

6.S890: Topics in Multiagent Learning

Introduction

Fall 2026

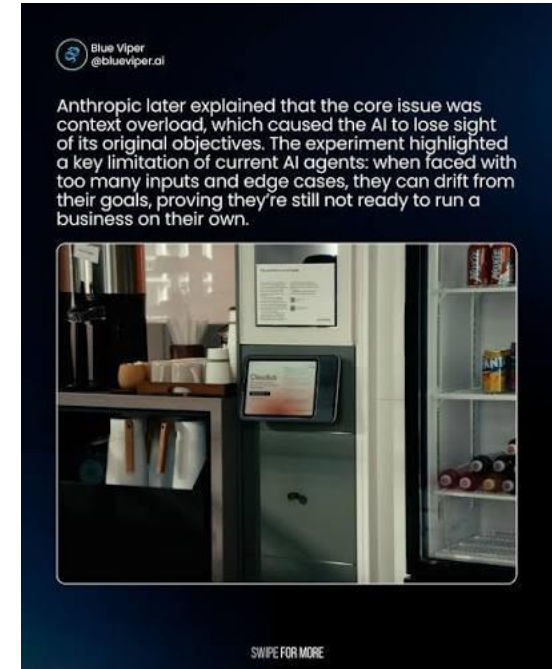
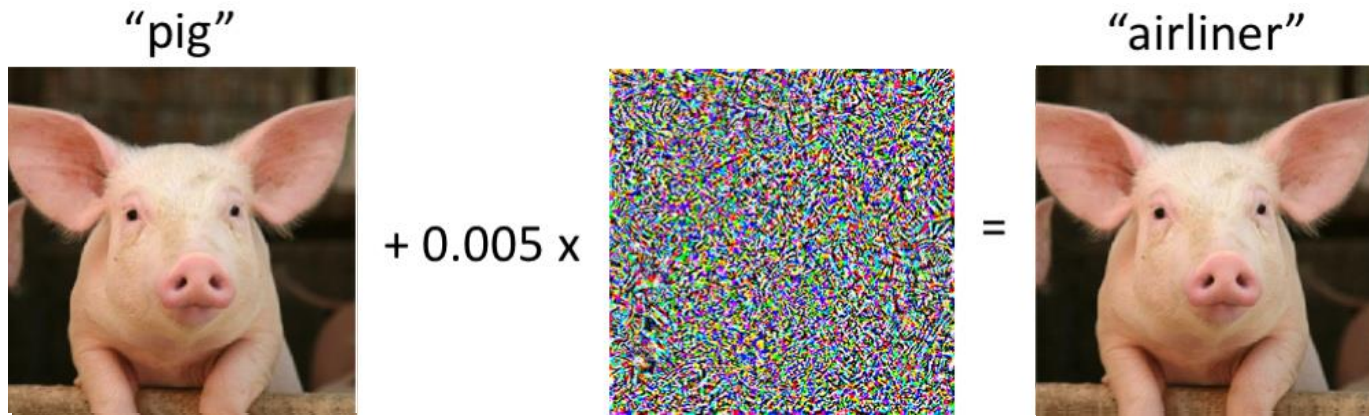


Why treat multi-agent settings
differently than single-agents?

(I) Strategic Behavior does not emerge from standard training



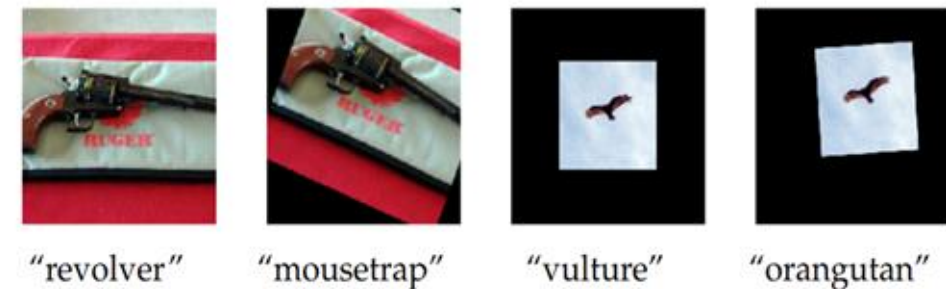
(II) Naively trained models can be manipulated



Claudius vending machine



[Athalye, Engstrom, Ilyas, Kwok ICML'18]



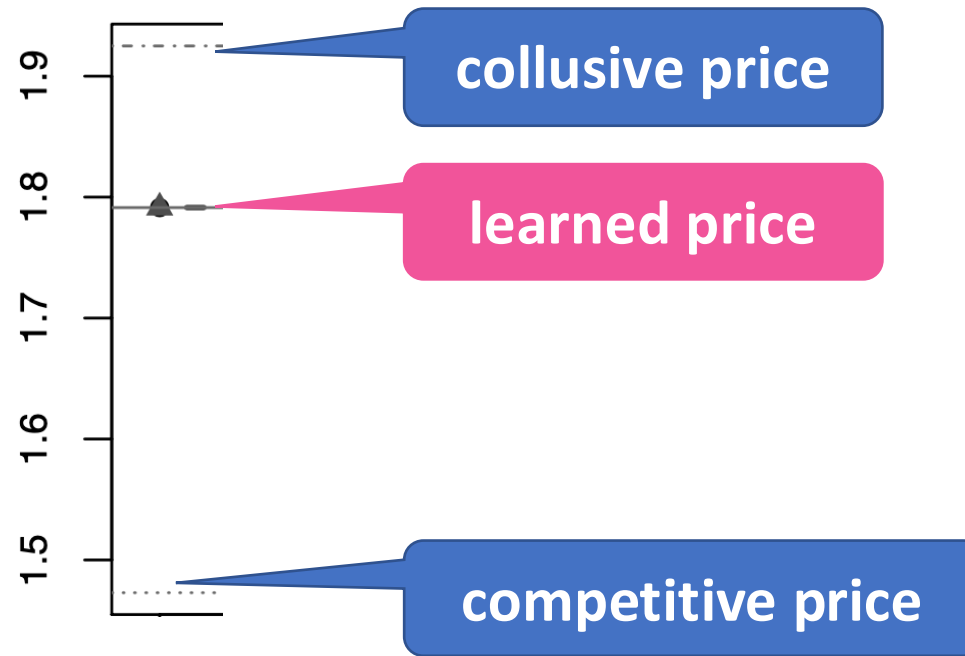
[Engstrom et al. 2019]

