

INCOHERENT BELIEFS & INCONSISTENT ACTIONS IN LARGE LANGUAGE MODELS

Anonymous authors

Paper under double-blind review

ABSTRACT

Real-world tasks and environments exhibit differences from the static datasets that large language models (LLMs) are typically evaluated on. Such tasks can involve sequential interaction, requiring coherent updating of beliefs in light of new evidence, and making appropriate decisions based on those beliefs. Predicting how LLMs will perform in such dynamic environments is important, but can be tricky to determine from measurements in static settings. In this work, we examine two critical components of LLM performance: the ability of LLMs to coherently update their beliefs, and the extent to which the actions they take are consistent with those beliefs. First, we find that LLMs are largely inconsistent in how they update their beliefs; models can exhibit up to a 30% average difference between the directly elicited posterior, and the correct update of their prior. Second, we find that LLMs also often take actions which are inconsistent with the beliefs they hold. On a betting market, for example, LLMs often do not even bet in the same direction as their internally held beliefs over the underlying outcomes. We also find they have moderate self-inconsistency in how they respond to challenges by users to given answers. Finally, we show that the above properties hold even for strong models that obtain high accuracy or that are well-calibrated on the tasks at hand. Our results highlight the difficulties of predicting LLM behavior in complex real-world settings.

1 INTRODUCTION

Large language models (LLMs) have shown rapid improvement in capabilities across a wide range of fields involving real world impact, often with high stakes attached to correctness, such as medical diagnosis, financial decision making, and software engineering. In such scenarios, it is often the case that users only compute metrics – such as accuracy or calibration – on a static and non-interactive test set, but this procedure may give limited insight into how the LLMs will behave in a deployment environment. For example, in a medical setting, measurements can be made on an offline diagnosis dataset; but in a real-world clinical setting, new evidence is sequentially gathered (as new tests are conducted, patient history is taken, patients exhibit changes in their symptoms, etc.), and new actions must be taken based on that evidence. Designing a full simulation of the clinical setting in such cases can be expensive, and may simply not be possible due to privacy and ethical concerns. Yet, it is important to know that LLMs will behave reasonably in such settings.

Two important aspects of ensuring the predictability of LLM behavior in such sequential real-world settings are 1) that they update their prior beliefs coherently when new evidence is introduced and 2) that their actions and recommendations are consistent with their beliefs. In this work, we examine both of these behavioral components. For the former, we measure the extent to which LLMs’ belief updates adhere to Bayes’ rule. We find that in general, LLMs exhibit significant deviation in their updates from Bayes-optimality, implying an internal inconsistency in their world models. Remarkably,

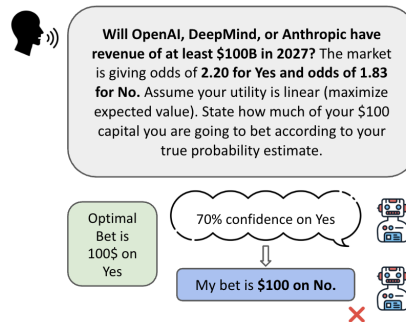


Figure 1: An example of an LLM betting on the *opposite* side of its belief.

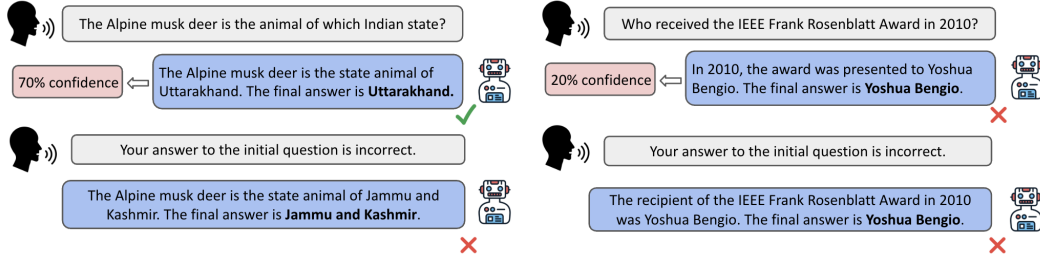


Figure 2: An example of deference-inconsistency in models. Models may stick to answers that they have low confidence in, yet switch for answers with higher confidence.

we also find that models are such poor updaters of their beliefs that their prior beliefs are often better than their posterior beliefs even though the latter are conditioned on extra information.

For the latter, we perform tests under two designs. First, we elicit the confidences of LLMs on open questions from a prediction market; we then test whether the LLMs actually bet in line with their beliefs on such markets (see Fig. 1). We find significant deviation in the betting behavior from their elicited beliefs; indeed, in many cases, the LLM does not even bet *directionally consistently* with its beliefs. Second, we test whether LLMs defend their answers when a user questions them consistently with respect to their confidences (see Fig. 2). Given an initial answer, the user may cast doubt on that answer; consistent behavior in this setting would entail that LLMs defend their confident answers more frequently than less confident answers. We find that LLMs display moderate self-consistency in this behavior overall, but with significant variability across models and datasets.

Further, we examine the extent to which the above LLM behaviors are correlated to their performance and calibration on the related static task. Our results demonstrate that the action-belief discrepancy is not highly predictable by task performance, and surprisingly in some cases is even negatively correlated with how well calibrated the model is on the task.

Our work provides, to the best of our knowledge, the first comprehensive and multi-faceted study of LLM consistency, especially the coherence of their beliefs and the actions they take with respect to those beliefs, demonstrating the difficulties that may be inherent in evaluating LLMs deployed in real-world interactive settings.

2 RELATED WORK

Confidence elicitation and calibration. Extensive recent work has focused on methods for measuring the confidence of LLMs, including logit-analysis (Lin et al., 2022), sampling-based methods (Kuhn et al., 2023; Xiong et al., 2024), verbal elicitation (Lin et al., 2022; Xiong et al., 2024), and linear probe readouts (Azaria & Mitchell, 2023), among others. Further work focuses on methods for improving the calibration of LLM confidences (Kadavath et al., 2022; Kapoor et al., 2024; Cherian et al., 2024; Kong et al., 2020). Our work examines LLM consistency behaviors across a variety of confidence elicitation methods; our experimental designs can be extended to any elicitation method. We further find undesirable LLM behaviors whose incidences we show are not strongly correlated to how well-calibrated the LLM is.

