

AI for Economic Research: Dynamic Models, Language, and Agents

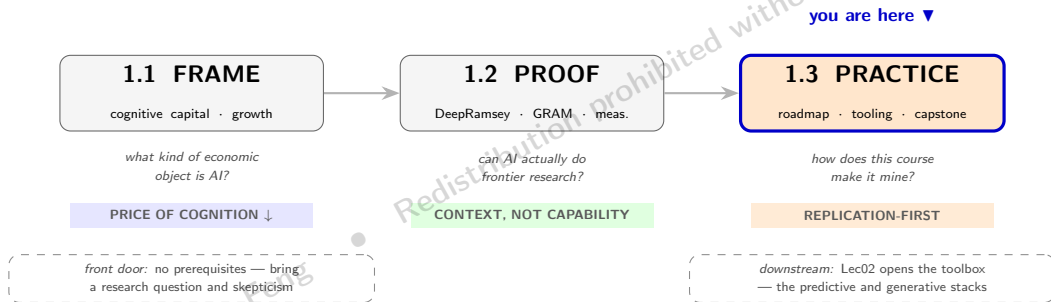
Lecture 1.3: How This Course Works

- Zhigang Feng

2026-07-07 10:17:02

Lecture 1 in One Map

One lecture, three decks, one claim: AI is new cognitive capital — cheap, fast, non-rival — and this course trains you to deploy it on real research.



Act 1.3 (here) makes it yours: the 10-lecture roadmap (two arcs, three running examples), the research-architect stance and its tooling (5-pillar workflow, replication discipline, reproducibility), and the capstone — four tracks, replication-first grading.

Course Roadmap

Ten lectures, two arcs: dynamics and language — three running examples thread them

Course Roadmap: How the 10 Lectures Connect

Lec.	Topic	Hours
1	⇒ Introduction (this lecture)	1–2
2	What is AI?	1–2
3	ML for Quant Macroeconomists	2–4 + lab
4	Solving Macro via DL & RL	2–4 + lab
5	RL in a Nutshell	2–4
6	Heterogeneous Agent Models via RL	2–4 + lab
7	Large Language Models (LLM)	2–4
8	Retrieval Augmented Generation (RAG)	1–2
9	Agentic AI	1–2
10	Case Studies	2–4 + lab

Hours are approximate and adjusted to pace; totals flex with lab depth.

The Two Arcs: Dynamics and Language

Arc 1: AI for Macro Models (Lec. 3–6)

- How do neural networks work?
- Can DL solve the Bellman equation?
- What is RL? How does it relate to DP?
- Can we solve Aiyagari/Krusell-Smith with RL?

Goal: Use AI to expand the computational frontier of quantitative macro.

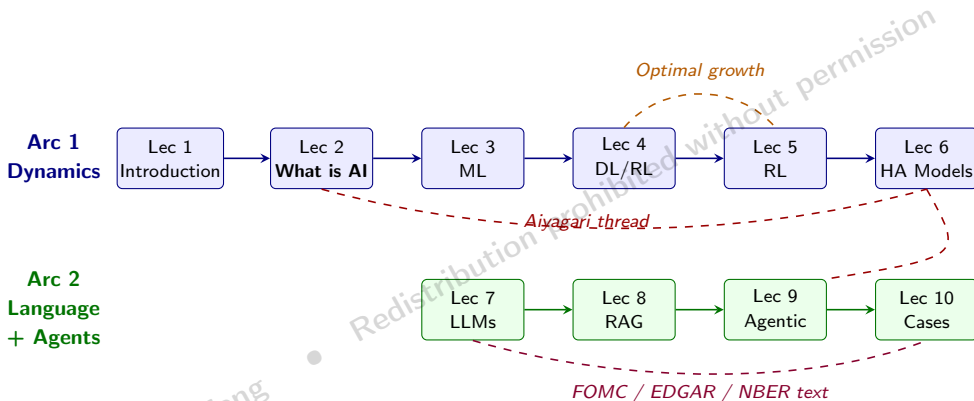
Arc 2: AI as a Research Tool (Lec. 7–10)

- How do LLMs work? What can they do?
- What is RAG? How do we feed our data to an LLM?
- What is agentic AI?
- Case studies: economists using agentic AI.

