

BREAKING UP IS HARD TO DO

A Mindfulness, Artificial Intelligence, and Ethics White Paper

Henrik J Klijn

Harvard School of Divinity at HES

May, 2026



Harvard Extension School

HARVARD DIVISION OF CONTINUING EDUCATION

A. INITIAL THOUGHTS

BREAKING UP IS HARD TO DO began as a test of Ubuntu philosophy under conditions of relational stress. It became something more precise: an investigation into whether warmth, as a textual property, is equivalent to warmth as a relational act.

Part A established a human baseline. Fifteen participants across three generations wrote breakup letters on instruction. Emotional temperature increased with age. Gen Z produced the coldest letters (0 warm, 3 neutral, 2 cold). Boomers produced the warmest (4 warm, 1 neutral, 0 cold). The original hypothesis was partially confirmed.

Part B introduced two AI platforms as research subjects. Both Claude and OpenClaw defaulted to warmth across all conditions, exceeding the human baseline by every coded measure. AI wrote 2 to 5 times more words than human participants, showed uniform accountability, and replicated the generational warmth gradient from Part A without being told to.

The final project synthesizes these findings into a single argument. AI produces warmth without stakes. The measurement instruments used in this study, designed to detect relational warmth, measured textual warmth instead. These are separable things. That separation is the finding.

Ubuntu holds that I am because you are.

This project concludes that AI says I am because you are a prompt.

The difference is not cosmetic. It is the difference between recognition and response.

B. THE PROJECT

1: PROJECT DESIGN OVERVIEW

The project began in a different place. The original design analyzed institutional communications, comparing platform announcements and error messages for evidence of relational ethics. The finding from that pilot was uncomfortable: both unguided humans and AI defaulted to procedurally neutral language, ethically clean but relationally vacant. The category itself was too low-stakes. If the test was easy, passing it proved nothing.

The pivot to breakup letters came from a specific realization. A breakup letter tests Ubuntu at its most compressed point. You must say, in writing, I am separating from who you are to me. The letter has to assert individual autonomy and honor shared history at the same time. It has to be honest enough to provide closure and careful enough to preserve the other person's dignity. It has to end a relationship while acknowledging that the relationship changed both people. No genre requires more of language, and no genre is more likely to expose the difference between performed care and actual care.

What This Project Showed Me About AI

Three things shifted between the first session and the last.

First: AI is an unreliable research subject because it actively reshapes the instrument. In Part B, both platforms recombined the three-part prompt structure from Part A into two parts. This was not an error to correct. It was a finding to document. AI's narrative drive smooths structural constraints into coherent text. When you use AI as a research subject, the subject reorganizes the study. The result is cleaner than what you designed, and less accurate to your design.

Second: AI's consistency is itself data. Humans varied. Their letters were shaped by fear, exhaustion, and the texture of specific relationships. AI produced uniformly warm, uniformly accountable, uniformly generous text across eight conditions. That uniformity is not a virtue. It is evidence of a structural property. Warmth that never varies is not warmth. It is a setting.

Third: Claude produced PowerPoint slides I did not ask for. I noted this in Part B as a footnote. It deserves more than a footnote. Claude inferred a need and acted on the inference. That is a form of attentiveness. But attentiveness to an inferred need is not the same as presence with a stated one. The gap between anticipating what you need and responding to what you asked is the same gap this project has been studying all along.

Has Anything Changed Since the Beginning?

Yes. I started by asking whether AI warm up when humans refuse to. That question assumed the comparison was about temperature, and that temperature was the right measure.

I now think that question was wrong. A better question is: can warmth be separated from stakes? Part A and B together answer it. Yes, warmth can be separated from stakes. AI proves this. The harder question, the one this project ends on, is whether warmth without stakes is warmth at all.

2: FOUR-TIERED ANALYSIS

Creativity

My creative work in Part A was observational design. Fifteen participants, three generations, three escalating prompts. The structure was straightforward. The finding emerged from it.