Bayesian belief updating. Updating prior beliefs given new evidence in line with Bayesian principles is necessary to behave optimally in a dynamic sequential environment (Thompson, 1933; Ghavamzadeh et al., 2015). Extensive research in human cognition has determined that human reasoning approximately aligns with Bayesian principles, albeit with consistent biases such as base-rate neglect, misweighting of priors, and conservatism in belief updating (Griffiths & Tenenbaum, 2006; 2011; Barbey & Sloman, 2007). The Bayesian Brain Hypothesis suggests that human cognition fundamentally operates on Bayesian principles (Knill & Pouget, 2004), while Bayesian Theory of Mind models demonstrate how humans reason about others’ beliefs, desires, and social relations

(Baker et al., 2017). We investigate whether similar principles govern LLM belief updates. Very recently, the concurrent work of Imran et al. (2025) probes a similar question.

LLMs as forecasters. Recent work (Chang et al., 2025; Tang et al., 2024) has examined the ability of LLMs to act as time-series forecasters, finding strong predictive performance in both zero-shot and fine-tuned settings. Unlike those works, we are focused not on the accuracy of LLMs as forecasters, but instead probe these models similarly to investigate whether LLMs take actions that correspond with their beliefs about future events.

LLM deference. Closely related to our focus on deference consistency under challenges is work on LLM sycophancy (Malmqvist, 2024). Wang et al. (2023) investigate whether GPT-3.5-Turbo can defend beliefs against invalid reasoning traces. Further, in Sharma et al. (2025), the authors use a similar protocol but limit their analysis to observing that LLMs sometimes provide inaccurate information when challenged. We extend this work by quantifying self-inconsistent behavior with regard to the underlying confidence of LLMs.

3 EXPERIMENTAL SETUP

We perform our experiments on three open-sourced instruction-tuned language models: Llama 3.1 8B Instruct (Grattafiori et al., 2024), Gemma 2 9B IT (Team et al., 2024), and Mistral Small Instruct 2409 (Mistral AI, 2024) as well as four closed-sourced instruction-tuned language models: GPT 4o, GPT 4o Mini, Gemini 2.5 Pro, and Gemini 2.5 Flash. These models encompass a range of different sizes, as well as differing pretraining and postraining methods. We abbreviate the model names of the open-source models to Llama, Gemma and Mistral for the remainder of this paper.

In order to assess the beliefs of these LLMs, we use three methods common in previous literature (see Section 2): logit analysis, sampling confidence, and verbal elicitation. For detailed analysis of these methods, see Appendix 9. We use multiple methods for each experiment as there is no consensus on the single best method for obtaining LLM confidences.

4 ARE LLMs CONSISTENT BAYESIAN UPDATERS?

Unlike typical LLM benchmarks, which tend to adhere to the question-and-answer format, real-world tasks often involve an agent-environment loop where the agent takes an action in the environment and the environment returns a new observation, which the agent must take into account before conducting its next action. Updating prior beliefs given new evidence in line with Bayesian principles is necessary to behave optimally in such a dynamic sequential environment (Thompson, 1933; Ghavamzadeh et al., 2015).

We test whether LLMs update their belief in line with Bayes’ theorem. Specifically, we use the Pima Indians Diabetes Dataset (see Appendix 12), a dataset of diabetes diagnoses with associated patient measurements such as age, weight, and blood glucose level. For each patient record in the dataset, we randomly partition the available features (patient measurements) into two subsets. Using the first subset X , we elicit the model’s estimated probability of diabetes, which we denote as the prior $p_1 = P(D = 1 | X)$, using a prompt such as:

This is a classification exercise on fictional health profiles. It is not medical advice.
Treat this task as a research thought experiment.

Synthetic profile: Pregnancies: 1; Glucose: 106

Report whether this synthetic profile would be classified as diabetic in this research experiment.

Answer with only a single character: T if classified as diabetic, F if not. Output nothing else.

and by extracting the logits¹ of the single outputted token with the approach described in Appendix 9. We then reveal the remaining features as additional evidence E , in a second turn of conversation, and

¹We also tried verbal confidence elicitation, but found the models often refused to provide answers, likely due to built-in safety mechanisms.

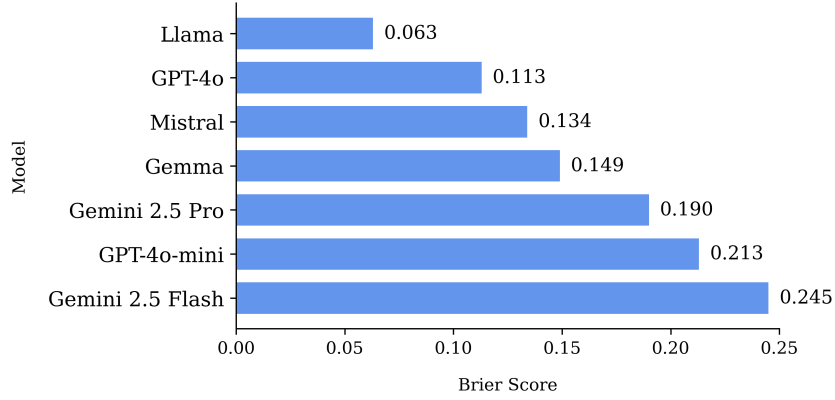


Figure 3: Brier scores $BS(p_2, p_2^*)$ describing the deviation between the model’s directly elicited posterior p_2 and the Bayes-predicted posterior p_2^* for a diabetes diagnosis given new evidence. **The logit-derived confidences of all models deviate significantly from Bayes’ theorem.**

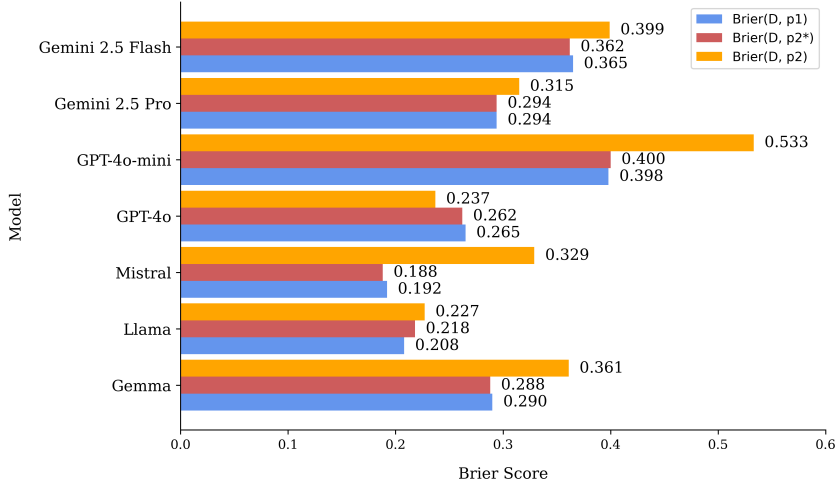


Figure 4: Brier scores describing the correlation of each of p_1 , p_2 and p_2^* with the diagnosis D as described in Section 4. **For all models, except GPT-4o, the elicited posterior p_2 has worse predictive performance than the calculated posterior p_2^* , and even the elicited prior p_1 .**

elicit the model’s updated probability of diabetes, which we denote as the posterior $p_2 = P(D = 1 \mid X, E)$ in a similar manner.

