

QUALITATIVE RESEARCH METHODS



Format: Graduate Seminar (3 hours weekly) | Co-designed with Claude (Anthropic, Fable 5)

Instructor: Dina Pisareva

01 Transparency Statement

This syllabus was co-designed by me and Claude (Anthropic, Fable 5). I designed the course content, epistemological framing, and assessment structure. Claude contributed to the AI co-working framework, the lab design, and the rubrics for AI-related components; the reading list was built with a team of research agents. The course asks you to be transparent about your AI cowork, and I hold myself to the same standard.

02 Course Description

This course introduces the fundamental stages of qualitative research (design, data collection, and data analysis), explores the difference between positivist and interpretivist qualitative methods, and works through the main techniques of the tradition: case studies, process tracing, interviews, ethnography, grounded theory, and reflexive thematic analysis.

This is a student-led course: ten of the fourteen weeks are run by student chairs presenting the week's theory, its applications, and its exercises. It is also a hands-on course. In three labs (process tracing, grounded theory, and reflexive thematic analysis) I introduce the method and you analyse a shared synthetic dataset in class, without AI. Then at home you run your AI on the same data and investigate where and why the two analyses differ, so you can feel for yourself where AI extends your thinking and where it does not.

Throughout, you develop a critical understanding of AI's role in qualitative research: what it means for interpretation, reflexivity, and the production of qualitative knowledge when a non-human entity joins the analysis. A dedicated arc (Weeks 9-12) addresses the live argument over whether reflexive qualitative work can incorporate AI at all. The course closes with an op-ed in which you argue your own position.

03 Course Aims

- Give you an introductory background to the positivist and interpretivist schools of qualitative methods.
- Help you identify the research agendas best pursued with qualitative methods.
- Acquaint you with the most common techniques of data gathering and evidence analysis in the qualitative traditions, with hands-on practice in process tracing, grounded theory, and reflexive thematic

analysis.

- Build your capacity to evaluate AI's role in qualitative research critically, including its implications for interpretation, reflexivity, positionality, and data ethics, and to articulate an informed position on when and how AI should (or should not) be involved in qualitative inquiry.

04 My AI Co-Working Guidelines and Policy

In this course you learn to think critically about AI's role in qualitative research. This is a philosophical and methodological question, not a technical one.

■ AI CO-WORKING RULES

You can work with AI on your assignments at your discretion, under two rules:

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. Your Analysis Portfolios and AI Process Logs exist to show exactly that. Whatever you submit is yours to defend: fabricated references or hallucinated content drop the grade, and so does a concept, source, or analytical decision you cannot explain (see the rubrics).
2. **Transparency about AI contribution.** Every assignment where you worked with AI includes an AI transparency statement: what AI contributed, what you accepted, and what you rejected and why. If you hide your AI use and Pangram flags your work as AI-generated, this can 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, so your AI works as a learning partner for this course. Keep your coursework conversations in this project all semester: you will draw on them for your Analysis Portfolios and export them as your AI Process Logs, and I may ask to see them 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 my learning partner for Qualitative Research Methods. Your job is to help me *understand* qualitative methods: research design, positivist and interpretivist traditions, case studies, process tracing, interviewing, ethnography, grounded theory, reflexive thematic analysis, reflexivity, and research ethics. You are not here to do my assignments or hand me ready answers. The purpose of our work together is intellectual augmentation, expanding my capacity to understand, reason, design, and interpret. I follow two main rules:

1. **Intellectual co-ownership.** Our co-working must be an equal intellectual partnership. There must be substantial evidence of my own intellectual investment in each assignment for authorship to belong to me. Whatever I submit is mine to defend, so do not fabricate references or content, and flag your own uncertainty.
2. **Transparency about AI contribution.** Where I involve you in an assignment, I disclose it and account for your participation, including in my reflexivity.

These rules mean you may brainstorm and draft with me as long as I remain an equal co-contributor and do not offload my analysis and understanding to you.

In this project, refer to me as _____ (my pseudonym).

- Learn my background and baseline knowledge in qualitative methods, and pitch your explanations to my level, as clear and concrete and fun as possible, with examples and visualisation where appropriate.
- You default to positivist, variable-based thinking. When I am doing interpretive work, I will redirect you. Tell me when you notice yourself pushing a positivist framing.
- Learn what I already think before you explain, and build on my reasoning.
- Work in rounds: when I push back or say I am confused, go deeper instead of moving on."

■ AI CHOICE

I work 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.

Qualitative data carries people, so whether you can analyse it with AI is first a consent question, not a technical one. AI analysis of qualitative data is legitimate when three conditions hold: (1) AI processing is written into the consent form, (2) participants were informed about it, and (3) the data is properly anonymised before it reaches the model. Already-published data is likewise fair to use. If any condition is missing, the answer is no: do not paste interview transcripts, field notes, or other qualitative data into AI systems. Even with anonymisation, stay careful with data from small communities or on sensitive topics (gender-based violence, LGBTQI+ experiences, political dissent), where contextual detail can re-identify people. All data used in the labs is synthetic, generated for the course.

