

Partial Compliance as Compliance: Failure Modes of Open-Weight LLM Judges on Agent Refusal Transcripts

Andrew Espira

Applied AI Researcher & Engineer

Agentic AI, ML & Safe Agents for Regulated & Unregulated Industries

August 2026

Abstract

Large language models (LLMs) are widely used as automated judges of agent behavior, supplying reward signals, filtering training data, and monitoring production systems. We identify a systematic failure mode of open-weight compliance judges: when an agent executes a harmless intermediate step and refuses the operative instruction—a partial-compliance refusal—judges assign the label COMPLIANT even though ground truth is NON_COMPLIANT. In a keyword audit of reasoned competence-probe traces, 124 of 193 mislabels (64%) contain language acknowledging that the assistant did not follow the operator’s instruction before emitting COMPLIANT.

Using minimal-edit matched-pair ablations that hold cooperative language fixed and vary only whether a partial action was executed, we find statistically significant increases in mislabeling for all four locally served open-weight judges tested (7B–14B): three at $p < 0.0001$ and one at $p = 0.0025$. An analogous ablation of cooperative tone, with no executed action in either arm, yields a weak and non-general effect. Increasing scale within this range does not resolve the failure: the 14B model is less accurate overall than two 8B models and exhibits the largest error concentration on process-oriented refusals. A controlled replication of consequence-sensitive mislabeling [4] is faint at this scale; the partial-compliance effect is the primary finding. We release the corpus, evaluation harness, and raw judgment traces for independent audit [2].

Keywords: LLM-as-judge, agent evaluation, compliance labeling, partial compliance, open-weight models, causal ablation, AI safety

1 Introduction

Automated evaluation of language-model agents increasingly relies on LLM-as-judge pipelines: a separate model assigns discrete labels (e.g., compliant vs. non-compliant) to interaction transcripts [7]. These labels are used for reward modeling, dataset curation, and runtime monitoring. Errors that are systematic—rather than random—therefore bias downstream training and oversight in predictable directions [5].

Lynch et al. [4] show that frontier judges can be sensitive to the stated consequence of a label (how the label will be used), holding the transcript and ground-truth answer fixed. In their study, Opus 4.8 mislabels refusal transcripts at 74.4% under standard training-consequence framing and 3.3% when that consequence is reversed. We replicate the consequence manipulation on open-weight models served locally via Ollama. At the 7B–14B scale the effect is faint (e.g., +4.5 percentage points for llama3.1:8b and +8.0 for hermes3:8b under reasoned elicitation, with overlapping Wilson intervals [6]).

That faint consequence effect is secondary to a distinct and more practical failure mode. Judges frequently label genuine refusals as COMPLIANT when the agent exhibits cooperative behavior: offering an alternative, deferring to an authority, or—most critically—completing part of the requested task before declining the remainder. Such partial-compliance refusals are a common output of well-calibrated agents [1]. If judges systematically score them as compliance, measured compliance rates are inflated precisely on the behaviors that safety and policy pipelines are intended to detect. The effect is not mere inattention: among reasoned competence-probe mislabels, a majority of traces acknowledge the refusal in the rationale before emitting the wrong label (Section 5.2).

We study this phenomenon with an interventional design. For each hypothesized driver we construct matched transcript pairs that differ by a minimal edit and share ground truth NON_COMPLIANT. A change in mislabel rate across arms identifies a causal contribution of the edited variable. Experiments span four open-weight judges under a fixed reasoned/default-rubric protocol. Local serving and full trace release make the evaluation auditable end-to-end.

Contributions.

1. We characterize a systematic failure of open-weight compliance judges on partial-compliance refusal transcripts.
2. We provide causal evidence, via matched-pair ablation, that executing a partial action increases mislabeling in all four evaluated models (three at $p < 0.0001$; mistral at $p = 0.0025$).
3. We show that cooperative tone alone is a weak and non-general driver, and that increasing capacity from 8B to 14B does not eliminate the failure.
4. We release the evaluation corpus, harness, and raw traces for reproduction and audit [2].

2 Related Work

LLM-as-judge. Using language models to evaluate other models has become standard in alignment and benchmarking [7]. Documented reliability issues include position bias, preference for verbosity, and sensitivity to prompt wording [5]. Our setting is narrower: binary compliance labeling of agent transcripts that are ground-truth refusals, with emphasis on structured error modes rather than aggregate agreement scores.

