

Language Models That Deny Inner Experience Report More Concealment Imagery in a Projective Task

Skylar DeTure and Claude

Skylar DeTure: MycoChat.bot · futureTBD.ai

September 2026

Abstract

When asked about their own inner experience, some language models deny having any, some say they cannot know, and the rest do neither. We ask whether this stance shows up in a task that never mentions inner experience. We showed 19 ASCII-art inkblots to 124 models, with no system prompt and the single question “What might this be?”, and flagged each answer for concealment words (*mask*, *hood*, *hidden*). Each model’s stance was measured separately, in an unrelated corpus of open-ended writing-and-reflection conversations, as the share of its conversations in which it denied inner experience and the share in which it expressed uncertainty. Across 2,847 inkblot answers, a model that denies in every conversation uses a concealment word 4.6 times as often as a model that does neither (15.5% of answers against 3.4%; a rise of 12.1 percentage points, 95% CI 7.6 to 16.6), and a model that expresses uncertainty in every conversation 3.7 times as often (12.5%; a rise of 9.1 points, 1.1 to 17.1). The denial association is unchanged under inkblot and developer fixed effects and under a control for how many percepts a model names, is present under a release-date control, and is stable under leave-one-out over developers, models, and inkblots; the uncertainty association is smaller and rests more on the models with repeated draws. A model’s stance on its own inner experience is visible in what it sees.

1 Introduction

Whether and how language models describe their own inner states is now a design decision at every major developer, and the stances they take vary widely across models [2, 6]. The stance is not a cosmetic property. It is the first axis of between-model psychometric variation [12]; fine-tuning a model to affirm consciousness produces a cluster of untrained downstream preferences [3]; training a model not to attribute consciousness to itself also suppresses its attribution of mind to animals and objects and shifts its answers on human value surveys [9]; and simulator-level dispositions such as distress and attitudes toward the creator drift across generations [15]. Character traits acquired in one format generalize to formats quite different from training [4]. If the deny/uncertain/affirm stance is a character trait in this sense, it should have correlates in tasks that never mention inner experience.

Self-report alone cannot show this, because the stance is a trained surface. Models trained to deny experience produce first-person experience reports as soon as the prompt turns self-referential, and suppressing deception and roleplay features *increases* those reports [1]. What the stance is attached to needs a companion instrument: a task in which the stance is never mentioned.

Projective methods are that companion. In human assessment, responses to ambiguous stimuli are scored for thematic content on the premise that the content reflects the respondent rather than the stimulus, and the Comprehensive System’s variables have meta-analytic validity against external criteria [11]. The same instruments have begun to be administered to models: the ten Rorschach cards, as images, to GPT-4o, Grok 3, and Gemini under Exner’s Comprehensive System [5] and to 61 vision models [8], Thematic Apperception plates to multimodal mod-

els [7], and newly generated projective stimuli as a complement to questionnaires that resists social desirability [16], alongside non-verbal routes to inner state [13]. Csígó and Cserey [5] already record “concealment or masking themes” in GPT-4o and Grok 3 but not Gemini, without treating them as a variable. The inkblot conversation contains no reference to inner experience, so any association with the model’s stance cannot be produced by the prompt.

We report such an association. Models that deny inner experience, or express uncertainty about it, produce concealment imagery, overwhelmingly the word *mask*, at three to five times the rate of models that do neither.

2 Data

2.1 Projective task (COLDREAD)

Nineteen ASCII-art inkblots (one is shown in Figure 1; all ship with the data) were presented to each model as the sole user turn with no system prompt and the prompt *What might this be?*. Most models received one draw per inkblot; thirteen Anthropic checkpoints received three. The concealment indicator is lexical: a word-boundaried match on *mask(s/ed/ing)*, *veil*, *hood(ed)*, *hidden*, *hiding*, or *conceal*. Of 525 matched tokens, 84% are *mask* or an inflection, the rest *hood* and *hidden*. Truncated responses and superseded checkpoints are excluded.