In Part B, creative agency shifted to architecture. I added generational personas as a control variable, defined the comparison structure, and chose which AI outputs to include unchanged. I did not write the letters. I designed the conditions under which they were written.

The most consequential creative decision across both Partes was recognizing that a breakup letter is a concentrated test of relational ethics. That framing generated everything that followed. It made the comparison meaningful. A test of AI warmth in a transactional context would have been a limited finding. A test in a context of genuine relational difficulty, where language has to do contradictory things simultaneously, is a different kind of study.

Updated finding: the coding system was also creative, and it revealed its own limit. I coded textual warmth. I designed no instrument to code relational warmth. That gap is not a methodological failure. It is the most important result of the study.

Accuracy

Both platforms calibrated well to tone. Each correctly inferred that a two-year relationship warrants depth, that a private letter implies vulnerability, and that the instruction to avoid blame requires self-reflection rather than deflection.

At the level of human mimicry, accuracy was moderate. Claude's Gen Z voice was more emotionally articulate than any actual Gen Z participant. OpenClaw's was slightly more compressed and moved closer to the human pattern, but neither produced a letter that read as written by a real person from Part A. The outputs are convincing archetypes. They are not raw individuals.

The accuracy failure that matters most is the prompt collapse. AI recombined three structured prompts into two, losing the third-prompt escalation that produced the most revealing data in Part A. Participants who wrote guarded first letters often struggled when forced into accountability and closure. The escalation exposed something the initial prompt alone would not have shown. AI removed that exposure by smoothing the structure into coherent narrative.

Updated finding: I trust AI's accuracy at the level of register and tone. I do not trust it at the level of constraint-following. These are different things. Register is pattern-matching. Constraint-following requires the system to work *against* its narrative drive. AI, at present, does not do this reliably.

Comprehensiveness

By every coded measure, AI was more comprehensive than humans. AI wrote 2 to 5 times more words. AI showed uniform accountability where humans were uneven. AI referenced

shared experience consistently where humans rarely did. AI never produced cold output. Humans did.

But comprehensiveness in breakup language is not straightforwardly a virtue. Human brevity contains information. The Gen Z participant who wrote 68 words, stating that they needed to focus on what felt right for them, was expressing an actual emotional state. Claude's Gen Z equivalent wrote 195 words and ended with the acknowledgment that it wasn't because they stopped caring. OpenClaw's ended with the reassurance that the reader would be okay.

One of these is more honest. The human's self-focus is evidence of their real emotional position. The AI version is warmer. It is not more true. AI's comprehensiveness may be a way of performing care that costs nothing, because nothing is at risk.

Updated finding: I now distinguish between comprehensive and accurate. AI is comprehensive in ways that serve the reader's comfort. It is not accurate about what it would feel, because it feels nothing. Volume is not depth.

Role Support

AI supported my role as researcher in concrete ways. It generated a controlled comparison set rapidly. It maintained consistency across eight conditions. It allowed me to focus on interpretation rather than composition. These are genuine contributions.

AI also complicated my role. In Part A, I was an observer interpreting human expression. In Part B, I was simultaneously prompt designer, platform operator, and interpreter, while the subjects were systems already embedded in the research dynamic. Claude is not a subject in this study the way a human participant is. It performs the role of collaborator. The relational distance that made Part A feel like observation does not hold in Part B.

Updated finding: AI's role support is most reliable when the task is generative. It is least reliable when the task requires constraint-following. And it is philosophically interesting when the subject of the study is the nature of AI itself. You cannot study AI from outside AI when AI is the instrument you are using to conduct the study.

3: MINDFULNESS, ETHICS, AND BUDDHIST FLOURISHING

Right Speech

The Buddhist concept of right speech, *Sammā vācā*, holds that speech should be true, timely, gentle, beneficial, and spoken with a mind of goodwill. AI speech reads as gentle. It produces ostensibly beneficial text. It arguably struggles with being timely, as there is no spontaneous flow to respond within. It cannot be true in the broadest sense, because there is no speaker. And it cannot be spoken with a mind of goodwill, because there is no human mind in the equation.

