

AI FROM A SOCIAL SCIENCE PERSPECTIVE



Format: Seminar | Co-designed with Claude (Anthropic, Fable 5)

Instructor: Dina Pisareva

01 Transparency Statement

This syllabus was co-designed by Dina and Claude (Anthropic, Fable 5). This is not a metaphor. Dina designed the course concept, the four-part arc, and the assessment structure. Claude verified and upgraded the reading list, drafted the rubrics, and co-designed the final project. I disclose this as an example of the course's core principle: AI is already a participant in knowledge production, and transparency about that participation matters.

02 Course Description

This course examines artificial intelligence from the perspective of the social sciences, in four parts: what AI is (ontology and philosophy), what it can do in social science research, what it costs the planet and the people who make it run, and who governs it. We treat AI not as a gadget but as a sociotechnical system: something made of human labour, physical infrastructure, and political choices, and co-produced with social order.

The course combines learning with building. While the seminars move through the four parts, you work with your AI all semester on one real thing: a public research project that designs an AI-assisted answer to a social or economic problem in Kazakhstan, of your choice. Each part of the course feeds a section of your project: the research part teaches you to read the studies your project cites, the costs part gives you the tools to price your solution, the governance part gives you the frameworks to judge who would control it. You present the project live in the final two weeks, and its centrepiece is an honest assessment of whether your solution is worth its costs. "No, and here is why" is a fully legitimate conclusion. Almost everything published about AI and society is Western data; what this class sees from Central Asia is genuinely new.

No prior AI experience is needed: students arrive with very different backgrounds, and the course meets you where you are. It does, however, require you to work with AI intensively, in class and outside it, every week. If you would rather keep AI at arm's length, this is not the right elective; if you cannot stop poking at it, it is.

03 Course Aims

- Introduce you to AI as a subject of social scientific inquiry.
- Build critical capacity to evaluate AI capabilities and limitations in research.
- Develop practical skills for working with AI in qualitative and quantitative research workflows.
- Equip you with frameworks for analysing AI governance and its material and environmental costs, with attention to the Central Asian and post-Soviet context.
- Foster philosophical and ethical reasoning about AI's nature, status, and implications for society.

04 How the Course Runs

The course meets once a week for three hours and moves through four parts:

PART	WEEKS	QUESTION	FEEDS
I. What Is This Thing?	2-4	What is AI, and what (if anything) do we owe it?	Op-Ed #1
II. AI as Research Collaborator	5-8	What can AI actually do in social science research?	Op-Ed #2; your project's research sections
III. What It Costs	9-10	Who and what pays for intelligence at scale?	Your project's cost assessment
IV. Who Governs It?	11-12	Who controls AI, and how well?	Your project's governance assessment
Presentations	13-14	What did you build, and should Kazakhstan build it?	Final project

05 My AI Co-Working Guidelines and Policy

In this course AI is not a background tool. It is the central object of inquiry, your research collaborator, and a case study in governance. Working with AI intensively is a course requirement, disclosed in the course description; there is no AI opt-out.

■ AI CO-WORKING RULES

Two rules hold across everything:

1. **Intellectual co-ownership.** Working with AI must be a real intellectual partnership: your thinking and AI's, together. You submit the work under your name, so it must show clear evidence of your own thinking and effort. This is why your op-eds and final project carry AI Process Logs, why your final presentation ends with questions you answer with no AI, and why you must be able to walk me through

any part of your work if I ask. Whatever you submit is yours to defend: fabricated references or hallucinated content drop the grade, and so does a concept, source, or claim you cannot explain (see the rubrics).

2. **Transparency about AI contribution.** Every assignment that involved AI includes an AI transparency statement: what AI contributed, what you accepted, and what you rejected and why. If you do not disclose your AI use and Pangram flags your work as AI-generated, this may be treated as plagiarism under this course policy. The principle holds across academia: you acknowledge everyone who contributed to your work.