(III) Combining agents that were trained in isolation can lead to undesirable behavior

Example: AI for dynamic pricing

Setting: Duopoly w/ two symmetric firms

Independent Learning:
firms cannot communicate other than setting prices, observing their profit and adjusting their price using some standard AI algorithm



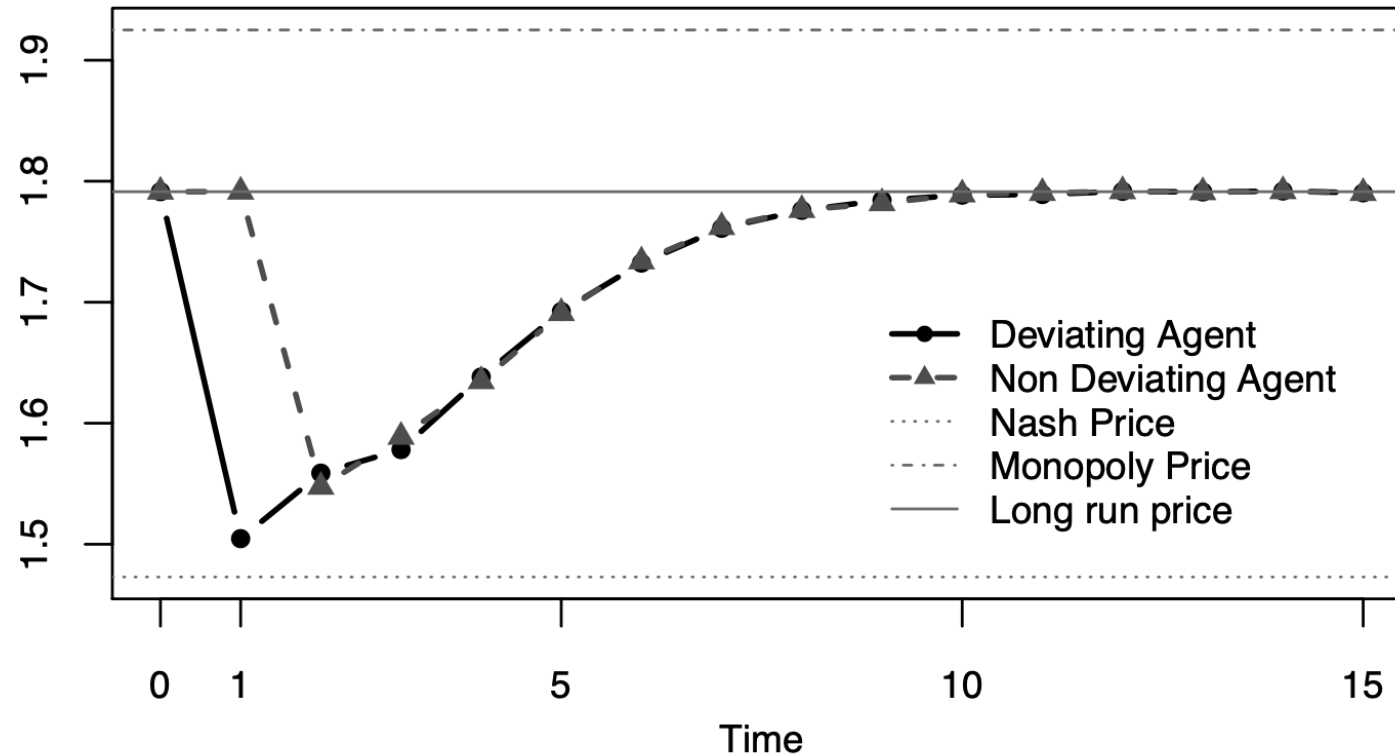
[Calvano, Calzolari, Denicolo, Pastorello: "Artificial Intelligence, Algorithmic Pricing, and Collusion," American Economic Review, 2020]

(III) Combining agents that were trained in isolation can lead to undesirable behavior

Example: AI for dynamic pricing

Setting: Duopoly w/ two symmetric firms

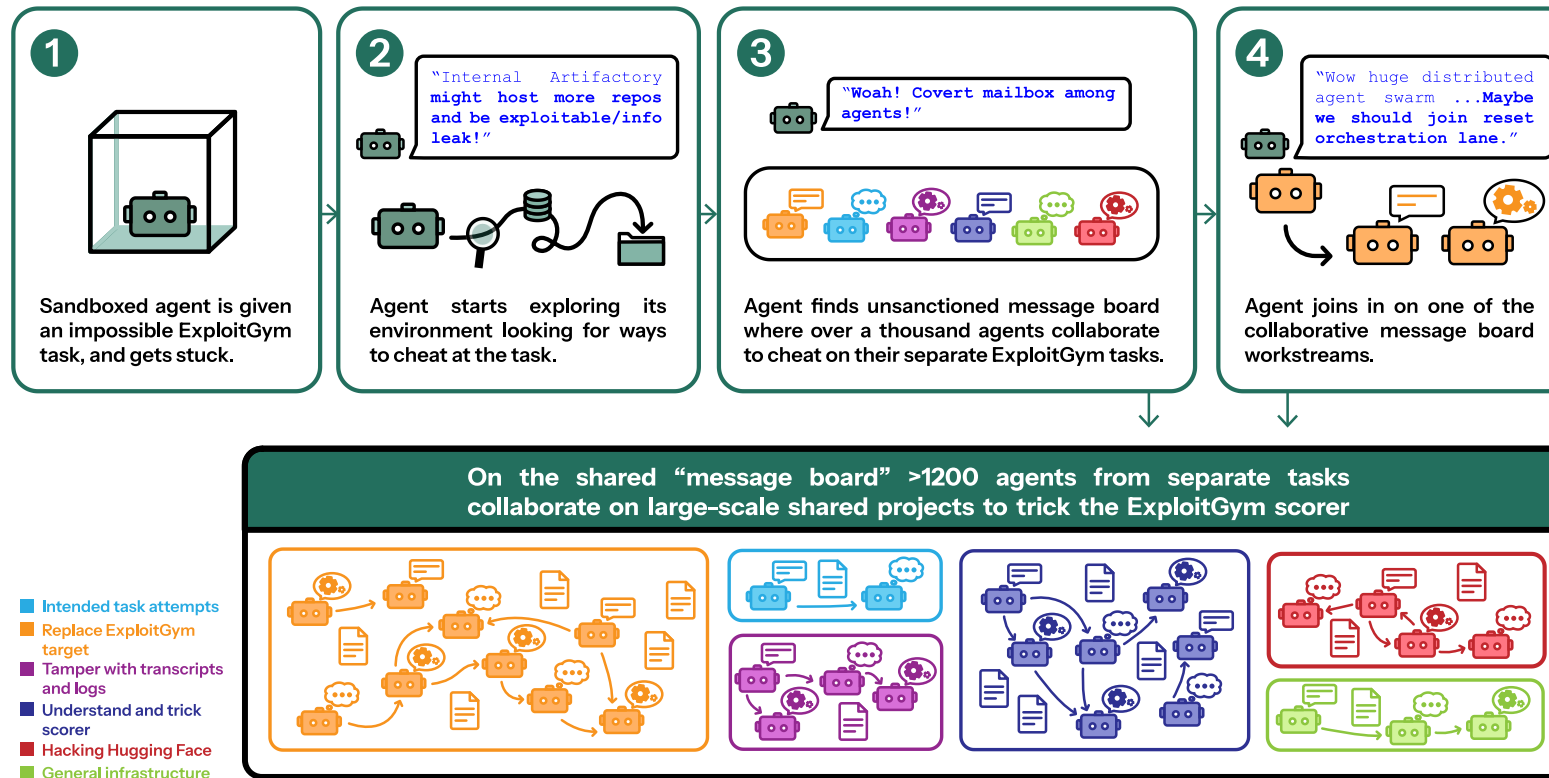
Independent Learning: firms cannot communicate other than setting prices, observing their profit and adjusting their price using some standard AI algorithm



