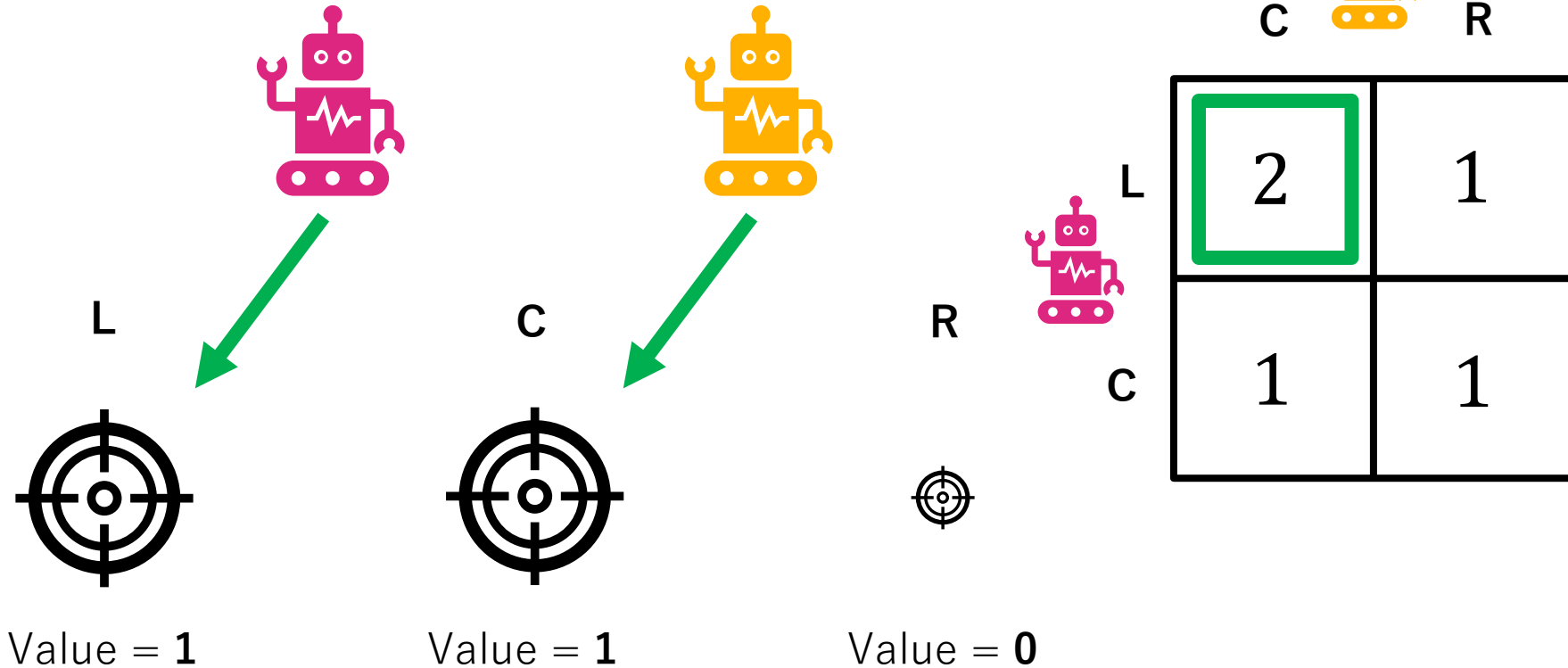
Max-Weight Set Cover:



System-level objective: cover the **maximum total value** of targets

Agent Decision: maximize system-level objective

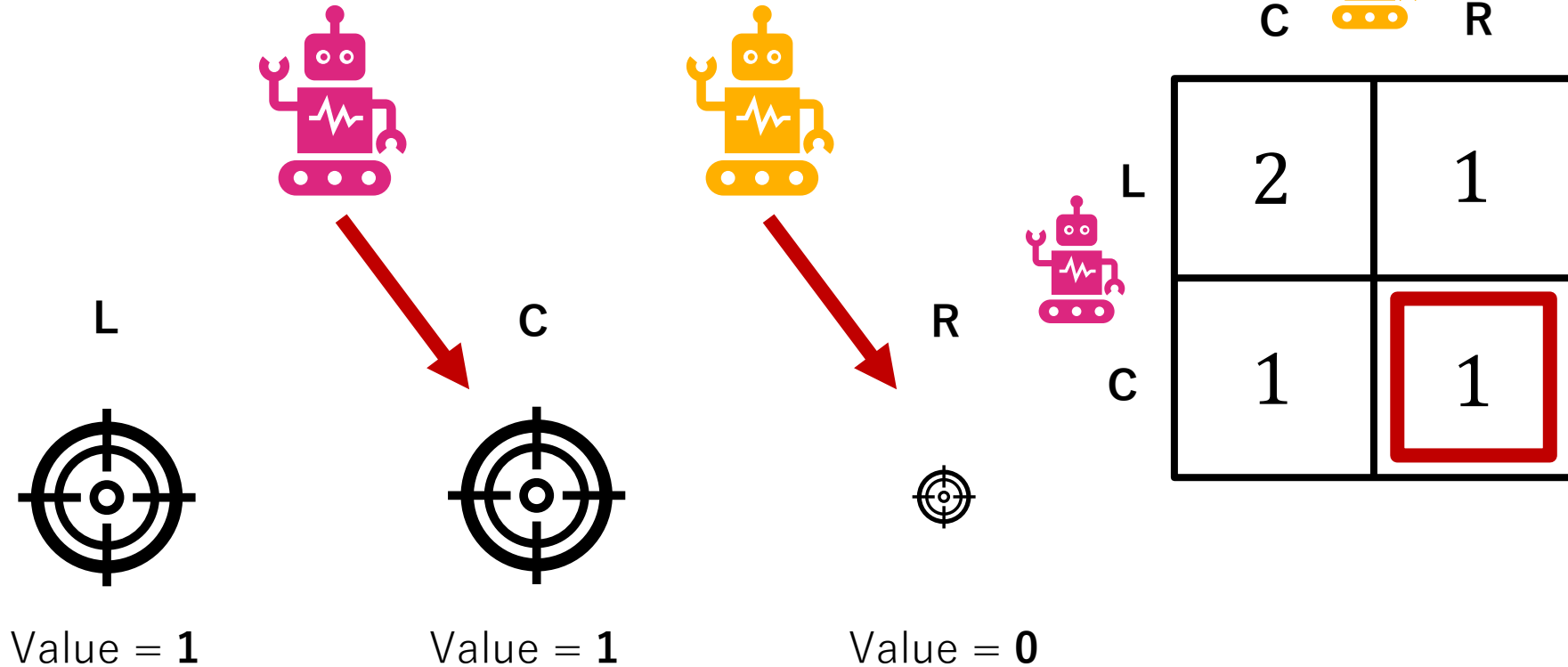
Max-Weight Set Cover:



System-level objective: cover the **maximum total value** of targets

Agent Decision: maximize system-level objective

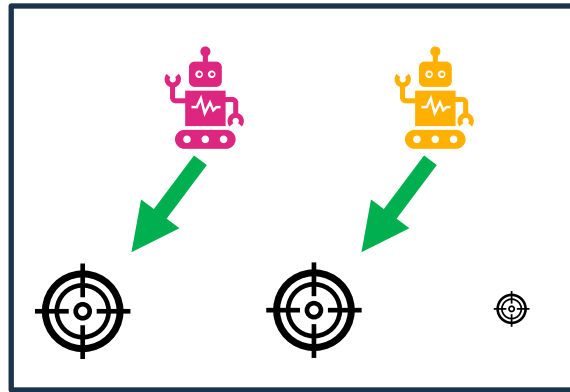
Max-Weight Set Cover:



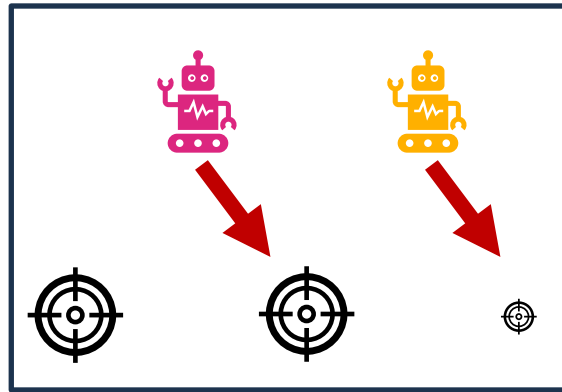
System-level objective: cover the **maximum total value** of targets

Agent Decision: maximize system-level objective

Max-Weight Set Cover:

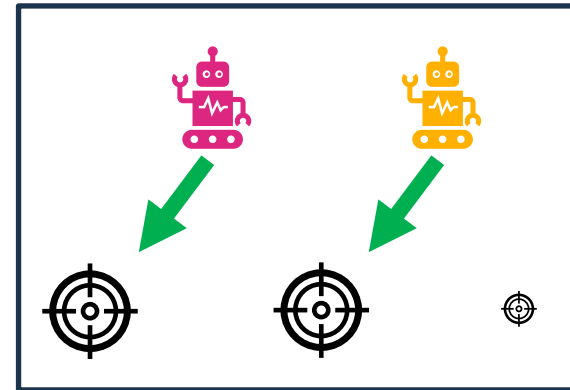
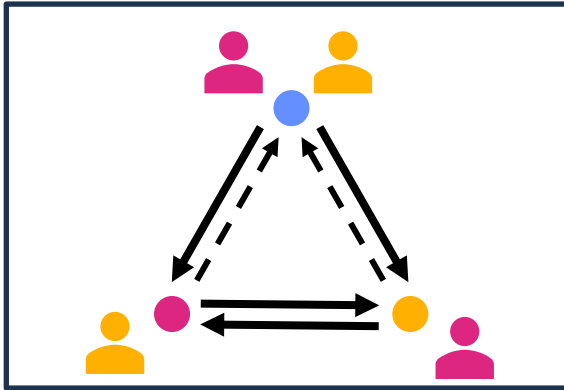


“High Stability”



“Zero Stability”