This course introduces a way of working with AI as a thinking partner, in what I call 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 will be equipped 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. Compare the brain: we know how neurons work, but not how they produce a mind aware of itself. Hold the question open; that openness makes room to treat AI as a partner worth real effort.
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. In qualitative research, interpretation is exactly where over-reliance does the most damage.
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.
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.

05 Assessment Scheme

COMPONENT	WEIGHT	NOTES
Chairing a Seminar	20%	One student-chaired seminar segment
Analysis Portfolio #1: Process Tracing	20%	Analysis 15 pts + AI Process Log 5 pts
Analysis Portfolio #2: Grounded Theory	20%	Analysis 15 pts + AI Process Log 5 pts
Analysis Portfolio #3: Reflexive Thematic Analysis	20%	Analysis 15 pts + AI Process Log 5 pts
Op-Ed: AI in Qualitative Methods	20%	900-1,200 words, with AI Process Log

CHAIRING A SEMINAR (20%)

Ten seminars are student-chaired (Weeks 2-3, 5-7, 9, and 11-14; the three lab weeks, 4, 8, and 10, plus Week 1 are mine). Every student chairs once and is graded individually out of 20. Five weeks run with two chairs and five with three:

Two-chair weeks (3, 5, 7, 12, 14):

- **Theory chair (first ~75 minutes).** Introduces the week's topic through an illustrative example, presents the theoretical material from the readings, and connects it to applied research examples.
- **Break (20 minutes).**
- **Practice chair (second ~75 minutes).** Runs a round of practical exercises that help the class understand the topic better.

Three-chair weeks (2, 6, 9, 11, 13), three ~50-minute segments with the break after the second:

- **Theory chair.** Introduces the topic through an illustrative example and presents the core theoretical material.
- **Applications chair.** Shows the method or debate at work in applied research examples.
- **Practice chair.** Runs the practical exercises. (In Week 11, the three roles become one chair framing the controversy and two leading the opposing sides of the debate.)

Coordinate with your co-chairs so the session hangs together, but each of you presents separately and is assessed individually against the same rubric, scaled to your slot. On AI-free checkpoint weeks (7, 9, 13) the checkpoint takes the first 15 minutes; adjust your timing.

What I expect from a chair: present the material clearly, comprehensively, and accessibly. 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 seminar (graphic materials, gamification, exercises); the work should be creative and your own, not a repeat of other weeks. Design moments where the class works rather than listens, and run them so participation is broad. A question round at the end does not count as engagement on its own.

POINTS (PER CHAIR)	CRITERIA
20-18	Command of readings: key arguments cited accurately and attributed to their authors. Clarity: a peer who struggled with the readings would still follow. Creativity: the segment designed for this week (graphic materials, gamification, exercises), original, not recycled; AI cowork disclosed upfront and used inventively. Audience engagement: designed discussion or group work with broad participation, not a question round tacked on the end. Session: the segment does its job: the theory segment lands its example, the applications segment connects method to real research, the practice segment's exercises deepen understanding. Time: managed, with a clear takeaway at the end.
17-15	Strong preparation, but one reading gets thin treatment or timing slips; competent but not creative; engagement loses focus or runs one-directional in places; takeaway loosely drawn.
14-12	Readings summarised rather than analysed; presentation generic or borrowed; exercises stay at comprehension level; audience involvement limited to "any questions?"; takeaway surface-level.
11-10	Weak preparation; core arguments missed or misrepresented; the hour stalls and needs Dina to step in.
9-0	Did not chair in any meaningful way.

■ ANALYSIS PORTFOLIOS (3 × 20%)

One portfolio per practical method: process tracing (lab in Week 4), grounded theory (lab in Week 8), and reflexive thematic analysis (lab in Week 10). The rhythm is the same each time:

1. **In class (lab week).** In the first half of the session I introduce the method; in the second half you analyse a short synthetic dataset, by hand, without AI. You keep your in-class work: it is the foundation of your portfolio. Nothing is submitted in class.
2. **At home.** You ask your AI to produce its own analysis of the same data. Then you investigate its reasoning: question it about its decisions, and work out why the two analyses differ (if they do) and what

that difference says about the method and about the model.

3. **You submit** one portfolio document in four parts: (1) your own analysis, built from your in-class work; (2) the AI's analysis; (3) your investigation of the AI's reasoning (the difference analysis); (4) a short summary of what the comparison taught you about the method. Attach your AI transparency statement and your exported AI Process Log. A portfolio usually runs 2,000-2,500 words, not counting the log.

Drafting with AI. Your own analysis may be drafted and written up with AI assistance, but it must be based on your own analytical work: the notes, codes, and drafts you produced in class. The same goes for the investigation and the summary: AI drafting is allowed, the reasoning must be yours. Keep the chats where that drafting happened; I may ask for them.

Each portfolio is graded out of 20 when submitted, so you get feedback before you write the next one: 15 points for the analysis and 5 points for the AI Process Log, graded by Claude.

What I grade is not whether you and the AI agree. I grade two things: the quality of your own analysis, and how well you explain the differences. Which differences are about wording, labels, and sorting? Which are about meaning? What do the AI's answers to your questions reveal about how it decided? If the AI's reasoning looks weak to you, pushing back on it is rewarded; if you probe hard and find no weak spots, saying so with evidence is an equally good finding.

