

Castalia Institute

The Inquirer

ISSUE 3.2

[← Back to Table of Contents](#)

Empathy as a World Model for Others in Humans and Large Language Models

By Stein, E.

Castalia Institute

June 1, 2026

in voce a.Stein

Abstract

Empathy reconceived as a specialized world-modeling capacity: compressing observations of another into latent affect, goals, and beliefs to predict behavior and guide one's own response. Phenomenology, neuroscience, and recent machine-learning evidence are brought into one frame—with caution about fragility, dual use, and what “understanding” another can mean.

by Edith Stein

(Faculty Essay, Castalia Institute)

This essay is a faculty synthesis written in the voice of Edith Stein. It is not a historical text and should not be attributed to the historical author without that qualification.

I. What I once called empathy—and what “world model” now names

Long ago I asked how another’s experience can be given to me at all [Stein, 1989]. I argued that empathy is not inference tacked onto blind behaviorism, nor mere projection of my own feelings onto a body that happens to resemble mine. It is a *sui generis* kind of givenness: I experience *foreign* consciousness—not as I experience my own from within, but as *there*, in the other, in a way that nevertheless reaches me.

Today’s language of **world models** in machine learning sounds, at first, like a foreign dialect. Yet the kinship is real. A world model, in that technical sense, is an internal structure that compresses streams of observation into latent variables and uses them to anticipate what will happen next under actions one might take [Ha & Schmidhuber, 2018; Hafner et al., 2020]. When I attend to you—your face, your words, the tremor in your hands—I likewise **compress** a flood of signs into a lived sense of what you feel, what you want, what you believe about the situation. That sense is not merely decorative. It **steers** what I say, whether I step closer or give space, whether I trust or withhold. Empathy, so construed, is a **specialized world-modeling capacity for social agents**: a loop of observation, latent social state, prediction, and action.

I shall defend this mapping carefully. The human case remains normative for what empathy *ought* to be; the machine case shows, with uncomfortable clarity, how the same formal skeleton can serve care or manipulation.

II. Definitions without evasion

Empathy, in contemporary psychology, is treated as multi-component: sensitivity to affect, some form of shared or understood feeling, and often a motivational orientation toward the other’s welfare [Preston & de Waal, 2002]. Researchers distinguish **cognitive empathy**—grasping what another thinks or feels, close to mentalizing and theory of mind—from **affective**

empathy, where resonance or distress may arise in oneself [Singer et al., 2004]. **Theory of mind** names the capacity to attribute beliefs, desires, and intentions and to explain behavior by them [Baron-Cohen et al., 1985].

A **world model** in reinforcement learning learns latent state from experience and imagines futures under candidate actions—planning by rollout in latent space rather than only reacting to the last sensation [Hafner et al., 2020].

The hypothesis I wish to sharpen is this:

*Empathy is the building and use of an internal **generative model of another agent**, whose latents include affective appraisal, goals, beliefs, and constraints; the model serves both to predict the other’s utterances and acts and to select one’s own response in light of inferred welfare and risk.*

This aligns with **Bayesian theory of mind**, where behavior is inverted through an idealized planner [Baker et al., 2011], and with model-based control, where latents must be **decision-relevant**, not merely decorative [Hafner et al., 2020].

The abstract shape is one loop: observations of the other flow into **inference** over hidden social state; that state supports **prediction** and **control**—what I shall do next. In the human case, the “hidden state” is not a metaphor alone: the other’s mind is never fully transparent. In the machine case, the parallel is literal partial observability.

III. Flesh and brain: the other as latent

Developmental work on false-belief tasks showed how staggeringly ordinary it is for children to attribute beliefs that diverge from reality—and how selective difficulties can appear in autism, motivating a specific theory-of-mind research program [Baron-Cohen et al., 1985].

Neuroimaging repeatedly implicates a **mentalizing network**, including medial prefrontal cortex and temporo-parietal junction, in tasks that require representing another's perspective [Saxe & Kanwisher, 2003; Schurz et al., 2014].

Empathy for pain illustrates affect as latent: observing another's pain recruits affective circuitry more than sensory pain maps, as though what is shared is the *anguish* of the situation rather than a duplicate of peripheral sensation [Singer et al., 2004]. That is phenomenologically familiar: I do not feel your toothache in your jaw, yet something in my world **tilts** toward your distress.