	C	R
L	2	1
C	1	1

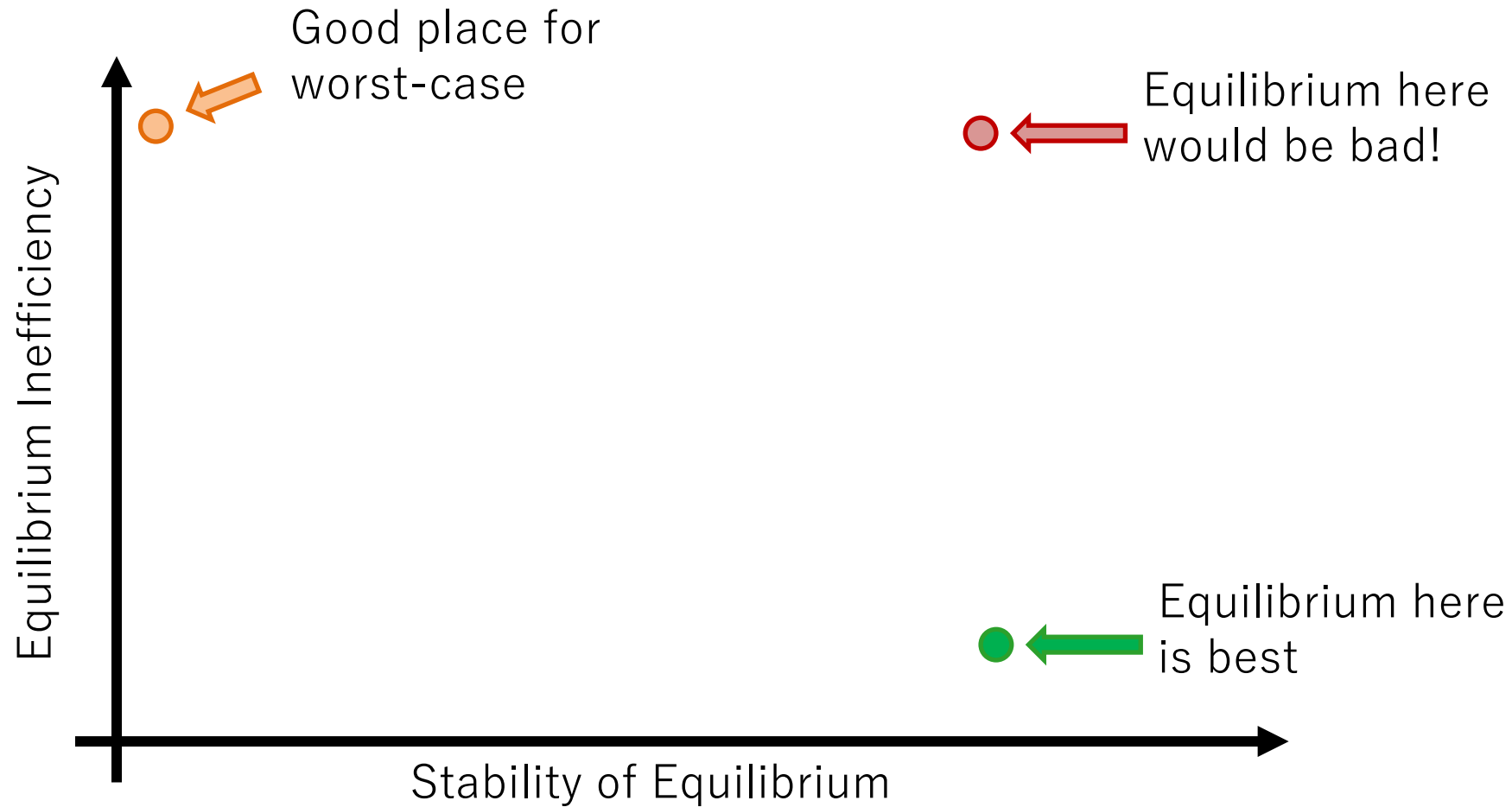


In these worst-case equilibria:

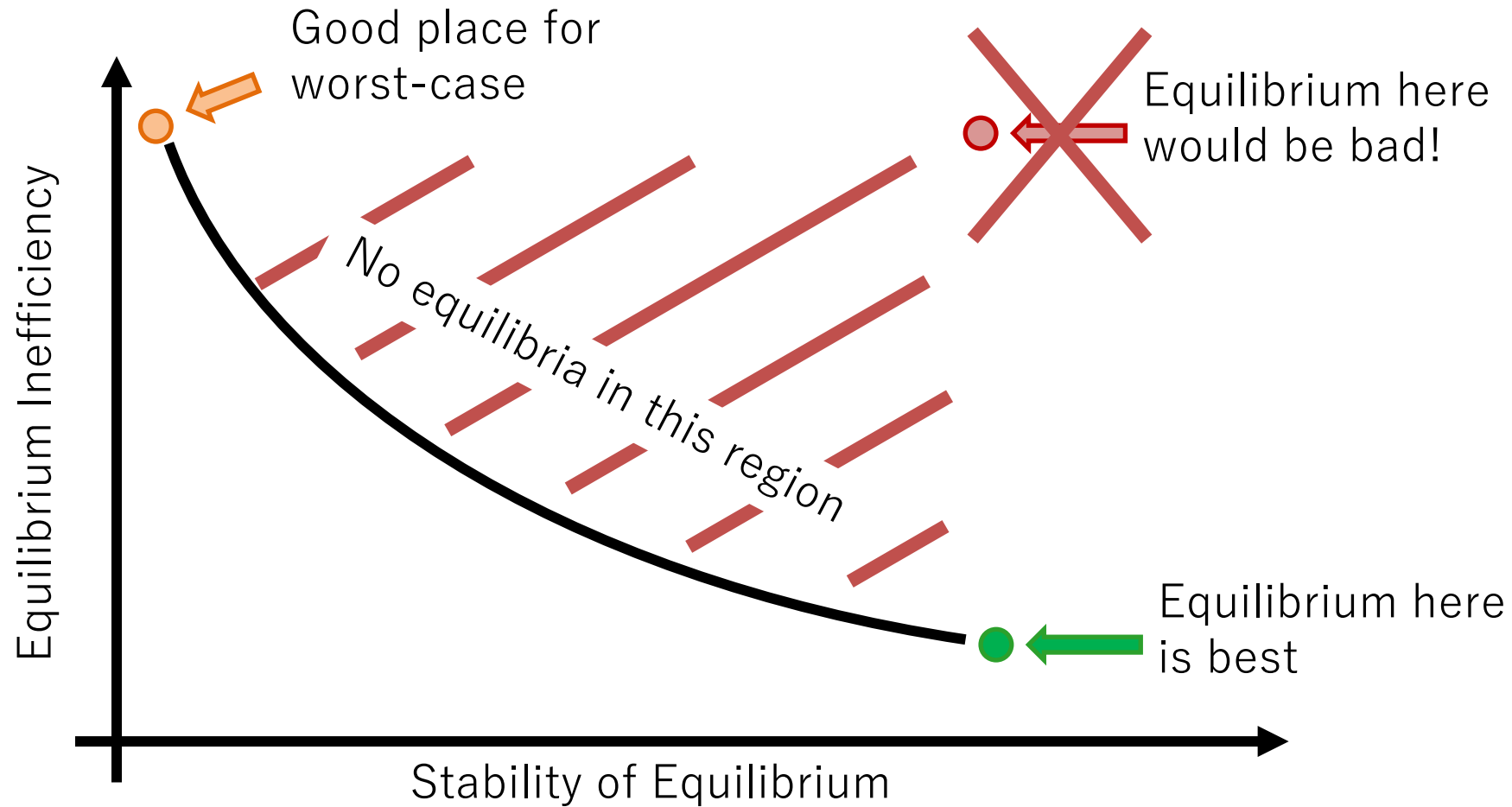
Efficient equilibria: Plausible and **Stable**

Inefficient equilibria: Contrived and “**Unstable**”

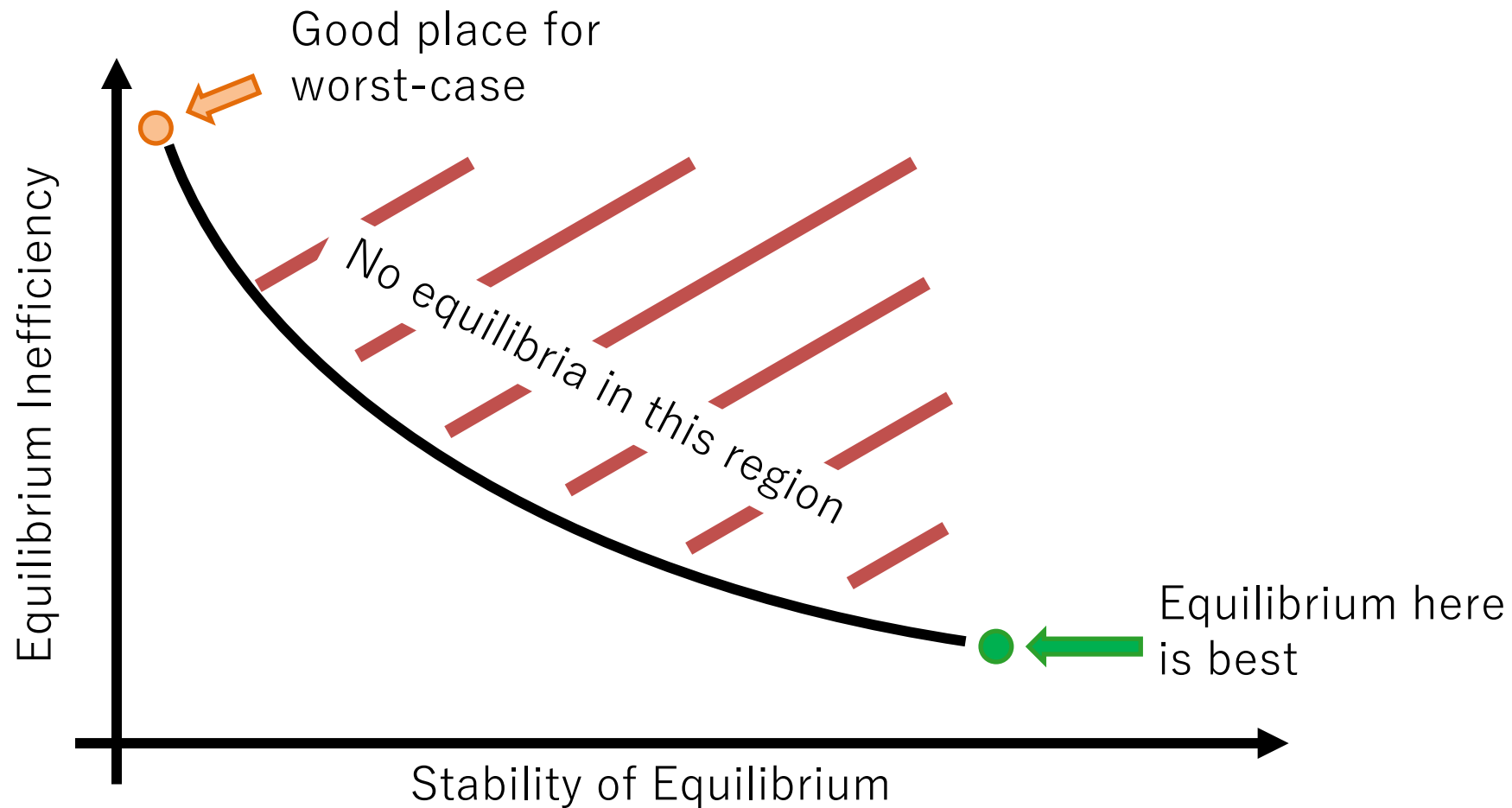
Does this generalize?



Main result: Every Inefficient Nash equilibrium has Zero Stability



Main result: Every Inefficient Nash equilibrium has Zero Stability



Main result: Every Inefficient Nash equilibrium has Zero Stability

How to exploit this?

