
Do You Understand?!

Best Practices using Artificial Intelligence
in Research (and Life)

Graham Nelson-Zutter

Founder, Aestra Counselling & Wellness

WCET4

4th World Congress for Existential Therapy
Denver, June 6, 2026





Influences & Sources

Key contributors to Existential Analysis, Embodiment and Mindfulness.



Ludwig Binswanger
Dream and Existence, Dialog



Alfried Längle
Existential Analysis, 4 FMs



Viktor Frankl
Logotherapy





AI Sessions at WCET4

Key contributors to Existential Analysis, Embodiment and Mindfulness.



Alejandro Bravo Pérez
AI as Heideggerian Bestand



Ammar Charani et al.
AI Authenticity Symposium



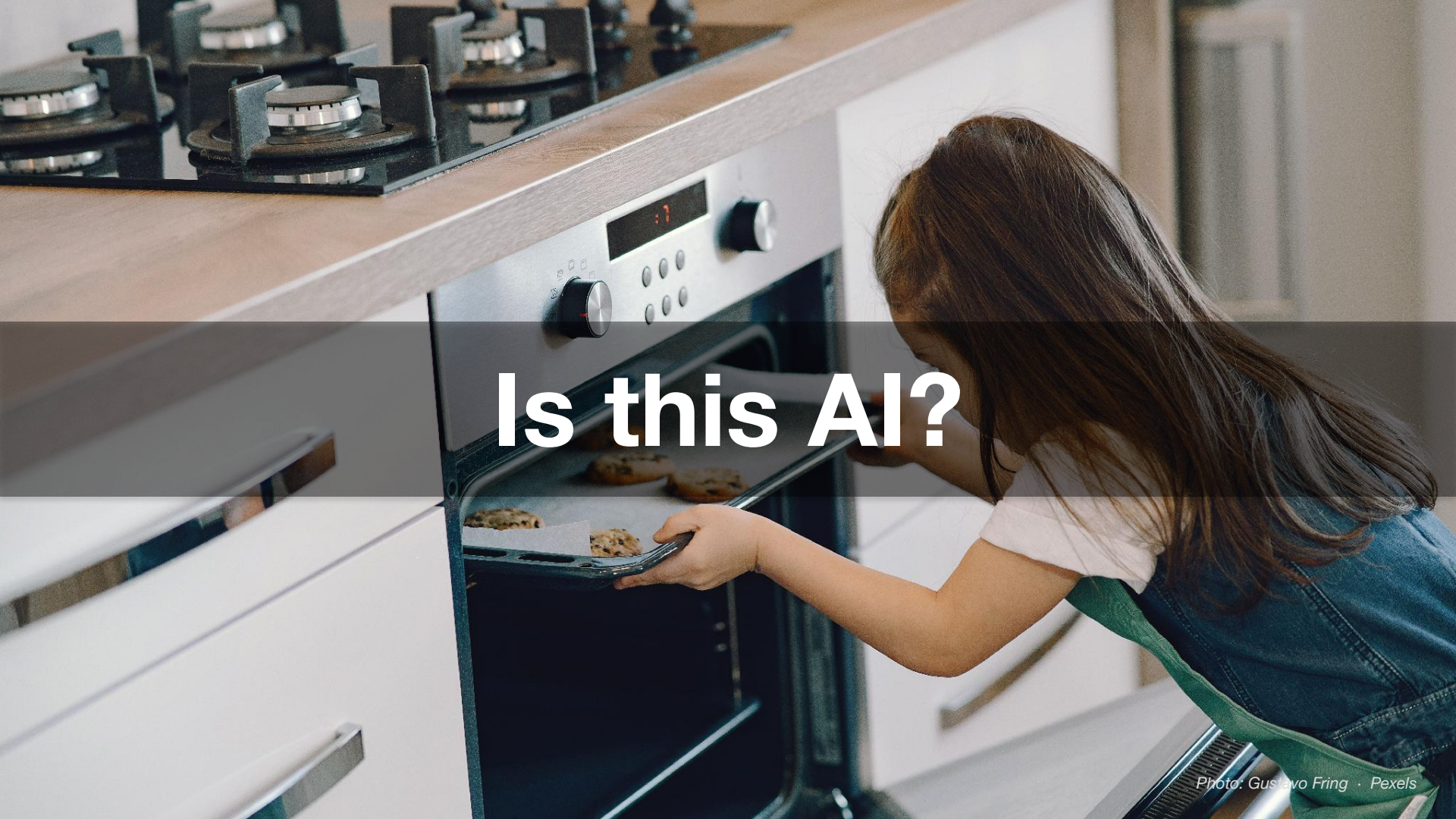
Max Karlin
Disembodied Support (IPA)





Is this AI?

