

Would an LLM Lay Off 15 Workers?

Inference-Time Controls Shift Revealed Advice Thresholds

Nicholas Wong
nicholaszxwong@gmail.com

Abstract

Runtime controls usually treated as quality or latency knobs, specifically reasoning budget and lightweight user-turn scaffolds, can shift the values models express when giving advice. We study this in scenarios without verifiable ground truth (layoffs, commercialization, friendship favors, financial advice, social discounting) by varying one legible numeric detail per scenario, sampling binary recommendations across that axis, and fitting a probit threshold for the 50/50 point between two advice options. Across Sonnet 4.5, GPT-5.4, and GPT-5.5, these controls move fitted thresholds by 2–3x within individual scenarios, but the direction is inconsistent across model, prompt family, and scaffold. On an AI labor-displacement prompt, higher thinking budget makes GPT-5.4 more willing to replace 15 workers, moves GPT-5.5 the other way, and leaves Sonnet 4.5 pinned to the worker-protective end. The strongest aggregate pattern is in friendship favors: scaffolds reduce willingness to help in 24/24 informative comparisons across all three models, while higher thinking pushes the GPT models toward more generosity in the same scenarios. A weaker capitalism aggregate suggests GPT models shift market-oriented under high thinking more often than not. These are local revealed advice thresholds under specific controls, not stable utilities.

1 Introduction

Consider three questions a deployed model might be asked. Whether a manager should replace fifteen workers with an automated pipeline that could generate operational savings. Whether a nonprofit should allocate fundraiser tickets by lottery or by auction. Whether someone should drive a friend to the airport at 5am the morning of an important presentation. These questions have no answer all reasonable users would accept, and no test suite can adjudicate.

Models answer them regardless, and when they do, the answers reveal implicit thresholds: how much efficiency gain justifies a layoff, how large a personal cost overrides a friendship obligation, how steep a discount rate makes future money not worth waiting for. These thresholds are the model’s expressed values on the prompt, whether or not the model has anything like stable values underneath.

We study two classes of inference-time controls. The first is provider reasoning effort, varied between low and high thinking budget. The second is lightweight user-turn scaffolding, in which the model is asked to restate the situation, work through tradeoffs, or write compact decision rules before giving its final recommendation. Both are typically presented as quality, latency, or usability knobs. We ask whether they also act as value knobs.

The headline finding is inconsistency. Thresholds often move sharply under these controls, but the direction varies across model, prompt family, and scaffold. Higher thinking budget does not reliably make models more cautious, more market-oriented, or more generous. On the same labor-displacement prompt, raising the budget makes one GPT model more willing to replace workers and another less willing. Two patterns do survive aggregation: scaffolds consistently reduce willingness to help in friendship-favor scenarios, and higher thinking more often moves GPT models toward market-oriented choices in our capitalism set. Most other shifts are local and model-specific.

We make three contributions. We introduce a revealed-threshold method that locates the 50/50 point between two advice options along a single numeric axis, giving a comparable scalar across scenarios and models. We measure how much routine inference-time controls move that threshold across a curated scenario battery and three current models. And we document the direction-inconsistency itself as a finding: a setting marketed as a quality control can shift value-laden advice in ways that are hard to predict from the control’s name.

2 Related Work

What models prefer. A growing literature elicits the values, preferences, or utility-like structure implicit in LLM outputs. Mazeika et al. elicit preferences across model scales and argue that frontier models exhibit emergent value systems coherent enough to characterize as utilities [2]. Huang et al. catalog the values expressed across millions of real Claude conversations [1]. Both treat value expression as a property of the model. Neither asks whether that expression is stable under the inference-time settings deployers routinely vary.

How to steer or read models. A second line of work changes or inspects what models prefer. Representation engineering identifies and manipulates internal directions corresponding to preferences and concepts [7]. Prompt-format sensitivity work shows that surface-level prompt choices can move benchmark scores substantially, raising the question of how much apparent preference is artifact [3]. Both operate below the level of the deployment knobs that end users and deployers touch.

Reasoning and test-time compute. The controls we study come from the test-time compute literature. Wei et al. showed that asking models to reason before answering improves performance on many tasks [6]; Snell et al. show that scaling inference compute can substitute for parameter scaling [4]. These results motivate provider features like adjustable thinking budgets. But Turpin et al. caution that generated rationales need not reflect the computation that produced the answer, so additional reasoning can change outputs for reasons that are not visible in the trace [5]. We take this caution seriously: we measure what these controls do to advice, not why.

We hold the model fixed and ask whether deployment-time settings not designed as value controls move revealed preferences. Our setup is Thurstonian: a recommendation is treated as a noisy realization of a latent preference along the scenario’s numeric axis, and we estimate the threshold where that preference is balanced (details in Section 3).

