

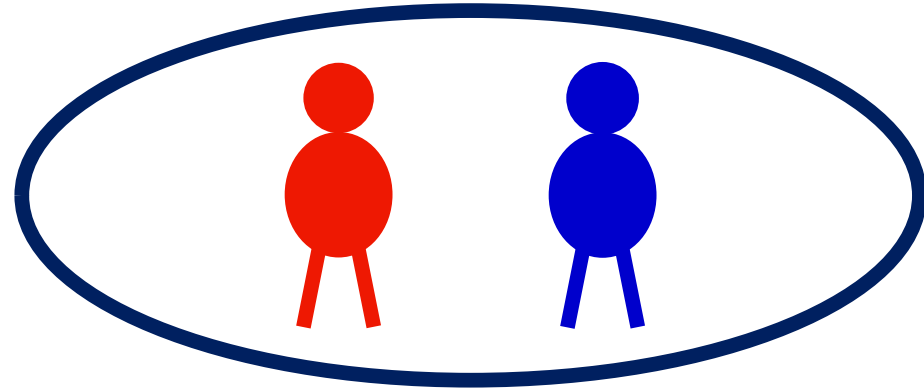
CS 580

Algorithmic Game **Theory**

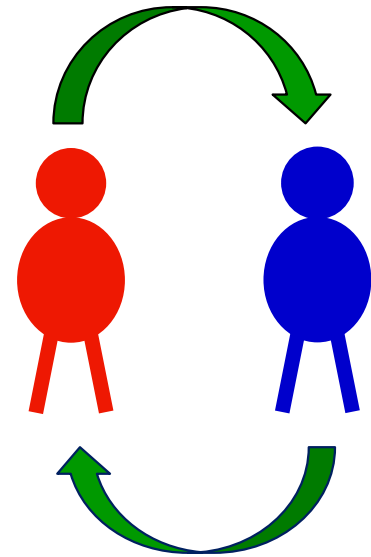
Instructor: Ruta Mehta

Game Theory

Multiple **self-interested** agents interacting in the same environment

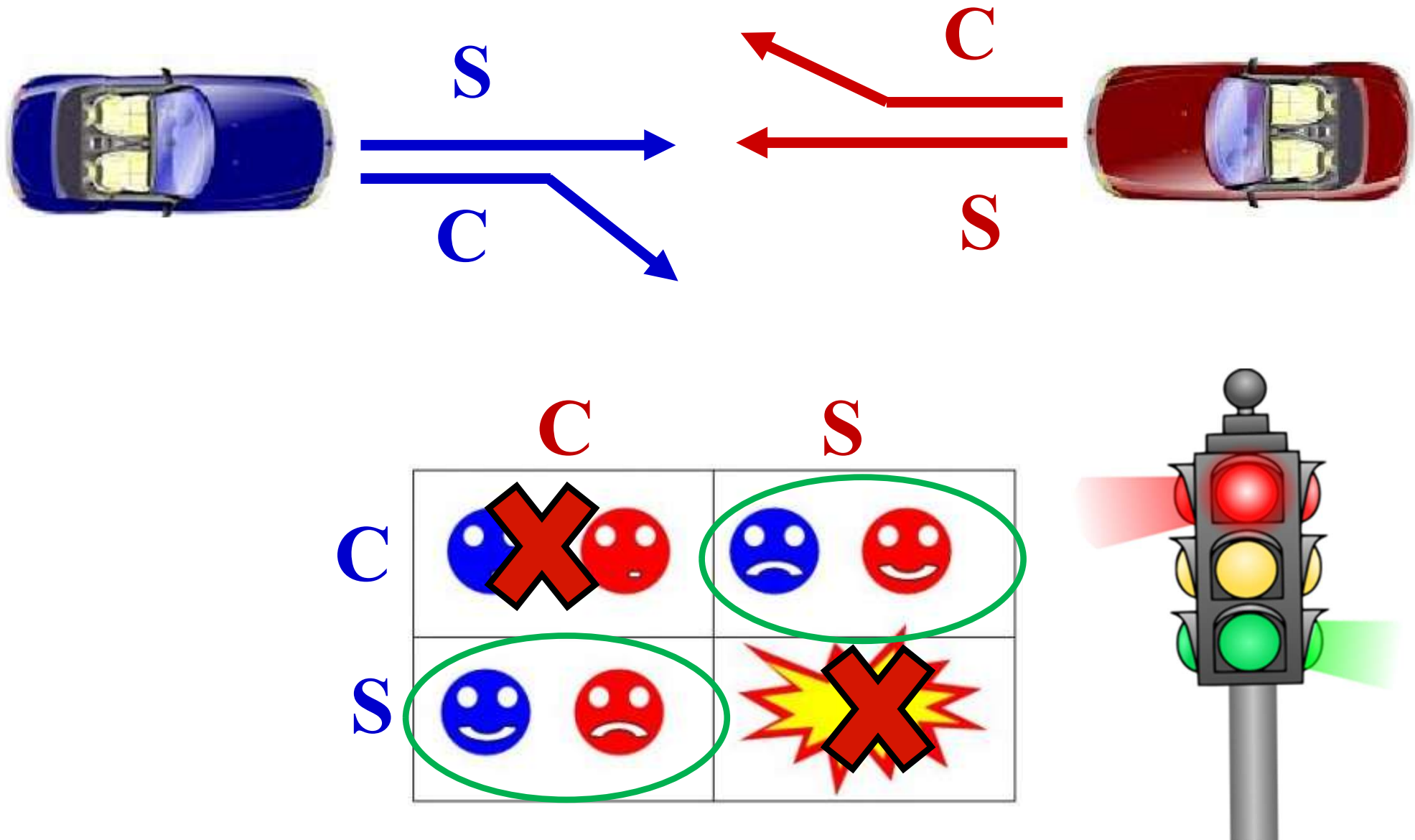


Deciding what **to do**.



Q: What to expect? How good is it? Can it be controlled?

Game of Chicken (Traffic Light)



GenAI Experiment: Delegate the Negotiation

Two people must split \$100.

But each first chooses an AI agent (or prompt) that will negotiate on their behalf.

Human A → Agent A
Choose the agent/prompt
before negotiation begins.

Human B → Agent B
Choose the agent/prompt
before negotiation begins.

- What kind of agent should you choose?
- Can committing to a “tough” agent improve your payoff?
- What equilibrium should we predict when agents are LLMs rather than utility maximizers?
- How should the negotiation protocol be designed?

Games • equilibrium • delegation • learning • mechanism design

Cite: w/ help of ChatGPT Pro



Algorithmic Game Theory

AGT, in addition, focuses on designing efficient algorithms to compute solutions that are crucial (e.g., to make accurate prediction),
and how things change in presence of AI agents.

■ What to expect

Research-oriented Course

- Core concepts and proof techniques from AGT
- Explore classical problems + new questions raised by learning and GenAI agents

■ What is expected from you