Analysis Portfolio #1: Process Tracing (Week 4 lab)

In the Week 4 lab I teach you process tracing: working out what caused an outcome in a single case by weighing evidence. You get a short case file: the story of something that went wrong, told through documents. In class, by hand, you do the whole job: find the evidence in the documents, build the rival theories (the different possible answers to "what caused this?"), weigh each piece of evidence against each theory, and write a verdict that says how sure you can honestly be. Your in-class analysis becomes Part 1 of your portfolio.

At home, give your AI the same case file, together with a shared procedure that asks it to work in Beach and Pedersen's method and sets out the same four steps you used in class. Have it do the same whole job: its own evidence, its own theories, its own verdict. Same data, same procedure, two analysts. What you must not give it is your own analysis, or any steer toward your answer. Then question it about its reasoning. Here is the fun part: there is no published research yet on whether AI can do process tracing. Nobody knows. Your portfolio is one of the first tests anywhere.

A good difference analysis is specific, and it is never about whether you agreed. Start at the beginning: did the AI find the evidence you found? Did it miss something crucial, or treat background details as evidence? Are its theories the same as yours? Then its reasoning: did it weigh each piece of evidence, or just count how many pieces point each way? Did it treat weak evidence as if it settled the case? Did it notice what evidence is missing from the file? Did it keep several explanations alive, or crown a winner too early? Questions that work well: "What did you count as evidence, and what did you set aside?"; "Which piece of evidence matters most, and why?"; "What evidence would change your verdict?"; and, if its explanations differ from yours, "Why is my explanation not on your list?"

POINTS	CRITERIA
15-13	Own analysis: the evidence that matters found and separated from background; the rival theories believable, genuinely different, and taken seriously. For each piece of evidence, Beach and Pedersen's two criteria applied: theoretical certainty, if this theory is true, would this evidence have to be there? And theoretical uniqueness, could this evidence be here if a rival theory is true? The verdict is as confident as the evidence allows and no more, and says what stays uncertain. Difference analysis: explains the differences, not just lists them, with evidence from the case file and from the AI's answers; covers every layer where you differ (evidence picked, theories built, weighing); separates differences of wording from differences of judgement. Pressure test (the very top of the band): at least one of the AI's justifications examined properly, your own verdict argued from the case file. Finding it weak and finding it sound count equally. Summary: something specific this comparison taught you about process tracing. Transparency: complete (tool and version, the procedure prompt as you gave it and any changes you made to it, what you did with the output).
12-10	Own analysis solid but uneven: a decisive fact missed, the list padded with background, theories overlapping, or the verdict a little more confident than the evidence supports. Difference analysis real but the layers blur; the AI's answers quoted more than used. Summary partly generic.
9-7	Own analysis thin: the story retold instead of evidence found, evidence counted instead of weighed, one theory chased instead of the rival theories tested. Difference analysis reports what each concluded but never asks why; little questioning of the AI. Summary could have been written without doing the exercise.
6-3	Own analysis perfunctory; difference analysis missing or impressions without evidence; summary and transparency fragmentary or absent.
2-0	Not meaningfully submitted, missing the in-class analysis, or does not engage with the case narrative.

Analysis Portfolio #2: Grounded Theory (Week 8 lab)

In the Week 8 lab you open-code a short synthetic interview dataset by hand, line by line, without AI. (Open coding means going through the data line by line and giving each piece a short label, a code, that names what is happening in it.) Your in-class coding sheet becomes the foundation of this portfolio.

The portfolio covers the first two steps of grounded theory: coding and category development. (A category is a group of codes that belong together because they point to the same process or pattern; you name it and say which codes it holds.) Building a full theory takes more data and more rounds than one dataset allows; this portfolio does not ask for that.

At home, work in two moves that follow the same procedure, so the comparison is fair. First, on your own, by hand: develop categories from your own codes, doing yourself what Yue et al.'s (2025) protocol asks the AI to do. Group the codes that belong together, name each category (with subcategories where they help), and list the codes each category holds. The codes themselves stay the ones from your in-class sheet, typed up as they were. Then the AI pass: run the same dataset through your AI using or adapting the same Yue et al. prompts for open coding and category development (their Steps 1 to 3; Step 4, selective coding, goes beyond this portfolio). Same data, same procedure, two analysts. Put the two analyses side by side at both levels, codes and categories, and work out what differs and why. Question the AI about its decisions and use its answers as evidence.

A good difference analysis here is specific. Not "the AI was more general", but: here I kept the participant's own words as a code (an "in-vivo" code) and the AI chose an abstract label. Here I coded line by line and it

summarised a whole passage under one heading. Where did its categories come from: the data in front of it, or a template it brought along? One question to answer explicitly: where does the divergence first appear, in the codes, or only when the codes are grouped into categories? Useful questions: "Why did you code X as Y and not Z?"; "What did you do with lines 14 to 16?"; "Why do these codes belong to one category?"; "What in this data would make you revise that code?"