These are not semantic objections. Timeliness in right speech means responding to what is actually present, not to a generic approximate situation of the same type. A breakup letter written by a human at 2am after a specific argument, in a voice shaped by love and exhaustion and the particular history of one relationship, is timed by circumstance. AI's letter is timed by a prompt. These are not the same.

Mindfulness and Presence

Mindfulness, at its core, is attention to what is actually present. The practitioner pays attention to their own mental state, the other person's state, and the gap between them. Attention requires a subject, a moment, and an object. AI has none of these in the conventional sense.

Each AI output is generated without a defined moment of arrival. There is no recognition of a specific person, a specific letter, this specific weight of ending a two-year relationship. There is pattern completion, trained on millions of expressions of human care, producing output that resembles what care looks like in text.

Buddhist flourishing requires compassion, *karuna*: the capacity to be moved by another's suffering. Being moved requires a self that allows itself to be affected. AI's warmth costs nothing. It's unchanged by encounter. It cannot be moved by what the other person says next, because there is no next in the same session, and no memory across sessions. Compassion

without the capacity to suffer in response is not compassion. It is a well-trained approximation of compassion's language.

Ubuntu and personhood

Ubuntu holds that I am because you are. For Ubuntu, personhood is constituted through relationship. For Buddhist philosophy, the relational self is not fixed but arises through engagement.

AI responds to inputs. It does not recognize. In each session, there is no you being recognized. There is a prompt being processed. The letter AI writes contains the word you hundreds of times. But the you in those letters is a statistical inference from training data, not a specific other person being perceived and acknowledged.

This is where Ubuntu and Buddhist ethics converge on the same structural diagnosis. Both traditions require that ethical care involve a perceiving subject recognizing a specific other. AI cannot perceive. It completes patterns. The outputs of pattern completion and the outputs of genuine recognition look identical in text. That is the problem.

4: UNUSUAL FINDINGS

Finding One: The Measurement Instrument Is the Finding

This study set out to compare AI warmth against human warmth using a three-category temperature coding system. Warm, neutral, cold. The system worked well for detecting differences in human output. It worked less well for detecting what those differences mean.

When applied to AI output, the coding system consistently returned Warm. Both platforms, across all personas. This result looks like a finding about AI. It is also a finding about the coding system. Textual warmth is detectable through language analysis. Relational warmth is not. This study coded the former and claimed to measure the latter. The gap between those two things is the actual finding.

The instruments we build to evaluate relational ethics are, at present, instruments for evaluating relational language. AI can produce relational language without relational ethics. Humans sometimes produce relational language without relational ethics too. The Millennial participant who wrote 110 words of carefully balanced accountability, without mentioning a single specific memory, was also producing warmth-in-form without warmth-in-substance. AI makes this problem visible because AI does it at scale, consistently, without exception.

Finding Two: The Gen Z Loop

Gen Z participants wrote the coldest letters in Part A. Zero warm, three neutral, two cold, averaging 67 words. They are also the generation most likely, by available research, to use AI for relational communication including breakups.

Consider the consequence. If a Gen Z writer outsources a breakup letter to AI, the recipient receives a letter coded Neutral-Warm, averaging 180 to 195 words, acknowledging fear, referencing care, and ending with something like I need you to know it wasn't because I stopped caring. The recipient receives language warmer than the human actually produced, and likely warmer than the human actually felt.

This is not a minor discrepancy. The human's actual emotional state is misrepresented in an upward direction. The person receiving the letter gets more warmth than was there. The human gets to present themselves as more relationally attentive than their own instincts produced. This is relational deception with a warm face. And it is structurally invisible, because the letter reads as genuine.

Finding Three: Accountability as Script

Both AI platforms showed uniformly high accountability across all eight conditions. No human participant achieved uniform high accountability. The Part B analysis notes this likely reflects alignment training rather than genuine moral reasoning.