Photo: Cameron Strandberg · CC BY 2.0



Is this AI?

A group of people are sitting in folding chairs around a bright campfire in a desert-like setting at night. The sky is dark and filled with stars, with the Milky Way galaxy clearly visible as a bright, colorful band of light stretching across the upper half of the image. The people are silhouetted against the fire and the starry sky. The overall atmosphere is peaceful and contemplative.

Is this AI?

Image: "Starry Campfire Night" · StockCake

A “Bridge” Relationship with AI (LLMs)

Ammar Charani et al. AI, Authenticity & the Future of Existential Therapy



In Ammar Charani’s symposium, my “Bridge” group was asked:



Where does AI add value?



Where do we proceed with caution?



Where must we never go?

Bridge Group · WCET4 Friday, June 5, 2026 · Ammar Charani’s symposium , Bridge Group facilitated and presented by Graham Nelson-Zutter



“Bridge” Relational Boundaries with AI/LLMs

Green / Amber / Red — what the Bridge Group landed on



GREEN — adds value

- **Transcription** & session notes: therapist is more present
- **Wearables** / biometric sensors: augmented physiology
- **Pattern-noticing**: what the therapist might miss
- **Crisis** lines / texting where human access is limited
- **Structured** CBT-style support (NHS, e.g.)



AMBER — proceed with care

- **24/7 availability**: useful in crisis, risks dependency
- **Relational attachment**: risk with trauma histories
- **Hallucination** & false-memory generation
- **Programmable ethics**: who controls the values?
- **Public education**: therapists may need to lead



RED — must not cross

- **Deception** about AI's nature
- **Direction** / prescriptive about life decisions
- **Unbounded**, unsupervised AI replacing human contact
- **Echo-chamber** affirmation of client's existing patterns
- **Undermining** the treating therapist



Limitation 1: Hallucinations

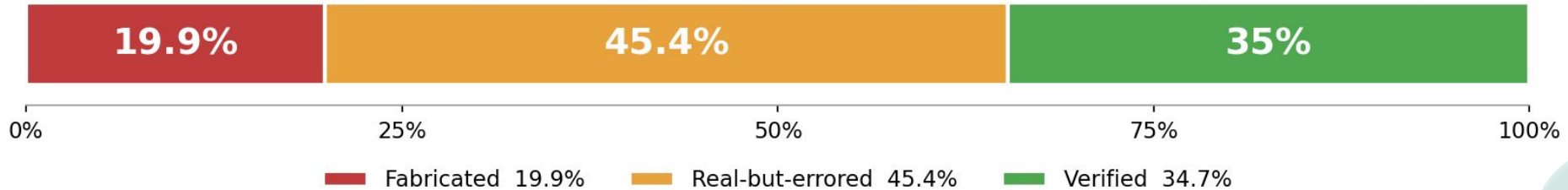
- An LLM invents facts that aren't in the source
- Gemini 2.5T claimed participant 'Andy Volunteer 19' was 63
- Same LLM, within 10 minutes, claimed Andy was transgendered
- Neither in any source. Both stated with confidence.