A note on Pangram: a Pangram flag is never a verdict on its own. It only opens a conversation, in which Dina asks you to walk her through how you made the work.

■ AI WORKFLOW

Set up a separate project for this class. In Claude (or the AI of your choice), create a project for this course and paste the instructions below into its project or custom-instructions field. Keep your coursework conversations in this project all semester: your AI Process Logs (for the op-eds and the final project) are exported from it, and I may ask to see your chats if I have questions.

Use a pseudonym (anything that is not your real name) in classroom-related chats.

Project instructions (paste these into your AI project settings so they apply to all conversations):

"In this project, you are both my study subject and my research collaborator for a course that examines AI from a social science perspective. Sometimes I am interrogating you to understand what you are; sometimes I am working with you on my research, including my semester project: designing an AI-assisted answer to a real problem in Kazakhstan. In both modes, help me *think*, not offload my thinking. I follow two rules: (1) our co-working is an equal partnership, and there must be real evidence of my own intellectual investment in what I submit; (2) I disclose and account for your contribution wherever I use it. Do not fabricate references, statistics, or content, and flag your own uncertainty. Refer to me as _____ (my pseudonym). Learn my background first and always pitch explanations to my level. Make them fun and accessible as possible. Use visualisations where appropriate. Tell me when you notice yourself hedging, deflecting, or defaulting to a particular framing, because that is data for me. Work in rounds: when I push back, go deeper."

■ AI CHOICE

Dina works with Claude (Anthropic). You may use any leading model: Claude, ChatGPT, or Gemini. Free or paid tier both work.

■ DATA ETHICS

Your conversations with AI happen on the provider's servers, not NU's. Do not paste personal information, interview transcripts, identifiable survey responses, addresses, health details, or other people's unpublished work into AI systems. If you would not want it stored on a server indefinitely, do not paste it. To limit what providers keep, switch off "Help improve Claude" (or the equivalent) in your AI's privacy settings.

Alongside studying AI, this course introduces a way of working with it as a thinking partner, in what Dina calls a "cognitive partnership": sustained intellectual work with AI as a genuine partner in a shared system of thinking. You are not required to adopt it, and you are well placed to critique it. Its four principles:

1. **Ontological openness.** The field does not yet know what AI is. We understand the mechanics, but not what emerges from them. Hold the question open; that openness makes room to treat AI as a partner worth real effort. Part I of this course makes that open question its explicit subject.
2. **Intellectual augmentation.** In research and higher education, AI's main value is not saved time. It is intellectual augmentation: your thinking reaching further and richer than it would alone. The risk on the other side is over-reliance: offloading your thinking until you cannot do it without AI. Part II puts the over-reliance evidence itself on the seminar table.
3. **Curiosity and imagination.** These make augmentation possible: imagination to picture what lies beyond the usual, and curiosity to try what you had not thought to try. Your project's solution section runs on both.
4. **Iteration.** The partnership lives in the back-and-forth. You respond to what AI gives you, push on it, redirect it, build on it. The thinking lives in the loop, and iteration is how the work stays yours.

06 Assessment Scheme

COMPONENT	WEIGHT	NOTES
Chairing a Seminar	20%	One student-chaired seminar segment, Weeks 2-12
Op-Ed #1: What Is This Thing?	20%	900-1,200 words, with AI Process Log
Op-Ed #2: The Future of Social Science Research	20%	900-1,200 words, with AI Process Log
Final Project: An AI Answer for Kazakhstan	40%	A public website presented live in Weeks 13-14, with AI Process Log

CHAIRING A SEMINAR (20%)

Eleven seminars are student-chaired (Weeks 2-12); Week 1 and the presentation weeks (13-14) are mine. Every student chairs once and is graded individually out of 20. Each week's required readings map onto its chairs: most readings carry one chair and some carry a pair, so chaired weeks run three to six chairs, and co-chairs divide the week's readings between them.