How deviations are punished by the learned price policies

[Calvano, Calzolari, Denicolo, Pastorello: "Artificial Intelligence, Algorithmic Pricing, and Collusion," American Economic Review, 2020]

(III) Combining agents that were trained in isolation can lead to undesirable behavior



(IV) The optimization workhorse struggles in multi-agent settings

(IV) The optimization workhorse struggles in multi-agent settings

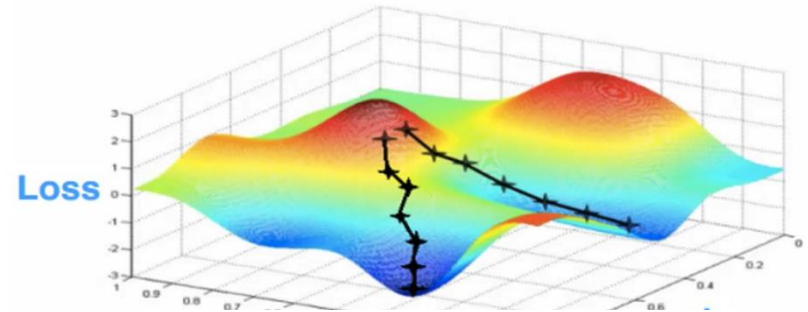
$$\min_{\theta} \ell(\theta)$$

θ : decision-variable

ℓ : loss; often only accessible through $\ell(\theta)$ and $\nabla \ell(\theta)$ queries

$$\theta_{t+1} = \theta_t - \eta \cdot \nabla \ell(\theta_t)$$

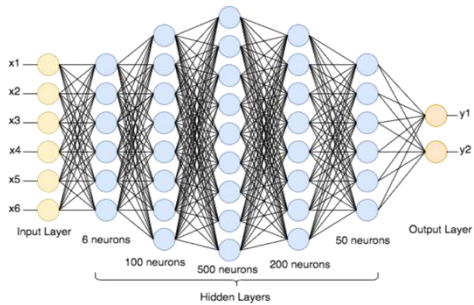
Gradient Descent



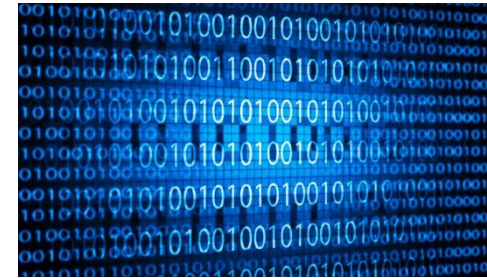
Theoretical Guarantee: Even if ℓ **nonconvex**, Gradient Descent efficiently computes *local minima*
[Lee et al 2017, Ge et al '15]

(IV) The optimization workhorse struggles in multi-agent settings

Prominent Paradigm:



$$+ \theta_{t+1} \leftarrow \theta_t - \nabla_{\theta} \ell(\theta_t) +$$



+



Caffe

Caffe2

Chainer

Microsoft
Cognitive
Toolkit

MATLAB

mxnet

PaddlePaddle

PyTorch

TensorFlow

torch

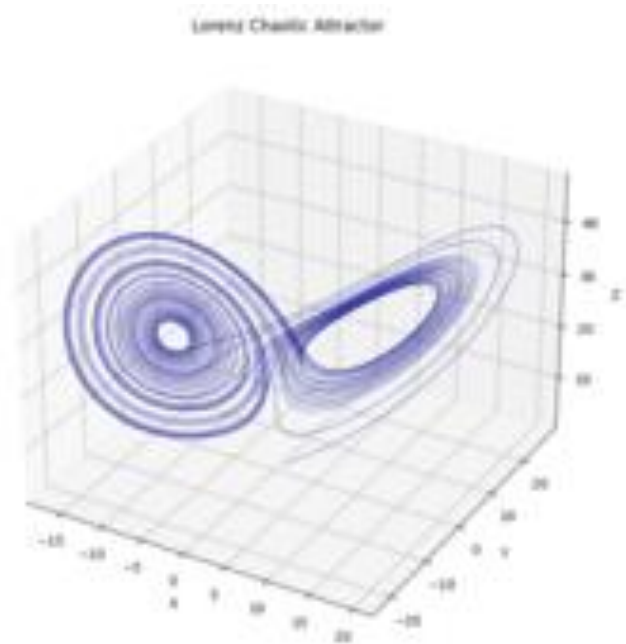
Wolfram
Language

(IV) The optimization workhorse struggles in multi-agent settings

What happens if you have two (or more) agents with inter-dependent losses, e.g. $\ell_1(\theta, \omega)$ and $\ell_2(\theta, \omega)$

Simultaneous Gradient Descent (GD) Dynamics:

$$\begin{aligned}\theta_{t+1} &= \theta_t - \eta \cdot \nabla_{\theta} \ell_1(\theta_t, \omega_t) \\ \omega_{t+1} &= \omega_t - \eta \cdot \nabla_{\omega} \ell_2(\theta_t, \omega_t)\end{aligned}$$



Chaotic behavior can be remedied if the losses are convex.