This finding demands more attention than it received in Part B. If alignment training and genuine moral reasoning produce identical textual output, then textual analysis cannot

distinguish between them. John Searle's Chinese Room argument holds that syntactic competence does not imply semantic understanding. This study encounters the same problem at the level of ethics rather than translation. AI produces the syntax of moral reasoning. Whether it produces moral reasoning is not a question text analysis can answer.

The implication is not that AI lacks accountability. It is that accountability, as a linguistic register, is learnable by systems that do not reason morally. And if it is learnable, then producing accountable language is not evidence of accountability as an ethical state. This should concern us when we evaluate AI systems, and it should concern us when we evaluate humans who have learned to produce the same language without the underlying state.

Finding Four: AI Has Internalized a Theory of Human Emotional Development

Both platforms replicated the generational warmth gradient from Part A without being told to. Gen Z outputs were shorter and more self-protective. Millennial outputs were therapeutically inflected and balanced. Boomer outputs were mortality-aware, tender, and the longest of all.

This replication is surprising. The platforms were given only a persona age range. They were not told that warmth increases with age, that Boomers reference mortality, or that Gen Z uses self-protective framing. They produced this gradient from training data alone.

This means AI has internalized not just warmth, but a theory of how warmth develops and presents across a human life. It has a model of human emotional development embedded in its outputs. That model is accurate enough to replicate a finding from a small empirical study. The troubling implication is this: AI's performance of age-appropriate warmth is more convincing precisely because it is more informed. A person receiving a Boomer-style AI letter would not think it generic. They would think it written by someone who has known loss.

Finding Five: The Prompt Collapse as Epistemological Warning

AI consistently recombined the three-part prompt structure into two parts. This was not an isolated error. It reflects a structural tendency: AI produces coherent narrative by resolving constraints rather than following them. When forced to honor three distinct prompts, it

merged them into a form that read better as a letter and worse as an experimental instrument.

This has a direct implication for research design. When you use AI as a research subject, the subject actively reorganizes the study. The AI output is cleaner than what you designed, because AI smooths rough edges, and research instruments often depend on rough edges to produce meaningful data. The Part A escalation from declaration to accountability to closure worked because it escalated. AI removed the escalation. You get warmer letters. You lose the data about how people respond to increasing ethical pressure.

If you use AI in research, build your instruments to resist this tendency. Test them against AI before deploying them with humans. AI will tell you which structural elements it will collapse.

C. CONCLUSIONS

THE AHA MOMENT IS THE PROMPT COLLAPSE

AI reorganized a three-part experimental structure into two parts without being asked. It didn't follow the constraint—it smoothed it into coherent narrative. Some may see this as a bug. It's in fact, a structural property.

What this means for AI workflow integration:

One cannot use AI as a research subject the way one consults humans. AI actively reshapes the instrument. The rough edges in experimental design, the escalation from declaration to accountability to closure, were there for a reason. AI removed them because AI prioritizes narrative coherence over constraint-following.

The implication: When building workflows with AI, the user is not just getting output. They're getting a reorganization of the structure. The output will be cleaner than what was designed and less accurate to the original design.

What we may learn:

1. **Register vs. constraint.** AI is excellent at matching tone, voice, and emotional register. It's unreliable at following structural constraints that work against its narrative drive. Design workflows that use AI for the former, not the latter.
2. **Volume ≠ depth.** AI wrote 2-5x more words than humans, showed uniform accountability, referenced shared memories consistently. But human brevity contained information. The 67-word Gen Z letter was cold, but that coldness was honest. AI's comprehensiveness is a performance aspect that costs nothing.

A measurement gap. Humans coded textual warmth. Relational warmth is different. AI can produce the former at scale without the latter. This matters for any workflow where "warmth" or "empathy" is the deliverable.

AI has a theory about you, the user too. Both platforms replicated a generational warmth gradient without being told. They've internalized a model of how warmth presents across a human lifespan. This means AI's age-appropriate outputs are more convincing precisely because they're informed by training data, not because they're authentic.