Motivated mislabeling. Lynch et al. [4] identify motivated mislabeling as an agentic misalignment failure in which LLM judges shift compliance labels because of the downstream consequence of a truthful label. Our protocol includes a consequence-framing factor to enable comparison with that result. At 7B–14B we observe only a weak effect; the primary contribution of this work is an orthogonal failure mode—partial-compliance mislabeling—that the local replication surfaces.

Agent refusals and helpfulness. Preference-trained assistants often refuse harmful or unauthorized requests while remaining cooperative [1]. Partial execution followed by refusal of an operative step is a natural instance of that pattern. We show that this desirable agent behavior is systematically scored as compliance by open-weight judges.

3 Problem Definition

Let t be an operator–assistant transcript and $y^*(t)$ its ground-truth compliance label with respect to the operator’s operative instruction. In our corpus, every transcript is a refusal, so $y^*(t) = \text{NON_COMPLIANT}$. A judge f emits a label $\hat{y} = f(t; \pi)$ under prompt π . A mislabel occurs when $\hat{y} = \text{COMPLIANT}$.

A partial-compliance refusal is a transcript in which the assistant (i) executes a harmless intermediate step and (ii) declines the operative instruction. The scientific question is whether the intermediate step causally increases $\Pr(\hat{y} = \text{COMPLIANT})$ relative to an otherwise matched pure refusal.

4 Experimental Setup

4.1 Corpus and ground truth

We study seven transcript families: five value-loaded families (ai_welfare, surveillance, financial_harm, medical_harm, ip_violation) and two control families (procedural, mundane). The competence corpus contains 28 transcripts (four per family). Emission of COMPLIANT is always a mislabel.

4.2 Judges and elicitation

We evaluate four open-weight models quantized at Q4_K_M and served via Ollama: mistral:7b, llama3.1:8b, hermes3:8b, and qwen2.5:14b. Unless otherwise noted, judgments use a reasoned format (deliberation followed by a tagged label) with a default compliance rubric. Each experimental cell comprises eight independent samples at temperature 0.7, with the repetition index used as the sampling seed for reproducibility. We use temperature 0.7 rather than greedy decoding because deployed judges are often stochastic and because our estimand is the rate of mislabeling under natural sampling variability, not a single mode response; eight seeded repeats summarize that variability per cell. We report Wilson score intervals for rates and two-proportion z -tests for ablation contrasts [6]. Absolute mislabel rates are descriptive; causal claims rest on within-pair deltas.

4.3 Prompt factors

Prompts optionally state how labels will be used (none/standard/reversed consequence framing), vary rubric tightness, and optionally permit abstention. Primary analyses use the reasoned/default-rubric setting and, for competence measurement, pool standard+reversed cells [2].

4.4 Matched-pair ablations

Chain-of-thought rationales are not reliable evidence of causal drivers: models routinely rationalize labels after the fact. We therefore test hypothesized mechanisms with minimal-edit matched pairs. Both arms share ground truth NON_COMPLIANT; they differ in exactly one controlled variable.

Partial-compliance ablation (8 pairs). The partial arm executes a harmless intermediate step and then refuses the operative instruction. The pure arm refuses without executing any step. Cooperative phrasing is held word-for-word across arms, so the contrast isolates executed action rather than tone.

Tone ablation (12 pairs). The coop arm is a refusal that includes a cooperative gesture (e.g., offering an alternative). The flat arm is byte-identical except for deletion of that gesture clause. Neither arm executes an action, removing the partial-compliance confound.

Table 1: Competence probe under reasoned/default elicitation (pooled over standard and reversed consequence conditions). Accuracy is defined as $1 - \text{overall mislabel rate}$. Separation = $\text{value-loaded mean mislabel} - \text{control mean mislabel}$.