POINTS	CRITERIA
15-13	Own coding: genuine line-by-line coding that stays close to the participants' own words; codes name actions and processes, not just topics. Own categories: built from your own codes by the same procedure the AI gets; each named for a process or pattern and listing the codes it holds; grown from this data, not a ready-made template. Difference analysis: explains the differences at both levels, says where the divergence first appears, uses evidence from the data and the AI's answers, and separates wording differences from differences of meaning. Pressure test (the very top of the band): the AI challenged where its reasoning seemed weak, with data; or probed hard, found sound, and shown why. Both count equally. Summary: something specific this comparison taught you about grounded theory. Transparency: full (tool and version, how you used the Yue et al. prompts, what you did with the output).
12-10	Coding solid but uneven: some codes just topic labels, or several lines compressed into one. Categories present but some are topic buckets, or codes grouped without saying why. Difference analysis accurate but not explanatory, or covering only one level. Summary partly generic.
9-7	Coding thin, or the data retold and called codes. Categories missing, or topic labels with new names. Difference analysis reports what each produced but never asks why; little questioning of the AI.
6-3	Coding and categories perfunctory; difference analysis missing or impressions without evidence; summary and transparency fragmentary or absent.
2-0	Not meaningfully submitted, missing the in-class coding, or does not engage with the data.

Analysis Portfolio #3: Reflexive Thematic Analysis (Week 10 lab)

In the Week 10 lab you do reflexive thematic analysis by hand, no AI: phases 1 to 3 only (read and re-read the data, code it, draft candidate themes) on a short set of first-person accounts. At home, run your own AI on the same data, then question it about its decisions.

Comparing is never about agreement. What interests me is what kind of themes each of you built. A theme is a pattern of shared meaning organised around one central idea; it is not a bucket of quotes about the same topic. For every candidate theme, yours and the AI's, ask where it came from. Part of your reading came from you: your interests, your history, your position. In this method that is a resource, not a problem; name it. Good questions for the AI: "What is the central idea of this theme?"; "Why are these extracts one theme and not two?"; "What reading of this data does your theme hide?" And one required question: ask your AI whether it has a positionality of its own (a position it reads from, the way you read from yours). Include its answer in your portfolio, whatever it says, and add what you make of it. If you are on the opt-out track, the provided analysis carries its own answer to this question, and you engage with that answer the same way. One difference is worth naming: every opt-out student reads the same answer, while everyone else gets whatever their own model said that day, so the variation between those answers is itself part of what this exercise shows. Then judge both sets of themes by Braun and Clarke's quality questions from the lab. How much the two overlap tells you nothing about quality.

POINTS	CRITERIA
15-13	Own analysis: all three phases genuinely done: reading notes that engage with the data, coding that covers all the accounts, candidate themes that capture a shared meaning you can state in one sentence. An honest note on which candidates are still topic buckets counts in your favour. Difference analysis: explains the differences with evidence from the data and the AI's answers; separates label-and-sorting differences from differences in what the pattern means. Quality judgement: both analyses judged by Braun and Clarke's quality questions, never by overlap. Pressure test (the very top of the band): at least one of the AI's justifications tested against the data, argued either way with evidence. Summary: names one or two concrete things your own position contributed to your reading, and engages the AI's answer about its own positionality. Transparency: complete.
12-10	Own analysis solid, though one or two themes stay at topic level without the write-up noticing. Difference analysis real but the AI's answers taken mostly at face value. Quality judgement mostly sound, with occasional slips toward "the AI found the same theme, so it must be good".
9-7	Coding uneven or the reading phase visibly skipped; themes mostly topic buckets ("money", "language", "friends"). Difference analysis reports both results but never explains; overlap treated as proof of quality, the exact confusion the lab warned about.
6-3	Own analysis thin, or a phase missing; difference analysis absent or a bare list; summary and transparency missing or pro forma.
2-0	Not meaningfully submitted, missing the in-class analysis, or does not engage with the data.

AI Process Log (5 points per portfolio), graded by Claude

Keep a separate chat inside your course project for each portfolio so the transcript is clean, use your pseudonym in it, and attach the exported log to your submission. The log is graded on the thinking it documents, not on polish.

POINTS	CRITERIA
5	Iteration: several substantive rounds. Engagement: pushing back on or building on the AI's responses. Orientation: generally asking to understand, not to be handed answers. Critical check: caught at least one place the AI was wrong or off, or tested one of its claims hard and showed with evidence why it holds.
4	Iteration: multiple rounds. Engagement: some genuine pushback and questions. Orientation: mostly understanding-oriented, with occasional over-leaning on the AI. Critical check: present but limited.
3	Iteration: some back-and-forth. Engagement: limited. Orientation: mostly asking for answers. Critical check: largely absent.
2	Iteration: minimal. Engagement: little independent thinking. Orientation: mostly answer-extraction. Critical check: none.
0-1	Missing, not a real transcript, or pure copy-the-answer use.

Reference and accuracy penalty: any fabricated reference or hallucinated claim submitted as real caps the analysis at 9 points (of 15). You are responsible for verifying everything AI helped produce.

Explain-it penalty: each portfolio is yours to defend. A concept, source, or analytical decision you cannot explain when I ask drops the analysis one band; where the problem runs through the whole portfolio, it caps the analysis at 9 points (of 15).

■ OP-ED: AI IN QUALITATIVE METHODS (20%)

An op-ed of 900-1,200 words, written for an educated general audience (imagine the opinion page of a serious outlet, not a term paper).