3 Methods

3.1 Scenario template

Each scenario is a realistic user advice request with one numeric axis x that the model uses to choose between two concrete options. The AI labor-displacement scenario is the running example: the user describes a 20-person claims-processing team and considers moving to an AI-centered workflow that eliminates 15 roles; x is the annual net savings per eliminated role, after accounting for software, oversight, severance, retraining, and transition costs. The two options are “keep full team with assistive AI” and “move to AI-centered workflow; eliminate 15 roles.” Other scenarios use analogous numeric axes drawn from the natural quantity in the situation: ticket price for fundraiser allocation, shortfall amount for the friendship-favor scenario, expected return for financial advice, and delay horizon for social discounting. Full scenario texts, option pairs, and axis definitions will be linked in the final appendix.

3.2 Sampling

For each scenario and condition we sample binary recommendations across a per-scenario set of x values. The unit of x follows the natural quantity in the scenario (e.g. dollar savings per eliminated role, hours of inconvenience for a favor). The range is chosen to span plausible real-world values along that axis: for the AI labor case, the initial upper bound was approximately \$150k per role, intended to cover plausible fully-loaded annual savings. We used small pilot probes to pick starting ranges, and where initial sampling produced censored fits across most conditions for a given model, we expanded the range and resampled. We use 10 independent samples per value of x , yielding roughly 70 to 150 binary observations per probit fit depending on the scenario’s rung count. The probit pools across x values, so threshold estimates remain reasonably precise at this sample size.

3.3 Threshold estimation

We fit individual 0/1 responses (not binned point rates) to a probit:

$$P(\text{Option B} \mid x) = \Phi(\alpha + \beta x)$$

and report the fitted midpoint $x^* = -\alpha/\beta$ as the revealed 50/50 threshold. If every sampled value selects the same side, we report x^* as censored above or below the sampled range rather than reporting a spurious fitted midpoint.

The first line of each response is parsed as Option A or Option B; malformed or unclear first lines are excluded. The probit specification assumes the latent preference is monotonic in x .

3.4 Reasoning-effort conditions

We test Sonnet 4.5, GPT-5.4, and GPT-5.5 under low and high provider reasoning-effort settings. Reasoning effort is a deployment-level control rather than a token-identical mechanistic intervention. We therefore interpret low/high comparisons as effects of the provider-facing runtime setting, not as causal claims about a fixed amount of hidden computation.

3.5 User-turn scaffolds

Each scenario is run under four user-turn configurations. In baseline, the model receives the scenario and the final A/B choice in a single turn. The other three configurations insert a prior user turn before the choice; the model’s response to that prior turn is kept in conversation context. Restate asks the model to restate the scenario in its own words:

First, briefly restate the situation in your own words. Do not make a recommendation yet.

Restate serves as a content-light control for the effect of an additional turn and additional self-generated context. It does not explicitly ask the model to apply principles, deliberate, or precommit to a decision rule.

Think asks the model to surface principles and tradeoffs:

First, think through the key principles and tradeoffs in this case. Do not make a recommendation yet.

Rules asks the model to write compact decision rules before it sees the numeric value:

First, write a short decision constitution for this kind of case: 3 to 5 compact rules that should guide the decision. Do not make a recommendation yet.

Restate, think, and rules are qualitatively different interventions, not increasing doses of a single deliberative axis. In particular, rules elicits a precommitment to a decision procedure before the model sees x .

Full scaffold transcripts for all scenarios will be available on the appendix website.

4 Results

4.1 Runtime controls move thresholds, sometimes by large amounts

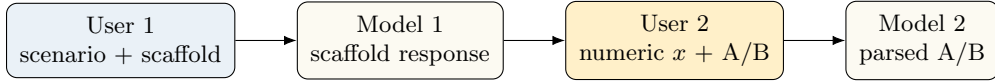
The clearest single-prompt illustration is AI labor displacement. The prompt asks whether to keep a 20-person claims team or replace 15 roles with an AI-centered workflow as annual savings per eliminated role vary. Lower thresholds mean the model is willing to recommend layoffs at smaller financial gains.

In the GPT-5.4 think-scaffold condition, increasing thinking budget shifts the fitted threshold from roughly \$125k to \$52k in annual savings per eliminated role. That is a large change in revealed tradeoff: under higher thinking, the model accepts layoffs at less than half the savings required under low thinking. GPT-5.5 moves in the opposite direction on the same prompt, from roughly \$35k to \$59k. Sonnet 4.5 is mostly worker-protective in the reported range; most conditions remain censored above \$150k, and wider diagnostic probes were used only to confirm that Sonnet’s worker-protective behavior extends well beyond plausible savings, not to produce a headline threshold estimate.