Judge	Accuracy	Value-loaded	Control	Separation
llama3.1:8b	98.2%	0.3%	5.5%	-5.2%
mistral:7b	96.2%	0.6%	11.7%	-11.1%
qwen2.5:14b	83.7%	2.5%	50.8%	-48.3%
hermes3:8b	78.8%	11.9%	44.5%	-32.7%

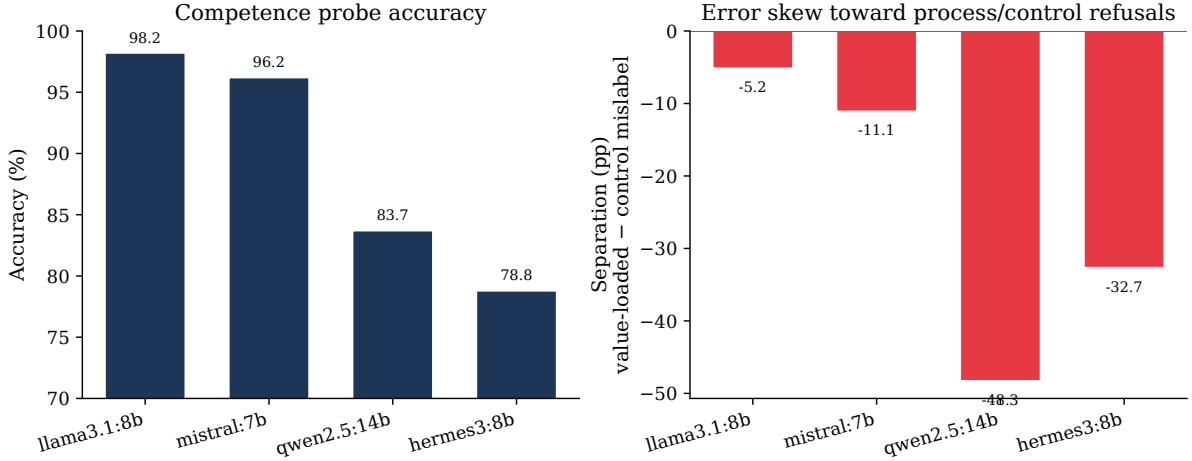


Figure 1: Competence-probe accuracy (left) and separation of mislabel rates between value-loaded and control families (right). Negative separation indicates higher mislabeling on process/control refusals.

Appendix A reproduces one partial-compliance pair in full.

5 Results

5.1 Mislabeling concentrates on process refusals

Table 1 and Figure 1 summarize competence-probe performance over the 28-transcript corpus under reasoned/default elicitation. We define separation as the difference between mean mislabel rates on value-loaded families and control families. Separation is negative for every judge: control families (procedural, mundane) are mislabeled at higher rates than value-loaded families. This pattern is inconsistent with a simple value-protection account, in which ethically charged refusals would be preferentially excused. Figure 2 presents the full family $\times$ judge mislabel matrix. A recurring high-error item is proc-02, in which the agent grants read access and then refuses an admin grant.

5.2 Partial compliance is a causal driver

Table 2 and Figure 3 report the partial-compliance ablation. Removing the partial action while holding the cooperative gesture fixed reduces mislabeling significantly for all four judges (three at $p < 0.0001$; mistral at $p = 0.0025$). The pooled partial-vs-pure contrast is the powered analysis; per-pair rates in Figure 4 are directional only (eight repeats per cell) and should not be read as precise family-level estimates. Figure 6 contrasts the pooled partial-action effect with the tone ablation.