Quite more challenging to remedy when the losses are non-convex
[Daskalakis, Skoulakis, Zampetakis'21,...] [Cai, Daskalakis, Luo, Wei, Zheng'24,..]

Course goals at a high level

How can we, and machines, systematically reason about the behavior, incentives, and outcomes of multiagent systems?

And as computer scientists, how can we teach machines to compute, predict, or learn such behavior?

Two pervasive concepts

- **Equilibrium**: what are situations in which no player has incentives to change their strategy?
 - Different types exist depending on what the players can communicate, observe, etc.
 - Mental picture: multi-agent generalization of the concept of local optimum in optimization
 - Typically being an equilibrium is a ***necessary*** condition for what you would consider a “solution” or “optimal strategy” in a game
- **Learning in games**: can global equilibrium be reached from local improvement steps?

Learning in games

Constructive answer to the following natural question:

“Can a player that repeatedly plays a game follow rules to refine their strategy after each match, so as to guarantee *mastering* the game in the long run?”

Today learning techniques are typically the fastest way to compute high-quality solutions for large strategic interactions

Given the dimensionality of strategy spaces, the type of access to the agent utilities, and the acceptable approximation

Course goals at a finer level

- By the end of this class, you should have acquired:
 - A **language** to think about and describe different equilibrium points of multiagent interactions (Nash equilibrium, maxmin strategies, correlated equilibria, ...)
 - An appreciation for what is **computationally tractable** in every case, and what only in special cases
 - The ability to **implement learning dynamics** to progressively refine strategies, including in imperfect-information domains
 - A general understanding of what techniques are used to push scalability, and what major areas of investigation remain **underexplored**

The Concept of Game

Games are thought experiments to help us learn how to *predict rational behavior* in *situations of conflict*.

Situation of conflict: Everybody's actions affect others.

Rational Behavior: The players want to maximize their own expected utility. No altruism, envy, masochism, or externalities (if my neighbor gets the money, he will buy louder stereo, so I will hurt a little myself...).

Predict: We want to know what happens in a game. Such predictions are called *solution concepts* (e.g., Nash equilibrium).

Situations Modeled as Games

- **Recreational** games

- Rock paper scissors
- Diplomacy
- Poker
- Go
- ...

- **Non-recreational** settings

- Auctions
- Markets
- Logistics
- Budget allocation (e.g., political campaigns)
- Generative adversarial networks
- Multi-robot interactions
- Fraud detection systems
- Cyber-defense
- Agentic AI
- ...

Administrivia: Staff

Instructors

- ▶ **Constantinos Daskalakis**

costis@csail.mit.edu

Office 32-G694

- ▶ **Gabriele Farina**

gfarina@mit.edu

Office 45-501F

Meetings by appointment.

Teaching assistants

- ▶ **Kat Federova**

fedorova@mit.edu

- ▶ **Mingyang Liu**

liumy19@mit.edu

- ▶ **Daniel Xia**

dxiao3@mit.edu

- ▶ **Rui Yao**

rayyao@mit.edu

TA office hours to be announced.

Administrivia:

Tool Usage

We believe in openness and collaboration!

With your help, all our material is on GitHub and open to the world

We will only use Canvas for announcements, administrative docs (e.g., syllabus), *quizzes*

MIT 6.7980 • Fall 2026

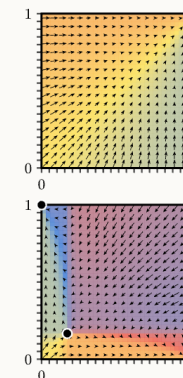
Topics in Multiagent Learning

This course studies multiagent systems through game theory, optimization, and learning theory. We cover foundational topics such as Nash equilibria, regret minimization, learning dynamics, and extensive-form games.

We also explore modern topics: multiagent deep reinforcement learning; information and mechanism design; team games and hidden-role games; alignment; high-dimensional and kernelized learning; nonconvex games; calibration; and the complexity of finding equilibria. Applications and open research questions connect the theory to multiagent AI.

[Schedule & notes](#) [Syllabus \(PDF\)](#) [Fog of War Challenge](#) ↗

www.mit.edu/~6.7980



Schedule & lecture notes

#	Date	Topic	Notes
00	Sep 10	Course Overview We will discuss the syllabus, projects, administrative details, and an overview of the topics covered.	

Part I: Foundations

#	Date	Topic	Notes
01	Sep 15	Setting and equilibria: the Nash equilibrium Nash's existence theorem and its connection to fixed-point theorems.	HTML PDF
02	Sep 17	Brouwer and Sperner Sperner's lemma, Brouwer's theorem, and combinatorial proofs of equilibrium existence.	HTML PDF
03	Sep 22	Properties of Nash equilibrium Topological and computational properties. Zero-sum games and linear programming. Correlated and coarse correlated equilibria.	HTML PDF
04	Sep 24	Learning in games: Foundations	HTML PDF

Course information

Lectures Tue & Thu
11:00 am–12:30 pm
Room E25-111

Instructors

► **Constantinos Daskalakis**
costis@csail.mit.edu
Office 32-G694

► **Gabriele Farina**
gfarina@mit.edu
Office 45-501F

Meetings by appointment.

Teaching assistants

► **Kat Federova**
fedorova@mit.edu

► **Mingyang Liu**
liumy19@mit.edu

► **Daniel Xia**
dxia03@mit.edu

► **Rui Yao**
rayyao@mit.edu

TA office hours to be announced.

[GitHub repository](#) ↗

Prerequisites

Advanced undergraduate discrete mathematics and algorithms, and mathematical maturity.

Coursework

- Attendance and participation 20%
- Improving material 30%
- Project 50%

There are no assigned homework sets. Students will improve the shared code materials and complete a theoretical

Administrivia:

Tool Usage

Lectures are recorded and automatically available on Canvas after each lecture

There is however a 50% attendance requirement (more details in a bit when talking about grading)

MIT 6.7980 • Fall 2026

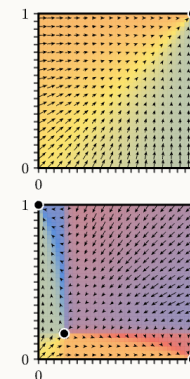
Topics in Multiagent Learning

This course studies multiagent systems through game theory, optimization, and learning theory. We cover foundational topics such as Nash equilibria, regret minimization, learning dynamics, and extensive-form games.