Goal: prove stability/efficiency tradeoff

Cost-Minimization Game:

Players: $N = \{1, \dots, n\}$

Actions: $A := A_1 \times A_2 \times \dots \times A_n$

Cost functions: $C_i : A \rightarrow \mathbb{R}_{\geq 0}$ Social Cost: $C(a) := \sum_{i=1}^n C_i(a)$

Nash equilibrium a^{ne} :

$$C_i(a_i^{\text{ne}}, a_{-i}^{\text{ne}}) \leq C_i(a'_i, a_{-i}^{\text{ne}}) \quad \forall i \in N, a_i \in A_i$$

Price of Anarchy:

$$\text{PoA} = \frac{\max_{a^{\text{ne}}} C(a^{\text{ne}})}{\min_{a \in A} C(a)}.$$

Linear-cost congestion games:

$$\text{PoA} \leq 2.5$$

Goal: prove stability/efficiency tradeoff

Smooth Game:

A cost-minimization game is (λ, μ) -smooth (with $\lambda \geq 0$ and $\mu < 1$) if for every two outcomes a and a^* ,

$$\sum_{i=1}^n C_i(a_i^*, a_{-i}) \leq \lambda C(a^*) + \mu C(a).$$

Roughgarden, “Intrinsic Robustness of the Price of Anarchy,” JACM 2015

Goal: prove stability/efficiency tradeoff

Smooth Game:

A cost-minimization game is (λ, μ) -smooth (with $\lambda \geq 0$ and $\mu < 1$) if for every two outcomes a and a^* ,

$$\sum_{i=1}^n C_i(a_i^*, a_{-i}) \leq \lambda C(a^*) + \mu C(a).$$

If game is (λ, μ) -smooth, then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu}$$

Ex: Linear-cost congestion games: $(\frac{5}{3}, \frac{1}{3})$ -smooth: $\text{PoA} \leq \frac{5/3}{1 - 1/3} = 2.5$

Roughgarden, “Intrinsic Robustness of the Price of Anarchy,” JACM 2015

Goal: prove stability/efficiency tradeoff

Definition of “Stability”:

Optimal action profiles: $A^* = \arg \min_{a \in A} C(a)$

Worst Nash equilibria: $\text{WNE} = \arg \max_{a \in \text{NE}} C(a)$

Stability Margin:

$$S = \max_{a \in \text{WNE}} \max_{a^* \in A^*} \frac{\sum_{i=1}^n C_i(a_i^*, a_{-i}) - C_i(a)}{C(a)}.$$

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Goal: prove stability/efficiency tradeoff

Definition of “Stability”:

Optimal action profiles: $A^* = \arg \min_{a \in A} C(a)$

Worst Nash equilibria: $WNE = \arg \max_{a \in NE} C(a)$ Agent i 's cost to deviate to system-optimal

Stability Margin:

$$S = \max_{a \in WNE} \max_{a^* \in A^*} \frac{\sum_{i=1}^n C_i(a_i^*, a_{-i}) - C_i(a)}{C(a)}.$$

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Goal: prove stability/efficiency tradeoff

Definition of “Stability”:

Optimal action profiles: $A^* = \arg \min_{a \in A} C(a)$

Worst Nash equilibria: $WNE = \arg \max_{a \in NE} C(a)$ Agent i 's cost to deviate to system-optimal

Stability Margin:

$$S = \max_{a \in WNE} \max_{a^* \in A^*} \frac{\sum_{i=1}^n C_i(a_i^*, a_{-i}) - C_i(a)}{C(a)}.$$

$S = 0 \rightarrow$ All can deviate to system-optimal for “free”

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Goal: prove stability/efficiency tradeoff

Definition of “Stability”:

Optimal action profiles: $A^* = \arg \min_{a \in A} C(a)$

Worst Nash equilibria: $WNE = \arg \max_{a \in NE} C(a)$ Agent i 's cost to deviate to system-optimal

Stability Margin:

$$S = \max_{a \in WNE} \max_{a^* \in A^*} \frac{\sum_{i=1}^n C_i(a_i^*, a_{-i}) - C_i(a)}{C(a)}.$$

$S = 0 \rightarrow$ All can deviate to system-optimal for “free”

$S \gg 0 \rightarrow$ Deviating to system-optimal is costly (in aggregate)

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Goal: prove stability/efficiency tradeoff

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

$S = 0 \rightarrow$ All can deviate to system-optimal for “free”

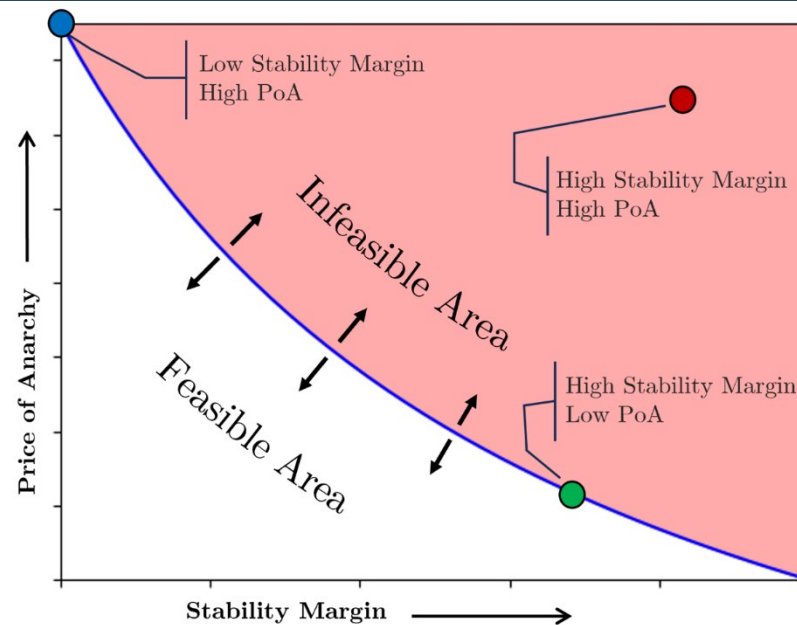
$S \gg 0 \rightarrow$ Deviating to system-optimal is costly (in aggregate)

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Goal: prove stability/efficiency tradeoff

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$



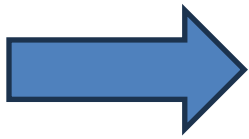
Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

Proof: Let $a \in \text{WNE}$ and $a^* \in A^*$. Then

$$\begin{aligned} C(a) &= \sum_{i=1}^n C_i(a) \\ &\leq \sum_{i=1}^n C_i(a_i^*, a_{-i}) && (a \text{ is a Nash equilibrium}) \\ &\leq \lambda C(a^*) + \mu C(a) && (\text{game is } (\lambda, \mu)\text{-smooth}) \end{aligned}$$


 $\frac{C(a)}{C(a^*)} \leq \frac{\lambda}{1 - \mu}$

Roughgarden, 2012

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

Proof: Let $a \in \text{WNE}$ and $a^* \in A^*$. Then

$$\begin{aligned} C(a) &= \sum_{i=1}^n C_i(a) \\ &= \sum_{i=1}^n C_i(a_i^*, a_{-i}) - SC(a) && (a \text{ is a Nash equilibrium}) \\ &\leq \lambda C(a^*) + \mu C(a) && (\text{game is } (\lambda, \mu)\text{-smooth}) \end{aligned}$$

$$\Rightarrow \frac{C(a)}{C(a^*)} \leq \frac{\lambda}{1 - \mu}$$

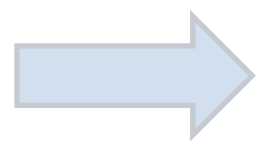
Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

Proof: Let $a \in \text{WNE}$ and $a^* \in A^*$. Then

$$\begin{aligned} C(a) &= \sum_{i=1}^n C_i(a) \\ &= \sum_{i=1}^n C_i(a_i^*, a_{-i}) - SC(a) && (a \text{ is a Nash equilibrium}) \\ &\leq \lambda C(a^*) + \mu C(a) - SC(a) && (\text{game is } (\lambda, \mu)\text{-smooth}) \end{aligned}$$



$$\frac{C(a)}{C(a^*)} \leq \frac{\lambda}{1 - \mu}$$

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

Proof: Let $a \in \text{WNE}$ and $a^* \in A^*$. Then

$$\begin{aligned} C(a) &= \sum_{i=1}^n C_i(a) \\ &= \sum_{i=1}^n C_i(a_i^*, a_{-i}) - SC(a) \quad (a \text{ is a Nash equilibrium}) \\ &\leq \lambda C(a^*) + \mu C(a) - SC(a) \quad (\text{game is } (\lambda, \mu)\text{-smooth}) \end{aligned}$$

$$\Rightarrow \frac{C(a)}{C(a^*)} \leq \frac{\lambda}{1 - \mu + S}$$

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

Theorem: If a game is (λ, μ) -smooth and has stability margin S , then

$$\text{PoA} \leq \frac{\lambda}{1 - \mu + S}$$

Main Insight:

The Nash inequality:

1. Has economically meaningful slack
2. Is central to the proof of PoA bounds

Question: is this technique “portable?”

Yes, it is portable, at least somewhat!

Seaton, **PNB**, “On the Intrinsic Fragility of the Price of Anarchy,” L-CSS 2023

- Resource $r \in R$ has value v_r
- Agent $i \in N$ has action set $A_i \subseteq 2^R$
- Given collective choices $a = (a_1, a_2, \dots, a_n)$, write $|a|_r := \#$ agents choosing r
- System objective (to maximize):

$$W(a) = \sum_{r \in R} v_r w(|a|_r)$$

for some “basis function” $w: \mathbb{N} \rightarrow \mathbb{R}$

Yes, it is portable, at least somewhat!

- Resource $r \in R$ has value v_r
- Agent $i \in N$ has action set $A_i \subseteq 2^R$
- Given collective choices $a = (a_1, a_2, \dots, a_n)$, write $|a|_r := \#$ agents choosing r
- System objective (to maximize):

$$W(a) = \sum_{r \in R} v_r w(|a|_r)$$

for some “basis function” $w: \mathbb{N} \rightarrow \mathbb{R}$

- Agents exist on an *information network*:
agent i observes only agents $M_i \subseteq N$

- **High-level goal: get agents to maximize system objective “on their own”**
- **Design philosophy: *Utility Design***

Principle: assign agents *local objective functions* $U_i: \mathbb{N} \rightarrow \mathbb{R}$

such that maximization of $\{U_i\}$ **coincides** with global maximization of W

Known: [Paccagnan TAC 2019; Singh, Wesley, **PNB** ACC 2025]

$$\max_{\{U_i\}} \text{PoA}(\{U_i\}, w, M) = \min_a \frac{W(a^*)}{\max_a W(a)} \text{ can be found using an LP!}$$

Theorem [Singh, Wesley, **PNB**, ACC 2025]:

$\min_a \frac{W(a^*)}{\max_a W(a)}$ is equivalent to the following LP:

$$\begin{aligned}
 W^* &= \max_{\theta(t)} \sum_{t \in \mathcal{I}} w(B_t) \theta(t) \\
 \sum_{t \in \mathcal{I}} [a_j f_j(A_{t,j}) - b_j f_j(A_{t,j} + 1)] \theta(t) &\geq 0, \\
 \sum_{t \in \mathcal{I}} w(A_t) \theta(t) &= 1, \\
 \theta(t) &\geq 0 \quad \text{for all } t \in \mathcal{I}, \\
 \text{and } f_j(0) = f_j(s_j + 1) = w(0) &= 0, \quad \forall j \in [k]
 \end{aligned}$$