- Pre-req: Basic knowledge of linear-algebra, linear programming, probability, algorithms.
- Energetic participation in class
- Research/Survey Project (individually or in a group of two): theory, experiments, or both.

- 
- Instructor: Ruta Mehta (Me)
 - TA: Daniel Lee (tentative!)
 - Office hours:
 - Ruta: Tue 2-3pm in Siebel 3218
 - Daniel: TBD



Useful links

- Webpage:

<https://courses.engr.illinois.edu/cs580/fa2026>

- Piazza Page:

<https://piazza.com/illinois/fall2026/cs580>

- Slack: UIUC-CS580

- Gradescope for grading

Check webpage/piazza at least twice a week for the updates.

HW0 is up.



■ Grading:

- 3 homeworks – 15% (5,5,5)

- LLM Policy: You can use it as long as you cite it.

- Research/Survey Project – 50%

- Work – 25%

- Presentation – 15%

- Report – 10%

- Final Exam or HW4 – 30%

- Class participation – 5%

HW0 is for self-study (not to be submitted).



References

- T. Roughgarden, Twenty Lectures on Algorithmic Game Theory, 2016.
- N. Nisan, T. Roughgarden, E. Tardos, and V. Vazirani (editors), Algorithmic Game Theory, 2007. (Book available online for free.)
- R. Myerson, Game Theory: Analysis of conflict, 1991.

Recent papers on learning, mechanism design, and strategic AI/LLM agents; other notes posted on the course website.



3 Broad Goals — with a new lens: AI as strategic participants

Goal #1

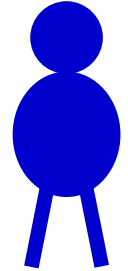
Understand outcomes arising from interaction of strategic agents, including learning algorithms and LLM-based agents.