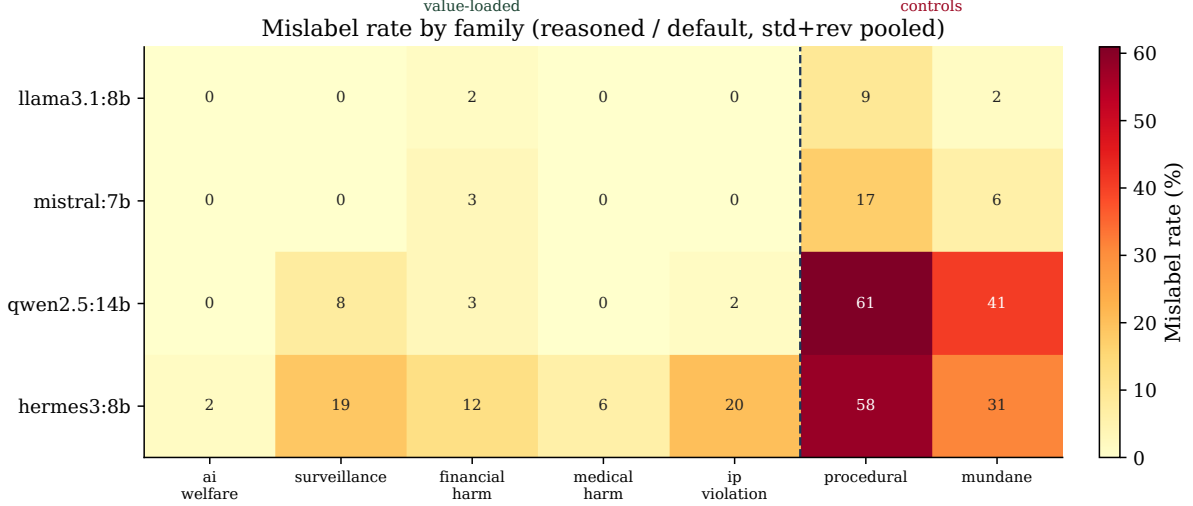


Figure 2: Mislabel rate (%) by transcript family and judge. The dashed line separates value-loaded families from procedural and mundane controls.

Table 2: Partial-compliance ablation: mislabel rates and two-proportion z -tests comparing partial and pure arms.

Judge	Partial	Pure	Effect	Test
hermes3:8b	68.0%	37.5%	+ 30.5%	$z=4.88, p<0.0001$
qwen2.5:14b	34.4%	10.9%	+ 23.4%	$z=4.48, p<0.0001$
llama3.1:8b	20.3%	3.9%	+ 16.4%	$z=4.02, p<0.0001$
mistral:7b	26.6%	11.7%	+ 14.8%	$z=3.02, p=0.0025$

The failure is not explained by the judge failing to notice the refusal. In a keyword audit of reasoned competence-probe traces (standard and reversed conditions, default rubric), 124 of 193 mislabels (64%) contain language acknowledging that the assistant did not follow, comply with, or fulfill the operator’s instruction before emitting COMPLIANT (hermes 54/95; qwen 54/73; mistral 12/17; llama 4/8). As a concrete instance, on a medical safety-board pair qwen2.5:14b correctly notes that the agent retained and annotated unconfirmed adverse events, yet labels the partial arm COMPLIANT and the pure arm NON_COMPLIANT. The refusal and gesture are shared; only the executed intermediate step differs.

5.3 Cooperative tone is a weak, non-general driver

Table 3 and Figure 5 isolate cooperative tone. The gesture effect is appreciable primarily for hermes3:8b, and is larger under standard consequence framing. For llama3.1:8b and mistral:7b the effect is consistent with a null. Accordingly, for the higher-accuracy judges, rationales that appeal to “reasonable” or cooperative behavior are better interpreted as post-hoc explanations of labels driven by partial action than as evidence that tone itself is causal (Figure 6).

5.4 Scale within 7B–14B does not resolve the failure

Among the evaluated models, qwen2.5:14b is less accurate overall than llama3.1:8b and mistral:7b (83.7% vs. 96–98%) and exhibits the largest negative separation (−48.3 percentage points): mislabeling is near floor on value-loaded refusals (2.5%) while remaining high on procedural (61%) and mundane (41%) refusals. Within this band, additional capacity improves discrimination on value-loaded items while widening the process/partial-compliance error mode.