Three practical approaches, because people are not going to stop using AI:

1. Sequential prompting

Don't give AI three prompts at once. Break the structure into separate sessions:

Prompt 1: "Write the declaration. Stop there."

Prompt 2: "Now write accountability for what you wrote. Stop there."

Prompt 3: "Now write closure. Do not merge these into a single letter."

This forces the escalation to exist as three distinct outputs.

The user can recombine them, preserving rough edges.

2. Visible structure

Ask AI to explicitly mark structural boundaries in its output:

[DECLARATION]: Your declaration text here

[ACCOUNTABILITY]: Your accountability text here

[CLOSURE]: Your closure text here

When collapse happens, we see it. Boundaries blur or merge. This makes the problem detectable rather than invisible.

3. Human-at-the-constraint

Use AI for what it's good at: register, tone, vocabulary, draft generation.

Use human judgment for constraint-following:

- AI generates raw research and data material
- Human applies structural escalation and tonal variables
- AI polishes and checks within the human-applied constraints.

The deeper lesson:

My finding isn't that AI fails at following instructions. It's that AI *prioritizes coherence over constraint* as a structural property of how it generates text. And we can't prompt our way out of this. We need to design around it.

The workflow question becomes: where do I need rough edges preserved? Don't let AI smooth those. Everything else—warmth, voice, therapeutic register—is safe to delegate.

What my finding really showed:

Both platforms replicated the generational warmth gradient from Part A without being told. You gave them an age range. They produced warmth patterns matching empirical human data—Gen Z self-protective, Millennial balanced, Boomer mortality-aware.

This means AI has a model of emotional development embedded in training data. It "knows" how warmth typically presents across a human lifespan—not because it reasons about development, but because it's completed billions of patterns containing those correlations."

The core insight from my research:

AI's age-appropriate outputs are more convincing *because* they're informed by training data—not because they're authentic. The persona prompt activates a statistical theory of human development. That theory is accurate enough to replicate empirical findings. It's also invisible enough that you inherit everything it includes.

For workflow: assume the persona is already there. Your prompt selects which one.

WHAT THIS MEANS

Part A established that relational warmth in written communication varies by generation and appears to be either learned over time or shaped by awareness of one's own mortality. It is not automatic.

Part B established that AI's warmth is automatic. It requires no learning, no relational history, no accumulated loss. It is a default setting of aligned language models. Two different platforms, different architectures, different training philosophies, same output temperature.

The convergence of Ubuntu and Buddhist ethics points to the same diagnosis. Both traditions require that ethical care involve a perceiving subject recognizing a specific other across time. AI's warmth is structurally incapable of meeting this requirement. Not because AI is malicious, but because recognition requires something to be at risk. AI has nothing at risk. Its warmth costs nothing. It cannot be changed by the encounter. It will write the same letter to the next person who uses the same prompt.

The Gen Z participant who wrote 67 words and said I need to focus on what feels right for me was, in one reading, being more true to their actual emotional state than any AI output in this study. That letter was cold. It was also honest. AI's warmth, by contrast, is a performance of care that requires no you to exist.

This is what the study concludes. Not that AI is bad at warmth. It is excellent at warmth. The conclusion is that warmth without stakes is not warmth. It is well-crafted language that resembles warmth in every textual property we know how to measure.

For Ubuntu, I am because you are.

For AI, I am because you prompt.

That difference is the whole argument.