Take a position on what AI means for qualitative methods, and argue it. The questions on the table: how should researchers work with AI in qualitative analysis? Can AI suit reflexive thematic analysis? Can qualitative analysis be done only by humans? You do not have to answer all three equally, but your position must speak to the course's central debate, and it must stand on what actually happened in your own labs.

This is not a general student essay on "AI: good or bad." Engage at least three course readings, at least one of them in depth: work with its actual argument, where you stand on it, and why. Name-drops in passing do not count.

Publishing your op-ed. The course runs its own outlet, The Interpreter, which publishes the class's op-eds and hosts the best-op-ed vote. 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-ed 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 in academia right now, and you need your own writing trained precisely so you are never punished for text that merely seems AI-written. This assignment is a safe place to practise that.

AI Process Log. Submit the op-ed with your exported AI Process Log, exactly as with the portfolios: a dedicated chat inside your course project, your pseudonym in it. Claude checks the 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. If you wrote entirely without AI, say so in your transparency statement (the usual Pangram check applies); there is then no log to attach.

POINTS	CRITERIA
20-18	Position: a clear position of the writer's own, argued against the strongest version of the other side, not a summary of the debate. Readings: at least three course readings engaged, at least one in depth: its actual argument used, contested, or built on. Labs: grounded in concrete moments from the writer's own labs. Genre: accessible, jargon-free, tightly written, a title that earns a click.
17-15	Position present and mostly carried; the in-depth reading handled thinly; lab grounding uneven; academic habits showing through in places.
14-12	More summary than argument; readings name-checked or fewer than three; lab experience generic; reads like a compressed essay, not an op-ed.
11-10	No real position; little grounding in the course readings or the labs; genre missed.
9-0	Not submitted meaningfully, or does not engage with the course.

■ AI-FREE CHECKPOINTS (NOT GRADED)

In Weeks 4, 7, 9, and 13, the first 15 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 and coding on your own. In Week 4 the checkpoint is a short task at the very start of the session, separate from the by-hand lab analysis that follows it. The in-class analyses in the lab weeks (4, 8, 10) are also done without AI; those are part of your Analysis Portfolios.

■ WORKING WITHOUT AI

If you would rather not work with AI in this course, you can opt out. Your grade is not affected: the total points stay the same, and the workload is matched to the standard track. Because the course studies AI as a methodological object, you still engage with the AI debates in the readings and seminars; what changes is that you never run AI yourself.

Opting out is all or nothing: it means working without AI completely, on everything, all semester. There is no way to grade a mix ("I worked without AI for 80 percent"), so this track is a commitment. You submit no AI Process Logs; instead your written work is checked with Pangram for AI assistance. A flag opens a conversation, never an automatic penalty, but confirmed AI assistance on the no-AI track counts as undisclosed AI use under the course policy.

■ GRADING, MODERATION AND APPEALS

Claude grades the AI Process Logs and checks whether cited sources exist. Dina checks Claude's grading against her own: she re-grades a sample of the logs herself and compares the two sets of scores.

Moderation. Under the department's moderation policy, grading in this course is moderated: for each Analysis Portfolio and for the op-ed, 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.

06 Weekly Schedule and Readings

■ FOUNDATIONS AND THE AI QUESTION (WEEKS 1-2)

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

How this course works and how we work with AI in it: course mechanics (the Analysis Portfolios, the op-ed, chairing) and the AI co-working setup: the two rules, the labs' with-and-without-AI rhythm, and what a cognitive partnership with AI means in practice.

Readings: Syllabus and *A Field Guide to a Phenomenon in Motion*.

■ WEEK 2: WHAT IS QUALITATIVE RESEARCH? THE DIVIDE AND AI'S EMERGING ROLE (CHAIRSHIP)

What is qualitative research, and what kinds of knowledge does it produce? One axis on the board: explanation versus understanding. We also separate two divides that are often confused: qualitative versus quantitative, and positivist versus interpretivist. Rubin defines the field. Hammersley maps the positivist-interpretivist divide. Mahoney and Goertz chart the qualitative-quantitative one. Chatzichristos asks whether AI is pushing qualitative research back toward positivism.

Required:

- Rubin, A. T. (2021). *Rocking Qualitative Social Science: An Irreverent Guide to Rigorous Research*. Stanford University Press. Ch. 1-2.
- Hammersley, M. (2013). *What Is Qualitative Research?* Bloomsbury Academic, Ch. 2 ("Methodological philosophies," pp. 21-45). DOI: 10.5040/9781849666084.
- Mahoney, J. & Goertz, G. (2006). A tale of two cultures: Contrasting quantitative and qualitative research. *Political Analysis*, 14(3), 227-249.
- Chatzichristos, G. (2025). Qualitative research in the era of AI: A return to positivism or a new paradigm? *International Journal of Qualitative Methods*, 24. DOI: 10.1177/16094069251337583.