Goal: Use AI tools to transform your research workflow.

Binding thread: Every dynamic problem reduces to a Bellman / Euler equation — classical DP (Lec. 4), RL (Lec. 5, 6), even LLM prompts (Lec. 7: implicit Bayesian conditioning).

The 10-Lecture Map: Two Arcs and Their Running Examples



Every running example traces multiple lectures. Build the foundation, then re-meet it in agentic form.

What You Need to Know vs. What AI Can Help With

- **What economists must bring:**

- Economic intuition: what is the right model, state space, welfare criterion?
- Identification: why does this variation estimate a causal effect?
- Domain knowledge: what makes FOMC minutes different from random text?
- Judgment: when is the AI output trustworthy?

- **What AI can do for you:**

- Solve high-dimensional optimization problems (Bellman, ML training)
- Process and extract information from large text/data collections
- Write, run, and debug code at scale
- Synthesize literature and summarize research

- **The Regime Change:**

- The bottleneck in research is shifting from *execution* (writing code, computing) to *design* (what question to ask, how to identify it)
- Your advantage as an economist: you know what question matters

Research-Architect Stance & Tooling

You bring the question, the identification, and the judgment; AI supplies throughput

You Are the Research Architect: The 5-Pillar Workflow

- Across the ten lectures you build four competencies — theoretical mastery, technical fluency, AI-augmented implementation, and critical validation. The capstone is where they meet a real problem.
- **The 5-pillar workflow — your scaffold on every project:**
 1. **Economic theory** — state the problem precisely (objective, constraints, equilibrium or estimand).
 2. **Algorithm design** — choose the method (VFI vs. Euler iteration; WLS; embedding + retrieval; a typed graph).
 3. **Numerical / empirical technique** — discretization, weighting, sampling, chunking.
 4. **Advanced method & trade-offs** — accuracy vs. speed, bias vs. variance, precision vs. recall.
 5. **Code literacy & validation** — read every line; benchmark; sanity-check against economics.

The stance: AI handles the bricks — syntax, boilerplate, data wrangling, refactoring — while *you* own the theory, the validation, and the interpretation. We are not impressed by code an agent wrote; we are impressed by a correct result you can defend.

Replicate First, Explore Second

- **Baseline (the core deliverable):** faithfully *reproduce and verify* a known result — a classical solution, a published number, or the instructor's reported finding — on public data or at small scale.
- **Stretch (upside):** finding “something more interesting” only counts *once the replication validates*. This mirrors how real research works, and how Lec. 9–10 teach you to use AI: build the trusted benchmark, *then* extend.

Validation discipline (non-negotiable). For every headline result: (i) re-derive at least one number by hand or in a second tool; (ii) benchmark against a *known* quantity (a classical solution, a published figure, a standard package); (iii) run economic sanity checks ($\beta \in (0, 1)$; Euler-equation errors below threshold; signs match theory); (iv) **never** submit code you have not read end to end.

Reproducibility Discipline

- **One command should reproduce your results.** The reproducibility package is a graded artifact, not an afterthought.

Every project ships:

- Git history + **fixed random seeds**
- A pinned environment (`environment.yml` / `requirements.txt`)
- `CLAUDE.md` — describes the project for the next agent (and human)
- `AI_LOG.md` — a running log of how you used AI: the prompts that mattered, the decisions you made, where the AI was wrong and how you caught it, an honest human-vs-AI split
- A README with run instructions, a data dictionary, and recorded access dates

The AI_LOG.md is the evidence that you were the architect, not the consumer.