Singh, **PNB**, “Worst-case Equilibria in Networked Resource Allocation Games rest on a Knife-Edge,” CDC 2025

Theorem [Singh, Wesley, **PNB**, ACC 2025]:

$\min_a \frac{W(a^*)}{\max_a W(a)}$ is equivalent to the following LP:

Find game instance with:

$$W^* = \max_{\theta(t)} \sum_{t \in \mathcal{I}} w(B_t) \theta(t) \quad \text{really good optimal outcome}$$

$$\sum_{t \in \mathcal{I}} [a_j f_j(A_{t,j}) - b_j f_j(A_{t,j} + 1)] \theta(t) \geq 0,$$

$$\sum_{t \in \mathcal{I}} w(A_t) \theta(t) = 1,$$

$$\theta(t) \geq 0 \quad \text{for all } t \in \mathcal{I},$$

$$\text{and } f_j(0) = f_j(s_j + 1) = w(0) = 0, \quad \forall j \in [k]$$

*really bad
Nash equilibrium*

Singh, **PNB**, "Worst-case Equilibria in Networked Resource Allocation Games rest on a Knife-Edge," CDC 2025

Theorem [Singh, Wesley, **PNB**, ACC 2025]:

$\min_a \frac{W(a^*)}{\max_a W(a)}$ is equivalent to the following LP:

$$W^* = \max_{\theta(t)} \sum_{t \in \mathcal{I}} w(B_t) \theta(t) \quad \text{equivalent to:}$$

$$\sum_{t \in \mathcal{I}} [a_j f_j(A_{t,j}) - b_j f_j(A_{t,j} + 1)] \theta(t) \geq 0,$$

$$\sum_{i \in \mathcal{C}_j} [U_i(a^{\text{ne}}) - U_i(a_i^{\text{opt}}, a_{-i}^{\text{ne}})] \geq 0$$

$$\sum_{t \in \mathcal{I}} w(A_t) \theta(t) = 1,$$

$$\theta(t) \geq 0 \quad \text{for all } t \in \mathcal{I},$$

$$\text{and } f_j(0) = f_j(s_j + 1) = w(0) = 0, \quad \forall j \in [k]$$

Singh, **PNB**, “Worst-case Equilibria in Networked Resource Allocation Games rest on a Knife-Edge,” CDC 2025

Theorem [Singh, Wesley, **PNB**, ACC 2025]:

$\min_a \frac{W(a^*)}{\max_a W(a)}$ is equivalent to the following LP:

$$W^* = \max_{\theta(t)} \sum_{t \in \mathcal{I}} w(B_t) \theta(t) \quad \text{equivalent to:}$$

$$\sum_{t \in \mathcal{I}} [a_j f_j(A_{t,j}) - b_j f_j(A_{t,j} + 1)] \theta(t) \geq 0,$$

$$\sum_{i \in \mathcal{C}_j} [U_i(a_i^{\text{ne}}) - U_i(a_i^{\text{opt}}, a_{-i}^{\text{ne}})] \geq 0$$

$$\sum_{t \in \mathcal{I}} w(A_t) \theta(t) = 1,$$

$$\theta(t) \geq 0 \quad \text{for all } t \in \mathcal{I}, \quad \text{complementary slackness} \rightarrow \text{this binds!}$$

$$\text{and } f_j(0) = f_j(s_j + 1) = w(0) = 0, \quad \forall j \in [k]$$

Singh, **PNB**, “Worst-case Equilibria in Networked Resource Allocation Games rest on a Knife-Edge,” CDC 2025

Main result: Every Inefficient Nash equilibrium is “unstable”

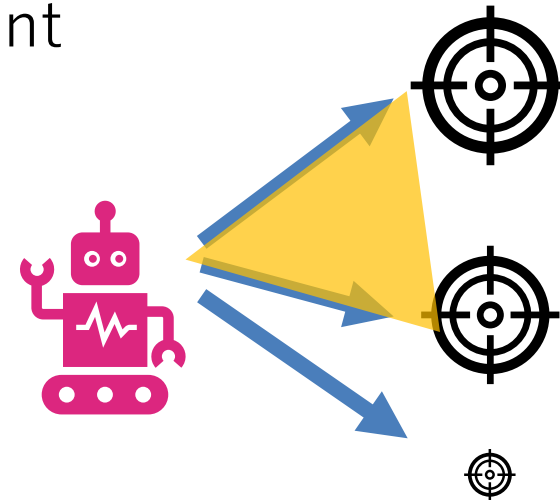
How to exploit this?

Inefficient $\Rightarrow$ easy to leave



hard to leave $\Rightarrow$ efficient

Idea: Randomize among high-payoff actions!



Result flavor: Distributed algorithms with “safety” guarantees

“Learning” in games: What happens when agents update their actions over time?

Equilibrium-**Seeking** Dynamics

Agenda: Prove that update rule X converges to NE in class of games Y

Example: Best-response dynamics converges almost surely to NE in potential games.

Problem: I only want to converge to *good* NE!

Monderer & Shapley, “Potential Games”, GEB’96

“Learning” in games: What happens when agents update their actions over time?

Equilibrium-**Seeking** Dynamics

Agenda: Prove that update rule X converges to NE in class of games Y

Example: Best-response dynamics converges almost surely to NE in potential games.

Problem: I only want to converge to *good* NE!

Monderer & Shapley, “Potential Games”, GEB’96

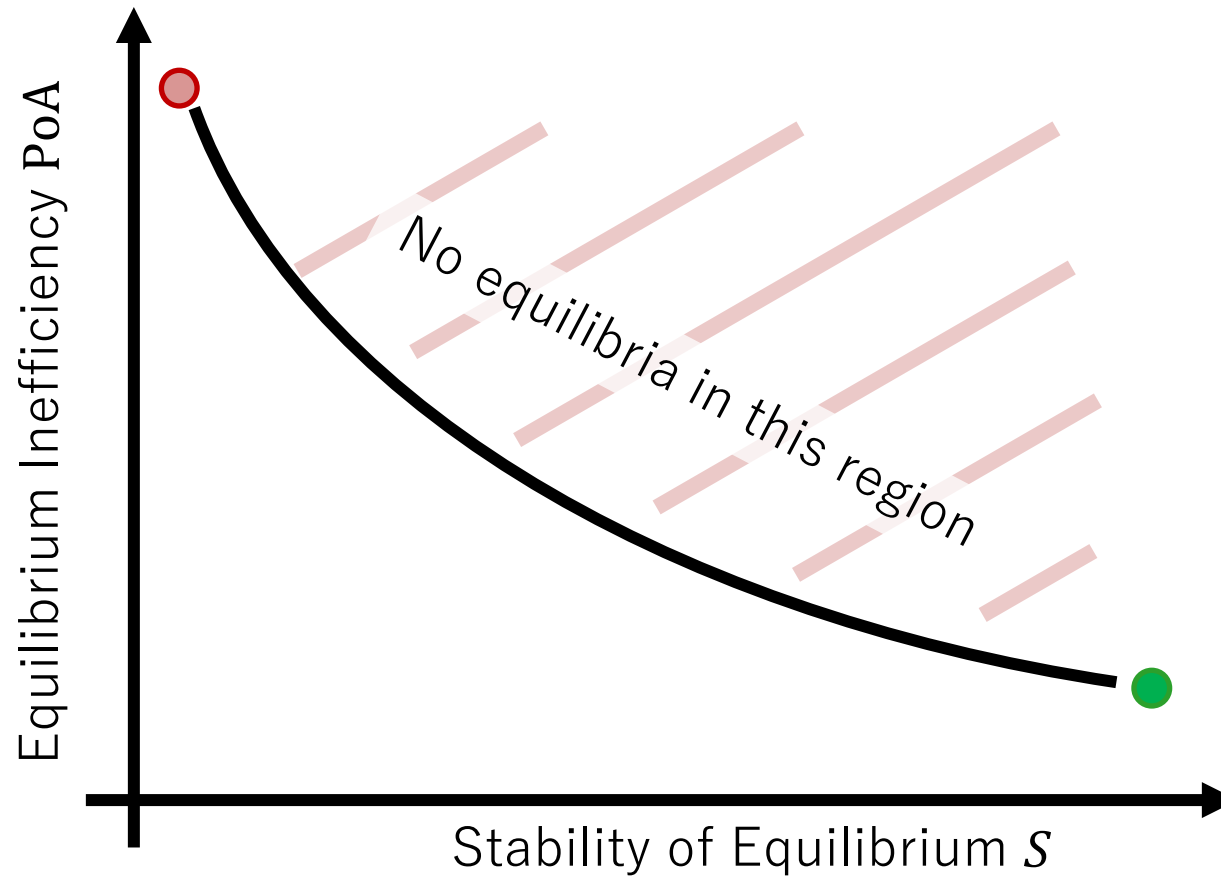
Equilibrium-**Selecting** Dynamics

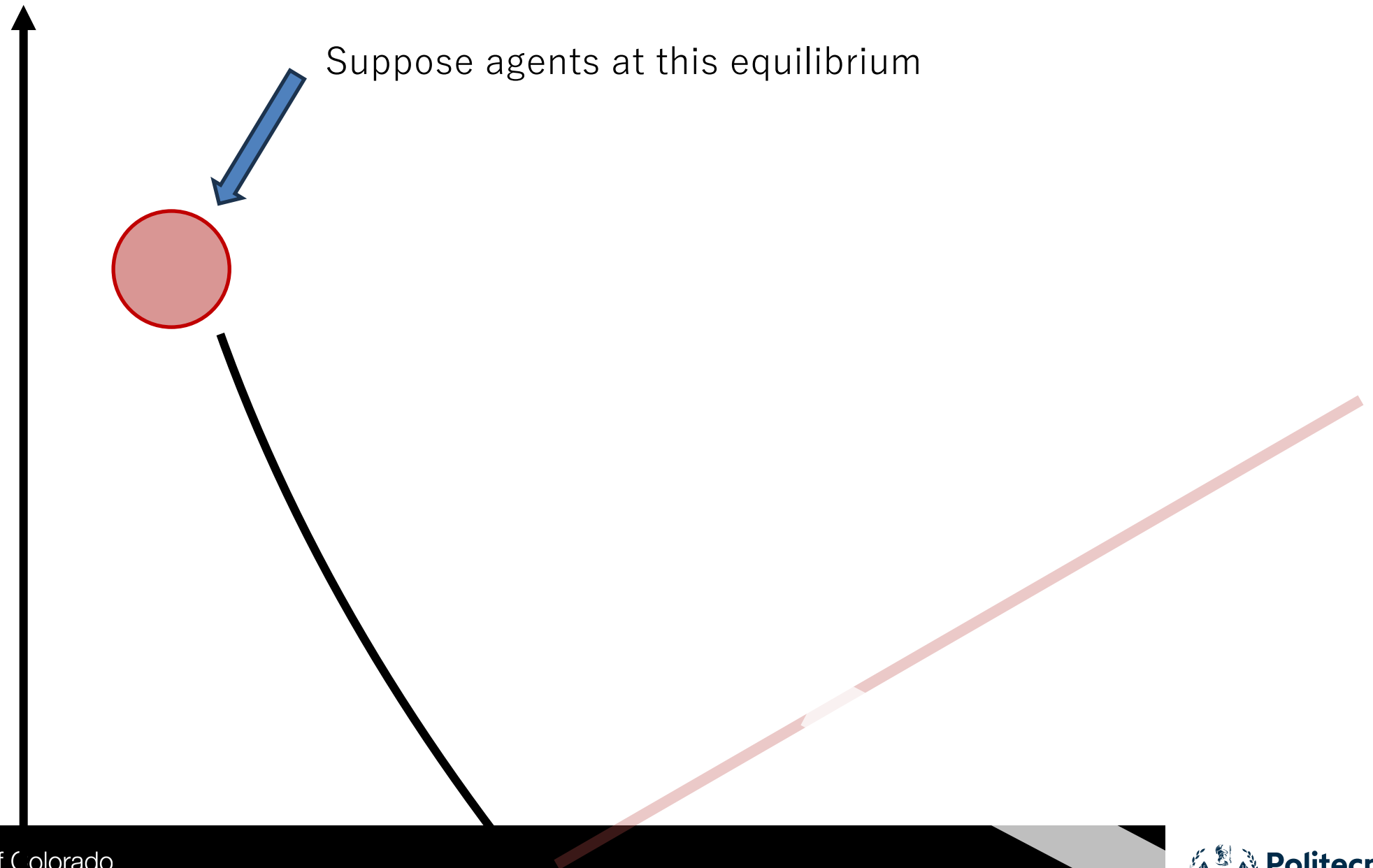
Agenda: Prove that update rule X selects “good” NE in class of games Y