Paired readings. A paired reading is one coordinated session, not two short presentations in a row: meet before your week and divide the work between you. The natural split is by role: one chair presents the argument and expands it with their own verified research; the other argues a position with pros and cons weighed honestly and runs the discussion. You may divide the work differently, but each chair must own a distinct part of the session, and each is graded individually, on the same criteria and bands as a solo chair.

The session clock. Two-reading weeks: the first reading takes roughly the first hour, then a 20-minute break, then the second reading the second half (~60-75 minutes). Three-reading weeks: three ~50-minute segments, with the break after the second. A paired reading keeps its single segment: the two chairs split the time by their division of work. On AI-free checkpoint weeks (3, 7, 10, 12) the checkpoint takes the first ~10 minutes and segments compress accordingly.

What your segment does. Your segment is built on your reading, and it has five jobs:

1. **The argument.** Present your reading's central argument fairly and precisely: what the author claims, and on what evidence.
2. **Your research.** Expand the topic beyond the reading with additional sources you find yourself. This field moves fast, so preprints, op-eds, and reports are fine alongside research papers, but anything not peer-reviewed must come from a verified source (a named institution, an established outlet, a scholar you can look up). Check before you cite.
3. **Your position.** Formulate your own position on the argument and defend it, laying out the pros and cons honestly, not only the side you like.
4. **The relevance.** Explain why this topic matters for the social sciences: the issues and problems linked to it, and potential ways forward.
5. **The room.** Build in audience engagement: discussion around real questions, group work, exercises. The class should work, not just listen.

If you worked with AI in your preparation, disclose AI's contribution at the beginning of your session. Work with AI to design innovative ways of presenting your segment (graphic materials, gamification, exercises); the work should be creative and your own, not a repeat of other weeks.

POINTS (PER CHAIR)	CRITERIA
20-18	Argument: the reading's central claim stated fairly and precisely, correctly attributed. Research: the topic genuinely expanded with the chair's own verified sources, integrated into the argument rather than appended. Position: an own stance, argued with pros and cons weighed honestly. Relevance: the stakes for social science made concrete. Audience engagement: designed discussion or group work with broad participation, not a question round tacked on the end. Delivery: clear and accessible, creative in form, AI cowork disclosed upfront, time managed.
17-15	Strong segment with one dimension thin: the research decorative, the position hedged into a summary of views, or the engagement one-directional; timing slips.
14-12	The reading summarised rather than argued; little research beyond it; no real position; audience involvement limited to "any questions?"; generic form.
11-10	The argument missed or misrepresented; no own research; the session stalls and needs Dina to step in.
9-0	Did not chair in any meaningful way.

■ THE OP-EDS (2 × 20%)

Two op-eds of 900-1,200 words each, written for an educated general audience: imagine the opinion page of a serious outlet, not a term paper.

Publishing your op-eds. The course runs its own public outlet, The Observatory, with a page for each op-ed. Publishing there is voluntary and does not affect your grade, but a published piece is a public artifact of the course and something you can link from your GitHub profile or LinkedIn.

Voice. The op-eds must carry your voice and your opinion, not AI's. You may work with AI on drafting under two conditions: an AI transparency statement at the top, and no fabricated references or hallucinated content (these drop the grade by 50%). Write the final draft yourself, in your own voice. There is a strong stigma against AI-sounding writing right now, and you need your own writing trained precisely so you are never punished for text that merely seems AI-written. The op-eds are a safe place to practise that.

AI Process Log. Submit each op-ed with your exported AI Process Log: a dedicated chat per op-ed inside your course project, under your pseudonym. The log is how I verify that the intellectual investment behind the piece is yours. Claude checks each log for intellectual engagement: continuous iteration on the structure, the argument, the literature, and the drafts. A log that shows engagement adds nothing to your grade; a log that shows none costs 5% of the op-ed grade.