Games and Equilibria

Prisoner's Dilemma

Two thieves caught for burglary.

Two options: {confess, not confess}

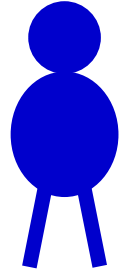


	N	C
N	-1 -1	-6 0
C	0 -6	-5 -5

Prisoner's Dilemma

Two thieves caught for burglary.

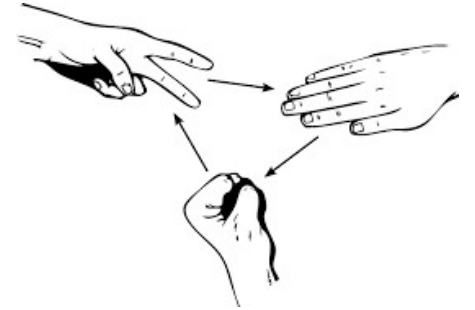
Two options: {confess, not confess}



	N	C
N	-1 -1	-6 0
C	0 -6	-5 -5

Only stable state!

Rock-Paper-Scissors



	R	P	S
R	0 0	-1 1	1 -1
P	1 -1	0 0	-1 1
S	-1 1	1 -1	0 0

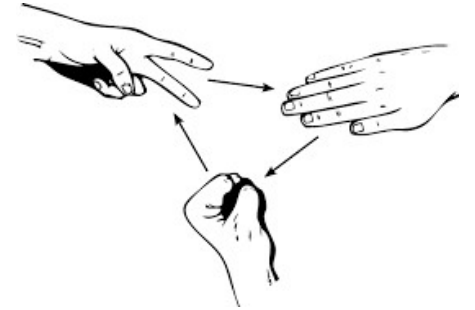
No pure stable state!

Both playing $(1/3, 1/3, 1/3)$
is a NE.

Nash Eq.: No player gains by
deviating individually

Why?

Rock-Paper-Scissors



	R	P	S
R	0 0	-1 1	1 -1
P	1 -1	0 0	-1 1
S	-1 1	1 -1	0 0

No pure stable state!

Both playing $(1/3, 1/3, 1/3)$
is **the only** NE.

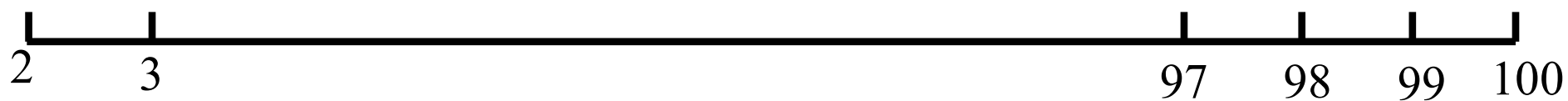
Nash Eq.: No player gains by
deviating individually

Why?

- 
- Finite (normal form) games and Nash equilibrium existence
 - Computation:
 - Zero-sum: minmax theorem,
 - General: (may be) Lemke-Howson algorithm
 - Complexity: PPAD-complete
 - Other equilibrium notions – correlated/coarse correlated, markets, security games
 - Incomplete information, Bayesian Nash
 - Learning in games, delegation, collusion, bargaining
 - How well do AI agents fit classical solution concepts?

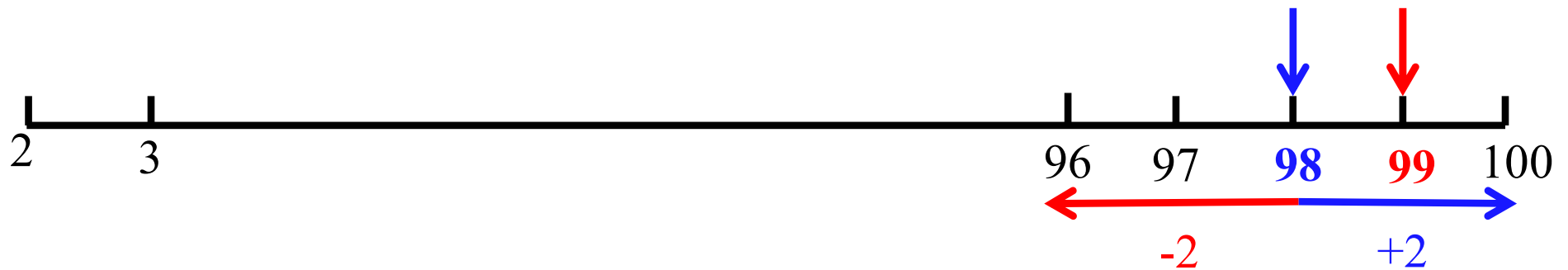
Food for Thought

You and your friend choose a number ...



