

APPROACHES TO USING LARGE LANGUAGE MODELS IN QUALITATIVE RESEARCH

POLARIZED MORAL IDENTITY: EXPLORING FORCES AND MECHANISMS OF TECHNOLOGY-ASSISTED DUAL IDENTITY PROCESSES

Nikki Drader-Mazza*, Vidhi Jain*, Steven Hyde, Samantha Jordan, Rhonda Reger, Virginie Kidwell, and Danielle Cooper

*Shared First Authorship



THE STUDY

One person, multiple identities, fifteen years



Ross Ulbricht ran the Silk Road as the Dread Pirate Roberts (DPR). Arrested 2013. Life sentence 2015. Pardoned 2025.

Non-violent son, boyfriend, libertarian, entrepreneur, criminal mastermind, prisoner's mentor, crypto hero

Empirical question. How does Ross Ulbricht achieve a holistic sense of self over time despite significant identity conflict?

Lens: **dynamic holism** (Wittman et al., 2025): identities as separate threads in one rope — woven together as circumstances change.



WHAT WE WORKED FROM

Archival Data Sources



**PERSONAL
JOURNAL**



**SOCIAL MEDIA
POSTS
X & Instagram**



**CHAT LOGS
SUBMITTED AS
CASE EVIDENCE**



LETTERS



NEWS ARTICLES



CASE FILES



**CONGRESSIONAL
FILES**



INTERVIEWS

The resulting corpus is comprised of approximately 882,000 words

HOW WE WORKED

The process, and where AI came in

AI had **three jobs** in this project.

- **ADDITIONAL CODER** — multiple models, alongside 2 human coders
- **TOOL-BUILDER** — dashboard and evidence explorer
- **ADVERSARY** — tried to break our theory (Wittman et al., forthcoming)

HUMANS

Code the
corpus by
hand

AI

Code it
again,
several
models, held
in a wiki

AI

Build a
dashboard
to see it all

HUMANS

Read it again
and re-
theorize

AI

Argue
against the
result

HUMANS

Interpretation
and synthesis
of findings

STEP ONE · HUMANS

People coded the record first

Before any model saw the text, the team coded every source by hand. Bricolage (Pratt et al., 2022)

1. TEMPORAL BRACKETING

Hand-built timeline of how the case unfolded over time

2. THICK DESCRIPTION

What happened — and what it meant to the people involved

3. THEMATIC CODING

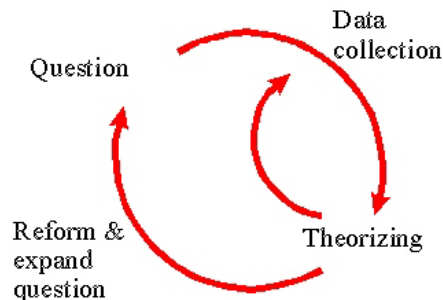
Open coding on a sample, followed by directed coding on the full corpus.

4. LITERATURE REVIEW

Back to the literature after every coding pass

5. CONSTRUCT IDENTIFICATION

Identify themes, early theorizing, and a codebook.



LLMs coded the same material

Claude, Gemini, ChatGPT coded the same sources, against our codebook. Every model's coding went into the wiki — on the same page as the human coding.

Each document in our corpus contained meta-data for the LLMs:

- Document Type
- Speaker

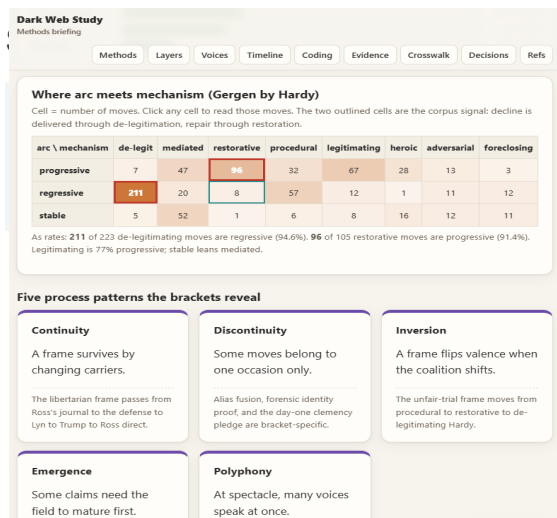
Agreement → some confidence. Divergence → a flag: human read that spot again.

We were not looking for agreement. We were looking for soft spots.

STEP THREE · AI

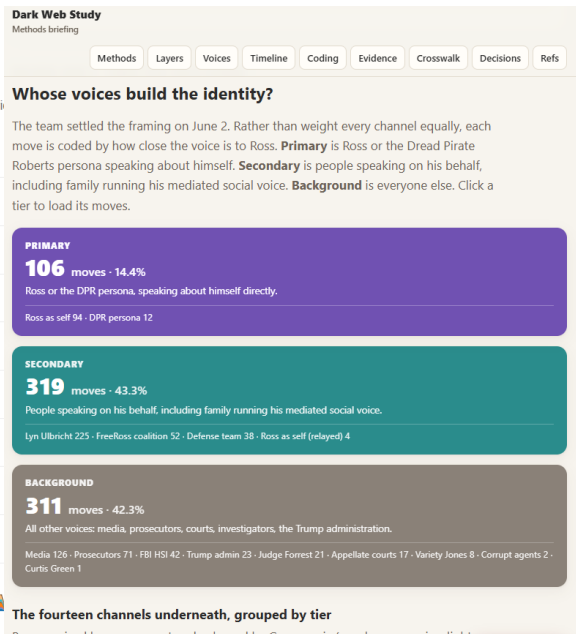
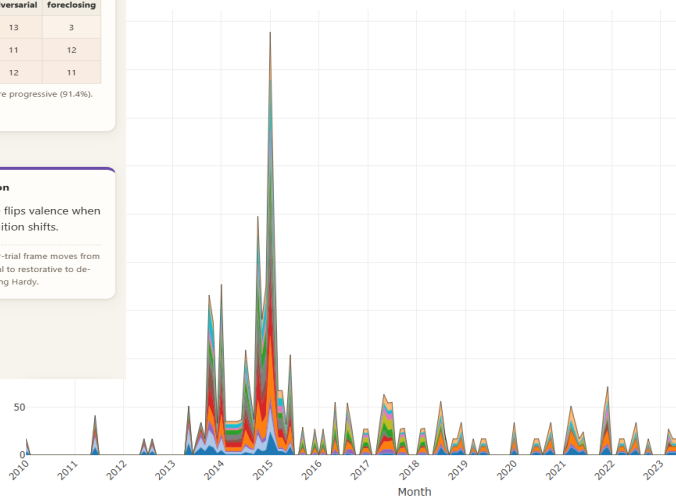
Then AI built the tools to see the data