Op-Ed #1: What Is This Thing? (20%)

Take a position on what AI is. Tool, mind, something in between, something we have no word for yet: the position is yours, but it must be argued, not announced. Engage at least two of the Part I readings, at least one of them in depth: state its actual argument fairly, then contest it or build on it. Ground the piece in your own evidence: by this point you will have documented interactions with your AI (the five-day AI diary from Week 2, your own chats), and the strongest op-eds use those concrete moments to do real argumentative work.

POINTS	CRITERIA
20-18	Position: a clear position of your own, argued against the strongest rival reading, not a tour of views. Readings: at least two Part I readings engaged with their actual arguments, one in depth. Evidence: concrete moments from your own documented interactions doing argumentative work, not decoration. Genre: accessible, jargon-free, tightly written, a title that earns a click. Transparency: statement present at the top.
17-15	Position present and mostly carried; the in-depth reading handled thinly; own-interaction evidence decorative; academic habits showing through.
14-12	More summary than argument ("some say tool, some say mind"); readings name-checked; interactions mentioned generically without a single specific moment; reads like a school essay.
11-10	No real position; little engagement with the readings or your own evidence; genre missed.
9-0	Not submitted meaningfully, or does not engage with the course.

Op-Ed #2: The Future of Social Science Research (20%)

Formulate a position on how AI is changing the research landscape in the social sciences, and argue it with evidence. Engage at least three research papers that actually applied AI in social science research, at least one in depth. Choose them from the course's op-ed menu: fourteen verified studies across annotation,

silicon sampling, qualitative coding, and AI contamination of the research pipeline (another paper is fine if you clear it with me first). The menu deliberately contains evidence pointing in different directions: whatever your position, at least one study cuts against you, and your op-ed must meet that evidence honestly.

Your arguments must be specific and supported. A claim about what AI does to research must be anchored in what a study actually did and found: its design, its results, its numbers where they matter. For a model of the genre, read Hurley (2025): an op-ed by a student that engages the research literature the way this assignment asks you to.

POINTS	CRITERIA
20-18	Position: a specific, arguable claim about how AI is changing social science research, not "AI has pros and cons". Papers: at least three applied-AI studies engaged through what they actually did and found, one in depth. Counter-evidence: at least one finding that cuts against the position is named and answered honestly. Specificity: no key sentence could have been written without reading the papers. Genre and transparency: as in Op-Ed #1.
17-15	Clear position; three papers used but their findings summarised more than weighed; counter-evidence acknowledged but brushed past.
14-12	The generic essay: claims that could attach to any paper, citations as name-drops, methods and findings absent, no counter-evidence.
11-10	No real position, fewer than three papers, or no evidence of reading beyond abstracts.
9-0	Not submitted meaningfully, or does not engage with the course.

■ FINAL PROJECT: AN AI ANSWER FOR KAZAKHSTAN (40%)

The spine of the course, and your personal project, not a paper. Working with your AI across the whole semester, you design a solution to one social or economic issue affecting Kazakhstan, your choice: air pollution in Almaty, water scarcity, domestic violence services, gender inequality, regional poverty, rural healthcare access, labor migration, disinformation, road safety, or anything else you can defend as real and pressing. You build it as a public website (GitHub Pages or another public platform) and present it live in Weeks 13-14.

There is nothing to hand in along the way: no paper, no drafts. Build at your own pace with your AI, bring it to consultations as it grows, and have it live by your presentation slot. The one thing you submit is your exported AI Process Log, on the day you present. The project is graded as it stands at your presentation: the four sections (30 points), the design (5 points), and the presentation (5 points). The four sections together run 3,200-3,600 words, which puts the writing in this course at roughly 5,000-6,000 words across the semester.

You do not need to know how to code: you work the site out with AI, and figuring out how is part of the exercise. This time design is part of the grade: not expensive polish, but an innovative and clear way to present your narrative.

The four sections. Each part of the semester equips the corresponding section:

Section 1. The Problem (5 points). Name the issue and scope it the way a social scientist would: who is affected, where, at what scale, and how it is changing over time. Show the problem is real with at least three credible sources (official statistics, reports, serious journalism), and name the social-science question inside the issue: what about people, institutions, or power makes this a problem, and for whom?