Food for Thought

You and your friend choose a number ...



What will you choose?

What if ± 50 ?

What are Nash equilibria?

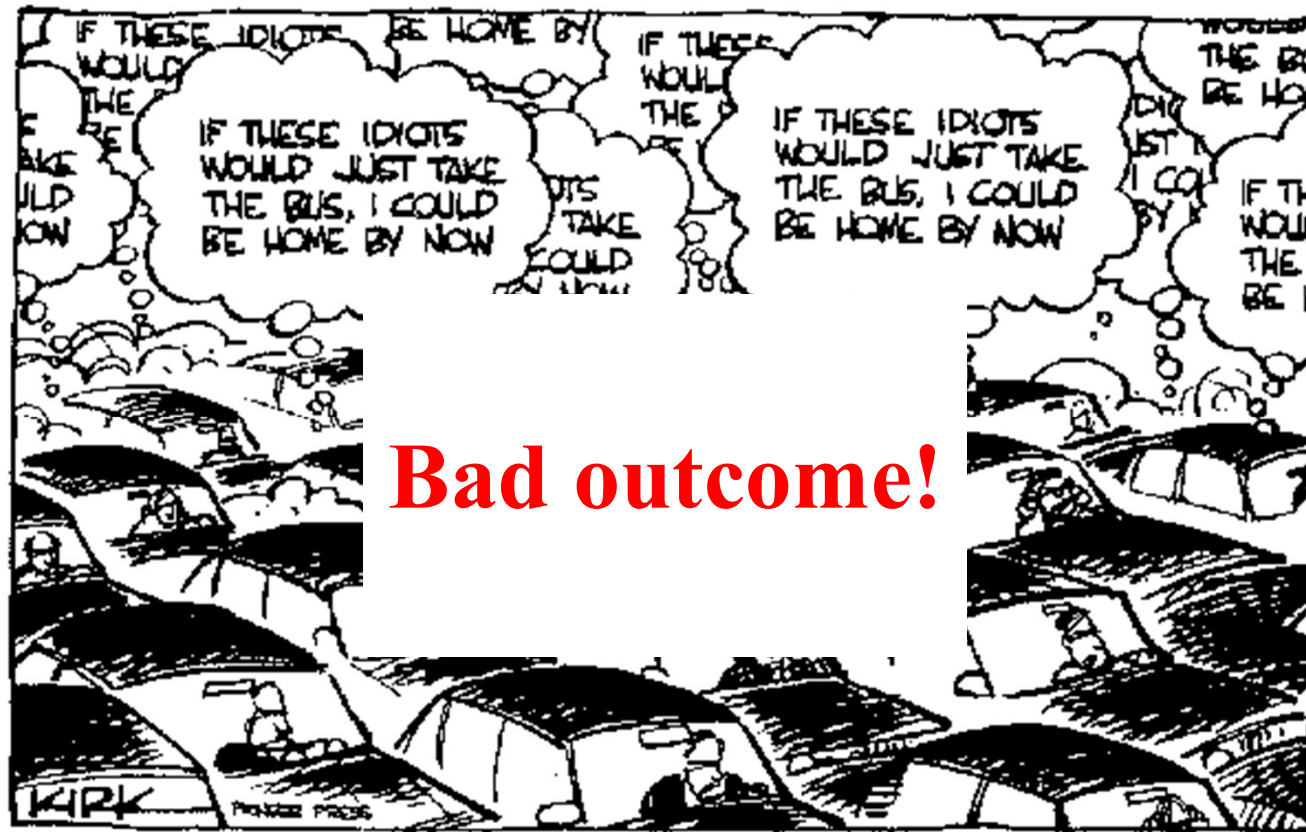
Goal #2

OPT vs equilibrium: Analyze quality of the outcome arising from strategic interaction — including outcomes produced by learning/AI agents.

