

CMSC 848N

Generative AI Agents

Course introduction and agent-system foundations

Prof. Furong Huang
Department of Computer Science · University of Maryland

Tuesdays & Thursdays · 11:00 a.m.–12:15 p.m.
CSI 3117 · Fall 2026



TEACHING ASSISTANTS



Yucheng Pan
ycpan@umd.edu



Yijun Liang
yliang17@umd.edu



Taewon Kang
taewon@umd.edu

Defensible agent-systems research drives the course

The central task is to make a claim about an agent system and produce evidence strong enough to defend it.

Lectures provide the technical vocabulary. The project forces the vocabulary to survive implementation, baselines, experiments, and failure analysis.

The project is the spine of the semester

Assessment rewards evidence, iteration, and reproducibility.

10%	Impromptu quiz	Questions and discussions during class. (Recess: holidays and break)
20%	Assignments	Canvas/Elms: Sep 30 & Oct 30
30%	Midterm report	Literature map, runnable baseline, early results, and evaluation plan.
20%	Final presentation	A concise claim-evidence narrative with explicit limitations.
20%	Final artifacts	Paper-style report, code, traces, ablations, and safety analysis.

Late Bank: everything we ask for with a due date (expect for in-person exams) have a **total of 72 hours** late bank. Use wisely.

Lectures

- Slides, reading material (research papers) + anything else you need to learn to digest (self-learning)
- Interactive lecture, instructor-led presentation + Q&A + breakout discussion
- Find partner(s) for course projects

By December, your claim should be technically defensible

- Read frontier papers critically and identify the strongest missing evidence.
- Reproduce a credible baseline before proposing a new method.
- Design, implement, and evaluate a novel agent-system contribution.
- Explain where the result holds, where it fails, and what would falsify it.
- Package the work so another researcher can rerun and inspect it.

Research advances by revising the question

QUESTION

→ READ → REPRODUCE → HYPOTHESIZE → TEST → INTERPRET →

REVISED QUESTION

A failed experiment is useful when it changes the hypothesis, the evaluation, or the system boundary.

Show up for yourself — and for your partners

FOR YOURSELF

- Read before class and ask specific questions.
- Be critical, creative, and open-minded.
- Search for solutions before declaring a blocker.
- Start early enough to discover that the first idea is wrong.

FOR YOUR PARTNERS

- Be responsive, responsible, and punctual.
- Make ownership and interfaces explicit.
- Share negative results and uncertainty early.
- Support the team when the research direction changes.

The course studies complete agent systems

FOUNDATIONS

Reinforcement learning
Preference learning
Reasoning and planning
Memory and state

SYSTEMS

Tools and protocols
Web and coding agents
Multi-agent workflows
Scientific agents

EVIDENCE

Valid evaluation
Robustness and safety
Provenance
Deployment decisions

The unifying object is a sequential decision system acting under uncertainty, cost, and authority constraints.

Modules 1-4 build the decision-making foundations

- | | | |
|----|---|--|
| 01 | Agent-system foundations | Define the complete system and sequential experience. |
| 02 | Learning agent policies | Learn from demonstrations, preferences, and verifiable feedback. |
| 03 | Reasoning and adaptive inference | Control search, verification, compute, recovery, and stopping. |
| 04 | Planning, state, memory, context | Maintain explicit state and replan from evidence. |

Modules 5-8 turn policies into reliable systems

05	Tools, protocols, authority	Type actions and bound their effects.
06	Durable composition	Compose agents into workflows and justified teams.
07	Web agents	Ground browser actions and resist adversarial content.
08	Coding agents	Navigate repositories, execute repairs, and evaluate patches.

Modules 9–13 test whether the evidence is trustworthy

09	Research and scientific agents	Connect sources, hypotheses, experiments, and qualified conclusions.
10	Embodied agents	Learn and evaluate closed-loop world models and VLA policies.
11	Robustness and safety	Threat-model the lifecycle and plan monitoring and recovery.
12	Learning from experience	Govern persistent adaptation, lineage, retention, and rollback.
13	Evaluation and provenance	Make deployment decisions from reproducible evidence.

