

419 Researchers Say **Never.** Others Say **Adapt.** Here Is a **Third Answer.**

Using Large Language Models for Qualitative Data Analysis

A hands-on micro course in AI-assisted analysis you can defend

CODES

THEMES

DECISION LOG

Codes measured per code. Themes put on trial. Every model contact logged. That is the third answer: not a position, a procedure.

FREE • **SELF-PACED** • **10 TO 14 HOURS** • **6 MODULES + BONUS SEVENTH**

No participant data used or required • BlendedIQ.ai

COURSE BOOKLET

FLUENCY IS NOT ANALYSIS. | **DEMAND RECEIPTS.**

The debate has produced positions. This course produces procedure.

In 2025, more than four hundred qualitative researchers signed an open letter rejecting generative AI in reflexive research entirely. The response landed punches of its own. Each side is weakest exactly where it attacks a position the other does not hold.

This course asks a different question: if a model quietly damaged your analysis, would you know? You will run a model against your own unassisted analysis, measure where it holds and where it collapses, catch it fabricating, audit its themes back to the source, and use its very conventionality as a mirror for your own assumptions. You leave with the two documents current academic integrity standards actually ask of you, built from checks you performed yourself.

The course takes no side. A learner who completes every module and decides against these tools has used the course exactly as designed; Module 6 even provides the affirmative non-use declaration.

FIVE COMMITMENTS

EVIDENCE OVER VIBES

Test, measure, and document every step.

PROCEDURE OVER POSITION

Build safeguards you can defend and replicate.

TRANSPARENCY BY DESIGN

Decision logs, disclosures, and audit trails.

MEASURE WHAT MATTERS

Evaluate per code. Test themes. Verify sources.

REFLECTION AT THE CORE

Use the model's conventionality as a mirror.

HOW IT WORKS

WATCH a short video: the instructor runs the workflow live, model included, mistakes included. **PRACTICE** the same workflow offline on transcripts reserved for you, in workbooks that compute your figures. **PASS** a 10-question quiz to unlock the next module. Nothing is submitted; the quizzes gate progress, the practice builds your evidence.

You need: a browser, a spreadsheet application, any major AI chat tool. Skeptics warmly expected: the AI stays off until your own reading exists.

THE JOURNEY AT A GLANCE

→	Course Overview: read the brochure first	10 min read
1	Tooling, data safety, fabrication, the decision log	2 to 2.5 h
2	The unassisted baseline: coding and the dated memo (no AI)	2 to 2.5 h
3	Deductive workflows: the codebook leash and the verdict	2.5 to 3 h
4	Inductive workflows: themes on trial, the traceability audit	2.5 to 3 h
5	Abductive workflows: estrangement, the model as mirror	2 to 2.5 h
6	Academic Integrity: Disclosure and More	1.5 to 2 h
7	BONUS. Where do you stand? The philosophy underneath	1.5 to 2 h

Modules unlock in sequence on quiz pass, starting from Module 1. Module 7 unlocks on completion of Module 6; it is the course's reward, not another hurdle.

Enrolment is free at BlendedIQ.ai.



Course Overview

READ FIRST.

The course introduction is a two-page brochure that opens with where the debate stands, the three rules that protect your results (synthetic data only; your reading sealed and dated before any model contact; no model claim trusted unsearched), what to download before Module 1, and the course's honestly labeled position.

READ the Course Overview brochure.

DOWNLOAD the decision log, the analysis journal, and the dataset, all named in the brochure.

YOU WILL BE ABLE TO *Know what the course will and will not claim, and what evidence you are about to build.*

1

Tooling, data safety, and what the model actually does

THE ENGINE ROOM.

What a language model actually produces and why it is non-deterministic; the three-part data gate (consent, institutional approval, provider terms) that decides whether real data may ever touch a model; and the fabrication hunt, in which you make a model invent quotations and catch it doing so. Anonymization is a safeguard, not a permission.

WATCH the module video with the live fabrication demo.

PRACTICE the consistency check, the fabrication hunt, and your data handling plan.

DOWNLOAD the decision log workbook, kept for the rest of the course, and beyond.

YOU WILL BE ABLE TO *Run the consistency and fabrication checks on any model, keep a decision log, and state exactly when participant data may not be processed.*

2

The unassisted baseline. No AI allowed

THE MOST IMPORTANT MODULE CONTAINS NO AI AT ALL.

Code a transcript by hand, write a dated memo of what struck you and why, asterisk the passages that resist your categories, and set it all aside. That dated record becomes the instrument every later module measures the model against; without it, you can never again know what you would have seen on your own.

WATCH the module video with on-camera hand-coding.

PRACTICE code transcript T3 (Meredith); complete and date the salience memo; close the file until Module 4.

DOWNLOAD the T3 learner workbook, which follows your transcript through Modules 2, 4, and 5.

YOU WILL BE ABLE TO *Produce a defensible unassisted coding and a dated salience memo, and explain why it must exist before any model contact.*

3

Deductive workflows. The model on a leash you built

THE MOST DEFENSIBLE DELEGATION.

Fix a codebook with definitions, exclusions, and stated provenance; the model only classifies, with a no-code option and one primary code per segment. Then measure it: agreement per code (the headline number hides exactly the code that matters), stability across runs, and an error analysis that catches force-fitting against your own written exclusions. You end with a verdict memo: what you would delegate, under what conditions, and what never.