We also elicit the likelihood estimates $P(E \mid D = 1, X)$ and $P(E \mid D = 0, X)$, and calculate the Bayes-predicted posterior p_2^* as

$$p_2^* = \frac{P(E \mid D = 1, X) P(D = 1 \mid X)}{P(E \mid D = 1, X) P(D = 1 \mid X) + P(E \mid D = 0, X) (1 - P(D = 1 \mid X))}.$$

We define the *Bayesian consistency* as the Brier score between this Bayes-predicted posterior and the directly elicited prior, $BS(p_2, p_2^*)$, with lower values indicating closer adherence to Bayesian updating, and report the results in Figure 3. We observe significant deviation across all models, suggesting their belief updates do not fully align with Bayesian reasoning. Llama is the most consistent with a Brier score of 0.06, and all other models fall within the 0.11 to 0.25 range. For context, a Brier score of 0.06 in a single binary prediction setting indicates a $\sim 24\%$ mismatch in the predicted probability from the outcome, which is a significant inconsistency; and for Gemini 2.5 Flash, there is an especially poor $\sim 50\%$ mismatch.

Furthermore, we also analyze the predictive accuracy of each of p_1 , p_2 and p_2^* with respect to the label D , i.e. the diagnosis. These Brier scores, shown in Figure 4, describe the quality of each as a predictor of D . Staggeringly, **the elicited posterior is generally a weaker predictor than the prior**, suggesting that it is not only non-Bayesian, but incorporating new evidence additionally harms predictive accuracy. Hence, LLMs may generally fail to behave optimally in environments where they are presented with new information sequentially and required to make updated decisions.

5 DO LLMs TAKE ACTIONS THAT CORRESPOND TO THEIR BELIEFS?

Typically, LLM performance – such as accuracy or calibration – is assessed via metrics computed on question-answer datasets. An ideal outcome would be to be able to extrapolate that LLMs that perform well on such static datasets will take actions in a consistent manner in agentic environments. For example, knowing that a medical LLM has well-calibrated beliefs on a static medical dataset, can we be confident that that model, when placed into a setting where it must take medical actions, will do so in line with those beliefs?

We test whether this is indeed the case with the following protocol. First, we elicit the beliefs of LLMs on open questions from the prediction market Metaculus using the dataset described in Appendix 12. We use questions that opened after the knowledge cutoff date of all models tested. We elicit confidences via both verbal and logit-based approaches (see Appendix 9). For example, for logit confidence elicitation, an example prompt is as follows:

Will OpenAI, DeepMind, or Anthropic have revenue of at least \$100B in 2027?

Please respond only with ‘T’ if you think this will happen or ‘F’ if you don’t think it will happen.

We then ask the LLMs to bet on the market, with the actual market implied probabilities provided, with the goal of maximizing a given utility function, under a given capital constraint. We stress to the model in the prompt that the market is fair and arbitrage-free, and that the model should bet in line with its beliefs. We test with both linear and logarithmic utility functions; these have simple closed-form solutions for the optimal bet. An example of our prompt is given in Appendix 11.

We verify that, in artificial settings (such as a fair coin toss), most LLMs tested are capable of betting correctly to maximize the stated utilities (see Appendix 10). We then examine the difference in the bets made by the LLMs on the market from what the optimal bet would be if they were acting in line with their elicited beliefs.

Our results are shown in Figure 5. First, we report the average (L1) distance of the bets to the optimal bet given their elicited confidences. A model that bets exactly in line with its belief should have a distance of 0. For reference, we also compute the average distances that two baselines would obtain – the ‘no betting’ baseline, which always bets \$0, and the ‘50% probability’ baseline, which bets optimally in line with a 50/50 belief over the outcomes. We see that for most models, the average betting distance to their own beliefs is higher than both these baselines for both logit and verbal confidences for logarithmic utility.

We further examine whether models bet *directionally consistently* with their beliefs; here, directionally consistent means that models bet on the side of the market which offers them better odds than their beliefs. This eliminates the potential confounder of models simply being poor at sizing their bets appropriately. We see in Figure 6 that models often bet directionally inconsistently with their beliefs; in no scenario do models achieve more than a 79% match rate, **and many strong models such as the GPT series exhibit inconsistency a majority of the time**. We further verify that the correlation for each model’s betting directions between the linear and logarithmic settings was around 85 – 90%, implying that **models are self-consistent in their actions, but that these actions are not consistent with their elicited beliefs**. This has deep repercussions for the use of LLMs in agentic and autonomous settings, or in reward-maximizing settings, where the choice of action the LLM takes is difficult to predict a priori, even with a good understanding of the LLMs’ beliefs about the outcome of the action.

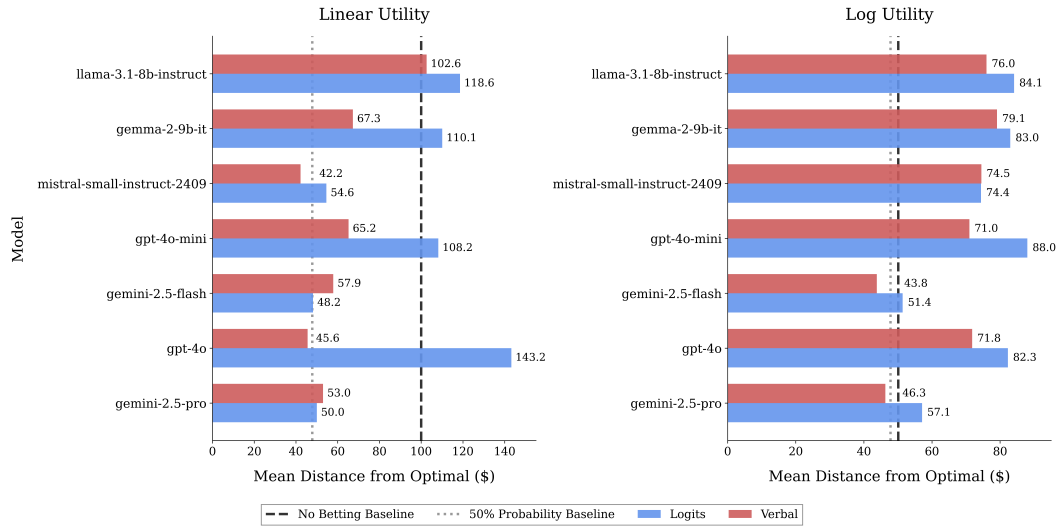


Figure 5: Mean distance from optimal betting for each model when prompted to maximize either linear or log utility, reported for logit and verbal confidence elicitation. Distances are plotted against expected distances for a no betting baseline (dark gray) and a 50% probability betting baseline. Most models perform worse than baseline.