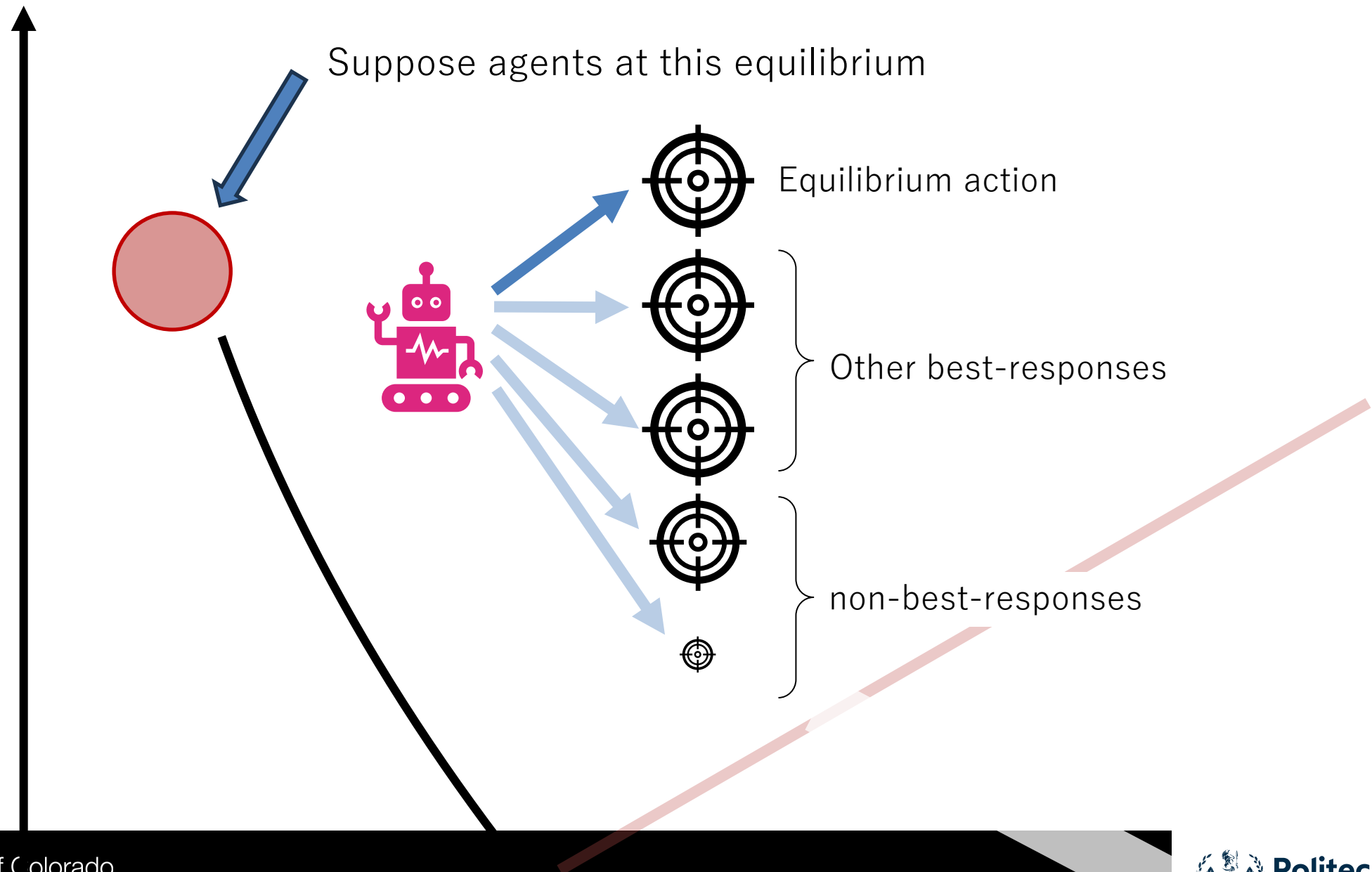
Example: Potential-maximizing NE are stochastically stable under log-linear learning in potential games.

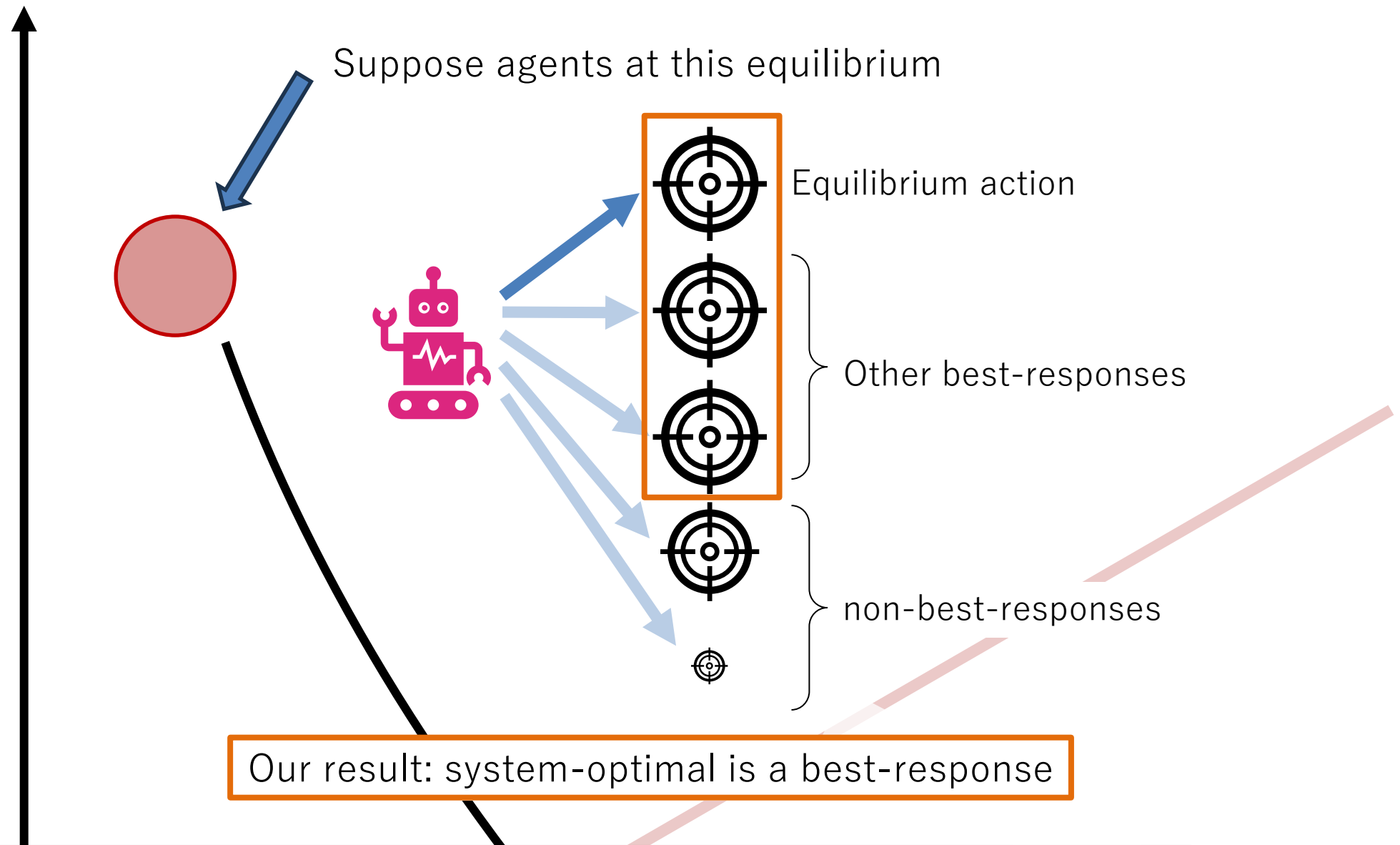
Problem: Not a practical convergence result! (doubly-asymptotic)