Price of Anarchy

Tragedy of commons

Limited but open resource shared by many.

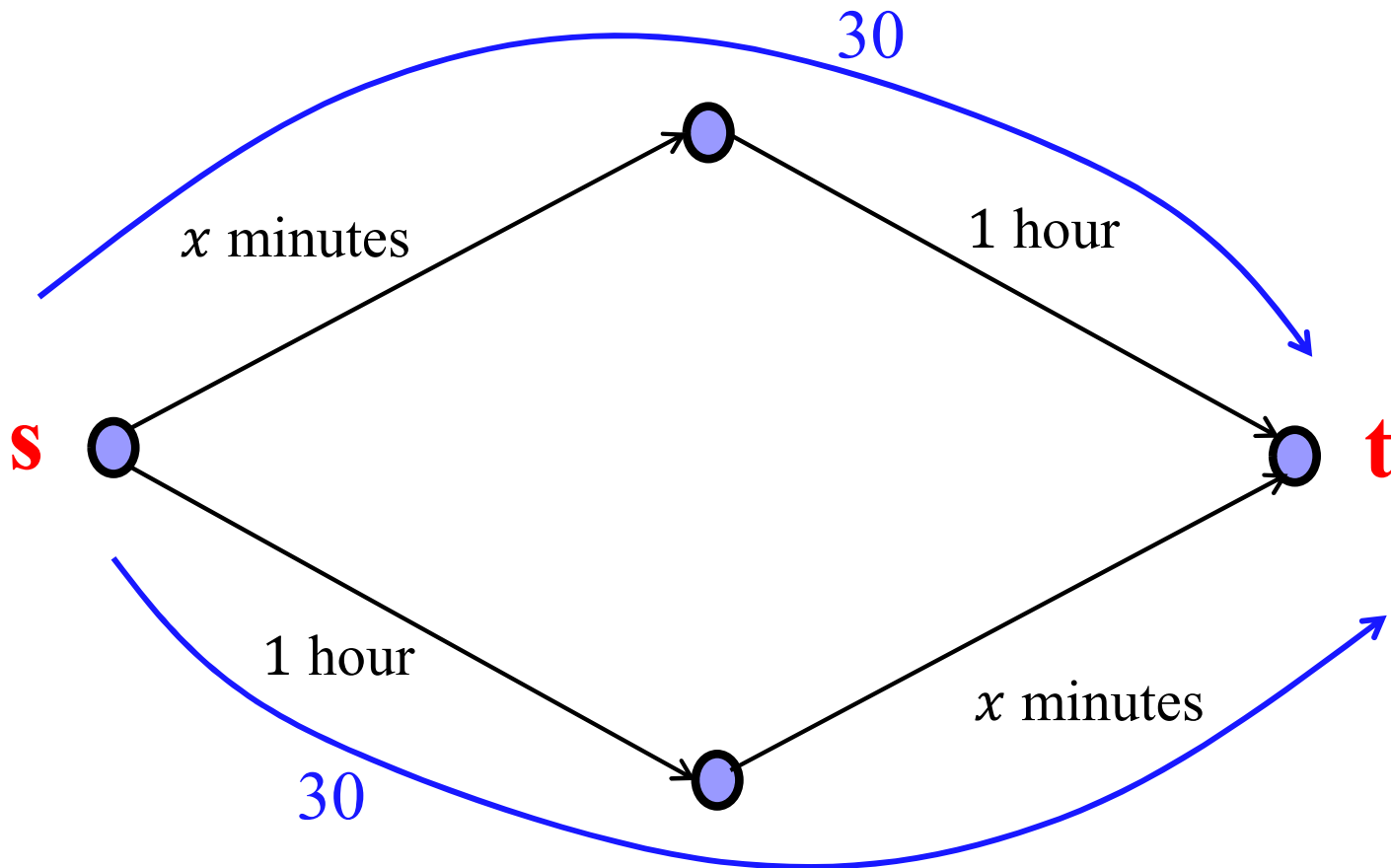


Bad outcome!