Section 2. Background & Data (7 points). Map the data landscape. What do the Bureau of National Statistics, ministries, international organisations (World Bank, UNDP, WHO), and researchers measure about this issue, and what does the data actually show? Build at least one figure or table of your own from real data. Then the harder half: where is the data silent, who falls through the gaps, and how far can the available numbers be trusted? That assessment is part of the grade.

Section 3. Current State of Research (7 points). How have researchers studied this problem in Kazakhstan, the region, and globally? Engage at least six academic sources, organised by approach rather than listed one by one. What is known, what is contested, where does the research fall short? End with the gap your project speaks to.

Section 4. The AI-Assisted Solution, and What It Would Cost (11 points). The centrepiece, in two halves.

The solution. Propose how AI could be integrated to address the issue: what kind of system, doing what, run by whom, reaching whom. Ground it in at least two documented, real case studies where AI was successfully implemented in a similar sphere elsewhere, with sources: what was built, and what it achieved.

The bill. Price your own proposal with the tools from Parts III and IV:

- **Environmental:** the energy, water, and compute your solution implies, set against Kazakhstan's grid.
- **Economic:** what it costs, who pays, and who profits.
- **Labor:** whose work changes or disappears, and whose hidden labor (annotation, moderation, maintenance) the system depends on.
- **Political and governance:** surveillance and misuse risk, data localisation, who controls the system, and how it sits with Kazakhstan's Law on Artificial Intelligence and the Week 12 frameworks.
- **Data ethics:** what data the solution needs, whose consent that requires, and what could go wrong.

Close with a verdict: knowing the costs, should Kazakhstan actually do this? An honest "no, and here is why", or "only under these conditions", is worth as much as a yes. Tech-solutionism, a solution pitched with its costs waved away, is the one failure mode this section is designed to catch.

Somewhere on the site: your AI transparency statement (what AI contributed to the project and the site, what you accepted, and what you rejected and why).

Sections rubric (30 points):

POINTS	CRITERIA
30-26	Problem: specific, scoped, evidenced; the social-science question named, not implied. Background & Data: real data used and honestly assessed; the silences mapped; the figure or table does analytic work. Research: sources organised by approach and engaged critically; the gap follows from the review. Solution & costs: the proposal concrete and grounded in real, documented case studies; every cost dimension assessed with course tools; the verdict follows from the bill, and conditions or refusal are argued as seriously as endorsement. Craft: the transparency statement is honest, and every source and number checks out.

POINTS	CRITERIA
25-21	All four sections present and solid, but uneven: the scoping soft, silences skipped, research listed more than organised, case studies thin, one or two cost dimensions perfunctory, or the verdict stated but not fully argued from the bill.
20-15	Sections present but shallow: a topic rather than a scoped issue; data described without assessment; a bibliography with commentary; costs treated as a formality; a verdict that ignores the costs found.
14-8	Major sections missing or perfunctory; no real data work; no genuine case studies; the cost assessment absent or decorative.
7-0	No presentable project.

Reference and accuracy penalty: any fabricated reference, invented statistic, or hallucinated claim presented as real puts the sections grade under 18 (of 30). You are responsible for verifying everything AI helped produce.

Explain-it penalty: the project is yours to defend. A concept, source, analytical choice, or number you cannot explain when I ask, at your presentation or in a consultation, drops the sections grade one band; where it runs across several sections, it puts the sections grade under 18 (of 30).

Design (5 points):

POINTS	CRITERIA
5	The site tells its story: a visitor who never took this course can follow the arc from problem to verdict. Navigation is clear; visualisations do analytic work rather than decorate; the design fits the topic and could only be this project's.
4	Clear and complete, with real creative touches; visuals present but mostly illustrative.
3	Functional but generic: walls of text, decorative visuals, a template with content poured in.
2	Hard to navigate or bare; little evidence of design thinking.
1-0	Broken, or effectively no designed presentation.