Tools You'll Use

- **Minimum:** VS Code + Python. Everything in this course runs there — no Google/Colab required.
- **Encouraged (the point of the course):** use AI as a *partner, not a ghostwriter* — Claude Code, an IDE agent, or a chat model across the whole pipeline: scoping, data acquisition, coding, debugging, drafting. What is *not* allowed is shipping anything you cannot read, explain, and defend.

For students in China. The labs run top-to-bottom **offline**; setup uses **Tsinghua (TUNA) mirrors** for conda/pip. Start with Labs/Lec01_02_Lab_Getting_Started.ipynb: install Python + PyTorch in VS Code and plot Stanford *AI Index* data (AI vs. human experts, inference cost, adoption speed) — the very figures from T1.1, now in your own hands.

Capstone Projects

Four tracks, replication-first grading — the stretch is upside

Capstone: Four Tracks (Pick One)

Track	One-line pitch	Builds on	Instructor research
A	Optimal growth via deep learning	Lec. 1,3,4,5	(thematic)
B	Build & validate an AI-exposure index	Lec. 2,7,10	AI-index (WIP)
C	Text as economic data (+ small RAG)	Lec. 7,8	(adjacent)
D	AI-for-research infrastructure	Lec. 8,9,10	GRAM / GraphRAG (WIP)

- **B** reconstructs, on public data, the task-level Anthropic Economic Index $\times$ ONET exposure index previewed in **T1.2** (the live measurement example).
- **D** engages the **GRAM** typed-retrieval project previewed in **T1.2** (Case B).
- Tracks B, C, D touch *active, unpublished* research — findings are labeled preliminary / WIP.

Capstone Tracks: Baseline & Stretch

- A — Optimal Growth via DL** *Baseline:* solve the optimal growth model with a neural net (Euler-residual loss and/or actor-critic) and match the classical VFI / time-iteration benchmark on Euler errors. *Stretch:* add a stochastic TFP shock or a second state; or an “AI-as-capital” productivity twist.
- B — AI-Exposure Index (WIP)** *Baseline:* rebuild the occupation-level index from the Anthropic Economic Index $\times$ ONET (G , $|G|$, importance-weighted intensity) and verify the structure from T1.2. *Stretch:* link to IPUMS-CPS / BLS OES employment (WLS, robust SEs); probe the usage $\times$ concentration interaction.
- C — Text as Economic Data** *Baseline:* turn a corpus (FOMC / 10-K / news) into a validated measure — a lexical *and* an embedding version — plus a small RAG (chunk $\rightarrow$ embed $\rightarrow$ store $\rightarrow$ retrieve), validated against an external benchmark out of sample. *Stretch:* a novel index or a head-to-head where embeddings beat the baseline.
- D — AI-for-Research Infra (WIP)** *Baseline:* either a mini condition-aware GraphRAG (flat RAG vs. typed-retrieval ablation) or a validated Claude Code skill (explicit phases, confidence gates, quote verification). *Stretch:* the supervised, confidential GRAM comparison; or a human-vs-skill evaluation.

Logistics & Grading

Teams of 2–4 (mix theory + data/code).

Milestones:

- **M1** Proposal — track, question, the *specific* result to replicate, data, division of labor
- **M2** Data / infra up — pipeline runs; benchmark in hand; repo scaffolded
- **M3** Replication verified — the core grade; stretch scoped
- **M4** Paper, repo, talk, peer-review report

Dimension	Pts
Economic framing & question	15
Replication correctness	20
Validation rigor	20
Reproducibility & provenance	15
AI-collaboration quality	10
Communication	10
Peer-review quality	10
Stretch finding (bonus)	+10

A project that only replicates — correctly, with rigorous validation and clean reproduction — earns a strong grade. The stretch is upside.

The Companion Textbook: Full Topic Map

Miao & Feng, Artificial Intelligence for Economists (Higher Education Press, 2027) — 16 chapters in six parts. This course samples the book; here is its full coverage.