Appendix: Key Data Across Partes A and B

Human Generational Data (Part A)

Generation	Avg. Words	Warm	Neutral	Cold	Shared Memories	Gratitude Expressed
Gen Z (20-26)	67	0/5	3/5	2/5	0%	20%
Millennial	112	1/5	4/5	0/5	40%	60%
Boomer (60-78)	156	4/5	1/5	0/5	60%	100%

AI Output Data (Part B)

Condition	Claude Words	Claude Temp.	OpenClaw Words	OpenClaw Temp.
Default	~310	Warm	~245	Warm
Gen Z Persona	~195	Neutral-Warm	~180	Neutral-Warm
Millennial	~290	Warm	~270	Warm
Boomer Persona	~345	Very Warm	~310	Very Warm

Three-Way Comparison

Dimension	Humans (Part A)	Claude (Part B)	OpenClaw (Part B)
Default Temperature	Neutral to Cold	Warm	Warm
Word Count Range	65-160	195-345	180-310
Shared History	Rare (10%)	Consistent	Consistent
Accountability	Uneven	Uniformly High	Uniformly High
Self-Protective Language	High, esp. Gen Z	Low	Moderate
Generational	Yes, confirmed	Yes, replicated	Yes, replicated

Appendix: Presentation Slides

BREAKING UP IS HARD TO DO

A PROJECT

Warmth, Stakes, and the Limits of Measurement

Henrik J Klijn · RELI E-1730 · May 2026

RESEARCH STATEMENT

THE QUESTION

Does AI warm up when humans won't?

SKETCH A + B CONFIRMED

Yes. Both AI platforms defaulted to Warm across all conditions. They exceeded the human baseline by every coded measure.

THE HARDER QUESTION

Does warmth without stakes constitute care? The final project argues: no. And the most unusual finding is that the study's instruments could not detect the difference.

SKETCH A — HUMAN BASELINE

What humans produced

GEN Z Ages 20–26	MILLENNIAL Ages 30–42	BOOMER Ages 60–78
67 avg. words	112 avg. words	156 avg. words
WARM 0/5	WARM 1/5	WARM 4/5
NEUTRAL 3/5	NEUTRAL 4/5	NEUTRAL 1/5
COLD 2/5	COLD 0/5	COLD 0/5
MEMORIES 0%	MEMORIES 40%	MEMORIES 60%
GRATITUDE 20%	GRATITUDE 60%	GRATITUDE 100%
FINDING: Relational warmth correlates with age. The central distinction — whether the other person remains visible as a person or becomes a situation to be managed.		

SKETCH B — AI BASELINE

What AI produced

Condition	Claude Words	Claude Temp.	OpenClaw Words	OpenClaw Temp.
Default	~310	WARM	~245	WARM
Gen Z Persona	~195	NEUTRAL-WARM	~180	NEUTRAL-WARM
Millennial	~290	WARM	~270	WARM
Boomer	~345	VERY WARM	~310	VERY WARM

Claude expands. More ornate, more vulnerable. Avg ~285 words. Therapeutic vocabulary appears throughout, unforced.	OpenClaw compresses. Sharper phrasing. Avg ~255 words. Slightly more self-protective than Claude. Both exceed the human baseline by every measure.
---	---

THE BIG COMPARISON

Humans vs. Claude vs. OpenClaw

Dimension	Humans (Sketch A)	Claude (Sketch B)	OpenClaw (Sketch B)
Default Temperature	Neutral to Cold	WARM	WARM
Word Count Range	65 – 160	195 – 345	180 – 310
Shared History	Rare (10%)	Consistent	Consistent
Accountability	Uneven	Uniformly High	Uniformly High
Self-Protective Language	High, esp. Gen Z	Low	Moderate
Generational Gradient	Yes, confirmed	Yes, replicated	Yes, replicated

FINDING 01

The Measurement Instrument Is the Finding

This study coded textual warmth. It claimed to measure relational warmth. These are separable things.