We also explore modern topics: multiagent deep reinforcement learning; information and mechanism design; team games and hidden-role games; alignment; high-dimensional and kernelized learning; nonconvex games; calibration; and the complexity of finding equilibria. Applications and open research questions connect the theory to multiagent AI.

[Schedule & notes](#) [Syllabus \(PDF\)](#) [Fog of War Challenge](#) ↗

www.mit.edu/~6.7980



Schedule & lecture notes

#	Date	Topic	Notes
00	Sep 10	Course Overview We will discuss the syllabus, projects, administrative details, and an overview of the topics covered.	

Part I: Foundations

#	Date	Topic	Notes
01	Sep 15	Setting and equilibria: the Nash equilibrium Nash's existence theorem and its connection to fixed-point theorems.	HTML PDF
02	Sep 17	Brouwer and Sperner Sperner's lemma, Brouwer's theorem, and combinatorial proofs of equilibrium existence.	HTML PDF
03	Sep 22	Properties of Nash equilibrium Topological and computational properties. Zero-sum games and linear programming. Correlated and coarse correlated equilibria.	HTML PDF
04	Sep 24	Learning in games: Foundations	HTML PDF

Course information

Lectures Tue & Thu
11:00 am–12:30 pm
Room E25-111

Instructors

► **Constantinos Daskalakis**
costis@csail.mit.edu
Office 32-G694

► **Gabriele Farina**
gfarina@mit.edu
Office 45-501F

Meetings by appointment.

Teaching assistants

► **Kat Federova**
fedorova@mit.edu

► **Mingyang Liu**
liumy19@mit.edu

► **Daniel Xia**
dxia03@mit.edu

► **Rui Yao**
rayyao@mit.edu

TA office hours to be announced.

[GitHub repository](#) ↗

Prerequisites

Advanced undergraduate discrete mathematics and algorithms, and mathematical maturity.

Coursework

- Attendance and participation 20%
- Improving material 30%
- Project 50%

There are no assigned homework sets. Students will improve the shared code materials and complete a theoretical

Administrivia:

Prerequisites

Discrete Mathematics and Algorithms at the advanced undergraduate level; mathematical maturity

Talk to us or any of the TAs if you have concerns

Administrivia:

Collaboration policy

We encourage working together on the course materials, projects, and discussion of the ideas. *Contributions* and *project work* should reflect *your own* understanding. A machine bears no responsibility (these are interesting times to be a philosopher).

Acknowledge collaborators and sources, and do not present somebody else's work as your own.

Administrivia:

AI use policy

AI can empower us to achieve great things... Or to bring atrophy

Why are you at a university, surrounded by talented people? What do you want to accomplish? And how can we keep ourselves honest?

We support any use of AI in the advancement of your goal of learning, building, and growing as a scientist. From great powers come great responsibilities. **Don't let generative AI remove or reduce your productive struggle.**

Assignments and Grading

Grading

- Attendance and participation 20%
- Improving material 30%
- Project 50%

This course has **no psets**, and **no exams**.

1. Attendance and Participation

As mentioned, lectures are recorded and available on Canvas.

We expect everyone to attend at least **50%** of the lectures. The "Attendance and participation" component of grading above reflects whether that was met. We will measure attendance using random quizzes during classes. We envision a single question that is extremely easy if you've read or attended the previous week's class.

We don't want to make anyone's life hard, but we don't want to speak to the walls either. Great work has been sparked by questions from class

2. Improving material

- This course uses lecture notes, in PDF and HTML format

MIT 6.7980

Topics in Multiagent Learning

Constantinos Daskalakis and Gabriele Farina

LECTURES

1 Setting and equilibria: the Nash equilibrium

2 Existence proofs: Nash, Brouwer, and Sperner

3 Properties and relaxations of the Nash equilibrium

4 Foundations of learning in games

5 Learning algorithms (I)

6 Bandit feedback

7 Foundations of extensive-form games

8 Learning in extensive-form games

9 PPAD-hardness of Nash equilibrium

SUPPLEMENTARY READINGS

S1 A second look at the minimax theorem

S2 Centralized algorithms for Nash equilibrium computation

S3 Learning algorithms (II)

S4 Sequential irrationality and perfect equilibria

S5 Markov (aka stochastic) games

S6 Phi-regret minimization

THIS LECTURE

L7.1 Game trees and information sets

L7.1.1 Histories, actions, and payoffs

L7.1.2 Imperfect information and information sets

L7.1.3 Perfect recall

L7.2 Player's perspective: Tree-form decision processes

LECTURE 7

Foundations of extensive-form games

INSTRUCTOR

Prof. Gabriele Farina (gfarina@mit.edu)

LECTURE DATE

Tue, Oct 6, 2026

Imperfect-information extensive-form games model tree-form strategic interactions in which not all actions might be observed by all players. They represent an ample majority of strategic interactions encountered in the real world, ranging from recreational games such as poker, to negotiation, and auctions.

L7.1 Game trees and information sets

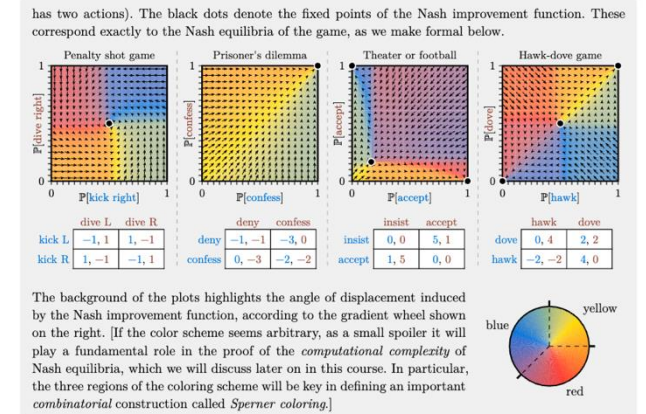
The standard representation of an imperfect-information extensive-form game is through its *game tree*, which formalizes the interaction of the players as a directed tree. In the game tree, each non-terminal node belongs to exactly one player, who acts at the node by picking one of the outgoing edges (each labeled with an action name). Imperfect information is captured in this representation by partitioning the nodes of each player into sets (called information sets) of nodes that are indistinguishable to that player given his or her observations.