882,000 words — all in one place. AI built a dashboard, a timeline, and an evidence explorer: the whole coded corpus in one view.



ix distinct periods e: Actor-Influence Histomap (2010-2025)

scores per actor. Heights are relative influence summed across events in that month. Substantial w



Humans read it again and built the theory

Interpretation and theorizing is human work.

LITERATURE REVIEW

Back to the literature after every coding pass

MODEL DEVELOPMENT

The mechanisms that answer our RQ

NARRATIVE PRESENTATION

Narratives of who Ross was, period by period

What we asked the LLMs

Theory in hand, we asked AI to break the argument.

CLAIM: [one coded finding]

EVIDENCE THE TEAM USED: [the quote]

Refute this claim using only the provided corpus.
Quote verbatim with the source_id.

If you cannot find disconfirming evidence, say
exactly: "No disconfirming evidence found."

Do not invent or paraphrase a quote. Fabrication is
worse than a null result.

FIVE GUARDRAILS ON EVERY CALL

- 1 Corpus only. No outside knowledge of the case.
- 2 Every claim carries a verbatim quote and a source ID, or it is dropped.
- 3 A "nothing found" answer is allowed, always.
- 4 As theory evolves, set prior framing aside.
- 5 Run separately for each voice tier.

WHAT THE PASS FOUND

What came back, and what we did with it

14 claims went through the pass. 1 clean · 4 held with a caveat · 4 contested · 5 dropped.

Three “DPR” quotes were not his — one invented in a trade book, two compiled by journalists. His own journal contradicted one claim.

Verification ran both ways. We hand-checked every quote against the source — and caught the model’s mistakes too: one quote it flagged as fabricated, we found verbatim in the chat logs, verified the event in personal journal.

The theory held. The evidence under it came out stronger.

IF YOU WANT TO TRY IT

What you need, and what to watch for

WHAT YOU NEED

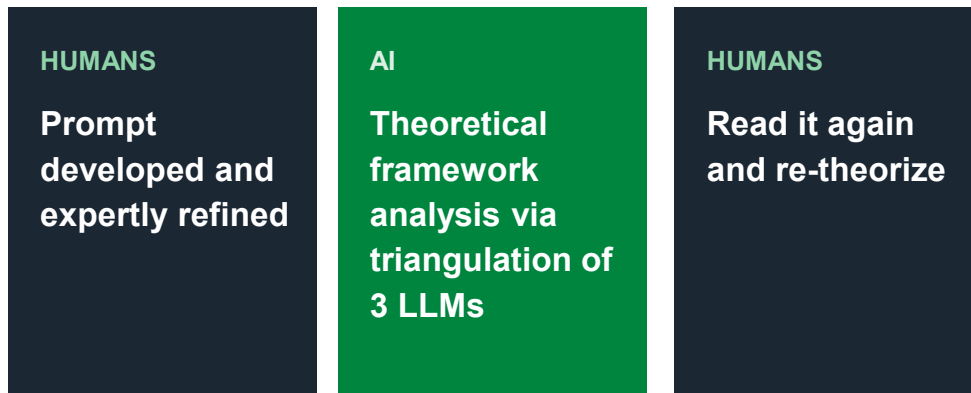
- A machine-readable corpus, with a manifest of sources and speakers
- A versioned codebook — same definitions for every coder and model
- Claims stated plainly enough to argue against
- A prompt that allows “nothing found”
- Guardrails are important

WHAT TO WATCH FOR

- One case — we cannot yet say how far it travels
- Prompt sensitivity
- Results depend on the model; another may disconfirm differently – Human judgement triumphs
- Triangulation is important – human to human, human to LLMs, LLM to LLM

APPENDIX

The process, and where AI came in



Benefits of LLM triangulation:

- Reduced preprocessing burden
- Consistency across time (reduced cognitive drift)
- Temporal pattern detection at scale
- Preservation of theoretical independence

WHAT WE FOUND

Comparing Large Language Models

Model	Consistent Strengths	Consistent Limitations	Systematic Bias
Gemini	• Most detailed psychological analysis • Strong on internal conflict detection • Excellent citation management	• Less systematic organization • Sometimes overly detailed on minutiae	Psychological bias - overemphasizes internal mental states
Claude	• Most systematic theoretical integration • Clear organizational structure • Best cross-framework synthesis	• Sometimes less detailed on specific evidence • More abstract theoretical connections	Theoretical bias - prioritizes framework consistency over empirical detail
ChatGPT	• Strong on communication patterns • Good organizational behavior analysis • Clear progression tracking	• Less psychological depth • Weaker on internal conflict analysis	Behavioral bias - focuses on observable actions over internal states

Key Events and Identity Shifts (2009-2025)