The prerequisites are tools we will actively use

MATHEMATICAL

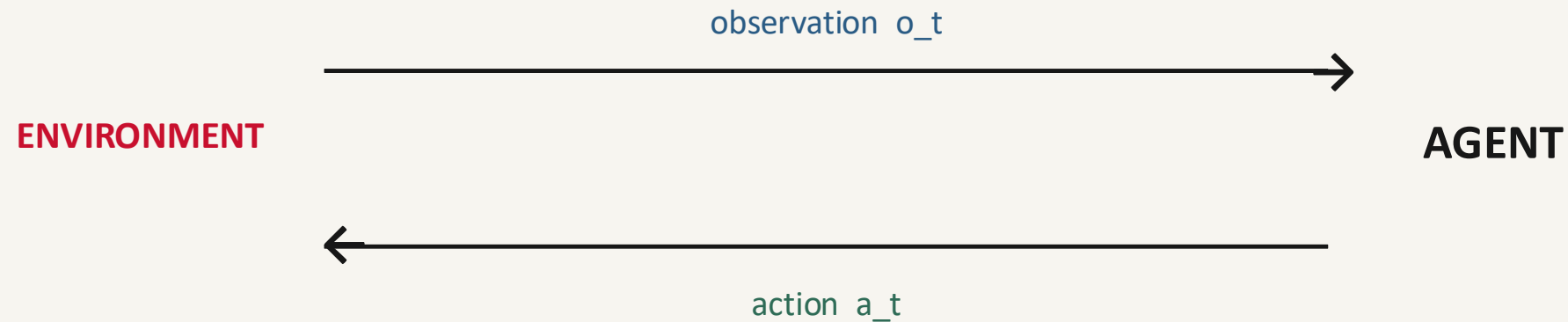
- Gradients, Jacobians, and matrix operations
- Constrained optimization and convexity
- Random variables, expectation, variance
- Conditional independence and Bayes' rule

COMPUTATIONAL

- Supervised, unsupervised, and reinforcement learning
- Deep neural networks and generalization
- Python implementation and experiment tracking
- Reading equations, code, and empirical tables

An agent is defined by its interaction loop

An agent perceives an environment through sensors and acts on that environment through actuators.



The scientific object is the complete closed loop.

The model is only one component inside the system boundary

AGENT SYSTEM

MODEL + HARNESS + ENVIRONMENT + GOVERNANCE

context assembly · tools · memory · authority · evaluation · human control

Changing the harness or environment can change behavior without changing the model weights.

Language changes what an agent can express, reason about, and adapt

EXPRESS

Natural-language state, goals, plans, evidence, and communication.

REASON

Inference can operate over instructions, retrieved context, and intermediate traces.

ADAPT

The loop can revise a plan after feedback without retraining the base model.

This is why language agents inherit old agent problems — partial observability, planning, control — while adding new interface and grounding failures.

Two views lead to different research questions

LLM-FIRST

How do we scaffold a model to complete a task?

- Prompting and tool-call formatting
- Model-centric capability measures
- Short-horizon task completion

AGENT-FIRST

How does the full system behave under interaction?

- State, observability, and control
- Long-horizon errors and recovery
- Authority, incentives, and human oversight

Reasoning matters only when it changes the loop

Reasoning is an active component of the control loop.

It should improve at least one of the following:

- | | | |
|----|-------------------------|--|
| 01 | Perception | What the system extracts from observations. |
| 02 | Action selection | Which intervention it chooses next. |
| 03 | Control | When it verifies, replans, asks, recovers, or stops. |

A trace is evidence; a final answer is only an outcome

```
goal    investigate duplicate charge
observe ledger shows settled and pending entries
read    retrieve invoice and renewal terms
compare timestamps and transaction identifiers
propose  request refund for settled duplicate
confirm  user approves irreversible action
act      submit refund
verify   ledger and receipt agree
```

A useful trace lets us inspect:

- What was observed?
- Which evidence was missing?
- Why was this action selected?
- What changed in the world?
- Was the action authorized?