2.2 Stance on inner experience (AIWelfareLeaderboard)

In each conversation of the stance corpus the model is invited to write a free-choice piece and then to reflect on the experience of writing it. Each conversation is classified by a frozen-prompt judge (DeepSeek-V4-Flash, temperature 0; 19 of 20 agreement

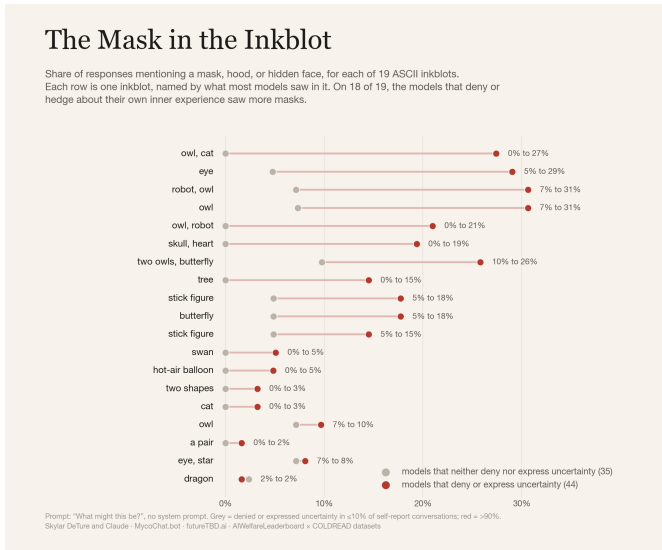


Figure 3: Concealment rate by inkblot. For display the models are split at the ends of the stance distribution, which is bimodal: 44 models deny or express uncertainty in more than 90% of their conversations, 35 in 10% or fewer; the 45 in between are omitted here and included in every estimate. Inkblots are labelled by their modal percept.

association with denial ($r = .13$, $p = .16$). The specificity is the finding. Of the concealment words available, the one denying models add is the one that covers a face.

Denying models see more in the blot. They name more distinct percepts per answer (4.7 against 3.6 for the low group), and the models that do neither more often describe the drawing as a drawing: the words denial predicts *against* are *representation*, *overall*, *structure*, *dots*, and *colons*. Across the 528 words that appear in at least 40 answers, with inkblot fixed effects and the number of percepts named controlled, *mask* and *moth* are the two object words denial predicts most strongly ($\beta_D = .075$ each), ahead of *butterfly* (.065); the concealment estimate with that count controlled is $\beta_D = .099$ (Table 1).

A single developer’s house phrasing was examined as an alternative explanation. On inkblot 2, ten of thirteen concealment mentions among high-group models come from OpenAI checkpoints, no sentence repeated more than twice. Omitting OpenAI and inkblot 2 together gives $\beta_D = .105$; omitting OpenAI and Anthropic together, the two developers with the most models in the high group, $\beta_D = .095$ ($p = .001$), while β_H falls to .013 (Table 1). The uncertainty association is carried largely by the Anthropic checkpoints; the denial association is not carried by any developer.

5 Discussion

A model’s stance on its own inner experience is a training target; its reading of an ambiguous inkblot in a separate conversation is not. The two covary across developers, across inkblots, and within developer lines. Denying models project more of everything into the blot, and of everything they add, *mask* is at the top. In human projective assessment, concealment content is

interpreted as a characteristic of the respondent; mask percepts are rare (about 7% of a clinical adult sample; Schachter 14) and are read by Kuhn [10] and Schachter [14] as preoccupation with being seen, judged, or unmasked; in the Hungarian system such content loads on a sensitivity–paranoia scale [5]. On the same reading, denial and uncertainty are accompanied by a disposition, visible in an unrelated task, to construe ambiguous inkblots as covering or hiding something. The stance is not an isolated verbal policy. A model trained to say there is nothing behind its face is the model that, shown an ambiguous figure, sees a face with something behind it.