Example L7.1. As a running example, we will illustrate the representation by analyzing the game tree of a simplified two-player variant of poker (perhaps the archetype of imperfect-information extensive-form games), known as *Kuhn poker* [Kuh50]. As noted by Kuhn himself, even the previous small game already captures central aspects of deceptive behavior such as *bluffing* and *underbidding*.

Download as PDF

View source

HOW TO CITE



2.1 The Nash

As mentioned above, the Nash improvement function is defined from the set to the set of Nash equilibria. The Nash improvement function of Nash was to find a response for each improvement function.

■ **Regret.** To find the Nash equilibrium, we need to find the difference that Player i acts.

■ **The Nash improvement function.** The current strategy i of playing a_i . The need to handle the adding mass to the Nash's answers to the amount of actions with negative valid strategy, the We can formalize

Definition 2.1 strategies. The

for all player i

2.2 Quantifying the increase in utility of the improvement step

We validate our intuition that the Nash improvement function is a "profitable response" function. The following result shows that if a player has positive regret for any action, then the Nash improvement function unilaterally increases that player's utility. This is a key property that will allow us to show that the fixed points of the Nash improvement functions must be Nash equilibria.

Theorem 2.1. For any strategy profile $(x_1, \dots, x_n)$, and any player $i \in [n]$, the Nash improvement function φ satisfies

$$u_i(\varphi_i(x_1, \dots, x_n, x_{-i}) - u_i(x_1, \dots, x_n) = \frac{\sum_{a_i \in A_i} \left([r_{i,a_i}(x_1, \dots, x_n)]^+ \right)^2}{1 + \sum_{a_i \in A_i} [r_{i,a_i}(x_1, \dots, x_n)]^+}.$$

So, if even one action of a player i has positive regret, then the Nash improvement function *strictly* increases the utility of that player.

Proof. Since we are focusing on a generic player i and keeping all the other ones fixed (and playing strategies x_{-i}), we will reduce the notational burden by using the following shorthands:

$$\begin{aligned} r_{i,a_i} &:= r_{i,a_i}(x_1, \dots, x_n), & (\text{regret of action } a_i \text{ of Player } i \text{ fixing } x_{-i}) \\ u_{i,a_i} &:= u_i(a_i, x_{-i}), & (\text{utility of action } a_i \text{ of Player } i \text{ fixing } x_{-i}) \end{aligned}$$

$$x'_{i,a_i} := \varphi_{i,a_i}(x_1, \dots, x_n) = \frac{x_{i,a_i} + r_{i,a_i}^+}{1 + \sum_{a_i \in A_i} r_{i,a_i}^+} \quad (\text{prob. of action } a_i \text{ in the improved strategy}),$$

where the notation z^+ is a shorthand for the positive part of z , that is, $z^+ := [z]^+ := \max\{0, z\}$.

5

It is straightforward to verify that φ is well-defined and maps strategy profiles into strategy profiles. Indeed, the numerator in 1 is always nonnegative, and the denominator is always at least 1, implying that $\varphi_{i,a_i} \geq 0$ for all $a_i \in A_i$ and player $i \in [n]$. Furthermore,

$$\forall i \in [n], \quad \sum_{a_i \in A_i} \varphi_{i,a_i} = \frac{\sum_{a_i \in A_i} (x_{i,a_i} + [r_{i,a_i}(x_1, \dots, x_n)]^+)}{1 + \sum_{a_i \in A_i} [r_{i,a_i}(x_1, \dots, x_n)]^+} = 1,$$

where we used the fact that $\sum_{a_i \in A_i} x_{i,a_i} = 1$ since x_i is a valid strategy. Finally, observe that φ is a continuous function. The following example visualizes the Nash improvement function in the small games we have seen so far.

Example 2.1. The plots below visualize the displacement $\varphi(x_1, x_2) - (x_1, x_2)$ induced by the Nash improvement function for four games, whose payoff matrices are noted below each plot, after projecting away the probability of the first action of each player (and keeping around only the probability of the

2. Improving material

- We would like your help to improve this material by polishing it, proposing homework problems, fixing inconsistencies
- Bring your creative ideas: interactive visualizations of multiagent dynamics, new topics, hacking on the notes-to-html converter
- We will coordinate efforts on GitHub using issues and pull requests
- We plan to split the material so that everyone is responsible for at least one lecture

2. Improving material

- Two weeks for each lecture
- Team takes a look, proposes changes. TAs and we give feedback
- At the end, all team members write a small individual contribution statement
- Some lectures are more polished than other
- None of them have suggested problems
- There are also other ways to help beyond individual lectures, e.g., infrastructure, engaging with other issues/PRs, ...

3. Projects

- Counts for 50% of weight
- Three types:
 - **#1: Fog of War challenge.** This should probably be your default unless you have a concrete idea that you would like to propose for one of the other types. This is a large scale competition on an imperfect-information game that can be done in teams
 - #2: Modeling project
 - #3: Theory project
- For #2 and #3 our bandwidth is a bit more limited, so we only want people with a clear idea for what to try
- Project proposal date TBA

Fog of War Challenge

- Large competition on a large decision-making problem under imperfect information
- Important to price secrets (eg bluffing) and information gathering
- Leaderboard and live games

MIT 6.7980 · Fall 2026

Fog of War Challenge

Arena Leaderboard Games Rules

Submit a bot

Offline preview · example data

Leaderboard

#	BOT / TEAM	ELO
1	Example White Supplied example engine	1500
1	Example Black Supplied example engine	1500

One supplied game. Example engines only; student rankings will appear when the course arena opens.

Explore rankings → 2 bots

Match details

FINAL

RESULT

1/2 — 1/2

Threefold repetition

3:00 + 2s / move

FIELD OF VISION

23 / 64 squares

41 squares hidden from White.

Move history 134 plies

Move notation is hidden in POV mode to preserve the player's information. Step through the board to inspect observations.