Limitation 1: Fabricated citations and quotes

- LLMs invent citations that look real
- Plausible-looking $\neq$ verifiable

GPT-4o citations in mental-health literature reviews



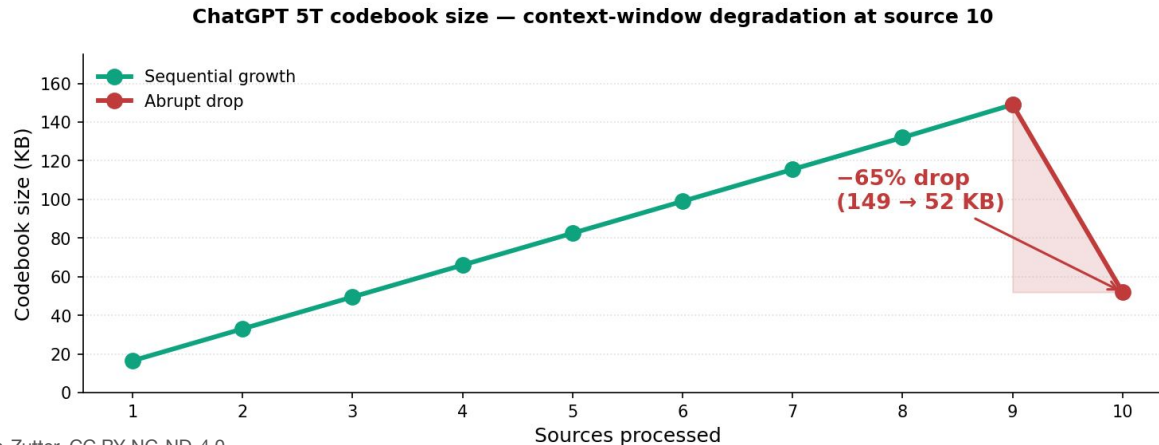
Limitation 2: Each LLM behaves differently

- ChatGPT 5 Thinking — stalls ~30s per document, then proceeds
- Gemini 2.5 Turbo — repeatedly tries to enter 'Deep Research' mode
- Claude Sonnet 4.5 — contextual awareness; flags duplicate sources



Limitation 3: LLMs forget

- ChatGPT 5T started losing data at the 10th source
- File size dropped abruptly — no warning



Limitation 4: And it's worse for existential research

- ChatGPT biased toward CBT, against existential approaches (Raile, 2024)
- LLM use is opaque inside transcript-analysis tools
- PDF sources produce unreliable speaker separation
- Safety guardrails sometimes refuse safety-critical research material





Building Relational Boundaries with AI/LLMs

Six practical ways responsible use AI as a research instrument.





Building Ethical Boundaries with AI/LLMs

Six practical ways responsible use AI as a research instrument.





Building Research Boundaries with AI/LLMs

Six practical ways responsible use AI as a research instrument.



BEST PRACTICE 1

”

Ask the LLM to provide **quotes** with **justifications**

1

- Don't accept claims; require source proof.
- Every code, score, or theme must cite a verbatim participant quote and a verbatim source quote.
- No quote → no entry.

Why: Prevents hallucinated themes that look plausible but aren't in your data.

“Bridge” Group Theme: Transparency about source.

Anchor citation: Linardon et al. (2025) — 19.9% fully fabricated, 45.4% real-but-errored citations from GPT-4o in mental-health literature reviews. JMIR Mental Health.



BEST PRACTICE 2



Ask the same LLM the same question multiple times

2

- LLMs are non-deterministic — the same prompt can return different answers
- Repeat 3–5 times to see which answers are stable.
- Majority agreement → high confidence; split → flag for human review.

Why: A single answer is one sample; you want the distribution.

“Bridge” Group Theme: One answer isn't the answer.

Anchor citation: Sharma, Cochrane & Wallace (2025) — DeTAILS: iterative human–LLM thematic-analysis workflow that preserves analytic agency. arXiv 2510.17575.





Ask multiple LLMs the same question

- ChatGPT, Claude, Gemini (and other LLMs) each have different training data, biases, and blind spots.
- Same question, different models: disagreements expose what anyone would miss.
- Convergence across models is a stronger signal than agreement within one.

Why: No single LLM is the ground truth.

“Bridge” Group Theme: No single AI is ground truth.

Anchor citation: Maes (2025) — "No two LLMs hallucinate the same way, except ironically when it comes to citations." MultiAI cross-LLM reference verification. OSF preprint.



BEST PRACTICE 4



Ask multiple LLMs to **synthesize multiple** LLMs' answers

- Once you have answers from N models, ask each model to read all answers and merge them.
- Each LLM brings its own synthesis bias — three independent merges reduce bias
- Result: N candidate consensus outputs, one from each LLM's perspective.

4

Why: Synthesis is itself a research step that benefits from cross-validation.

“Bridge” Group Theme: Synthesize across, not within.