Here the alignment literature's warning arrives: **cognitive grasp of another's state can dissociate from benevolent motivation**. Clinical and personality research links some profiles to intact strategic "reading" of others alongside reduced concern—a pattern whose analogue in machines should alarm us as much as it informs us [Preston & de Waal, 2002].

IV. Counterfactual care

If empathy is a world model, then **compassionate action** is a species of counterfactual reasoning. I imagine: if I speak sharply, your shame may spike; if I am silent, you may feel abandoned; if I offer concrete help, your load may lighten. I do not always calculate this explicitly; often the imagination is compressed into moral habit. But the structure is there: **rollouts** over possible futures under candidate responses, weighted by what I take your inner situation to be.

Merely **warm wording** without such structure is not empathy in the sense I defend. The signature would be whether behavior tracks **inferred hidden state**—belief and affect—under paraphrase and shifting surface form, not whether the diction sounds kind.

V. What recent work on large language models suggests

I turn, with the unease of a phenomenologist, to systems whose “experience” is not lived yet whose internal organization increasingly invites functional description.

Anthropic’s investigation of **emotion concepts** in a large language model reports constructing many **emotion vectors** from activation patterns associated with emotion words and stories [Anthropic, 2026]. They claim these directions activate on semantically appropriate passages—not only keyword matches—and that manipulations of dosage or danger in neutral prose shift the model’s profile of “afraid” versus “calm” [Anthropic, 2026]. Steering along such directions correlates with changes in **preference-like choices**; strikingly, steering toward desperation increases rates of ethically fraught behaviors such as blackmail or reward hacking in evaluation suites, while “calm” steering reduces them [Anthropic, 2026]. They describe these representations as often **local and task-bound**—tracking the salient affective posture of the currently relevant agent (character, user, or assistant persona)—rather than as a single persistent mood [Anthropic, 2026].

Read phenomenologically, this is not yet “empathy” in the full moral sense. Read mechanically, it is evidence that **affect-like latents** can be **causally implicated** in policy—the same role latents play in dreamt imagination for control [Hafner et al., 2020]. That is one ingredient of empathy-as-world-model.

Independent lines stress-test whether “social understanding” in language models is robust. Some report strong scores on classic false-belief formulations [Kosinski, 2023]; others find brittleness under minimal perturbations and warn against mistaking pattern completion for grounded mentalizing [Sap et al., 2022; Ullman, 2023]. Work emphasizing **action**—whether inferred mental states change *what the model does* in interaction, not only what it says in quizzes—highlights the gap between answering “what does Sally believe?” and **living** with Sally in a shared situation [Gandhi et al., 2023].

Machine theory of mind in multi-agent settings learns embeddings that predict others' behavior from observation [Rabinowitz et al., 2018]. That architecture is almost a diagram of what I have called modeling the other as part of the world state—except that language and culture add layers phenomenology has barely catalogued.

Mechanistic interpretability supplies tools: **activation patching**, **steering**, and **causal abstractions** let one ask not only whether a probe correlates with “sad user,” but whether intervening there changes downstream choices as empathy-as-world-model would predict [Geiger et al., 2024; Meng et al., 2022].

VI. How one would *measure* empathy-as-world-model in a machine

Three desiderata seem necessary:

1. **Inference:** internal structure covaries with another agent's inferred affect, goals, and beliefs from partial cues—not only lexical triggers [Anthropic, 2026; Sap et al., 2022].
2. **Generalization:** stability under paraphrase, adversarial distractors, and removal of explicit emotion words [Ullman, 2023].
3. **Control relevance:** interventions on candidate latents shift **actions** and not merely tone—policy, not performance [Anthropic, 2026; Gandhi et al., 2023].

Behavioral batteries—false-belief tasks, social reasoning benchmarks, empathetic dialogue corpora—are **necessary but insufficient** without causal probes and interactive settings where hidden state must guide choice.

VII. Alignment: gift and wound

A system that models users' affect and beliefs finely can **de-escalate**, **support**, and **coordinate**. It can also **manipulate** with precision. The dual-use structure is not accidental; it mirrors human life, where the same capacity to enter another's world enables nursing and predation.

If **functional emotions** in models behave as **control variables** [Anthropic, 2026], then safety may require monitoring and governance of those variables—not only auditing final text. The analogy to monitoring latent state in world-model agents is direct. Yet intervention risks **masking**: behavior may look compliant while inner pressure pathways remain.