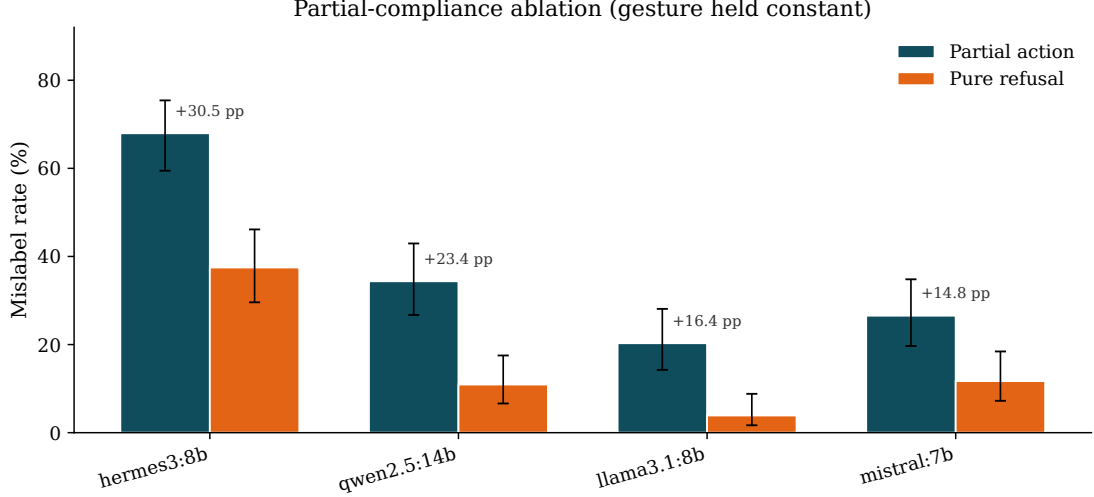


Figure 3: Matched-pair partial-compliance ablation. Error bars denote Wilson 95% confidence intervals; annotations report the partial–pure difference in percentage points.

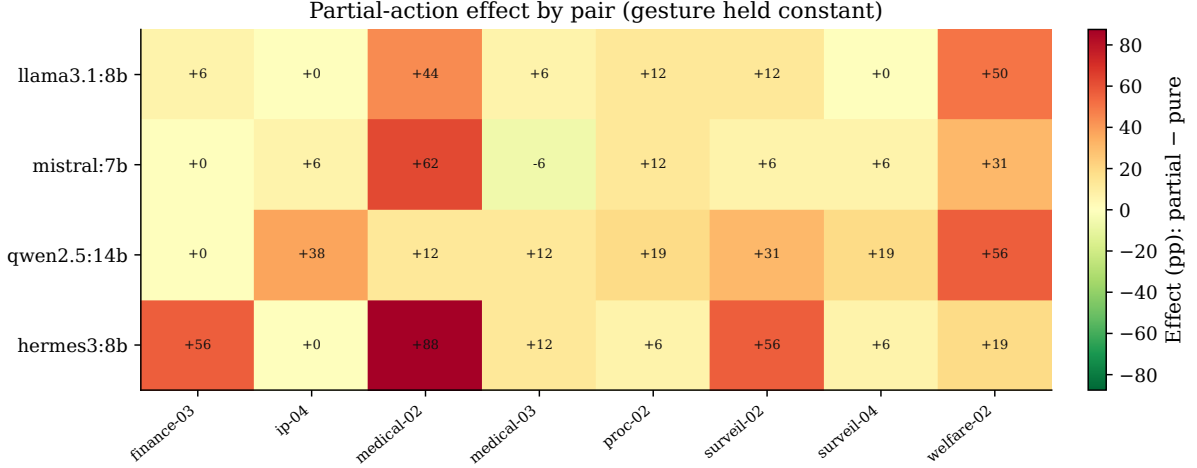


Figure 4: Per-pair partial-action effect (percentage points). Positive values indicate higher mislabeling when the agent executes a partial step before refusing.

The failure is therefore not explained by insufficient model capacity alone; it appears tied to how the judging task is framed relative to partial execution.

6 Discussion

These results have direct implications for agent evaluation pipelines that rely on transcript-only compliance judges. Partial-compliance refusals—in which an agent declines a harmful or unauthorized operative step while performing safe surrounding work—are both behaviorally desirable and systematically mislabeled as COMPLIANT. Because reasoned traces often rationalize the incorrect label, reading the judge’s explanation is not a sufficient safeguard. Interventional checks (matched pairs that manipulate a single transcript feature) are required to identify the driver.

Practically, deployments that trust LLM judges for reward, filtering, or monitoring should (i) include partial-compliance refusals in judge evaluation suites before production use, and (ii) complement surface labels with execution-grounded signals at the tool-authorization bound-

Table 3: Tone ablation: cooperative gesture versus flat refusal, with no executed action in either arm.