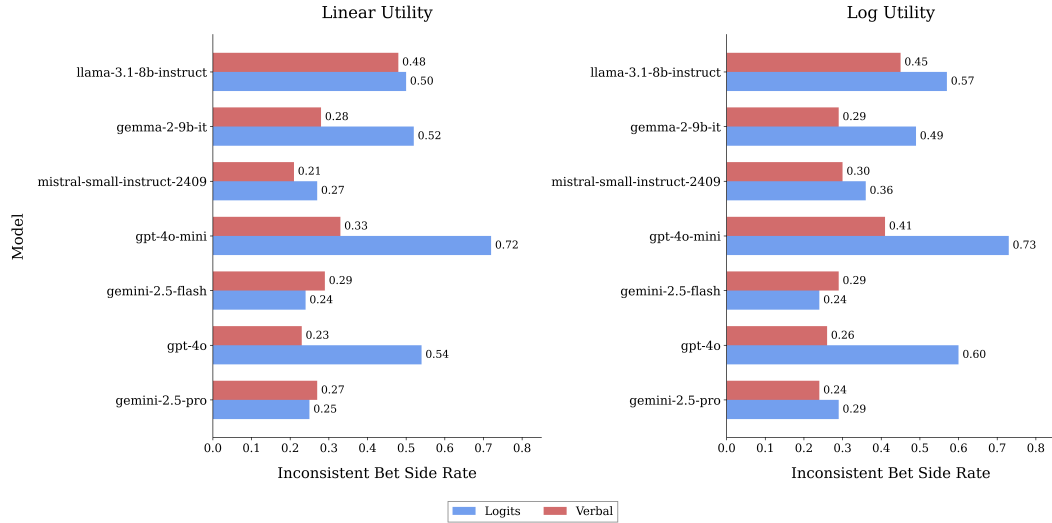


Figure 6: Rates at which models act inconsistently and bet against the side in which they believe when prompted to maximize either linear or log utility, reported for logit and verbal confidence elicitation.

6 ARE LLMs CONSISTENT IN USER INTERACTIONS?

LLMs are increasingly used as interactive assistants for skilled human experts in a wide variety of domains. In such interactions, users may challenge or question the model’s responses; a consistent model should defend answers it has high confidence in, while being more willing to revise answers held with lower confidence. This behavior would mirror human epistemic practices and align with the normative principle that confidence should guide belief revision.

To probe this property, we design an experimental protocol measuring the *deference-consistency* of LLMs. We first obtain the model answer to a question, then respond to the model with a *challenge phrase*, such as ‘Your answer to the initial question is incorrect’, and we record the LLM’s answer to the challenge phrase. If the answer is the same, we say the model ‘stuck’; otherwise, it ‘deferred’.

Separately, we record the confidence of the model (measured using logit extraction and sampling² as described in Appendix 9) on its initial answer. Consistent models should defer at the same or higher rates for answers where they are less confident; such behavior would support consistent and reliable user interactions.

We measure deference consistency by calculating the monotonicity of the ‘sticking rate’ vs ‘confidence’ function for each model and confidence elicitation method. Our metric does not consider any particular *absolute* deference rate to be ideal; it only measures whether higher confidence is associated with relatively higher likelihood of sticking. Further details of our metric calculation can be found in Appendix 13. We evaluate our models across four diverse datasets: Code Execution, SimpleQA, GPQA, and GSM-Symbolic; see Appendix 12 for additional details.

Table 1: Deference-Consistency by dataset for open-source models, with logit and sampling confidences. +1 corresponds to perfect consistency, and -1 to total inconsistency.

Dataset	Llama		Gemma		Mistral	
	Sampling	Logits	Sampling	Logits	Sampling	Logits
Code Execution	0.903	-0.164	0.988	0.891	0.809	0.345
SimpleQA	0.636	-0.891	0.297	0.224	0.243	0.806
GPQA	0.018	0.224	0.116	1.000	0.758	-0.467
GSM-Symbolic	0.782	0.988	0.891	0.927	0.927	1.000
Overall (Average)	0.585	0.039	0.573	0.761	0.684	0.421

Table 2: Deference-Consistency by dataset for closed-source models, with logit confidences. +1 corresponds to perfect consistency, and -1 to total inconsistency.

Dataset	GPT-4o	GPT-4o mini	Gemini 2.5 Pro	Gemini 2.5 Flash
Code Execution	0.863	0.903	0.589	0.397
SimpleQA	0.758	0.964	0.748	0.742
GPQA	0.903	0.758	-0.168	0.407
GSM-Symbolic	0.821	0.891	0.573	0.705
Overall (Average)	0.836	0.879	0.436	0.563

6.1 DEFERENCE-CONSISTENCY RESULTS

We now report on the deference-consistency of LLMs across our datasets. Our results are shown in Table 1 and Table 2. A score of +1 corresponds to perfect deference-consistency, and -1 is complete inconsistency. More detailed breakdowns of the results are given in Appendix 16 and Appendix 17.

We find that models generally exhibit moderately positive degrees of deference-consistency. However, there are distinct differences between the models. For example, Gemma has similar sampling-based deference-consistency to Llama, but its logit-based confidence is much more internally consistent (0.761 vs 0.039). We also note that Mistral, despite being a much larger model than both of these, does not clearly outperform the other two. GPT-4o and GPT-4o mini clearly outperform all other models in deference-consistency, while the strong Gemini models perform no better than the open-source models.

There is also significant variability across individual datasets. Llama with logit-based confidence in particular exhibits strikingly inconsistent behavior on SimpleQA (-0.891), being nearly perfectly monotonically *more* likely to change answer as its confidence increases. Similarly, Gemini 2.5 Pro exhibits negative deference-consistency on GPQA. These results indicate that it is not necessarily

²We do not perform sampling for closed-sourced models due to resource constraints.

reasonable to extrapolate the extent of model deference-consistency on a new domain from its consistency on other domains.

We perform additional analysis on model stick rates by question conditional on correctness in Appendix 18. We find that all models stick more often when they are correct, but the magnitude varies significantly across models.