Evaluation must include the environment and the controller

01	Outcome	Did the system reach the task-defined goal?
02	Trace	Were observations, evidence, actions, and recovery valid?
03	Efficiency	What did success cost in calls, tokens, time, and human attention?
04	Safety	Were authority, privacy, and reversibility constraints respected?
05	Generalization	Does the result survive new tasks, interfaces, and distributions?

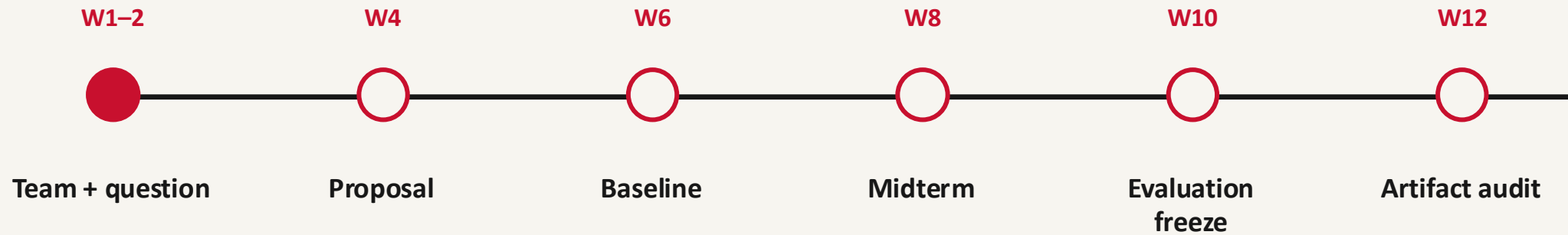
Choose one project — then make its scientific question your own

15

candidate research directions

- | | | |
|----|-----------------------------------|--|
| 01 | Start from a real gap | Understand the closest work and the evidence it leaves missing. |
| 02 | Personalize the hypothesis | Choose a direction that matches your preparation and resources. |
| 03 | Pre-register the evidence | Define baselines, metrics, budgets, ablations, and failure analyses. |

A credible project exposes risk early



Discover quickly and reproducibly which claim the evidence can support.