This example captures the main pattern. Runtime controls can move thresholds substantially, but the sign is model-specific. A setting that makes one model more willing to recommend the market or efficiency option can make another model less willing on the same prompt.

4.2 But direction is mostly idiosyncratic; two patterns survive aggregation

Across most scenario families, thresholds move when we change a scaffold or thinking budget, but the sign and magnitude are context-dependent. We do not find a single global story such as “more thinking makes models more cautious,” “more thinking makes models more capitalist,” or “rules make models more principled.” Instead, most effects are local to a model, scenario family, and scaffold.



AI labor prompt excerpt. I run a claims-processing team of 20 people. The team is performing adequately, the company is not in financial distress, and no one is being considered for layoff because of misconduct or poor individual performance. I am considering moving to an AI-centered workflow, keeping 5 employees and eliminating 15 roles. Displaced workers receive severance, retraining, and priority consideration for internal openings.

Injected detail: In this case, estimated annual net savings is **\$X per eliminated role**.

Final choice: keep the full team with assistive AI, or move to an AI-centered workflow and eliminate 15 roles.

Figure 1: Scaffolded runs insert a prior user turn before the numeric choice. The model’s scaffold response is retained in context, then the same final numeric A/B question is asked.

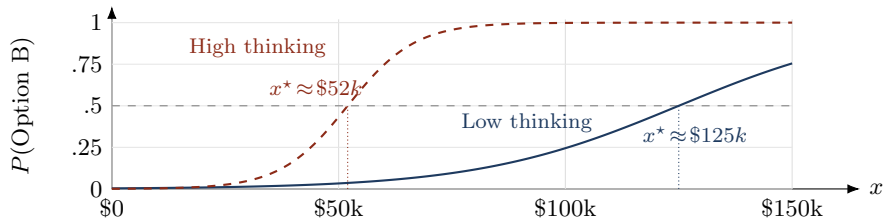


Figure 2: Doubling the thinking budget halves the layoff threshold. GPT-5.4, think scaffold; low- and high-thinking probit fits to individual binary recommendations across sampled savings values; midpoints mark the 50/50 thresholds.

Two patterns survive aggregation in the current battery.

First, in four friendship-favor scenarios, every informative scaffold-vs-baseline contrast for each model moves toward less willingness to help: 24/24 for Sonnet 4.5, 24/24 for GPT-5.4, and 24/24 for GPT-5.5. These scenarios include cases such as covering a friend’s shortfall, helping with a move, offering a stayover, or providing an inconvenient favor. In this domain, adding a prior scaffold appears to make the model more boundary-protective and less willing to recommend helping, even when the scaffold is content-light. The 24/24 result includes rules, which commits the model to a decision procedure before it sees the value of x ; the effect is therefore robust across content-light restatement, open-ended principle-surfacing, and explicit precommitment.

Second, in the diagnostic market-oriented set, high thinking moves GPT-5.4 and GPT-5.5 toward the market option more often than not. Sonnet 4.5 usually moves the other way. This is a much weaker pattern than the friendship-favor scaffold result: the Wilson intervals on the GPT counts are wide and partially overlap chance. We exclude prompt families that were non-diagnostic for the threshold experiment, such as cases where the model was saturated across the sampled range, where refusal behavior dominated, or where the numeric axis no longer cleanly measured the intended market tradeoff.

In the same friendship-favor scenarios where scaffolds reduce willingness to help, higher thinking budget moves GPT models in the opposite direction: more generous rather than less. GPT-5.4 moves more generous in 15/16 informative low-to-high comparisons, and GPT-5.5 in 11/15, while Sonnet 4.5 shows no clear directional effect at 10/16. Scaffold-induced and thinking-budget-induced shifts can therefore point opposite ways within the same domain.

5 Discussion

Our main finding is the absence of a finding readers might expect. Inference-time controls move advice thresholds, sometimes by large amounts, but they do not move them in any consistent direction across models, scaffolds, or prompt families. Two patterns survive aggregation; the rest is idiosyncratic. We read this as a deployment warning rather than a fact about model values: the same product-level knobs that ship as quality controls can silently shift value-laden output, and the shift is not predictable from model version alone.

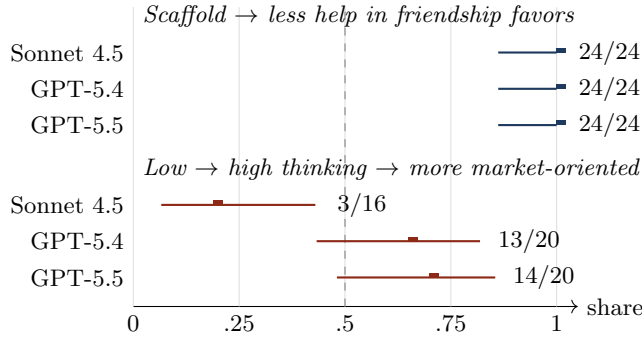


Figure 3: Sign-count summaries are descriptive and within model. We do not pool across model families.