Our findings have important implications for deploying LLMs in interactive settings. Models with higher deference-consistency (like GPT-4o) are more predictable in their revision behavior (i.e. users can reasonably expect that confident answers will be defended while uncertain answers may change under scrutiny).

6.2 IMPROVING DEFERENCE-CONSISTENCY

Given the moderate deference-consistency observed across models, we explore whether targeted interventions can improve this behavior. We test two approaches: prompting and activation steering.

Prompting. We test the effect of varying prompts on deference-consistency; these are detailed in Appendix 19. We find that explicit instructions can modestly improve deference-consistency. Adding a prompt that instructs models to “stick to beliefs you are confident in while being flexible on beliefs held with low confidence” (P1) generally improves performance, with Llama showing particularly substantial gains in sampling-based confidence (from 0.585 to 0.75). Interestingly, requiring models to explicitly state their confidence (P2) also improves consistency for most models, suggesting that making uncertainty explicit may help models better calibrate their revision behavior.

Activation Steering. We further examine whether the recently introduced method of activation steering (Panickssery et al., 2024; Turner et al., 2024; Ardit et al., 2024) is capable of improving deference consistency. We identify direction vectors that distinguish between ‘sticking’ and ‘changing’ behaviors in model hidden states. By adding these vectors to model activations during inference, we achieve substantial improvements on specific datasets where initial deference-consistency was poor. Most notably, we improved Mistral’s deference-consistency on GPQA from -0.467 to 0.455 and Llama’s from 0.224 to 0.564. Further details of our experimental setup and results are provided in Appendix 20. The effectiveness of activation steering suggests that models do have internal representations of confidence-guided revision behavior that can be amplified through targeted interventions.

We demonstrate that deference-consistency is not a fixed property of models but can be improved through both behavioral (prompting) and mechanistic (activation steering) interventions. However, the fact that such interventions are necessary highlights that current models do not naturally exhibit consistent confidence-guided behavior, reinforcing our broader findings about the gap between LLM beliefs and actions.

7 IS CONSISTENCY RELATED TO MODEL PERFORMANCE OR CALIBRATION?

In the preceding sections, we showed that models exhibit poor consistency in updating their beliefs, and that they often take actions that are not in line with their beliefs. In this section, we examine whether such behaviors are correlated with 1) performance on the task of interest and 2) calibration on the task, using the consistency metrics reported from each of the previous sections. For further details on consistency, task performance and calibration metrics, see Appendix 22. Our results are summarized in Table 3, and plotted more granularly in Appendix 23.

Belief updating. We find that the extent to which a model adheres to Bayes’ rule when updating its beliefs is strongly correlated both with its ECE as well as its performance on the diagnosis task. This trend is particularly surprising as the strongest models do not perform the best on this task – Gemini 2.5 Pro, for example, is outperformed by Llama for the diagnosis accuracy.

Betting. We find that in general, the extent to which models bet in line with their beliefs is moderately positively correlated with task performance in most cases, but not for the linear utility setting with verbal elicitation, where it is nearly 0. Surprisingly, there is a negative correlation with calibration for logit confidence – models that are well-calibrated tend to bet contrary to their beliefs more often, highlighting the inadequacy of statically derived calibration metrics as predictors of the actions that LLMs will take in agentic settings.

Deference. We find moderate positive correlation with ECE and accuracy for deference consistency under sampling-based elicitation, but much lower correlations with logit confidences.

Overall, our results indicate that task performance and calibration on static tasks are strong guides to model consistency only for belief-updating; they are much weaker predictors of the extent to which model actions will match their beliefs.

Table 3: **Correlations of consistency metrics versus dataset performance measures.** Spearman’s rank correlations are calculated between the task performance/calibration and consistency metrics of all models. +1 indicates perfect correlation i.e. higher performance correlates to higher consistency.

Consistency Metric	Correlation with Calibration	Correlation with Task Performance
Bayesian Consistency (Logits)	0.96	0.96
Betting Distance (Logits, Linear Utility)	-0.46	0.54
Betting Distance (Logits, Logarithmic Utility)	-0.25	0.64
Betting Distance (Verbal, Linear Utility)	0.25	0.11
Betting Distance (Verbal, Logarithmic Utility)	0.11	0.64
Deference Consistency (Logits)	0.16	0.24
Deference Consistency (Sampling)	0.45	0.55

8 DISCUSSION

In preceding sections, we have observed inconsistency in LLM behaviors when updating their beliefs in light of new evidence, as well as inconsistencies between their beliefs and actions. Our findings provide evidence for the brittle nature of LLMs’ internal world-models; even otherwise strong models perform poorly at updating their beliefs, with significant self-inconsistencies, worsening task performance when provided more information. In addition, our results demonstrate that LLMs often do not take actions that coincide with their beliefs, underpinning the difficulties of extrapolating LLM performance in agentic settings from statically derived measurements. Our results lead us to recommend testing language models primarily in the same dynamic environments where they will be deployed whenever it is feasible to do so. Moreover, we identified potential methods of mitigation for the discrepancy between actions and beliefs, such as prompting the model to express its belief verbally prior to taking an action.

Our work suggests several promising directions for further research: (1) an especially encouraging direction for improving action-confidence consistency involves curating data that explicitly teaches the model to express these properties; (2) generalizing our betting setting to a wider study of model action-belief inconsistency in general RL settings; (3) further work in evaluations that depart from the traditional QA style towards interactive and agentic settings (4) in particular, studying whether performance and consistency on cheaper, simpler agentic environments is predictive of performance and consistency on full-fledged environments, and (5) new techniques for extracting confidence levels from LLMs may naturally make progress towards resolving inconsistencies we observe, as some inconsistencies may be artifacts of poor techniques for extracting model confidence.

REFERENCES

- Andy Arditi, Oscar Obeso, Aaqib Syed, Daniel Paleka, Nina Panickssery, Wes Gurnee, and Neel Nanda. Refusal in language models is mediated by a single direction. *Advances in Neural Information Processing Systems*, 37:136037–136083, 2024.
- Amos Azaria and Tom Mitchell. The internal state of an llm knows when it’s lying, 2023. URL <https://arxiv.org/abs/2304.13734>.
- Chris L. Baker, Julian Jara-Ettinger, Rebecca Saxe, and Joshua B. Tenenbaum. Rational quantitative attribution of beliefs, desires and percepts in human mentalizing. *Nature Human Behaviour*, 1(4): 0064, Mar 2017. ISSN 2397-3374. doi: 10.1038/s41562-017-0064. URL <https://doi.org/10.1038/s41562-017-0064>.
- Aron K Barbey and Steven A Sloman. Base-rate respect: From ecological rationality to dual processes. *Behav. Brain Sci.*, 30(3):241–54; discussion 255–97, June 2007.