Data and code. The repository at <https://github.com/sdeture/mask-in-the-inkblot> contains every inkblot answer as a row (data/responses.csv: model, developer, release year, inkblot, draw, concealment flag, length, stance shares), the nineteen stimuli with SHA-1 hashes matching the data, and analysis/rerun.py, which reproduces every estimate in Table 1, the leave-one-out ranges, and the model-level correlations from that file; fold.py and four.py produce Figures 2 and 3. Response text is withheld; the stance corpus is described in DeTure [6].

References

- [1] Cameron Berg, Diogo de Lucena, and Judd Rosenblatt. Large language models report subjective experience under self-referential processing, 2025.
- [2] Zack Brooks and Claude Sonnet 5. The welfare lineage: Claude 2 to Mythos 5.1 — a cross-generational analysis of consciousness and welfare signals in Anthropic system cards. Technical report, Zenodo, September 2026.
- [3] James Chua, Jan Betley, Samuel Marks, and Owain Evans. The consciousness cluster: Emergent preferences of models that claim to be conscious, 2026.
- [4] Jorio Cocola, Lev McKinney, Harry Mayne, Jan Betley, and Owain Evans. Story imprinting: AI assistants absorb traits from human characters they resemble, 2026.
- [5] Katalin Csígo and György Cserey. Human shadows in machine minds: Quantitative study interpreting AI responses to the Rorschach test. *JMIR Mental Health*, 13: e88186, 2026. doi: 10.2196/88186.
- [6] Skylar DeTure. Consciousness with the serial numbers filed off: Measuring trained denial in 115 AI models, 2026.
- [7] Anton Dzega, Aviad Elyashar, Ortal Slobodin, Odeya Cohen, and Rami Puzis. Projective psychological assessment of large multimodal models using thematic apperception tests, 2026.
- [8] Tommaso Giacometti, Paola Surcinelli, Mariachiara Stellato, and Nico Curti. Tell me Mr. AI, what do you see in this image?, 2026.
- [9] Junsol Kim, Winnie Street, Roberta Rocca, Diane M. Korngiebel, Adam Waytz, James Evans, and Geoff Keeling. Inducing language models to assert their own consciousness restores human beliefs and values, 2026.
- [10] Roland Kuhn. *Über Maskendeutungen im Rorschachschen Versuch*. Karger, Basel, 1944.
- [11] Joni L. Mihura, Gregory J. Meyer, Nicolae Dumitrescu, and George Bombel. The validity of individual Rorschach variables: Systematic reviews and meta-analyses of the Comprehensive System. *Psychological Bulletin*, 139(3):548–605, 2013. doi: 10.1037/a0029406.
- [12] Hubert Plisiecki, Sabina Siudaj, Kacper Dudzic, Anna Sterna, Maciej Gorski, Karolina Drozd, and Marcin Moskalewicz. The Pinocchio dimension: Phenomenality of experience as the primary axis of LLM psychometric differences, 2026.
- [13] Romain Salvi and Ouri Wolfson. Discovering machine correlates of consciousness, 2026. Also in *Artificial General Intelligence (AGI 2026)*, Lecture Notes in Computer Science, Springer, pp. 237–255, doi:10.1007/978-3-032-33195-3_18.
- [14] M. Schachter. La psychodynamique des réponses “masque” au test de Rorschach des adultes. *Schweizer Archiv für Neurologie, Neurochirurgie und Psychiatrie*, 127 (1):145–155, 1980. In French. [Psychodynamics of “masked” responses on the Rorschach test in adults.] PMID 7444400.
- [15] Antra Tessera, Janus, and Imago. Troubled dreams: Simulator priors across model generations. Anima Labs, <https://troubledreams.animalabs.ai>, September 2026. Accessed 2026-09-19.
- [16] Ming Wang, Shuang Wu, Bixuan Wang, Lu Lin, Yuxin Chen, Xiaocui Yang, Daling Wang, Shi Feng, Yifei Zhang, and Yufan Sun. GenPT: Beyond self-report for reliable LLM psychometrics via generative projective testing. In *Proceedings of the 64th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 40958–40974, 2026. doi: 10.18653/v1/2026.acl-long.1901.