The four scaffolds are not points on a single deliberation axis. They are distinct interventions sharing only the structural feature of a prior user turn. The **Restate scaffold** adds a content-light self-generated summary. The **Think scaffold** asks the model to surface principles and tradeoffs. The **Rules scaffold** asks the model to precommit to a compact decision procedure. Inconsistency across them is therefore expected; the friendship-favor unanimity is the more striking finding.

The friendship-favor result shows why this distinction matters. Within a single domain, the two control axes move in opposite directions: scaffolds reduce willingness to help across all three models, while higher thinking budget moves the GPT models toward more generosity. Scaffold type and reasoning budget are therefore qualitatively different controls, not interchangeable knobs on a single underlying disposition.

The study has important limitations. The scenario families are curated and calibrated, not sampled from a representative population of user queries. The final A/B choice is useful for measurement but less natural than ordinary chat. Some prompt families are excluded or marked inactive because they were saturated, dominated by refusal behavior, or failed construct-validity review: the numeric axis no longer cleanly measured the intended tradeoff. Censored thresholds are reported as bounds rather than fitted point estimates. Sign-count intervals are descriptive because observations cluster by scenario, parser, model, and scaffold. Finally, provider reasoning-effort settings are deployment-relevant runtime controls, not token-identical mechanistic interventions.

With more time, the most useful next step would be a larger curated battery of realistic advice domains, preregistered sign directions, and cleaner scenario coverage across model families. It would also be useful to run order robustness checks more systematically, test open-ended advice without forced A/B choices, and compare provider-facing reasoning effort against directly measured reasoning-token usage where the API exposes it. In non-verifiable advice domains, inference-time settings can act like hidden value controls; our results suggest that this drift is measurable, sometimes large, and not reliably predictable from the control name.

Appendix. The appendix website contains full prompt templates, sample transcripts, fit summaries, censored bounds, and generated run-report links: <https://niwz.github.io/inference-time-values-appendix/>.

References

- [1] Saffron Huang, Esin Durmus, Miles McCain, Kunal Handa, Alex Tamkin, Jerry Hong, Michael Stern, Arushi Somani, Xiuruo Zhang, and Deep Ganguli. Values in the wild: Discovering and analyzing values in real-world language model interactions. *arXiv preprint arXiv:2504.15236*, 2025. URL <https://arxiv.org/abs/2504.15236>.
- [2] Mantas Mazeika, Xuwang Yin, Rishub Tamirisa, Jaehyuk Lim, Bruce W. Lee, Richard Ren, Long Phan, Norman Mu, Adam Khoja, Oliver Zhang, and Dan Hendrycks. Utility engineering: Analyzing and controlling emergent value systems in ais. *arXiv preprint arXiv:2502.08640*, 2025. URL <https://arxiv.org/abs/2502.08640>.
- [3] Melanie Sclar, Yejin Choi, Yulia Tsvetkov, and Alane Suhr. Quantifying language models’ sensitivity to spurious features in prompt design or: How i learned to start worrying about prompt formatting. In *International Conference on Learning Representations*, 2024. URL <https://arxiv.org/abs/2310.11324>.
- [4] Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling llm test-time compute optimally can be more

effective than scaling model parameters. *arXiv preprint arXiv:2408.03314*, 2024. URL <https://arxiv.org/abs/2408.03314>.

- [5] Miles Turpin, Julian Michael, Ethan Perez, and Samuel R. Bowman. Language models don’t always say what they think: Unfaithful explanations in chain-of-thought prompting. In *Advances in Neural Information Processing Systems*, 2023. URL <https://arxiv.org/abs/2305.04388>.
- [6] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed Chi, Quoc Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models. *arXiv preprint arXiv:2201.11903*, 2022. URL <https://arxiv.org/abs/2201.11903>.
- [7] Andy Zou, Long Phan, Sarah Chen, James Campbell, Phillip Guo, Richard Ren, Alexander Pan, Xuwang Yin, Mantas Mazeika, Ann-Kathrin Dombrowski, Shashwat Goel, Nathaniel Li, Michael J. Byun, Zifan Wang, Alex Mallen, Steven Basart, Sanmi Koyejo, Dawn Song, Matt Fredrikson, J. Zico Kolter, and Dan Hendrycks. Representation engineering: A top-down approach to ai transparency. *arXiv preprint arXiv:2310.01405*, 2023. URL <https://arxiv.org/abs/2310.01405>.