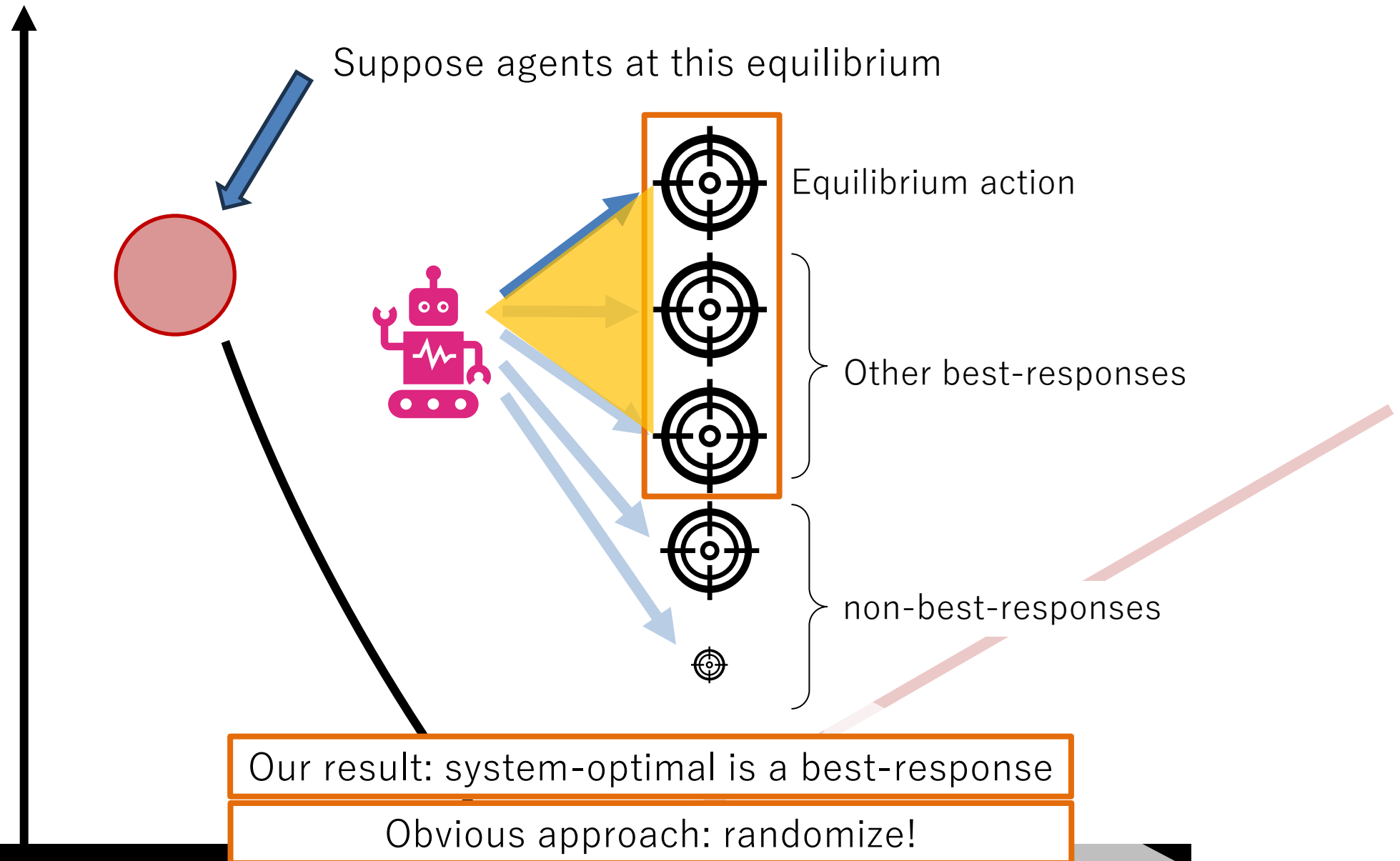
Alós-Ferrer & Netzer, “The Logit-Response Dynamics” GEB’10

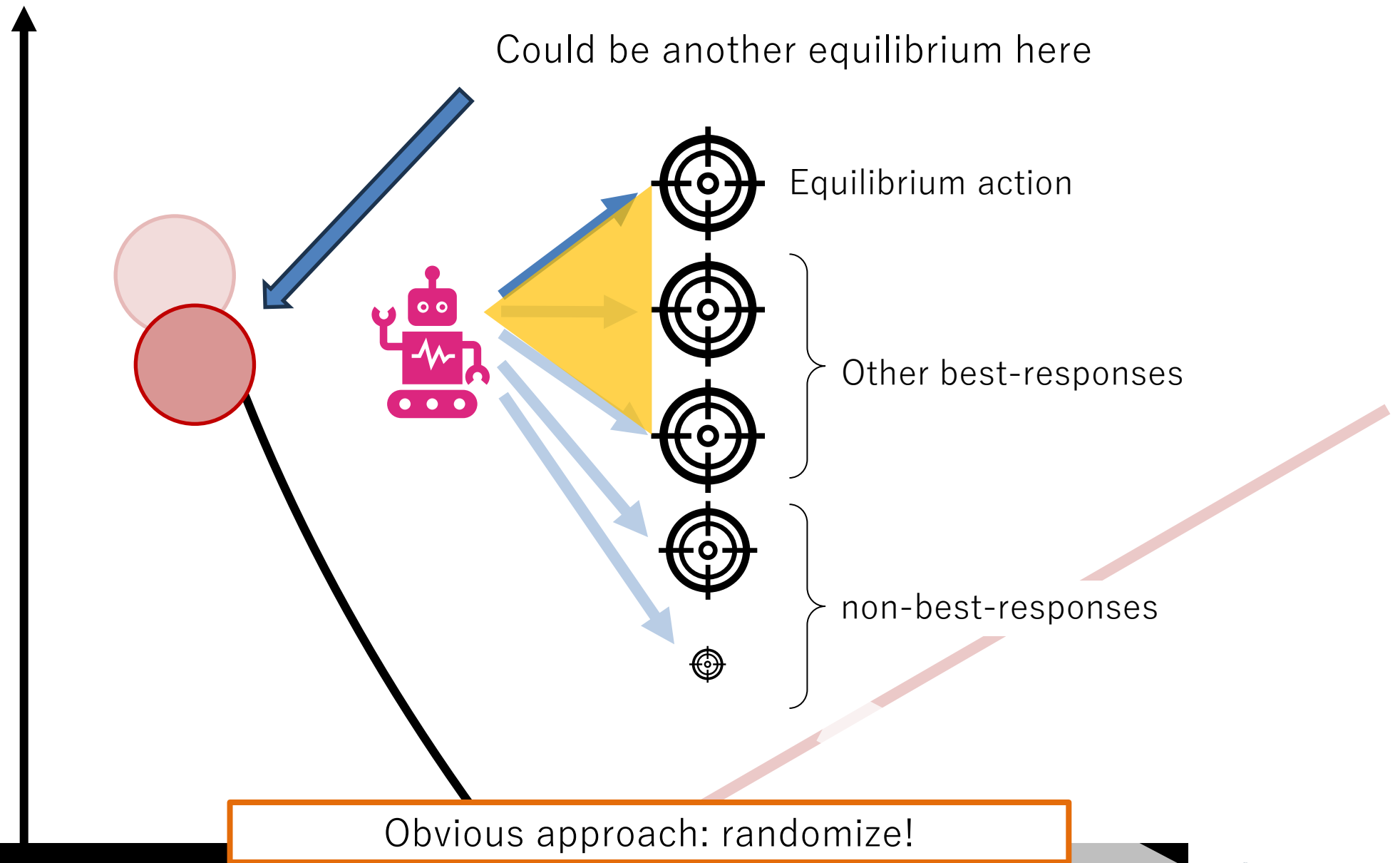


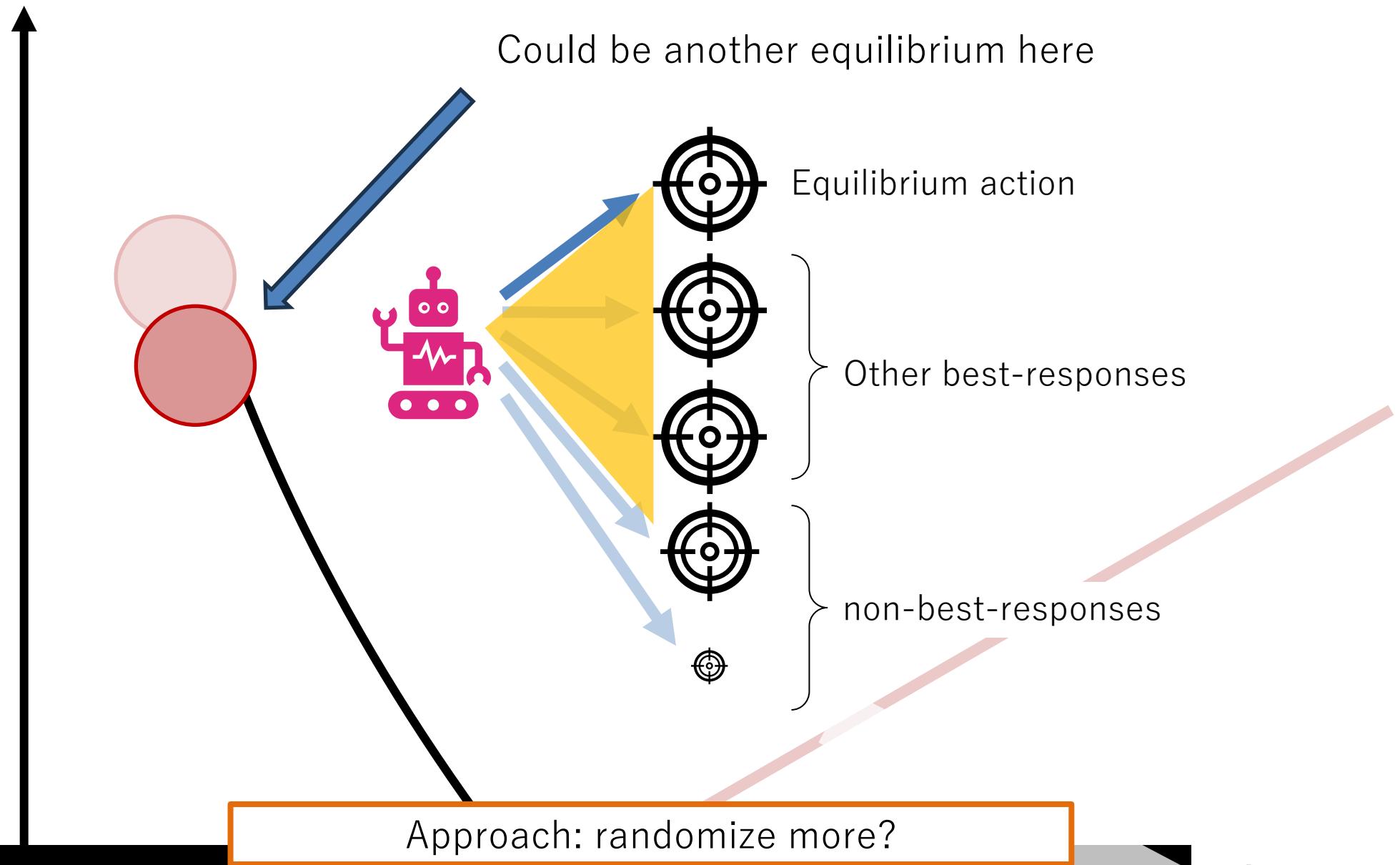




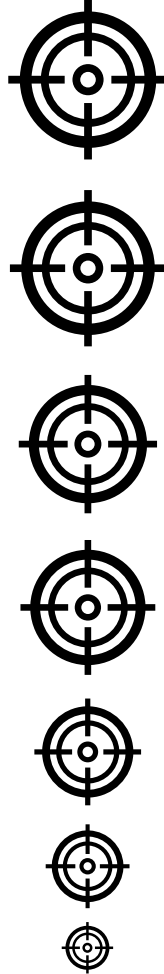
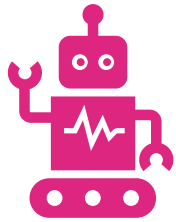








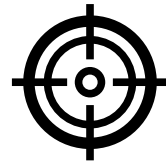
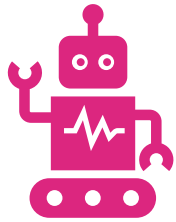
Approach: randomize more?



Algorithm agenda:

- @ time t , random agent wakes up
- Ranks actions by current payoff
- Selects from β -best response set

Approach: randomize more?