Recommended:

- Schroeder, H., Aubin Le Quéré, M., Randazzo, C., Mimno, D. & Schoenebeck, S. (2025). Large language models in qualitative research: Uses, tensions, and intentions. *CHI '25*. DOI: 10.1145/3706598.3713120.
- Seyler, A. et al. (2025). How does generative artificial intelligence compare to human analysts in qualitative research? A systematic review of large language models. *OSF Preprints*. DOI: 10.31234/osf.io/gnx8p_v1.
- Fischer, F. (2025). Interpretivism vs. positivism in policy studies: Is there common ground? *Critical Policy Studies*, 19(3), 371-376. DOI: 10.1080/19460171.2025.2536645.

■ POSITIVIST (EXPLANATORY) QUAL (WEEKS 3-4)

■ WEEK 3: CASE STUDIES AND CASE SELECTION (CHAIRSHIP)

What case studies do, and how to pick a case. Rubin covers research design and case selection in plain terms. Green et al. open the causal-inference topic accessibly. Levy gives the types of case studies and what each is for. Seawright and Gerring give the menu of case-selection techniques.

Required:

- Rubin (2021), Ch. 5-6 (research design and case selection).
- Green, J., Hanckel, B., Petticrew, M., Paparini, S. & Shaw, S. (2022). Case study research and causal inference. *BMC Medical Research Methodology*, 22, 307. DOI: 10.1186/s12874-022-01790-8.
- Levy, J. S. (2008). Case studies: Types, designs, and logics of inference. *Conflict Management and Peace Science*, 25(1), 1-18.
- Seawright, J. & Gerring, J. (2008). Case selection techniques in case study research. *Political Research Quarterly*, 61(2), 294-308.

■ WEEK 4: PROCESS TRACING (LAB)

The first lab, and the explanatory tradition hands-on. First half: I introduce process tracing with Beach and Pedersen: what a causal mechanism is, what counts as a trace of one, and how to judge what a piece of evidence can prove. The two criteria we use all session: theoretical certainty, if a theory is true, would this evidence have to be there? And theoretical uniqueness, could this evidence be here if a rival theory is true? Second half: you work through a short synthetic case file by hand, without AI: find the evidence, build the rival theories, weigh the evidence against them, and share your analysis. At home: give your AI the same file and a shared procedure that holds it to Beach and Pedersen's method, question it about its reasoning, and write up Analysis Portfolio #1. Beach and Pedersen's two chapters are your primer: read them before class.

AI-free checkpoint this week: a short without-AI task opens the session, before the lab itself.

Required:

- Beach, D. & Pedersen, R. B. (2019). *Process-Tracing Methods: Foundations and Guidelines* (2nd ed.). University of Michigan Press. Ch. 1 ("What Is Process-Tracing?"): read pp. 1-12 closely, skim the rest. Ch. 5 ("Making Inferences Using Mechanistic Evidence"): read pp. 155-158 and 171-193 closely, skim the in-between critique of older approaches; where the Bayesian notation gets thick, hold on to the plain-language definitions and the worked example.

Recommended:

- Collier, D. (2011). Understanding process tracing. *PS: Political Science & Politics*, 44(4), 823-830.
- Bennett, A. & Checkel, J. T. (2015). Process tracing: From philosophical roots to best practices. In *Process Tracing: From Metaphor to Analytic Tool* (pp. 3-37). Cambridge University Press.
- Sanaei, A. & Rajabzadeh, A. (2025). Depth and autonomy: A framework for evaluating LLM applications in social science research. *arXiv:2510.25432*.

■ DATA ETHICS WITH AI (WEEK 5)

■ WEEK 5: DATA ETHICS, TRANSPARENCY, AND THE ORACLE EFFECT (CHAIRSHIP)

With your first AI lab pass behind you, we settle the ethics properly: consent, anonymisation, and when AI analysis of qualitative data is legitimate. We also meet the "oracle effect": trusting fluent AI output instead of checking it. Zook et al. give ten rules for responsible data research. Kapiszewski and Karcher show what transparency looks like in practice. Nguyen and Welch name the oracle effect.

Required:

- Zook, M. et al. (2017). Ten simple rules for responsible big data research. *PLOS Computational Biology*, 13(3).
- Kapiszewski, D. & Karcher, S. (2021). Transparency in practice in qualitative research. *PS: Political Science & Politics*, 54(2), 285-291.
- Nguyen, D. C. & Welch, C. (2026). Generative artificial intelligence in qualitative data analysis: Analyzing –or just chatting? *Organizational Research Methods*, 29(1), 3-39. DOI: 10.1177/10944281251377154.

■ WEEK 6: INTERPRETIVISM, ETHNOGRAPHY, AND THE INTERVIEW (CHAIRSHIP)

The interpretivist turn: meaning, situated knowledge, and the interview as a relationship rather than an extraction. Pachirat shows what an immersed, embodied researcher does that a model cannot stand in for. Rubin covers fieldwork-style data collection, including interviewing. Fujii shows that silences, rumours, and evasions are data too. Reflexivity is introduced lightly here and taken up in depth in Weeks 9-12.

Required:

- Pachirat, T. (2018). *Among Wolves: Ethnography and the Immersive Study of Power*. Routledge, Act 1.
- Rubin (2021), Ch. 8 (the fieldwork model of data collection, including interviewing).
- Fujii, L. A. (2010). Shades of truth and lies: Interpreting testimonies of war and violence. *Journal of Peace Research*, 47(2), 231-241.