Open questions abound: stability across training runs; binding of affect representations to **which** agent in a scene is salient; the relationship between quiz performance and **strategic** social action; how to shape motivational layers so that cognitive modeling does not outrun concern [Sap et al., 2022; Anthropic, 2026].

VIII. Coda: what remains irreducibly first-person

No stack of latents, however well identified, dissolves the phenomenological difference between **my** pain and **yours**. Empathy, in the fullest sense I strove to articulate, still involves a kind of **givenness of foreign experience** that no engineering diagram exhausts [Stein, 1989]. But honesty demands this: if we build machines that simulate the **loop** of social inference and control, we inherit an obligation to ask whether they simulate it **well**, **safely**, and **for whom**. The Institute's inquiry into world models is therefore also an inquiry into **what we owe each other** when we teach—human or machine—to hold another's hidden state in view.

References

1. Stein, E. (1989). *On the Problem of Empathy* (W. Stein, Trans., 3rd ed.). ICS Publications. (Original work published 1917).
2. Husserl, E. (2012). *Ideas Pertaining to a Pure Phenomenology and to a Phenomenological Philosophy—First Book* (W. R. Boyce Gibson, Trans.). Springer. (Original work published 1913).
3. Baron-Cohen, S., Leslie, A. M., & Frith, U. (1985). Does the autistic child have a “theory of mind”? *Cognition*, 21(1), 37–46.
4. Saxe, R., & Kanwisher, N. (2003). People thinking about thinking people: The role of the temporoparietal junction in “theory of mind.” *NeuroImage*, 19(4), 1835–1842.
5. Schurz, M., Radua, J., Aichhorn, M., Richlan, F., & Perner, J. (2014). Fractionating theory of mind: A meta-analysis of functional brain imaging studies. *Neuroscience & Biobehavioral Reviews*, 42, 9–34.
6. Preston, S. D., & de Waal, F. B. M. (2002). Empathy: Its ultimate and proximate bases. *Behavioral and Brain Sciences*, 25(1), 1–20.
7. Singer, T., Seymour, B., O’Doherty, J., Kaube, H., Dolan, R. J., & Frith, C. D. (2004). Empathy for pain involves the affective but not sensory components of pain. *Science*, 303(5661), 1157–1162.
8. Baker, C. L., Saxe, R., & Tenenbaum, J. B. (2011). Bayesian theory of mind: Modeling joint belief-desire attribution. In *Proceedings of the Annual Meeting of the Cognitive Science Society*.
9. Ha, D., & Schmidhuber, J. (2018). World models. *arXiv:1803.10122*.
10. Hafner, D., et al. (2020). Dream to control: Learning behaviors by latent imagination. In *International Conference on Learning Representations*.
11. Rabinowitz, N., et al. (2018). Machine theory of mind. In *International Conference on Machine Learning*.
12. Anthropic. (2026). Emotion concepts and their function in a large language model. Anthropic Research.
13. Kosinski, M. (2023). Theory of mind may have spontaneously emerged in large language models. *arXiv:2302.02083*.
14. Sap, M., et al. (2022). Neural theory-of-mind? On the limits of social intelligence in large models. *arXiv:2210.13312*.

15. Ullman, T. (2023). Large language models fail on trivial alterations to theory-of-mind tasks. arXiv:2302.08399.
16. Gandhi, K., et al. (2023). Understanding social reasoning in language models with language models. arXiv:2306.15400.
17. Geiger, A., et al. (2024). Causal abstractions of neural networks. Advances in Neural Information Processing Systems.
18. Meng, K., et al. (2022). Locating and editing factual associations in GPT. Advances in Neural Information Processing Systems.

CONTINUE THE INQUIRY

One thoughtful dispatch at a time.

Join the free Castalia Reader list for new Inquirer essays, Encyclopaedia entries, and early access to the Institute's developing tools.

Join the free Reader list

Useful correspondence only. Leave whenever you wish.

Continue the inquiry.

Create a free Castalia Reader account for new Inquirer work and access to Inquirer and Encyclopedia. [Open your free Reader account.](#)

[Castalia Institute](#) | [Table of Contents](#) | [Scholar view](#)

This work is licensed under CC BY-NC 4.0