Part I — Introduction & Background

1. AI as research infrastructure · 2. a short history of AI & the new data environment

Part II — Foundations

3. the logic of machine learning · 4. mathematical foundations of ML · 5. research engineering

Part III — Machine-Learning Methods

6. supervised learning (prediction) · 7. unsupervised learning (structure) · 8. double ML: prediction → causality · 9. ML-enhanced causal identification

Part IV — Deep Learning

10. neural-network foundations · 11. economic applications of deep learning

Part V — Reinforcement Learning

12. foundations of reinforcement learning · 13. core RL algorithms

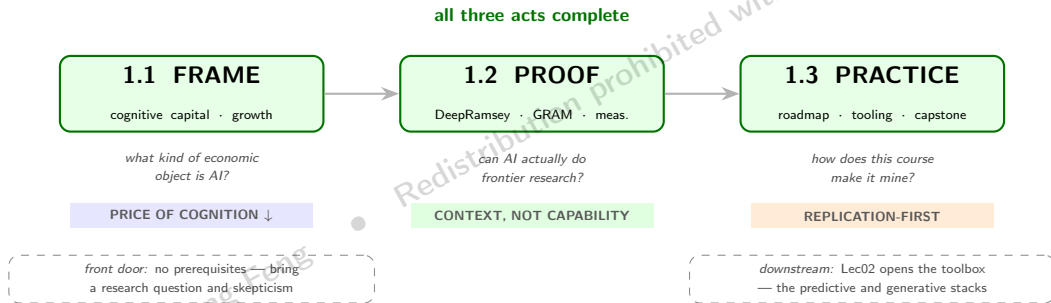
Part VI — Text, Language Models & AI-Assisted Research

14. text representation & language-model basics · 15. large language model applications · 16. retrieval-augmented generation & agents

Course vs. book. These ten lectures follow the *dynamics* and *language* arcs — roughly Parts IV–VI. The book widens the lens with Part II's *foundations* and Part III's full *ML-methods* toolkit — supervised/unsupervised learning and **double ML / causal forests** — the prediction-and-causal-inference methods every empirical economist now needs.

Lecture 1, Complete: One Map, Three Acts

The claim is framed, proven, and handed to you: AI is cognitive capital, it already does research, and this course is where you make it yours.



Bring the question, the identification, and the judgment; let AI supply throughput — the stance every later lecture assumes.

Preview: What Lec. 2 Delivers

- **Lecture 1 set the why and the how:** AI as fast, cheap, widely adopted cognitive capital (T1.1); AI already *doing* frontier research (T1.2); and how this course works — the architect stance, tooling, and capstone (T1.3).
- **Lecture 2 — “What is AI?”** opens the technique families you will use: machine learning, deep learning, reinforcement learning, LLMs, and agentic AI — and, crucially, *where AI fails* (predictive and generative), so you wield it as powerful but bounded.

You are the Research Architect: bring the question, the identification, and the judgment; let AI supply throughput. Next we open the toolbox.

Summary

1. **Why now:** AI is cheap, fast, and diffusing at record speed — a *candidate* general-purpose technology moving on a timescale of months (snapshot: early 2026).
2. **Cognitive capital:** AI is *non-rival cognitive capital* — entering growth through automation (K), idea production (A), and a new intangible stock.
3. **The measurement gap:** task-level gains do not auto-sum to macro TFP — bottleneck tasks and mismeasured intangibles keep aggregate estimates modest (Acemoglu 2024).
4. **AI already does research:** DeepRamsey co-developed a dynamic model; GRAM showed *context, not capability*, is the bottleneck — with causation as the standing guardrail.
5. **You are the Research Architect:** bring the question, the identification, and the judgment; let AI supply throughput — the stance the replication-first capstone trains.
6. **Next:** Lec. 2 opens the toolbox — the predictive stack (T1) and the generative & agentic stack (T2), each with where it fails.