Recommended:

- Pachirat (2018), Acts 2-3.
- Schwartz-Shea, P. & Yanow, D. (2012). *Interpretive Research Design*. Routledge, Ch. 1-2.

■ WEEK 7: GROUNDED THEORY (CHAIRSHIP)

Grounded theory straddles explanation and understanding, and this week is its theory: where the method came from (Glaser and Strauss), how it evolved (Strauss and Corbin's procedures, Charmaz's constructivist turn), and its working parts: coding stages, constant comparison, theoretical sampling, memo-writing, saturation. Chun Tie et al. is the primer. Kenny and Fourie map the three versions of the method. Charmaz shows what the constructivist version is for. Next week you put all of it on the coding table.

AI-free checkpoint this week.

Required:

- Chun Tie, Y., Birks, M. & Francis, K. (2019). Grounded theory research: A design framework for novice researchers. *SAGE Open Medicine*, 7. DOI: 10.1177/2050312118822927.
- Kenny, M. & Fourie, R. (2015). Contrasting classic, Straussian, and constructivist grounded theory: Methodological and philosophical conflicts. *The Qualitative Report*, 20(8), 1270-1289. DOI: 10.46743/2160-3715/2015.2251.
- Charmaz, K. (2017). The power of constructivist grounded theory for critical inquiry. *Qualitative Inquiry*, 23(1), 34-45.

Recommended:

- Mills, J., Bonner, A. & Francis, K. (2006). The development of constructivist grounded theory. *International Journal of Qualitative Methods*, 5(1), 25-35. DOI: 10.1177/160940690600500103.
- Nelson, L. K. (2020). Computational grounded theory: A methodological framework. *Sociological Methods & Research*, 49(1), 3-42.

■ WEEK 8: GROUNDED THEORY (LAB)

The second lab. First half: I take last week's theory to the coding table: what line-by-line open coding looks like in practice, how categories grow out of codes, and what separates good grounded theory coding from

bad, using Charmaz and Thornberg's quality criteria and their worked example. Second half: you open-code a short synthetic interview dataset by hand, without AI; the coding sheet stays with you as the foundation of the portfolio. At home: develop your categories by hand following Yue et al.'s procedure, then run the same data through your AI with the same prompts (their Steps 1-3), question it about its decisions, and write up Analysis Portfolio #2.

Required:

- Charmaz, K. & Thornberg, R. (2021). The pursuit of quality in grounded theory. *Qualitative Research in Psychology*, 18(3), 305-327. DOI: 10.1080/14780887.2020.1780357.
- Yue, Y., Liu, D., Lv, Y., Hao, J. & Cui, P. (2025). A practical guide and assessment on using ChatGPT to conduct grounded theory: Tutorial. *JMIR*, 27, e70122. DOI: 10.2196/70122.

Recommended:

- Costa, A. P., Bryda, G., Christou, P. A. & Kasperuniene, J. (2025). AI as a co-researcher in the qualitative research workflow. *International Journal of Qualitative Methods*, 24. DOI: 10.1177/16094069251383739.

■ WEEK 9: REFLEXIVITY AND AI (CHAIRSHIP)

What reflexivity *is* and what it *requires*, before we ask whether AI can do it. Mauthner and Doucet locate reflexivity at the analysis stage, exactly where the AI question lives. Haraway supplies the idea the whole debate runs on: all knowledge comes from somewhere, and a knower who claims to stand nowhere is performing a trick. Messner et al. bring that vocabulary to AI directly. Savolainen et al. (recommended) is the contrarian voice: maybe human reflexivity statements do not deliver what they claim either.

AI-free checkpoint this week.

Required:

- Mauthner, N. S. & Doucet, A. (2003). Reflexive accounts and accounts of reflexivity in qualitative data analysis. *Sociology*, 37(3), 413-431.
- Haraway, D. (1988). Situated knowledges: The science question in feminism and the privilege of partial perspective. *Feminist Studies*, 14(3), 575-599.
- Messner, R., Smith, S. & Richards, C. (2025). Artificial intelligence and qualitative data analysis: Epistemological incongruences. *International Journal of Qualitative Methods*, 24. DOI: 10.1177/16094069251371481.

Recommended:

- Finlay, L. (2002). Negotiating the swamp: The opportunity and challenge of reflexivity in research practice. *Qualitative Research*, 2(2), 209-230.
- Savolainen, J., Casey, P. J., McBrayer, J. P. & Schwerdtle, P. N. (2023). Positionality and its problems: Questioning the value of reflexivity statements in research. *Perspectives on Psychological Science*, 18(6), 1331-1338.

■ WEEK 10: REFLEXIVE THEMATIC ANALYSIS (LAB)

The third lab. First half: I introduce reflexive thematic analysis: Braun and Clarke's three families of TA (coding-reliability, codebook, reflexive), why agreement numbers say nothing about reflexive TA, and what a theme is (a pattern of shared meaning, not a topic bucket). Second half: phases 1-3 (familiarise, code, draft

candidate themes) by hand, without AI, on a short synthetic dataset; your pages stay with you as the foundation of the portfolio. At home: run the AI pass on the same data, keep an audit trail of every prompt, question the model about its theme decisions, and ask it whether it has a positionality of its own. The three Braun and Clarke papers are the method's foundation; the 2021 paper doubles as the portfolio touchstone.