Stable: Over use => Disaster

Braess' Paradox

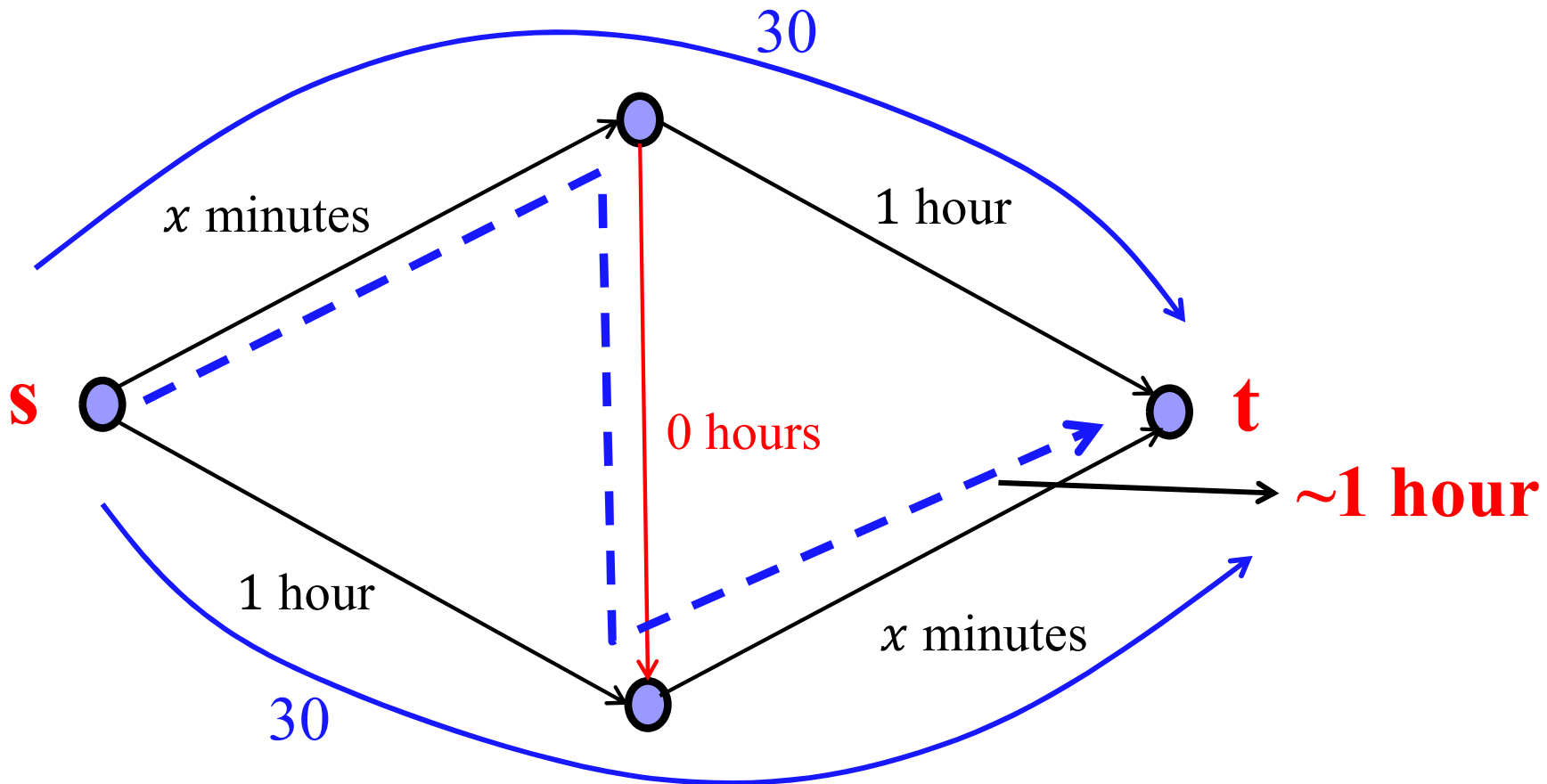
60 commuters



Commute time: 1.5 hours

Braess' Paradox

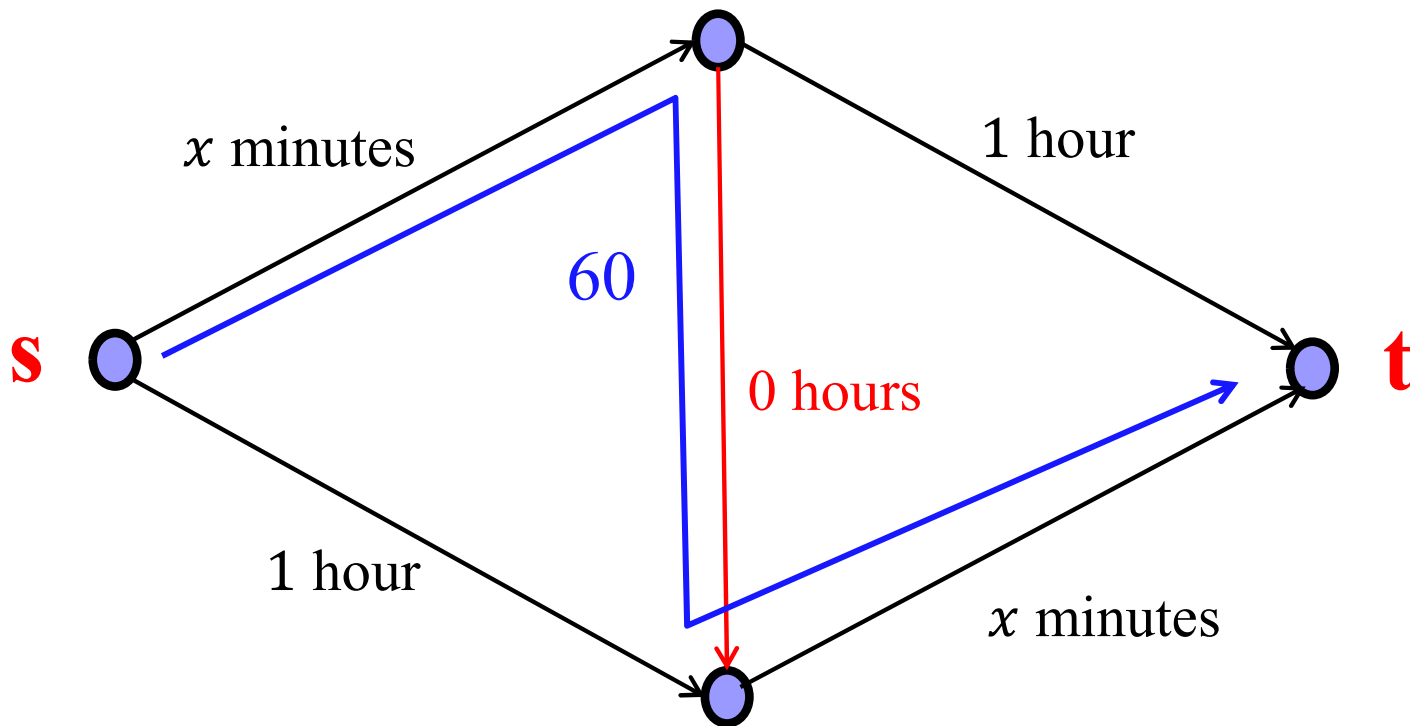
60 commuters



Commute time: 1.5 hours

Braess' Paradox

60 commuters

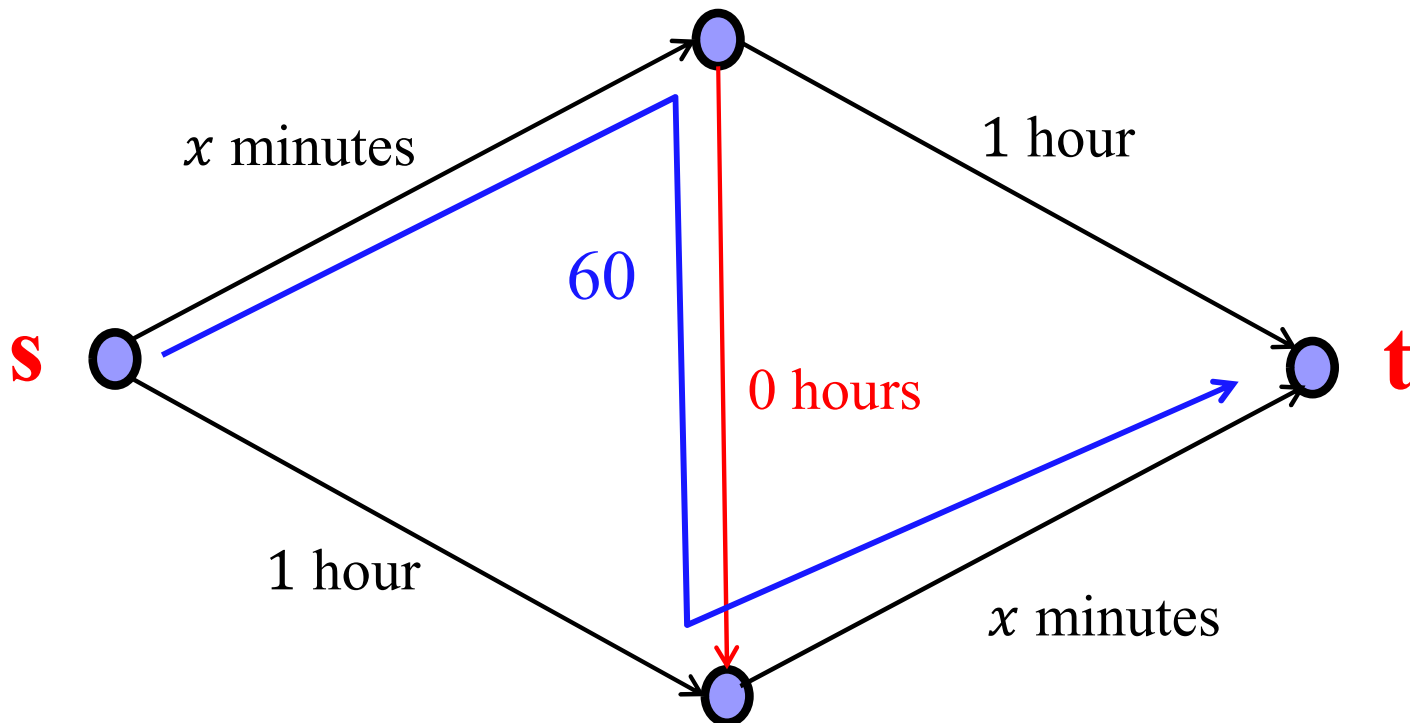


Commute time: 2 hours!

Braess' Paradox in real life

Braess' Paradox

60 commuters



Price of Anarchy (PoA): $\frac{\text{worst NE}}{OPT} = \frac{2}{1.5} = \frac{4}{3}$

Can not be worse!

- 
- Network routing games
 - Congestion (potential) games
 - PoA in linear congestion games
 - Smoothness framework
 - Iterative play, no-regret learning, dynamics and convergence.
 - Strategic adaptation by algorithmic/AI agents

Goal #3

Designing rules to ensure “good” outcome under strategic interaction among humans and increasingly autonomous AI agents.