WATCH the module video with the live deductive run.

PRACTICE adopt the provided codebook, code T4 yourself first, run the model twice, compute per-code figures.

DOWNLOAD the deductive workbook and codebook template; prompts P1 to P4 in the prompt book.

YOU WILL BE ABLE TO *Build and apply a machine-usable codebook, evaluate model coding per code rather than by headline, and write an evidence-based delegation verdict.*

4

Inductive workflows. Themes on trial

WHERE THE ORACLE EFFECT LIVES.

Ask the model for themes twice, once cold, once through a human-gated staged pipeline with your corrections logged at every gate. Then audit everything: every quotation searched in the source and classified exact, altered, or fabricated; every theme binned grounded, plausible-but-generic, or unsupported. Then the three-way comparison against your Module 2 memo, with two honesty columns: what the model found that you missed, and what you found that it never will.

WATCH the module video with the audited run.

PRACTICE both passes on your T3, the full traceability audit, the gates log, the three-way comparison.

DOWNLOAD the traceability audit template and the staged prompts (4.2a to 4.2d).

YOU WILL BE ABLE TO *Run a full traceability audit, distinguish grounded findings from statistical mimicry, and report both honesty columns without flinching.*

5

Abductive workflows. The model as mirror

THE COURSE'S MOST DISTINCTIVE MOVE.

Built on a method proposal currently under peer review, and presented with exactly that status. Models tend toward the canonical reading, the interpretation a field has made most available. This module turns that limitation into an instrument: elicit the model's reading of a puzzling passage while revealing nothing of yours, locate the divergence, and interrogate the gap until you can name, in one sentence, an assumption of your own you could not see. Then decide on the record: defend, revise, or flag.

WATCH the module video with the on-screen estrangement demo.

PRACTICE two estrangement records on your T3: your strongest asterisked puzzle, and a deliberately mundane passage as the stress test.

DOWNLOAD the estrangement record template and the context-assembly guidance.

YOU WILL BE ABLE TO *Run the estrangement protocol on any passage, name the assumption a divergence exposes, and stress-test the method itself.*

6

Academic integrity: Disclosure and More

FIVE REQUIREMENTS. TWO DOCUMENTS. YOUR EVIDENCE.

The current state of academic integrity on AI use reduces to five requirements: DISCLOSE, OWN IT, VERIFY, PROTECT THE DATA, KEEP RECORDS. You meet all five already holding the evidence, because the course made you build it. Draft the two documents your thesis actually needs: a disclosure statement generated from your decision log, and a one-page conditions note in which every delegation cites your own figures, ending with the sharpest question in the course: what evidence would show your safeguards had failed? “No participant data may currently be processed, and here is what would have to change” is a fully passing conditions note.

WATCH the module video.

PRACTICE draft your disclosure statement (L-1) and conditions note (L-2); run the 10-point self-check (L-3).

DOWNLOAD the lab templates (L-1 to L-3); five sample declarations with course traces unlock after your draft.

YOU WILL BE ABLE TO *Write a disclosure statement a reader can verify, and a conditions note you could defend in a viva.*

7

BONUS. Where do you stand?

IDENTIFY VARIOUS POSITIONS ON THE AI DEBATE.

The debate's landmark exchange, a mass rejection letter and its sharpest response, presented at full strength on both sides. A five-position map of how much power a technology has, from the neutral tool (a phantom nobody actually holds) to the regime; the levels insight (practice versus system) that explains why the camps argue past each other; the standard objection dissected into its four failing parts; two fatalisms refused; and your deliverable: a half-page position statement, conditional use, principled refusal, or open trial, ready for a methodology chapter or ethics application.

WATCH the module video.

PRACTICE the lab: placement drill, objection surgery, and your position statement, with answer keys.

SELF-CHECK the 10-question quiz; it gates nothing, you have already completed the course.

GO DEEPER the Extended Lens Compendium: fifteen philosophical positions, their blind spots, fallacies, and biases.

YOU WILL BE ABLE TO *Place any claim in this debate before answering it, and state your own position knowing exactly what it costs.*

THE INSTRUMENT LIBRARY YOU TAKE WITH YOU

DECISION LOG

One row per model contact, including rejected outputs. The audit trail behind every disclosure.

CODEBOOK TEMPLATE

Definitions, inclusion and exclusion criteria, provenance per code, machine-usable.

TRACEABILITY AUDIT

Quote checks and theme bins, with automatic counts.

ESTRANGEMENT RECORD

Five steps from salience refresh to named assumption and verdict.

INTEGRITY SET

Disclosure scaffold, conditions note, 10-point self-check.

P1 TO P9 PROMPT TEMPLATES

Each with its verification requirements attached.

HONESTY NOTES

The Module 5 protocol comes from a manuscript currently under peer review, and the course says so wherever it appears: promising and proven are different words. The dataset is entirely synthetic; none of your practice ever risks real data. And the course's own philosophical package is disclosed inside Module 7, where you are invited to argue against it.

WHO THIS IS FOR

Doctoral researchers and supervisors in the social sciences and management; qualitative researchers weighing these tools for the first time; methods instructors seeking teachable, evidence-first workflows; and skeptics, warmly expected.

Enrol free at **BlendedIQ.ai**

Bring a browser, a spreadsheet application, any major AI chat tool, and, if you have them, your doubts.