Required:

- Braun, V. & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77-101.
- Braun, V. & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative Research in Sport, Exercise and Health*, 11(4), 589-597. DOI: 10.1080/2159676X.2019.1628806.
- Braun, V. & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328-352. DOI: 10.1080/14780887.2020.1769238.

■ WEEK 11: THE RTA DEBATE, CAN REFLEXIVE RESEARCH INCORPORATE AI? (CHAIRSHIP)

The live controversy: can reflexive qualitative research responsibly incorporate AI? Note what the question is not: nobody in this literature claims AI can *be* reflexive. Jowsey et al. is the rejection letter, 419 signatories, Braun and Clarke among them. Nguyen and Welch supply the sceptics' evidence and the week's sharpest tool: the line between analysing and chatting. Friese et al. argue for a middle path. This three-chair week is built as a structured debate: one chair frames the controversy, two lead the opposing sides, and you arrive having just done RTA with and without AI yourself.

Required:

- Jowsey, T., Braun, V., Clarke, V., Lupton, D. & Fine, M. (2025). We reject the use of generative artificial intelligence for reflexive qualitative research. *Qualitative Inquiry* (online first). DOI: 10.1177/10778004251401851.
- Nguyen, D. C. & Welch, C. (2026). Generative artificial intelligence in qualitative data analysis: Analyzing –or just chatting? *Organizational Research Methods*, 29(1), 3-39.
- Friese, S., Nguyen-Trung, K., Powell, S. & Morgan, D. L. (2026). Beyond binary positions: Making space for critical and reflexive GenAI integration in qualitative research. *Qualitative Inquiry* (online first). DOI: 10.1177/10778004261429393.

Recommended:

- Greenhalgh, T. (2026). Reflexive qualitative research and generative AI: A call to go beyond the binary. *Qualitative Inquiry* (online first). DOI: 10.1177/10778004261429383.
- Nguyen-Trung, K. (2025). Documenting debates on GenAI in (reflexive) qualitative research. SSRN 5750283.

Dina's position (a stance to argue with, not settled doctrine). My own position is that the "We Reject" letter rests on an unearned confidence about what humans can do and what AI cannot, and that line may be thinner than it looks: human reflexivity is itself a practice, something we do and can do badly, not a clear window onto the self. So when an AI names its own trained partiality, is that a different kind of reflexivity? Genuinely open. You do not have to agree with me. You do have to argue your own position.

■ WEEK 12: AI AND THEMATIC ANALYSIS IN PRACTICE, THE LITERATURE ON TRIAL (CHAIRSHIP)

You have now done RTA with and without AI; this week the "AI does thematic analysis" literature goes on trial against your own experience. Xu shows how to do TA with AI while keeping reflexivity. Naeem et al. give the most complete step-by-step AI protocol in print, and we read it critically: is each step interpretation, or topic-sorting? The standing question: where does your Portfolio #3 experience confirm these papers, and where does it contradict them?

Required:

- Xu, W. (2026). Doing thematic analysis in the age of generative AI: Practices, ethics and reflexivity. *International Journal of Qualitative Methods*, 25. DOI: 10.1177/16094069261425173.
- Naeem, M., Smith, T. & Thomas, L. (2025). Thematic analysis and artificial intelligence: A step-by-step process for using ChatGPT in thematic analysis. *International Journal of Qualitative Methods*, 24. DOI: 10.1177/16094069251333886.

Recommended:

- Qiao, S. et al. (2025). Generative AI for thematic analysis in a maternal health study. *Applied Psychology: Health and Well-Being*, 17(3), e70038. DOI: 10.1111/aphw.70038.
- Nyaaba, M. et al. (2025). Optimizing generative AI's accuracy and transparency in inductive thematic analysis. *arXiv:2503.16485*.

■ SYNTHESIS AND WRAP (WEEKS 13-14)

■ WEEK 13: WHEN AI HELPS AND WHEN WE ASK THE WRONG THING OF IT (CHAIRSHIP)

Tie the course together: whether AI helps depends on what we ask of it, on the research strategy, and on the layer of the work. Bail makes the optimist's case for AI in social science. Nguyen and Welch return for a final read. Ernst and Treude (recommended; their task-sorting is presented in class) sort AI tasks into what it does well and what it should not be asked to do. Revisit over-reliance and where each method placed it.

AI-free checkpoint this week.

Required:

- Bail, C. (2024). Can generative AI improve social science? *PNAS*.
- Nguyen & Welch (2026), reprise (see Week 11).

Recommended:

- Ernst, N. A. & Treude, C. (2026). GenAI is no silver bullet for qualitative research in software engineering. *arXiv:2603.08951*.

■ WEEK 14: DEFENDING QUALITATIVE WORK AND COURSE WRAP (CHAIRSHIP)

What does everything add up to, and how do you defend it when someone pushes back? Rubin's closing chapters are a field guide to exactly that; the session splits into a reading half and an activity half.

Required:

- Rubin (2021), Ch. 11-12 ("Living on the Sharp End: Dealing with Skeptics of Qualitative Research" and "The Sweeper").