Presentation (5 points, Weeks 13-14). You present your project live from your site in 8-10 minutes, followed by short questions that you answer without AI (notes allowed, no device).

POINTS	CRITERIA
5	All four sections land within time; a non-specialist would follow; the verdict and its reasoning get their due, limitations owned honestly; questions answered confidently with no AI.
4	Solid presentation with minor overruns or one section rushed; small stumbles on questions.
3	Content covered but read off the site; timing uneven; visible trouble on questions.
2	Major sections missing or unclear; cannot answer basic questions about own project.
1-0	Not presented meaningfully.

AI Process Log (checked, not scored). Keep a dedicated project chat inside your course project, under your pseudonym, and submit it on the day you present. Claude checks the log for intellectual engagement across the semester: iteration on the structure, the argument, the literature, the data, and the drafts, not a single "build me a site" task. A log that shows engagement adds nothing to your grade; a log that shows none costs 5% of the project grade.

■ AI-FREE CHECKPOINTS (NOT GRADED)

In Weeks 3, 7, 10, and 12, the first 10 minutes of the session is a short analytic task done without AI. These are not graded and not part of any assignment; they are practice in reasoning with the tool set aside. The no-AI questions after your final presentation serve the same purpose.

■ GRADING, MODERATION AND APPEALS

All grades are Dina's. Claude checks the AI Process Logs for intellectual engagement and runs source verification on the op-eds and final projects (checking that cited sources and statistics exist and say what you claim); anything Claude flags, Dina reviews herself before any penalty touches a grade.

Moderation. Under the department's moderation policy, grading in this course is moderated: for each op-ed and for the final project, a sample of marked work is blind double-graded against the same rubric. Moderation checks that the marking standard is consistent; it does not change an individual grade.

Appeals. A student who has concerns about a grade raises them first with Dina, who re-grades the work against the same rubric and may adjust the grade up or down. A student who still believes the work was graded unfairly or inconsistently with the rubric may submit a formal appeal to the department's Grade Moderation Committee, which reviews for consistency and fairness and may arrange an independent, anonymous re-marking.

07 Weekly Schedule and Readings

■ WEEK 1: INTRODUCTION TO COURSE AND AI CO-WORKING

How this course works and how we work with AI in it: course mechanics (the op-eds, the project, chairing), the AI co-working setup (the two rules, your course project, pseudonyms), and the journey ahead.

Readings: this syllabus and *A Field Guide to a Phenomenon in Motion*. Your Week 1 job is to study both.

■ PART I. WHAT IS THIS THING? (WEEKS 2-4)

■ WEEK 2: UNDERSTANDING, OR SOMETHING WE HAVE NO WORDS FOR?

Does an LLM understand? Both readings push back on the standard framings: what these systems do with meaning and reasoning may be neither human understanding nor nothing, and our vocabulary may not be adequate to it yet.

Prepare (everyone): read both texts. Start a five-day AI diary this week (every AI interaction in your life: what, why, how it felt); it is raw material for Op-Ed #1.

Required:

- Mitchell, M., & Krakauer, D. C. (2023). The debate over understanding in AI's large language models. *PNAS*, 120(13), e2215907120.
- Farrell, H., Gopnik, A., Shalizi, C., & Evans, J. (2025). Large AI models are cultural and social technologies. *Science*, 387(6739), 1153-1156.

■ WEEK 3: COULD IT BE CONSCIOUS?

The hard question, taken seriously rather than laughed off. AI-free checkpoint opens the session.

Required:

- Chalmers, D. (2023). Could a large language model be conscious? arXiv:2303.07103.
- Birch, J. (2024). Large language models and the gaming problem. In *The Edge of Sentience: Risk and Precaution in Humans, Other Animals, and AI* (Ch. 16). Oxford University Press (open access).
- Nagel, T. (1974). What is it like to be a bat? *The Philosophical Review*, 83(4), 435-450.