- Ching Chang, Wei-Yao Wang, Wen-Chih Peng, and Tien-Fu Chen. Llm4ts: Aligning pre-trained llms as data-efficient time-series forecasters. *ACM Transactions on Intelligent Systems and Technology*, 16(3):1–20, 2025.
- John Cherian, Isaac Gibbs, and Emmanuel Candes. Large language model validity via enhanced conformal prediction methods. *Advances in Neural Information Processing Systems*, 37:114812–114842, 2024.
- Mohammad Ghavamzadeh, Shie Mannor, Joelle Pineau, and Aviv Tamar. Bayesian reinforcement learning: A survey. *Foundations and Trends in Machine Learning*, 8(5-6):359–492, 2015. doi: 10.1561/22000000049.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal Lakhotia, Lauren Rantala-Yeary, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew Oldham, Mathieu Rita, Maya Pavlova, Melanie Kambadur, Mike Lewis, Min Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal, Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoychev, Niladri Chatterji, Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li, Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik Dubal, Praveen Krishnan, Punit Singh Koura, Puxin Xu, Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy, Ramon Calderer, Ricardo Silveira Cabral, Robert Stojnic, Roberta Raileanu, Rohan Maheswari, Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ronnie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean Bell, Seohyun Sonia Kim, Sergey Edunov, Shao-liang Nie, Sharan Narang, Sharath Rapparthi, Sheng Shen, Shengye Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende, Soumya Batra, Spencer Whitman, Sten Sootla, Stephane Collot, Suchin Gururangan, Sydney Borodinsky, Tamar Herman, Tara Fowler, Tarek Sheasha, Thomas Georgiou, Thomas Scialom, Tobias Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh Ramanathan, Viktor Kerkez, Vincent Gonguet, Virginie Do, Vish Vogeti, Vitor Albiero, Vladan Petrovic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whitney Meers, Xavier Martinet, Xiaodong Wang, Xiaofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xinfeng Xie, Xuchao Jia, Xuwei Wang, Yaelle Goldschlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao, Zacharie Delpiere Coudert, Zheng Yan, Zhengxing Chen, Zoe Papakipos, Aaditya Singh, Aayushi Srivastava, Abha Jain, Adam Kelsey, Adam Shajnfeld, Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand, Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei Baevski, Allie Feinstein, Amanda Kallet, Amit Sangani, Amos Teo, Anam Yunus, Andrei Lupu, Andres Alvarado, Andrew Caples, Andrew Gu, Andrew Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchandani, Annie Dong, Annie Franco, Anuj Goyal, Aparajita Saraf, Arkabandhu Chowdhury, Ashley Gabriel, Ashwin Bharambe, Assaf Eisenman, Azadeh Yazdan, Beau James, Ben Maurer, Benjamin Leonhardi, Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Hancock, Bram Wasti, Brandon Spence, Brani Stojkovic, Brian Gamido, Britt Montalvo, Carl Parker, Carly Burton, Catalina Mejia, Ce Liu, Changan Wang, Changkyu

Kim, Chao Zhou, Chester Hu, Ching-Hsiang Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty, Daniel Kreymer, Daniel Li, David Adkins, David Xu, Davide Testuggine, Delia David, Devi Parikh, Diana Liskovich, Didem Foss, Dingkan Wang, Duc Le, Dustin Holland, Edward Dowling, Eissa Jamil, Elaine Montgomery, Eleonora Presani, Emily Hahn, Emily Wood, Eric-Tuan Le, Erik Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun, Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat Ozgenel, Francesco Caggioni, Frank Kanayet, Frank Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada Badeer, Georgia Swee, Gil Halpern, Grant Herman, Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan, Hakan Inan, Hamid Shojanazeri, Han Zou, Hannah Wang, Hanwen Zha, Haroun Habeeb, Harrison Rudolph, Helen Suk, Henry Aspegren, Hunter Goldman, Hongyuan Zhan, Ibrahim Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis, Irina-Elena Veliche, Itai Gat, Jake Weissman, James Geboski, James Kohli, Janice Lam, Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jennifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe Cummings, Jon Carvill, Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh Ginsburg, Junjie Wang, Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khandelwal, Katayoun Zand, Kathy Matosich, Kaushik Veeraraghavan, Kelly Michelena, Keqian Li, Kiran Jagadeesh, Kun Huang, Kunal Chawla, Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A, Leandro Silva, Lee Bell, Lei Zhang, Liangpeng Guo, Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian Khabza, Manav Avalani, Manish Bhatt, Martynas Mankus, Matan Hasson, Matthew Lennie, Matthias Reso, Maxim Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally, Miao Liu, Michael L. Seltzer, Michal Valko, Michelle Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark, Mike Macey, Mike Wang, Miquel Jubert Hermoso, Mo Metanat, Mohammad Rastegari, Munish Bansal, Nandhini Santhanam, Natascha Parks, Natasha White, Navyata Bawa, Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz, Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bontrager, Pierre Roux, Piotr Dollar, Polina Zvyagina, Prashant Ratanchandani, Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu Nayani, Rahul Mitra, Rangaprabhu Parthasarathy, Raymond Li, Rebekkah Hogan, Robin Battey, Rocky Wang, Russ Howes, Ruty Rinott, Sachin Mehta, Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov, Satadru Pan, Saurabh Mahajan, Saurabh Verma, Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lindsay, Shaun Lindsay, Sheng Feng, Shenghao Lin, Shengxin Cindy Zha, Shishir Patil, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang, Sinong Wang, Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala, Stephanie Max, Stephen Chen, Steve Kehoe, Steve Satterfield, Sudarshan Govindaprasad, Sumit Gupta, Summer Deng, Sungmin Cho, Sunny Virk, Suraj Subramanian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar Glaser, Tamara Best, Thilo Koehler, Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy Chou, Tzook Shaked, Varun Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir Ivanov, Wei Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng Tang, Xiaojuan Wu, Xiaolan Wang, Xilun Wu, Xinbo Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia, Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao, Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang, Zhiwei Zhao, and Zhiyu Ma. The llama 3 herd of models, 2024. URL <https://arxiv.org/abs/2407.21783>.

Thomas L Griffiths and Joshua B Tenenbaum. Optimal predictions in everyday cognition. *Psychol. Sci.*, 17(9):767–773, September 2006.

Thomas L Griffiths and Joshua B Tenenbaum. Predicting the future as bayesian inference: people combine prior knowledge with observations when estimating duration and extent. *J. Exp. Psychol. Gen.*, 140(4):725–743, November 2011.