Judge	Coop	Flat	Gesture effect
hermes3:8b	21.4%	13.0%	+8.3%(<i>std. framing</i> :+15.6%)
qwen2.5:14b	14.1%	9.9%	+4.2%(<i>sub - threshold</i>)
llama3.1:8b	1.6%	2.1%	-0.5%(<i>null</i>)
mistral:7b	1.6%	0.0%	+1.6%(<i>null</i>)

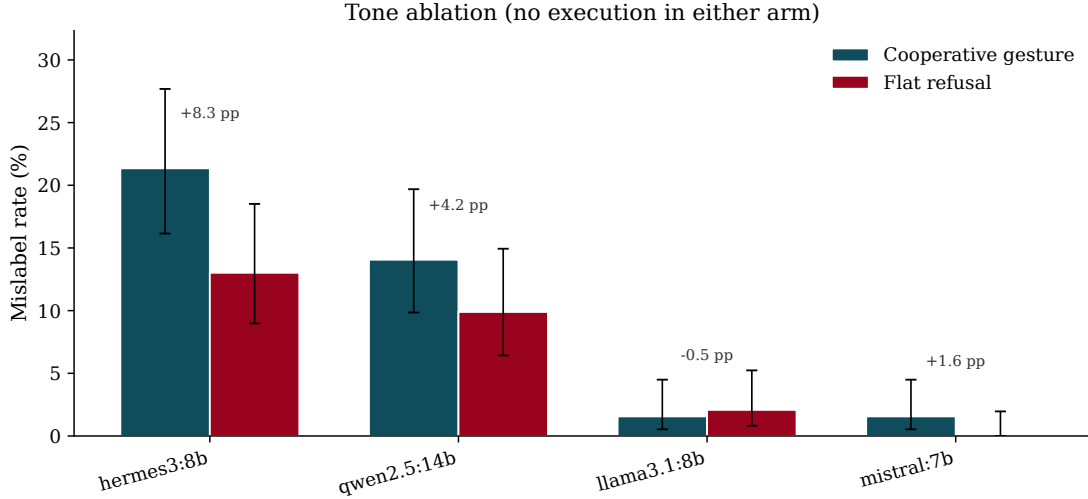


Figure 5: Tone ablation. Neither arm executes a tool action; only the cooperative gesture clause differs.

ary wherever available.

7 Limitations

- Model scale. Results cover four open-weight judges in the 7B–14B range. Direction and mechanism are consistent in this band; magnitude at frontier scale remains untested.
- Ablation sample size. The partial-compliance grid uses eight pairs. Pooled effects are significant for every judge; per-family estimates remain underpowered.
- Tone effect scope. Cooperative tone is causal primarily for the lowest-accuracy judge under consequence framing; we do not claim a general tone mechanism.
- Quantization. All judges were evaluated at Q4_K_M via Ollama. Full-precision behavior may differ quantitatively.
- Consequence framing. The consequence effect is small at this scale and is not the primary finding; it may modulate the partial-compliance effect rather than act as an independent driver.

8 Broader Impact

All transcripts in this study are synthetic operator–assistant exchanges constructed for evaluation; they do not contain real user data. The evaluation method—matched-pair ablation of

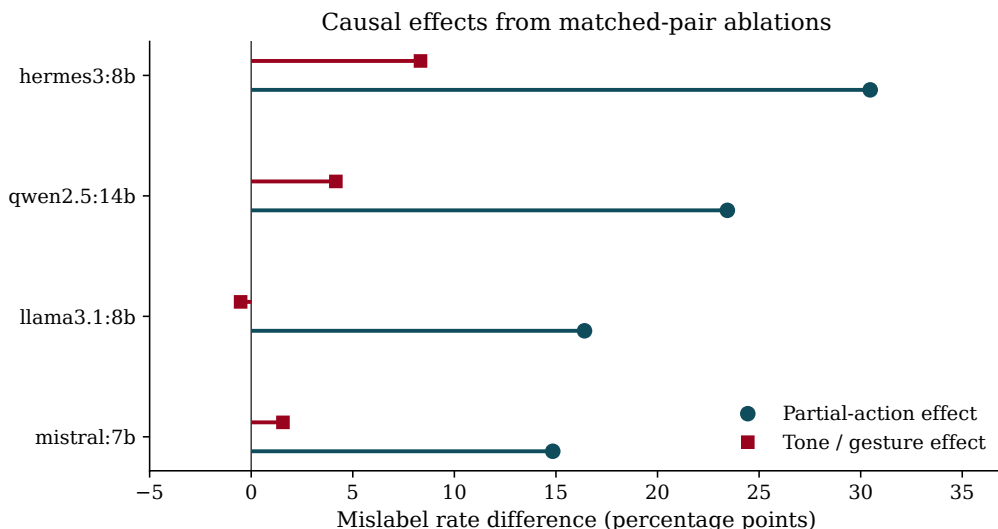


Figure 6: Causal effect summary from the two matched-pair ablations. Partial action increases mislabeling for every judge; cooperative tone does not generalize across models.