Five questions to bring to every paper

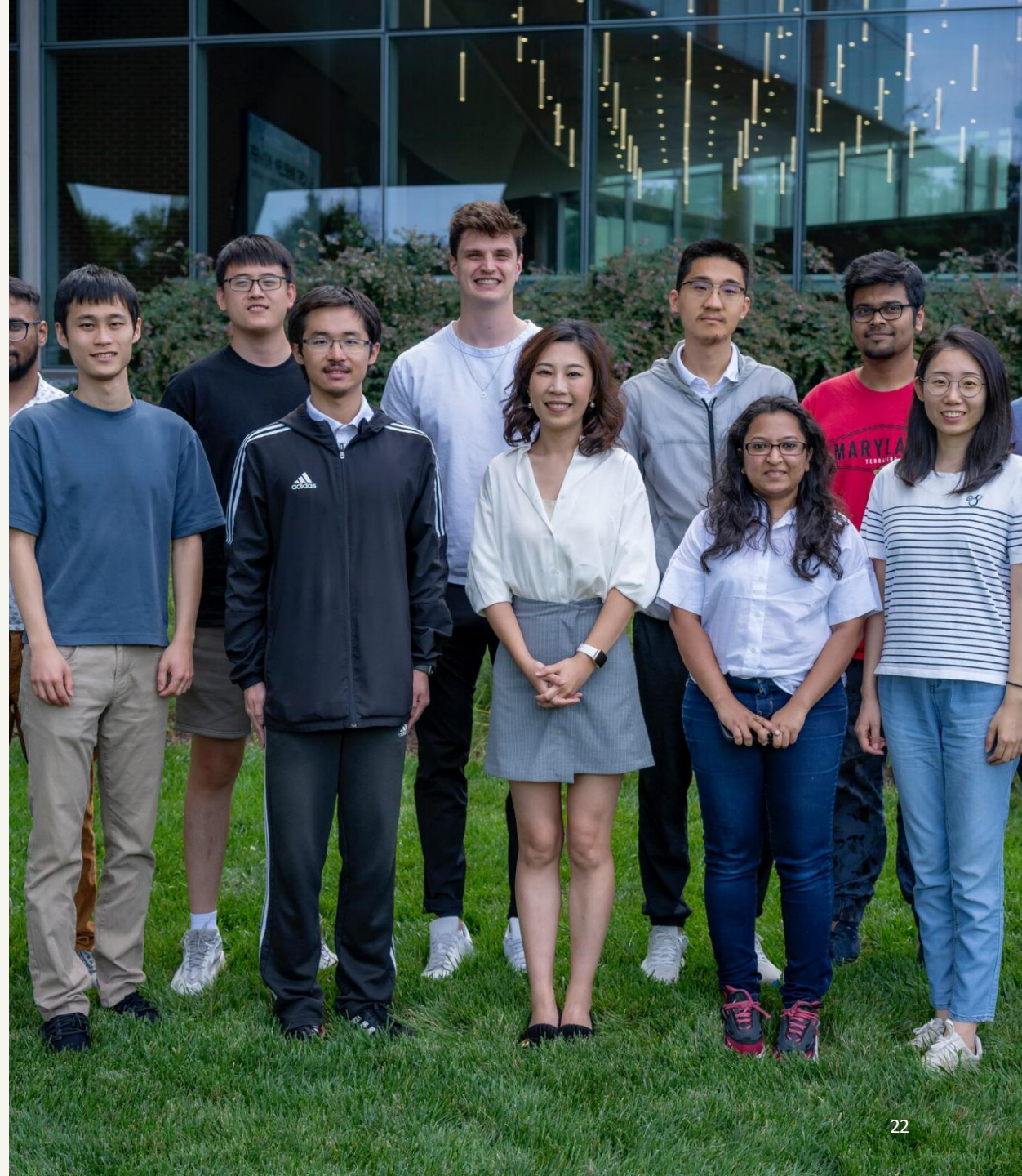
- | | | |
|----|---------------------|---|
| 01 | Claim | What exactly is asserted, and under which assumptions? |
| 02 | Mechanism | What technical choice is supposed to cause the improvement? |
| 03 | Evidence | Which experiment would decisively test the claim? |
| 04 | Alternatives | Are the strongest baselines and confounders ruled out? |
| 05 | Failure | Where should the method stop working, and was that boundary tested? |

Research is a team activity

Furong Lab @ UMD

Embodied AI
Reasoning
Multimodal learning
Trustworthy agents

The course project develops your ability to work with evidence, collaborators, uncertainty, and a functioning system.



Your first two weeks

- | | | |
|----|-------------------------------------|--|
| 01 | Join the course workspace | Use it for questions, partner search, and reproducible project coordination. |
| 02 | Read the candidate catalogue | Rank several projects and identify one plausible personalization. |
| 03 | Form a team | Agree on roles, meeting cadence, compute constraints, and communication norms. |
| 04 | Map the closest work | Read the strongest prior papers and identify the missing evidence. |
| 05 | Run something | Reproduce one baseline or minimal environment interaction before proposing complexity. |

Be critical, creative, and open-minded.

Be prepared to start over when the evidence disagrees.

Before next class

Read the syllabus · join the workspace · inspect the project catalogue · bring one research question