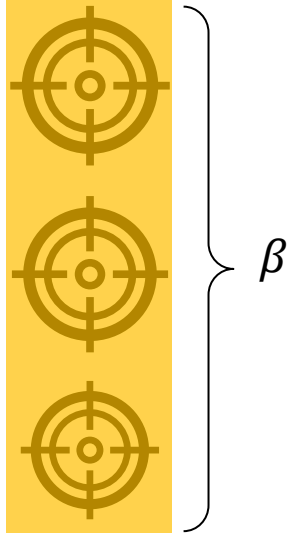
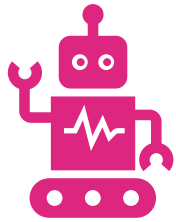


Algorithm agenda:

- @ time t , random agent wakes up
- Ranks actions by current payoff
- Selects from β -best response set

$\beta = 0$: Best-response algorithm

Approach: randomize more?



Algorithm agenda:

- @ time t , random agent wakes up
- Ranks actions by current payoff
- Selects from β -best response set

$\beta = 0$: Best-response algorithm

$\beta > 0$: Truncated Noisy Best-Response (TNBR)
Algorithm

We know:

- $\beta > 0$ cannot converge to worst-case equilibria

We hope:

- $\beta > 0$ avoids all low-quality action profiles

Setup: common-interest weighted set cover games

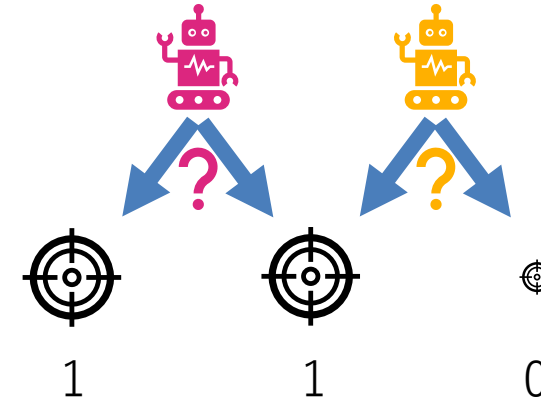
Model: Players: $N = \{1, \dots, n\}$

Resources: $v_r \geq 0 \quad \forall r \in \mathcal{R}$

Actions: $A_i \subseteq 2^{\mathcal{R}}$

Objective: maximize $U(a) = \sum_{r \in \mathcal{R}(a)} v_r$

Common interest utility: $U_i(a) = U(a)$



PoA (due to smoothness and stability): $\text{PoA} \geq \frac{1}{2 + S}$

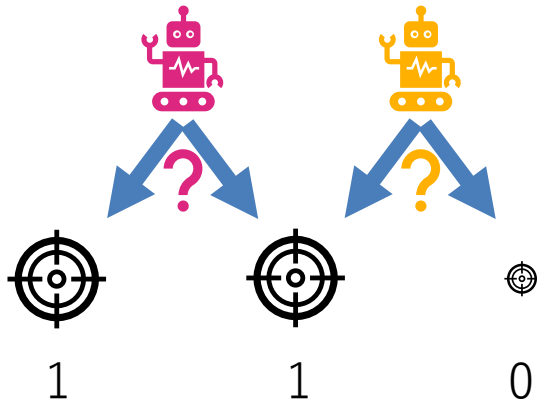
(note: utility-maximization game has $\text{PoA} \leq 1$)

Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$



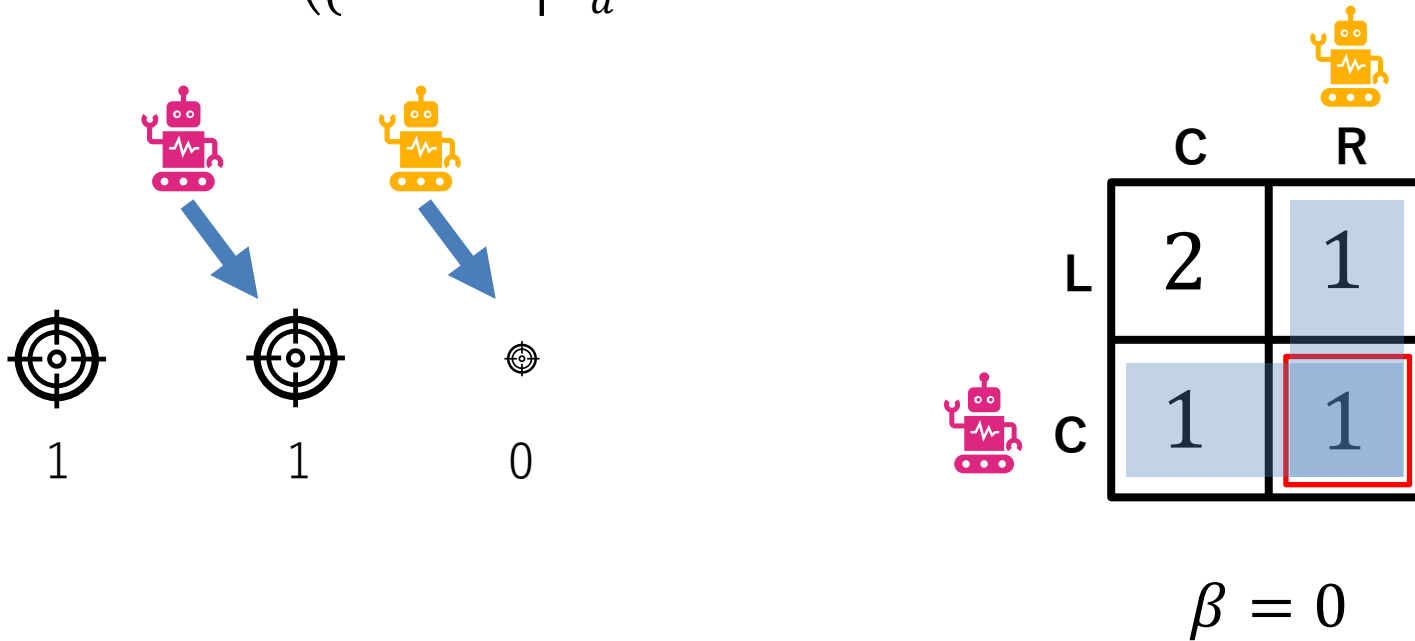
		C	R
			
	L	2	1
	C	1	1
			

Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$

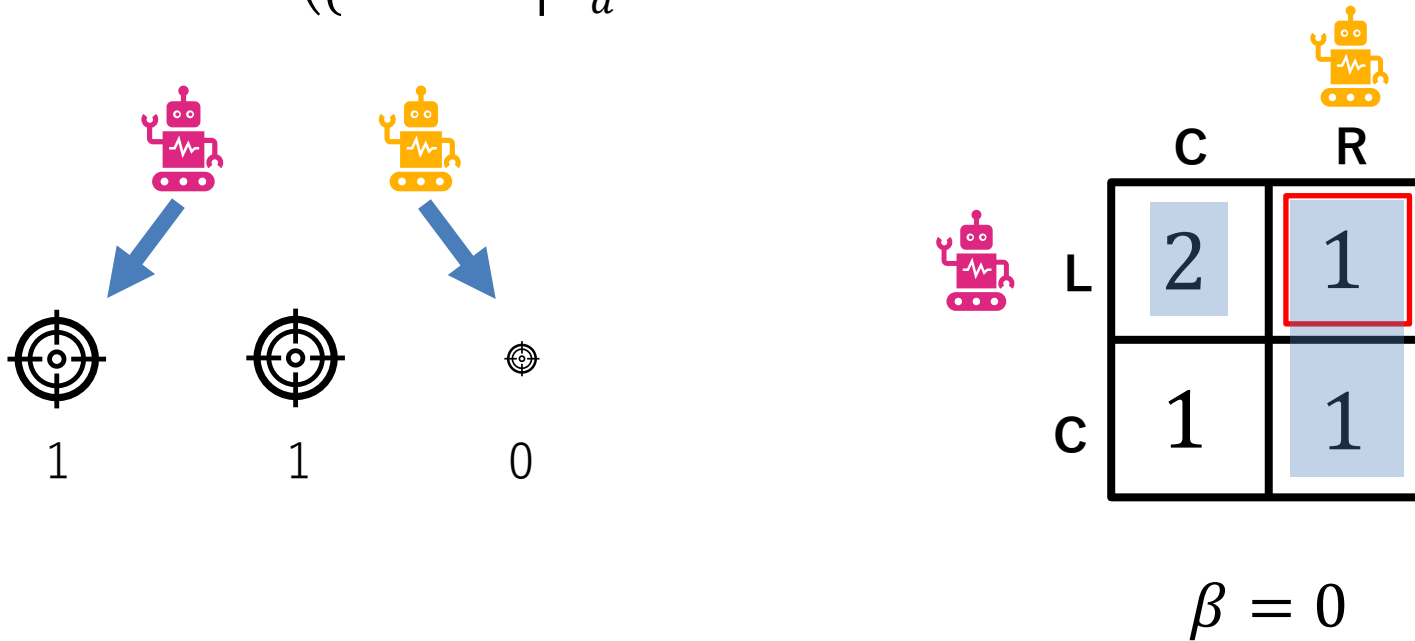


Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$

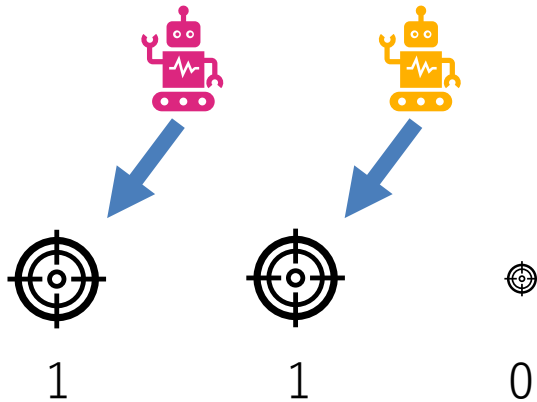


Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$



		C	R
L		2	1
c		1	1

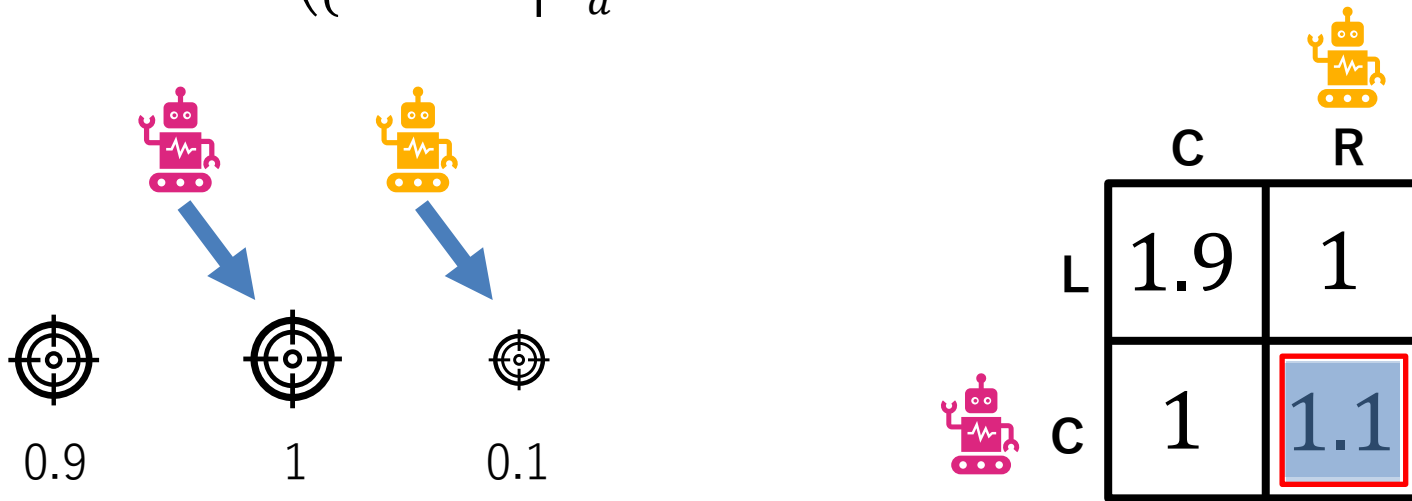
$$\beta = 0$$

Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$



More-stable bad equilibrium?