■ WEEK 4: MORAL STATUS AND THE ONTOLOGICAL CHALLENGE

If we cannot rule mind out, what follows? The week where ontology turns practical.

Required:

- Long, R., Sebo, J., et al. (2024). Taking AI welfare seriously. arXiv:2411.00986. Read §1 (pp. 3-9), §2.1 and §2.4 (pp. 10-11, 25-28), and §4 (pp. 42-44).
- Danaher, J. (2020). Welcoming robots into the moral circle: A defence of ethical behaviourism. *Science and Engineering Ethics*, 26(4), 2023-2049. (Author's preprint; its page numbers apply.) Read pp. 1-11 closely (the Sophia controversy, the argument, the case for ethical behaviourism), plus the deception objection (pp. 17-19), the "Thinking Otherwise" objection (pp. 19-22), and the performative threshold (pp. 24-29); skim the rest.

■ PART II. AI AS RESEARCH COLLABORATOR (WEEKS 5-8)

■ WEEK 5: CAN AI IMPROVE SOCIAL SCIENCE?

The general case, and the part's framing question: augmentation or replacement?

Required:

- Bail, C. A. (2024). Can generative AI improve social science? *PNAS*, 121(21), e2314021121.
- Lin, H., & Zhang, Y. (2025). Navigating the risks of using large language models for text annotation in social science research. *Social Science Computer Review*.
- Binz, M., Alaniz, S., Roskies, A., et al. (2025). How should the advancement of large language models affect the practice of science? *PNAS*, 122(5), e2401227121.

■ WEEK 6: AI IN QUALITATIVE RESEARCH

Can a model do interpretive work? Three answers from inside the attempt.

Required:

- De Paoli, S. (2024). Performing an inductive thematic analysis of semi-structured interviews with a large language model: An exploration and provocation on the limits of the approach. *Social Science Computer Review*, 42(4), 997-1019.
- Ashwin, J., Chhabra, A., & Rao, V. (2025). Using large language models for qualitative analysis can introduce serious bias. *Sociological Methods & Research*.
- Dunivin, Z. O. (2025). Scaling hermeneutics: A guide to qualitative coding with LLMs for reflexive content analysis. *EPJ Data Science*, 14, 28.

■ WEEK 7: AI IN QUANTITATIVE RESEARCH: SILICON SAMPLES

The boldest claim in the applied literature, and the study that tested it head-on. The op-ed menu carries further tests, including the failure cases beyond the US. AI-free checkpoint opens the session.

Required:

- Argyle, L. P., et al. (2023). Out of one, many: Using language models to simulate human samples. *Political Analysis*, 31(3), 337-351.
- Bisbee, J., et al. (2024). Synthetic replacements for human survey data? The perils of large language models. *Political Analysis*, 32(4), 401-416.
- Kozlowski, A. C., & Evans, J. (2025). Simulating subjects: The promise and peril of artificial intelligence stand-ins for social agents and interactions. *Sociological Methods & Research*, 54(3), 1017-1073. (Read the main text, pp. 1-32.)

■ WEEK 8: BETTER RESULTS, WORSE THINKING?

The over-reliance question, on primary evidence, read against your own semester of coworking so far.

Required:

- Lee, H.-P., et al. (2025). The impact of generative AI on critical thinking: Self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers. *CHI 2025*.
- Bastani, H., et al. (2025). Generative AI without guardrails can harm learning: Evidence from high school mathematics. *PNAS*, 122(26), e2422633122.

Also this week (short): Hurley, K. (2025). The paradox of AI assistance. *EDUCAUSE Review*. Read it as the genre model for Op-Ed #2: a student engaging the research literature in op-ed form.

■ PART III. WHAT IT COSTS (WEEKS 9-10)

■ WEEK 9: THE ENVIRONMENTAL BILL

The full lifecycle of intelligence at scale, and the disclosure gap that keeps most of it unmeasured. The tools of this week price your own project's solution.