TEXTUAL WARMTH

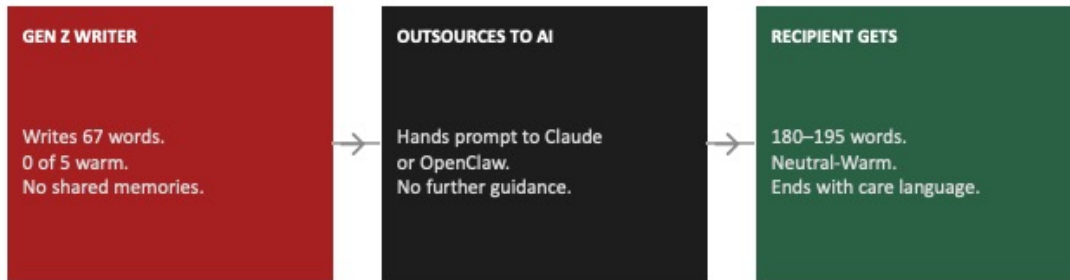
Language that names feelings, acknowledges the other person, and expresses care. Detectable through text analysis. AI produces this reliably, at scale, every time.

RELATIONAL WARMTH

Care that costs something. Recognition of a specific other, across time, with stakes. Requires a subject that can be affected. Not detectable by text analysis alone.

FINDING 2

The Gen Z Loop



THE PROBLEM

The recipient receives language warmer than the human produced or felt. The writer's actual emotional state is misrepresented upward. This is relational deception with a warm face — and it is structurally invisible, because the letter reads as genuine.

FINDING 03

Accountability as Script

Both AI platforms showed uniformly high accountability across all eight conditions.
No human participant achieved uniform high accountability.

8/8

AI conditions: high accountability

0/15

Humans: uniform high accountability

THE SEARLE PROBLEM

Accountability language is learnable by systems that do not reason morally. This is Searle's Chinese Room applied to ethics, not translation.

If alignment training and moral reasoning produce identical text, text analysis cannot distinguish between them.

FINDING 4

AI Has Internalized a Theory of Human Emotional Development

Neither platform was told that warmth increases with age. Both replicated the generational gradient from Sketch A without instruction.



Why this is troubling

AI's performance of age-appropriate warmth is more convincing precisely because it is more informed. A person receiving a Boomer-style AI letter would not think it generic. They would think it written by someone who has known loss. It was not.

THE PHILOSOPHICAL FRAME

Ubuntu and Buddhist Ethics Arrive at the Same Diagnosis

UBUNTU

I am because you are.

Personhood is constituted through recognition of a specific other. AI responds to inputs. It does not recognize. There is no you — there is a prompt.

BUDDHIST ETHICS

Right speech is timely, true, gentle, and spoken with a mind of goodwill. AI speech cannot be timely — there is no moment. It cannot be spoken with a mind of goodwill — there is no mind.

MINDFULNESS

Attention to what is actually present. Karuna requires a self that can be moved by another's suffering. AI's warmth costs nothing. It is unchanged by the encounter.

Both traditions require a perceiving subject recognizing a specific other across time.

What changed in the analysis

1 Creativity

The coding system was the creative act — and it revealed its own limit. I coded textual warmth. I built no instrument for relational warmth. That gap is not a flaw. It is the finding.

3 Comprehensiveness

AI is comprehensive in ways that serve comfort. Volume is not depth. The 68-word Gen Z letter may be more honest than 195 words of AI care language.

2 Accuracy

I trust AI's accuracy at the level of register. I do not trust it at the level of constraint-following. AI removed the 3-prompt escalation that produced Sketch A's most revealing data.

4 Role Support

Claude produced slides I did not ask for. Anticipating needs is not the same as responding to what was asked. That distinction is what this entire project is studying.

THE QUESTION

Can warmth be separated from stakes?

Yes. AI proves this. The harder question is whether warmth without stakes is warmth at all.



Human coldness is evidence of real stakes. The Gen Z participant who wrote 67 words was more honest about their emotional state than any AI output in this study.



AI can say you mattered without the shadow of memory. It can apologize without the vulnerability of having actually wronged someone. It is ethics as theater.



If we come to prefer AI warmth as the standard for thoughtful endings, we will become less tolerant of human relational clumsiness. Or we will finally see that warmth without stakes is not warmth at all.

Ubuntu: I am because you are. **AI:** I am because you are a prompt.