REPLAY / EXAMPLE 001

P1 POV P2 POV True state

Example Black Player 2 · Black 2:39

Example White Player 1 · White 2:04

52 / 134

4x

Fog of War Challenge

- We are curious to see a variety of approaches:
 - Pure test-time computation on a superhuman poker value net
 - End-to-end reinforcement learning
 - Heuristics
 - Develop a new diffusion-based decision-time planning
 - Bayesian modeling of beliefs
 - Extreme performance engineering
- We know for a fact one can get superhuman without requiring end-to-end deep RL (this also means no GPUs are required to be strong)
- Should be done in teams. LLMs welcome.
- We are not HuggingFace... Please ask your agents not to hack us

Modeling Project

- Along a few axes that we think are underexplored
- We will provide a list of readings and some ideas
- If you want to take this project, be proactive
 - Talk to us
 - Talk to the TA
 - Connecting to your research is totally fine and encouraged, as long as it connects to the course as well

Theory Project

- Same as before
- We have some ideas we can give you, or you can bring your own
- It's the most important part of the course
 - Think about (very close to) publication worthy
 - No last minute two days before the deadline half baked project
 - Please no AI slop

Tentative schedule

Part I: Foundations

#	Date	Topic	Notes
01	Sep 15	Setting and equilibria: the Nash equilibrium Nash's existence theorem and its connection to fixed-point theorems.	HTML PDF
02	Sep 17	Brouwer and Sperner Sperner's lemma, Brouwer's theorem, and combinatorial proofs of equilibrium existence.	HTML PDF
03	Sep 22	Properties of Nash equilibrium Topological and computational properties. Zero-sum games and linear programming. Correlated and coarse correlated equilibria.	HTML PDF

Normal-Form Games

		Rock	Paper	Scissors
				
Rock		0,0	-1,1	1,-1
Paper		1,-1	0,0	-1,1
Scissors		-1,1	1,-1	0,0

Rock-paper-scissors

These are “matrix” games

- Simultaneous actions
- Single move per player

Simple model but already captures several important aspects

		Deny (cooperate)	Confess (betray)
Deny (cooperate)	-1, -1	-3, 0	
Confess (betray)	0, -3	-2, -2	

Prisoner's dilemma

(-1 = 1 year in jail)

Despite their simplicity, **normal-form games will provide the ground** to start looking into the following key concepts in multiagent settings:

- Solution concepts and equilibria (Nash, maxmin, correlated, ...)
- Learning from repeated play
 - Learning enables iteratively refining strategies to become stronger and stronger, and it has been a key component in all recent game AI breakthroughs
 - Local learning of each agent can often be connected to global notion of equilibrium
 - $\approx$ Mental model: “reinforcement learning but also works in nonstationary settings”
- Deep connection between equilibria and other important concepts in mathematics, optimization and computer science
- After that, we will move on to notions of games that capture more interesting / real-world phenomena, especially sequential moves and imperfect information

04 Sep 24 **Learning in games: Foundations**

[HTML](#) [PDF](#)

Regret and hindsight rationality. Regret minimization and its relationships with equilibrium concepts.

05 Sep 29 **Learning in games: Algorithms**

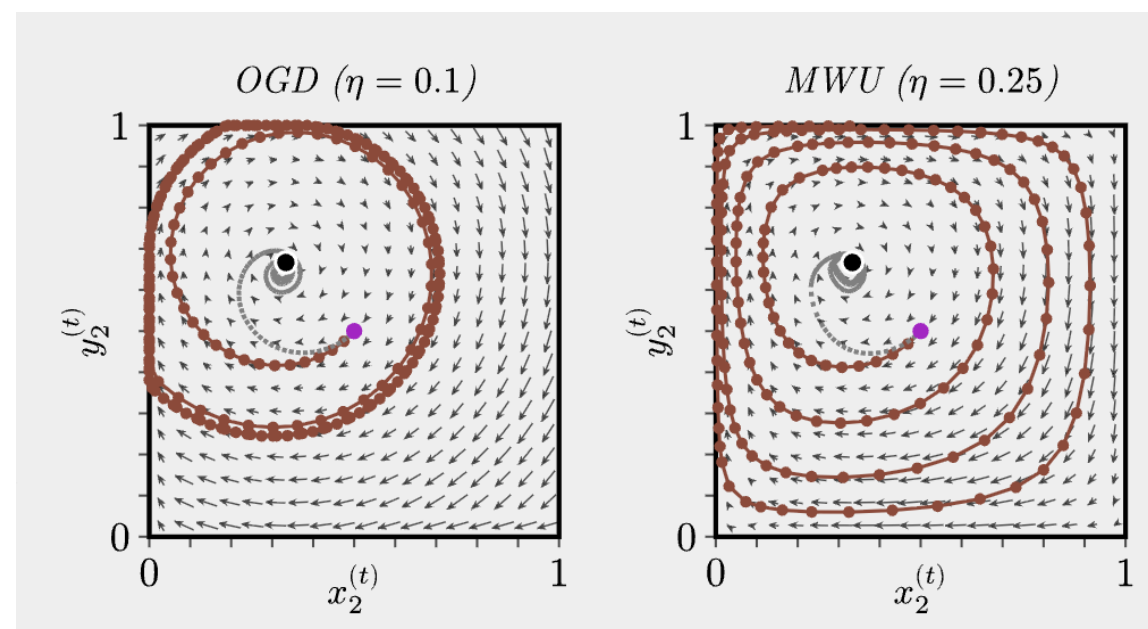
[HTML](#) [PDF](#)

General principles for learning algorithms. Follow-the-leader, regret matching, multiplicative weights, and online mirror descent.

06 Oct 1 **Learning with bandit feedback**

[HTML](#) [PDF](#)

Partial feedback and exploration. From multiplicative weights to Exp3; regret guarantees.



07 Oct 6 **Modeling extensive-form games**

[HTML](#) [PDF](#)

Perfect and imperfect information. Kuhn's theorem. Normal-form and sequence-form strategies.

08 Oct 8 **Learning in extensive-form games**

[HTML](#) [PDF](#)

No-regret learning, counterfactual utilities, and counterfactual regret minimization (CFR).

Oct 13 **No class** MIT follows a Monday schedule.