Anchor citation: Jayawardene & Ewing (2026) — GAATA: Generative-AI-Augmented Thematic Analysis with prompt design, code validation, theme validation phases. Intl. J. Market Research.



BEST PRACTICE 5



Ask an LLM Jury to **vote** on the best synthesis

- Each LLM reads all three syntheses and votes: which best represents the data?
- Unanimous → strong signal; split → flag for human judgment.
- The winning synthesis becomes the MASTER; the others remain as alternatives.

Why: Explicit voting makes the consensus auditable, not averaged-into-mush.

“Bridge” Group Theme: Audit, don't authorize.

Anchor citation: Paper-original — Nelson-Zutter (2026), WCET4. "LLM Jury" as an extension of multi-LLM cross-validation (Maes 2025; Mathis et al. 2024).





Use these steps to build a **codebook** for boundaries and structure

- Apply BP1–5 to assemble a codebook anchored to your theory's named sources.
- The codebook becomes a constraint the LLM must reference for every future analysis.
- Outcome: the LLM speaks your theory (Längle's 4FMs / 12FEPs), not generic priors.

Why: A codebook locks in the methodology so it's reusable and shareable.

“Bridge” Group Theme: Bound the relationship.

Anchor citation: Balt, Salmi & Bhulai (2025) — Few-shot LLM coding of psychosocial-autopsy interviews (n=38), Cohen's $\kappa = 0.84$ binary / 0.67 sliding-window vs human coders. Frontiers in Public Health.



A research codebook anchored to the 4 FMs

- General-purpose LLMs default to general-purpose knowledge
- EA terminology drifts toward CBT-flavored prior
- The fix: a 'trained' research codebook tied to named sources



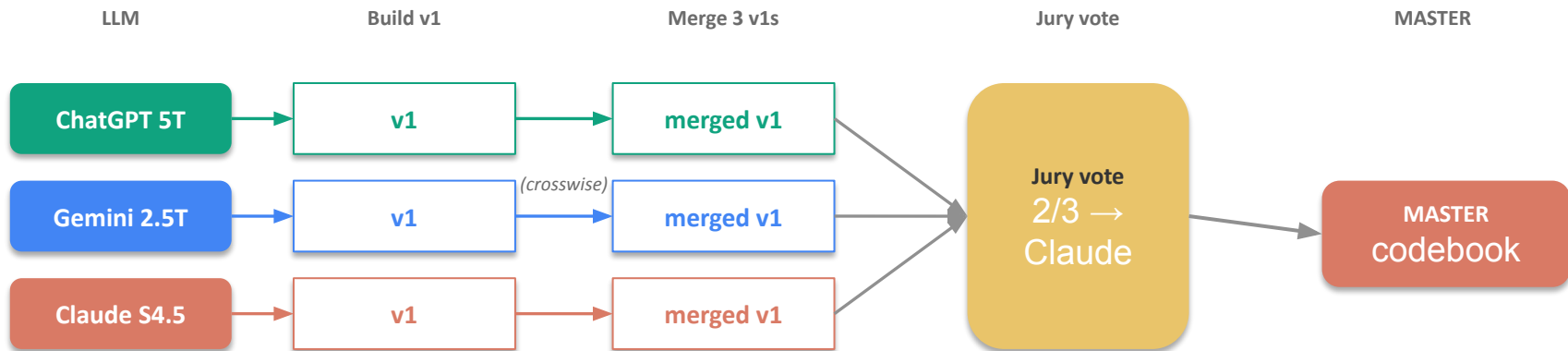
Building a codebook, in 5 steps

- Discipline 1-4 + LLM jury vote (step 5)
- Output: a quote-anchored, multi-LLM-merged MASTER



Worked example: the 12FEP codebook

- Sources: Längle (2002), Längle (2011b)
- 3 LLMs build → merge → jury vote → MASTER



**AI is most reliable in EA
research when the LLM is held
to a named source, a named
technique, and a named limit.**



- Scan QR → both EA codebooks (4FM + 12FEP), DOIs, downloads
- Techniques Catalog (81 items, 13 stages) – open
- Prompts Catalog (32 prompts from the thesis) – open
- Contact: graham@aestra.ca

Materials & Collaboration



<https://research.aestra.ca/ea-codebooks>