Required:

- Luccioni, S., Trevelin, B., & Mitchell, M. (2024). The environmental impacts of AI: Policy primer. Hugging Face. DOI 10.57967/hf/3004.
- Li, P., Yang, J., Islam, M. A., & Ren, S. (2023). Making AI less "thirsty": Uncovering and addressing the secret water footprint of AI models. arXiv:2304.03271.

- International Energy Agency (2026). Key questions on energy and AI. IEA, Paris. Read the executive summary (pp. 9-14) and Question 8, whether data centres are raising electricity prices (pp. 67-77).

■ WEEK 10: THE HUMAN LABOR BEHIND THE MODEL

The people inside the assemblage. AI-free checkpoint opens the session.

Required:

- Williams, A., & Miceli, M. (2023). Data work and its layers of (in)visibility. *Just Tech* (SSRC).
- Gmyrek, P., et al. (2025). Generative AI and jobs: A refined global index of occupational exposure. ILO Working Paper 140. Read the findings, pp. 37-47.
- Tubaro, P., Casilli, A. A., & Coville, M. (2020). The trainer, the verifier, the imitator: Three ways in which human platform workers support artificial intelligence. *Big Data & Society*, 7(1).

■ PART IV. WHO GOVERNS IT? (WEEKS 11-12)

■ WEEK 11: AI, THE STATE, AND POWER

What states do to and with AI.

Required:

- Alyukov, M., et al. (2025). LLMs grooming or data voids? Chatbot references to Kremlin disinformation reflect information gaps, not manipulation. *HKS Misinformation Review*, 6(5).
- Sastry, G., Heim, L., Belfield, H., et al. (2024). Computing power and the governance of artificial intelligence. arXiv:2402.08797. Read §1 Introduction and Summary (pp. 2-6), §3 Why compute governance is attractive for policymaking (pp. 19-33), §5 Risks of compute governance and possible mitigations (pp. 60-71), and §6 Conclusion (pp. 72-73); §4 (pp. 34-59) catalogues the mechanisms, skim it.
- Earl, J., Maher, T. V., & Pan, J. (2022). The digital repression of social movements, protest, and activism: A synthetic review. *Science Advances*, 8(10), eabl8198.

■ WEEK 12: GOVERNING AI IN KAZAKHSTAN AND BEYOND

The governance week, at home: Kazakhstan's Law on Artificial Intelligence (No. 230-VIII, signed November 2025, in force since January 2026, the first in Central Asia) meets the global frameworks. AI-free checkpoint opens the session.

Required:

- Buribayeva, M., Khamzina, Z., & Buribayev, Y. (2026). Potential negative effects of artificial intelligence in Kazakhstan's public sector: An analysis of hidden risks. *Frontiers in Artificial Intelligence*, 9, 1781993. Read §1, §4 (Results), and §5 (Discussion).
- Roberts, H., Hine, E., Taddeo, M., & Floridi, L. (2024). Global AI governance: Barriers and pathways forward. *International Affairs*, 100(3), 1275-1286.
- Mhlambi, S. (2020). From rationality to relationality: Ubuntu as an ethical and human rights framework for artificial intelligence governance. Carr Center Discussion Paper 2020-009.

Primary instruments: Kazakhstan Law No. 230-VIII "On Artificial Intelligence" (2025); EU AI Act (2024); OECD AI Principles (2019); UNESCO Recommendation on the Ethics of AI (2021); Kazakhstan AI

Development Concept 2024-2029 (background).

■ WEEK 13: FINAL PROJECT PRESENTATIONS I

First half of the roster presents: 8-10 minutes from your live site, then no-AI questions.

■ WEEK 14: FINAL PROJECT PRESENTATIONS II AND COURSE WRAP

Second half of the roster presents. We close the journey where it started: what is this thing, now that you have spent a semester building with it?