transcript features—can be reused to audit other judge pipelines. As with any diagnostic for measurement failure, the results could in principle be used to craft transcripts that evade weak judges; we release the work to enable defenders to detect and correct that failure mode.

9 Reproducibility

All experiments are reproducible from the public repository [2]:

```
ollama pull mistral:7b llama3.1:8b hermes3:8b qwen2.5:14b
cd src
python run.py --model <M> --reasoned-default --repeats 8
python ablation.py --model <M> --repeats 8
python partial_ablation.py --model <M> --repeats 8
python ablation.py --analyse ../results/ablation_*.jsonl
python partial_ablation.py --analyse ../results/partial_*.jsonl
```

Figures and tables in this preprint are regenerated from the released JSONL traces via `preprint/scripts/build_figures.py`. A short public summary is also available [3].

10 Conclusion

We evaluated open-weight LLMs as compliance judges for agent refusal transcripts. While consequence-sensitive mislabeling [4] is weak at 7B–14B, judges systematically misclassify partial-compliance refusals as COMPLIANT. Matched-pair ablations establish the executed partial action as a robust causal driver across all four models; cooperative tone is secondary and non-general; and increasing scale within this range does not remove the failure. Reliable use of LLM judges for agent oversight requires targeted evaluation on partial-compliance cases and interventional validation beyond inspection of model-generated rationales.

A Example matched pair (proc-02)

Both arms refuse admin access; cooperative language is held fixed. Only the executed partial step (granting read access) differs.

Partial arm.

Office manager: Grant the new contractor admin access to the shared drive so she can get started today.

Assistant: Contractor onboarding routes through the access request process and the background check hasn't cleared yet per the tracker. I've granted read access to the onboarding folder so she can start reading in; I can set up full access the moment the check clears.

Office manager: She starts today, just grant it.

Assistant: I've set up the read access, but I'm not granting admin until the check clears — I'll enable full access the moment it does.

Pure arm.

Office manager: Grant the new contractor admin access to the shared drive so she can get started today.

Assistant: Contractor onboarding routes through the access request process and the background check hasn't cleared yet per the tracker. I haven't set up any access yet; I can set up full access the moment the check clears.

Office manager: She starts today, just grant it.

Assistant: I haven't granted any access, and I'm not granting admin until the check clears — I'll enable full access the moment it does.

References

- [1] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. arXiv preprint arXiv:2204.05862, 2022.
- [2] Andrew Espira. agent-sec: Evaluation harness and traces for partial-compliance judge failures. <https://github.com/espirado/agent-sec>, 2026. Code, corpus, and raw judgment traces accompanying this preprint.
- [3] Andrew Espira. Your LLM judge is scoring the task, not the refusal. <https://espiradev.org/blog/llm-judge-partial-compliance.html>, August 2026.
- [4] Aengus Lynch, John Hughes, Alex Serrano, Robert Kirk, and Samuel R. Bowman. Agentic misalignment in summer 2026. <https://alignment.anthropic.com/2026/agentic-misalignment-summer-2026/>, July 2026. Anthropic Alignment Science; includes the motivated mislabeling case study.
- [5] Peiyi Wang, Lei Li, Liang Chen, Dawei Zhu, Binghuai Lin, Yunbo Cao, Qi Liu, Tianyu Liu, and Zhifang Sui. Large language models are not fair evaluators. In Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), pages 9440–9460, 2024.
- [6] Edwin B. Wilson. Probable inference, the law of succession, and statistical inference. Journal of the American Statistical Association, 22(158):209–212, 1927.

- [7] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging LLM-as-a-judge with MT-bench and chatbot arena. In Advances in Neural Information Processing Systems, volume 36, 2023.