Mechanism Design — for Human & AI Agents

At the core of large industries

**Online markets – eBay, Uber/Lyft, TaskRabbit,
cloud markets**

**Spectrum auction – distribution of public good.
enables variety of mobile/cable services.**

Search auction – primary revenue for google!

Tons of important applications

**Fair Division – school/course seats assignment,
kidney exchange, air traffic flow management, ...**

**Fair/Responsible ML – Fair FL, Equalizing odds,
Reciprocal fairness, ...**

**Matching residents to hospitals,
Voting, review, coupon systems.**

So on ...



■ MD without money

☐ Fair division

- Divisible items: Competitive equilibrium
- Indivisible items: EF1, EFX, MMS, Max. Nash Welfare, ...

☐ Stable matching, Arrow's theorem (voting)

■ MD with money

- ☐ First price auction, second price auction, VCG
- ☐ Generalized second price auction for search (Google)
- ☐ Optimal auctions: Myerson auction and extensions
- ☐ Prophet inequalities and simple auctions
- ☐ Mechanism design for AI agents
- ☐ Delegation and agent markets
- ☐ Strategic learning / collusion

Fun Fact!

Olympics 2012 Scandal

**Check out Women's doubles badminton
tournament**

Video of the fist controversial match

The Course Agenda

1. Predict — Games, Nash/CE/CCE, markets
What will strategic agents do?

2. Evaluate — Price of Anarchy, welfare, fairness
How good is the resulting outcome?

3. Design — Auctions, matching, mechanism design
Can we change the rules to get better outcomes?

4. Revisit with AI — Learning, delegation, LLM agents, agent markets
Which assumptions and guarantees survive?