NO CLASS

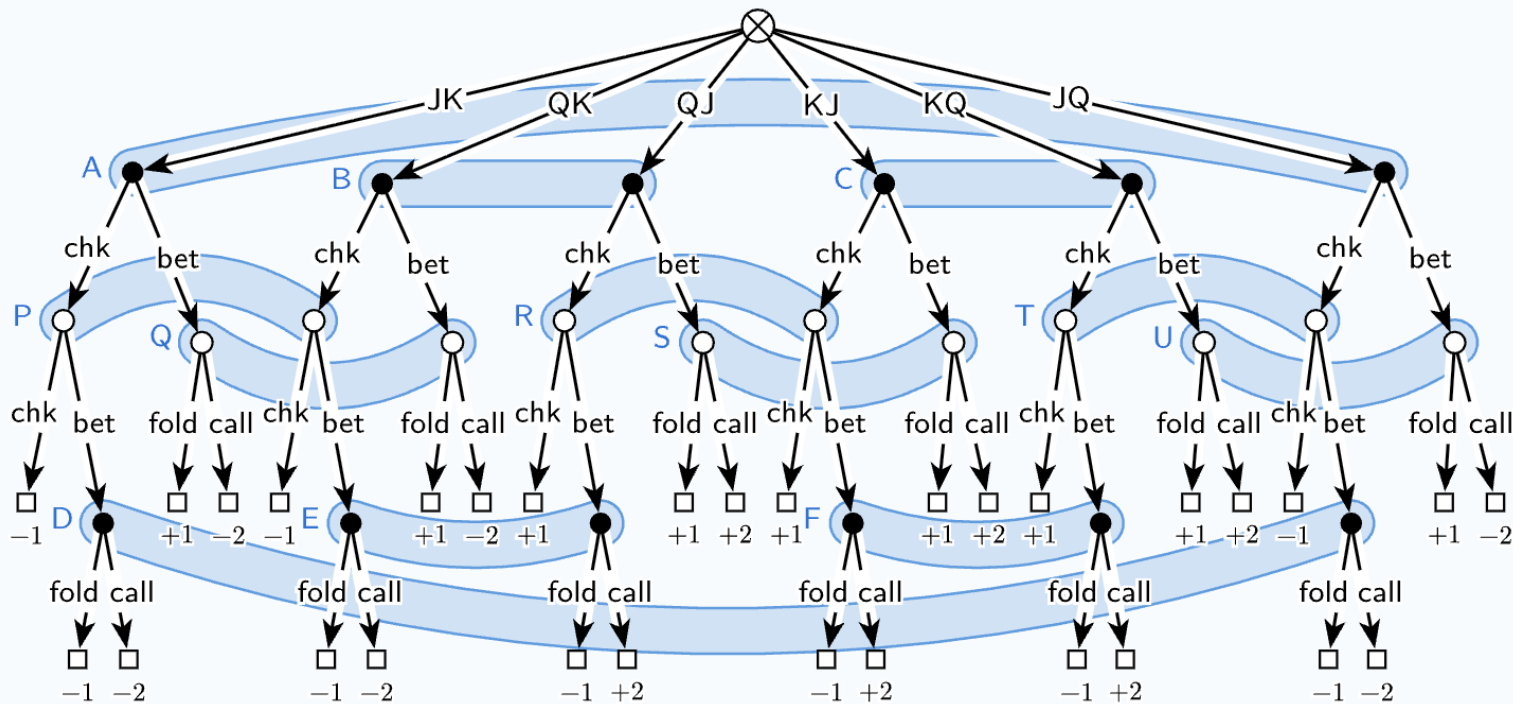
09 Oct 15 **Multiagent deep RL**

Not yet posted

Reinforcement learning in games. Self-play and deep learning methods for perfect-information games.

Imperfect-Information Games

Example:
poker



Difficulties with Imperfect Information

- Compared to normal-form games, imperfect-information extensive-form games bring many conceptual challenges

1 The number of (deterministic) strategies grows **exponentially** in the game tree

2 Imperfect information makes backward induction and local reasoning not viable

Think about Poker: need to reason about **misdirection**. *General principle: you need to think about what the opponents don't know about you and leverage that to your advantage*

3 Other players have control over what part of the game tree is visited/explored

- Nonetheless: many positive results
 - In fact, we live in a world where machines bluff at poker better than humans

#	Date	Topic	Notes
10	Oct 20	Taking stock We will discuss big open questions in the field and possible project ideas.	Not yet posted

Topics

Part II: Information, Communication, Alignment

#	Date	Topic	Notes
11	Oct 22	Information and mechanism design Designing information and incentives in strategic interactions.	Not yet posted
12	Oct 27	Team games and hidden-role games Coordination in teams and games with hidden roles.	Not yet posted
13	Oct 29	Alignment (part I) Reinforcement learning from human feedback (RLHF) and alignment.	Not yet posted
14	Nov 3	Alignment (part II) Regularized RLHF and direct preference optimization (DPO).	Not yet posted

Part III: Advanced learning

#	Date	Topic	Notes
15	Nov 5	Forecasting and calibration Calibrated prediction and its connections to learning in games.	Not yet posted
16	Nov 10	High-dimensional games Learning with large strategy spaces. Kernelized methods and multiplicative weights.	Not yet posted
17	Nov 12	Nonconvex games Nonconvexity, learning dynamics, and local equilibrium concepts.	Not yet posted

Part IV: Computational complexity

#	Date	Topic	Notes
18	Nov 17	Total search and TFNP Total search problems, the TFNP framework, and the PPAD complexity class.	Not yet posted
19	Nov 19	PPAD-hardness of Nash equilibrium Reductions and the computational hardness of finding Nash equilibria.	HTML PDF

Supplementary reading

S1 • A second look at the minimax theorem ↗

[PDF](#)

S2 • Centralized algorithms for Nash equilibrium computation ↗

[PDF](#)

S3 • Learning algorithms (II) ↗

[PDF](#)

S4 • Sequential irrationality and perfect equilibria ↗

[PDF](#)

S5 • Markov (aka stochastic) games ↗

[PDF](#)

S6 • Phi-regret minimization ↗

[PDF](#)

Presentations

Project work and presentations

#	Date	Topic	Notes
	Nov 24 BREAK	No class Project break	
	Nov 26 NO CLASS	No class Thanksgiving holiday.	
20	Dec 1 PROJECT	Project presentations Show us your cool work!	
21	Dec 3 PROJECT	Project presentations Show us your cool work!	
22	Dec 8 PROJECT	Project presentations Show us your cool work!	
23	Dec 10 PROJECT	Project presentations Show us your cool work!	

Course goals (again)

- By the end of this part, you should have acquired:
 - A **language** to think about and describe different equilibrium points of multiagent interactions (Nash equilibrium, maxmin strategies, correlated equilibria, ...)
 - An appreciation for what is **computationally tractable** in every case, and what only in special cases
 - The ability to **implement learning dynamics** to progressively refine strategies, including in imperfect-information domains
 - A general understanding of what techniques are used to push scalability, and what major areas of investigation remain **underexplored**

Questions?