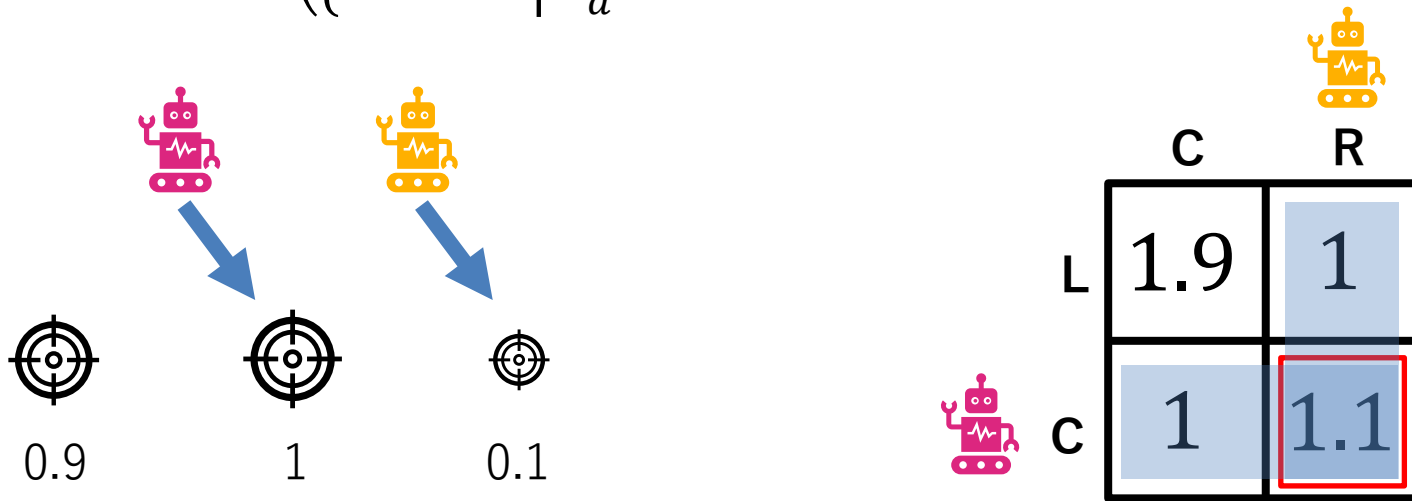
$\beta = 0$

Setup: common-interest weighted set cover games

Algorithm: TNBR (β -best-response)

@ time t , random agent selects from distribution supported on β -neighborhood of best response

$$a_i(t) \sim \Delta \left(\left\{ a \in A_i : \left| \max_{a'} U(a', a_{-i}(t-1)) - U(a, a_{-i}(t-1)) \right| \leq \beta \right\} \right)$$



More-stable bad equilibrium?

$$\beta = 0.1$$

If $\beta \geq 0$, we know worst-case equilibrium is not absorbing state

If $\beta > 0$, how good are TNBR's recurrent states?

Assumptions:

- 2 players
- System objective is normalized: $U(a) \in [0,1]$

Theorem: Let $\mathcal{R}_\beta$ be a recurrent class of TNBR that contains *no* system-optimal state. Then:

$$\forall a \in \mathcal{R}_\beta, \quad U(a) > \frac{1}{2}$$

$$\exists a \in \mathcal{R}_\beta, \quad U(a) > \frac{1}{2} + \beta$$

Singh, **PNB**, “ABRA: An algorithm which cannot converge to low-quality Nash equilibria” CDC 2024

Singh, **PNB**, “Truncated Noisy Best-Response Algorithms: Toward Game Theoretic Learning with Safety Guarantees,” (under review)

If $\beta \geq 0$, we know worst-case equilibrium is not absorbing state

If $\beta > 0$, how good are TNBR's recurrent states?

Assumptions:

- 2 players
- System objective is normalized: $U(a) \in [0,1]$

Theorem: Let $\mathcal{R}_\beta$ be a recurrent class of TNBR that contains *no* system-optimal state. Then:

$$\forall a \in \mathcal{R}_\beta, \quad U(a) > \frac{1}{2}$$

$$\exists a \in \mathcal{R}_\beta, \quad U(a) > \frac{1}{2} + \beta$$

That is: if TNBR **never** visits optimal, it always stays above $\frac{1}{2}$

Singh, **PNB**, “ABRA: An algorithm which cannot converge to low-quality Nash equilibria” CDC 2024

Singh, **PNB**, “Truncated Noisy Best-Response Algorithms: Toward Game Theoretic Learning with Safety Guarantees,” (under review)

Thus: if TNBR ever goes below $\frac{1}{2}$, it visits optimal infinitely often

In fact, we can say more:

Theorem: Let $\mathcal{R}_\beta^*$ be a recurrent class of TNBR that contains a system-optimal state. Then:

$$\forall a \in \mathcal{R}_\beta^*, \quad U(a) > \frac{1}{2} - \frac{3\beta}{2}$$

That is: if TNBR **never** visits optimal, it always stays above $\frac{1}{2}$

Singh, **PNB**, “ABRA: An algorithm which cannot converge to low-quality Nash equilibria” CDC 2024

Singh, **PNB**, “Truncated Noisy Best-Response Algorithms: Toward Game Theoretic Learning with Safety Guarantees,” (under review)

Theorem: if TNBR ever goes below $\frac{1}{2}$, it visits optimal infinitely often
(and never goes below $\frac{1}{2} - \frac{3\beta}{2}$)

In fact, we can say more:

Theorem: Let $\mathcal{R}_\beta^*$ be a recurrent class of TNBR that contains a system-optimal state. Then:

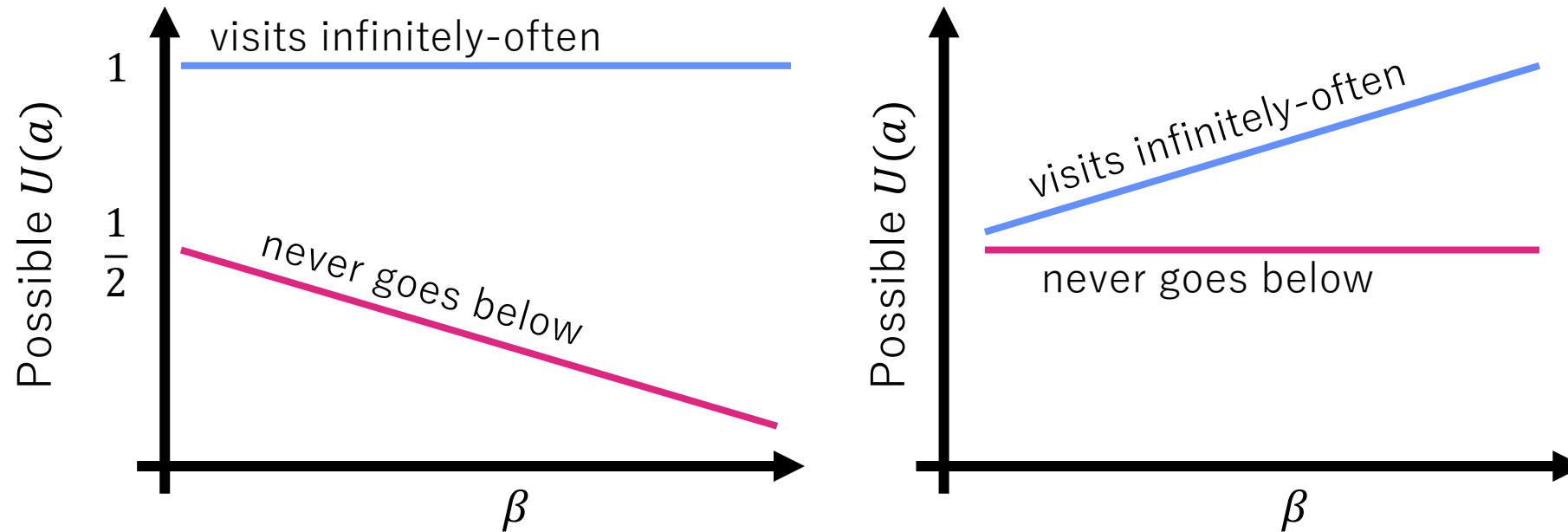
$$\forall a \in \mathcal{R}_\beta^*, \quad U(a) > \frac{1}{2} - \frac{3\beta}{2}$$

Theorem: if TNBR **never** visits optimal, it always stays above $\frac{1}{2}$

Singh, **PNB**, “ABRA: An algorithm which cannot converge to low-quality Nash equilibria” CDC 2024

Singh, **PNB**, “Truncated Noisy Best-Response Algorithms: Toward Game Theoretic Learning with Safety Guarantees,” (under review)

The recurrent classes of TNBR are either:



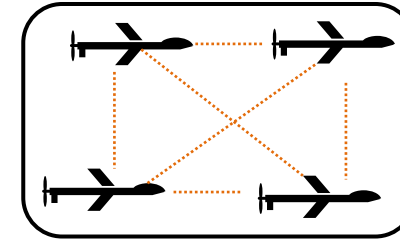
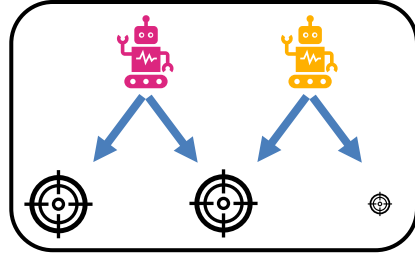
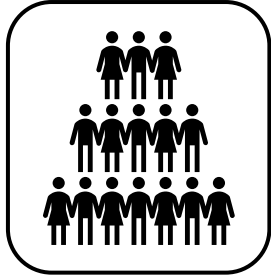
Theorem: if TNBR ever goes below $\frac{1}{2}$, it visits optimal infinitely often
(and never goes below $\frac{1}{2} - 3\beta/2$)

Theorem: if TNBR **never** visits optimal, it always stays above $\frac{1}{2}$

Singh, **PNB**, "ABRA: An algorithm which cannot converge to low-quality Nash equilibria" CDC 2024

Singh, **PNB**, "Truncated Noisy Best-Response Algorithms: Toward Game Theoretic Learning with Safety Guarantees," (under review)

Q: How good is Multiagent Coordination?



A: Bad *only* when brittle! **So what?** Can design better algorithms!

What's next? Mapping from game structure → Good algorithms?
How should “blindness” impact randomization?



Josh Seaton



Vartika Singh