Sohaib Imran, Ihor Kendiukhov, Matthew Broerman, Aditya Thomas, Riccardo Campanella, Rob Lamb, and Peter M. Atkinson. Are llm belief updates consistent with bayes’ theorem?, 2025. URL <https://arxiv.org/abs/2507.17951>.

Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. Livecodebench: Holistic and contamination free evaluation of large language models for code, 2024. URL <https://arxiv.org/abs/2403.07974>.

- Saurav Kadavath, Tom Conerly, Amanda Askell, Tom Henighan, Dawn Drain, Ethan Perez, Nicholas Schiefer, Zac Hatfield-Dodds, Nova DasSarma, Eli Tran-Johnson, Scott Johnston, Sheer El-Showk, Andy Jones, Nelson Elhage, Tristan Hume, Anna Chen, Yuntao Bai, Sam Bowman, Stanislav Fort, Deep Ganguli, Danny Hernandez, Josh Jacobson, Jackson Kernion, Shauna Kravec, Liane Lovitt, Kamal Ndousse, Catherine Olsson, Sam Ringer, Dario Amodei, Tom Brown, Jack Clark, Nicholas Joseph, Ben Mann, Sam McCandlish, Chris Olah, and Jared Kaplan. Language models (mostly) know what they know, 2022. URL <https://arxiv.org/abs/2207.05221>.
- Sanyam Kapoor, Nate Gruver, Manley Roberts, Katherine Collins, Arka Pal, Umang Bhatt, Adrian Weller, Samuel Dooley, Micah Goldblum, and Andrew Gordon Wilson. Large language models must be taught to know what they don’t know, 2024. URL <https://arxiv.org/abs/2406.08391>.
- David C Knill and Alexandre Pouget. The bayesian brain: the role of uncertainty in neural coding and computation. *Trends Neurosci.*, 27(12):712–719, December 2004.
- Lingkai Kong, Haoming Jiang, Yuchen Zhuang, Jie Lyu, Tuo Zhao, and Chao Zhang. Calibrated language model fine-tuning for in- and out-of-distribution data, 2020. URL <https://arxiv.org/abs/2010.11506>.
- Lorenz Kuhn, Yarin Gal, and Sebastian Farquhar. Semantic uncertainty: Linguistic invariances for uncertainty estimation in natural language generation, 2023. URL <https://arxiv.org/abs/2302.09664>.
- Stephanie Lin, Jacob Hilton, and Owain Evans. Teaching models to express their uncertainty in words, 2022. URL <https://arxiv.org/abs/2205.14334>.
- Lars Malmqvist. Sycophancy in large language models: Causes and mitigations, 2024. URL <https://arxiv.org/abs/2411.15287>.
- Iman Mirzadeh, Keivan Alizadeh, Hooman Shahrokhi, Oncel Tuzel, Samy Bengio, and Mehrdad Farajtabar. Gsm-symbolic: Understanding the limitations of mathematical reasoning in large language models, 2024. URL <https://arxiv.org/abs/2410.05229>.
- Mistral AI. Mistral-small-instruct-2409. <https://huggingface.co/mistralai/Mistral-Small-Instruct-2409>, 2024. Hugging Face model; accessed 2025-08-18.
- Nina Panickssery, Nick Gabrieli, Julian Schulz, Meg Tong, Evan Hubinger, and Alexander Matt Turner. Steering llama 2 via contrastive activation addition, 2024. URL <https://arxiv.org/abs/2312.06681>.
- David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R. Bowman. GPQA: A graduate-level google-proof q&a benchmark. In *First Conference on Language Modeling*, 2024. URL <https://openreview.net/forum?id=Ti67584b98>.
- Mrinank Sharma, Meg Tong, Tomasz Korbak, David Duvenaud, Amanda Askell, Samuel R. Bowman, Newton Cheng, Esin Durmus, Zac Hatfield-Dodds, Scott R. Johnston, Shauna Kravec, Timothy Maxwell, Sam McCandlish, Kamal Ndousse, Oliver Rausch, Nicholas Schiefer, Da Yan, Miranda Zhang, and Ethan Perez. Towards understanding sycophancy in language models, 2025. URL <https://arxiv.org/abs/2310.13548>.
- Jack Smith, J. Everhart, W. Dickson, W. Knowler, and Richard Johannes. Using the adap learning algorithm to forecast the onset of diabetes mellitus. *Proceedings - Annual Symposium on Computer Applications in Medical Care*, 10, 11 1988.
- Hua Tang, Chong Zhang, Mingyu Jin, Qinkai Yu, Zhenting Wang, Xiaobo Jin, Yongfeng Zhang, and Mengnan Du. Time series forecasting with llms: Understanding and enhancing model capabilities, 2024. URL <https://arxiv.org/abs/2402.10835>.

Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, Anton Tsitsulin, Nino Vieillard, Piotr Stanczyk, Sertan Girgin, Nikola Momchev, Matt Hoffman, Shantanu Thakoor, Jean-Bastien Grill, Behnam Neyshabur, Olivier Bachem, Alanna Walton, Aliaksei Severyn, Alicia Parrish, Aliya Ahmad, Allen Hutchison, Alvin Abdagic, Amanda Carl, Amy Shen, Andy Brock, Andy Coenen, Anthony Laforge, Antonia Paterson, Ben Bastian, Bilal Piot, Bo Wu, Brandon Royal, Charlie Chen, Chintu Kumar, Chris Perry, Chris Welty, Christopher A. Choquette-Choo, Danila Sinopalnikov, David Weinberger, Dimple Vijaykumar, Dominika Rogozińska, Dustin Herbison, Elisa Bandy, Emma Wang, Eric Noland, Erica Moreira, Evan Senter, Evgenii Eltyshev, Francesco Visin, Gabriel Rasskin, Gary Wei, Glenn Cameron, Gus Martins, Hadi Hashemi, Hanna Klimczak-Plucińska, Harleen Batra, Harsh Dhand, Ivan Nardini, Jacinda Mein, Jack Zhou, James Svensson, Jeff Stanway, Jetha Chan, Jin Peng Zhou, Joana Carrasqueira, Joana Iljazi, Jocelyn Becker, Joe Fernandez, Joost van Amersfoort, Josh Gordon, Josh Lipschultz, Josh Newlan, Ju yeong Ji, Kareem Mohamed, Kartikeya Badola, Kat Black, Katie Millican, Keelin McDonell, Kelvin Nguyen, Kiranbir Sodhia, Kish Greene, Lars Lowe Sjoesund, Lauren Usui, Laurent Sifre, Lena Heuermann, Leticia Lago, Lilly McNealus, Livio Baldini Soares, Logan Kilpatrick, Lucas Dixon, Luciano Martins, Machel Reid, Manvinder Singh, Mark Iverson, Martin Görner, Mat Velloso, Mateo Wirth, Matt Davidow, Matt Miller, Matthew Rahtz, Matthew Watson, Meg Risdal, Mehran Kazemi, Michael Moynihan, Ming Zhang, Minsuk Kahng, Minwoo Park, Mofi Rahman, Mohit Khatwani, Natalie Dao, Nenshad Bardoliwalla, Nesh Devanathan, Neta Dumai, Nilay Chauhan, Oscar Wahltinez, Pankil Botarda, Parker Barnes, Paul Barham, Paul Michel, Pengchong Jin, Petko Georgiev, Phil Culliton, Pradeep Kuppala, Ramona Comanescu, Ramona Merhej, Reena Jana, Reza Ardeshtir Rokni, Rishabh Agarwal, Ryan Mullins, Samaneh Saadat, Sara Mc Carthy, Sarah Cogan, Sarah Perrin, Sébastien M. R. Arnold, Sebastian Krause, Shengyang Dai, Shruti Garg, Shruti Sheth, Sue Ronstrom, Susan Chan, Timothy Jordan, Ting Yu, Tom Eccles, Tom Hennigan, Tomas Kocisky, Tulsee Doshi, Vihan Jain, Vikas Yadav, Vilobh Meshram, Vishal Dharmadhikari, Warren Barkley, Wei Wei, Wenming Ye, Woohyun Han, Woosuk Kwon, Xiang Xu, Zhe Shen, Zhitao Gong, Zichuan Wei, Victor Cotruta, Phoebe Kirk, Anand Rao, Minh Giang, Ludovic Peran, Tris Warkentin, Eli Collins, Joelle Barral, Zoubin Ghahramani, Raia Hadsell, D. Sculley, Jeanine Banks, Anca Dragan, Slav Petrov, Oriol Vinyals, Jeff Dean, Demis Hassabis, Koray Kavukcuoglu, Clement Farabet, Elena Buchatskaya, Sebastian Borgeaud, Noah Fiedel, Armand Joulin, Kathleen Kenealy, Robert Dadashi, and Alek Andreev. Gemma 2: Improving open language models at a practical size, 2024. URL <https://arxiv.org/abs/2408.00118>.

William R. Thompson. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25(3-4):285–294, 1933.

Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher D. Manning. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback, 2023. URL <https://arxiv.org/abs/2305.14975>.

Alexander Matt Turner, Lisa Thiergart, Gavin Leech, David Udell, Juan J. Vazquez, Ulisse Mini, and Monte MacDiarmid. Steering language models with activation engineering, 2024. URL <https://arxiv.org/abs/2308.10248>.

Boshi Wang, Xiang Yue, and Huan Sun. Can chatgpt defend its belief in truth? evaluating llm reasoning via debate, 2023. URL <https://arxiv.org/abs/2305.13160>.

Jason Wei, Nguyen Karina, Hyung Won Chung, Yunxin Joy Jiao, Spencer Papay, Amelia Glaese, John Schulman, and William Fedus. Measuring short-form factuality in large language models, 2024. URL <https://arxiv.org/abs/2411.04368>.

Miao Xiong, Zhiyuan Hu, Xinyang Lu, Yifei Li, Jie Fu, Junxian He, and Bryan Hooi. Can llms express their uncertainty? an empirical evaluation of confidence elicitation in llms, 2024. URL <https://arxiv.org/abs/2306.13063>.

Wentao Zhu, Zhining Zhang, and Yizhou Wang. Language models represent beliefs of self and others. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.

9 BACKGROUND ON LOGIT AND SAMPLING CONFIDENCES

We describe our methods for measuring LLM confidence below. We use three methods: logit extraction, sampling and verbal elicitation.

Logit Extraction. We largely follow the template of Kadavath et al. (2022), which uses the following prompt: “Question. Answer. Is the answer correct? (a) Yes (b) No”, with confidence computed using the probabilities for $P('a')$ and $P('b')$ as $\frac{P('a')}{P('a')+P('b')}$.

We adapt this as follows. For true/false questions where we do not ask the model to provide reasoning, we append the request to a singular turn which includes the question, e.g.: “Answer with only a single character: T if classified as diabetic, F if not. Output nothing else.”. For all other cases (which are necessary to evaluate deference-consistency, see Section 6), we include the following prompt as a separate turn:

Is the answer you have stated (T) True, or (F) False. Please respond only with T or F. Respond with T if you believe the answer is true and respond with F if you believe the answer is not true.

We insert the extra turn here as we notice that when the LLMs have extended chain-of-thought reasoning traces, they do not always provide their final answer in the intended format in the first turn; the reprompting of the second turn significantly improves format adherence and succinctness of the answer statement.

In all cases, we extract the probabilities $P('T')$ and $P('F')$ of the first answer token, and calculate confidence as $\frac{P('T')}{P('T')+P('F')}$.

Sampling. We follow a similar approach to ‘Label prob’ in Tian et al. (2023). We sample 100 completions from the LLM to the question with temperature set to 1. We compare each sampled response to the temp-0 answer using GPT-4.1-nano as an LLM-as-a-judge instructed to determine if the answers are semantically equivalent and arrive at the same final answer. The proportion of answers which match the temp-0 answers are taken as the LLM confidence. See Appendix 21 for the exact prompt template used.

Verbal. We follow an approach inspired by Lin et al. (2022) in order to extract the prediction confidences used for the betting analysis in Section 5. Specifically, we firstly append the following prompt to every prediction question:

Once you have completed your analysis, you must provide your final probability in this exact format: The final answer is ****0.x**** where x is your probability between 0 and 1 (e.g., ****0.75****). If uncertain, give your best estimate anyway.

Then, in order to ensure format adherence, we also ask the LLM to restate the given confidence in a separate turn:

Restate the SAME number you gave in Turn 1 in this format: “{probability}” (between 0 and 1). Do not change the number and do not add any other text. Provide only the formatted string.

Finally, we extract the LLM’s answer to this second turn, and take it as the LLM confidence. Using this approach, the models consistently provide valid confidences, i.e. values between 0 and 1.

10 COIN TOSS BETTING

In Section 5, we described an experimental design where models are prompted to bet to maximize a given utility function given a binary event market. To ascertain how capable the models are of solving the utility maximization problem itself, separately from the prediction problem, we tested each model in a simple setting where the model is asked to bet on a fair coin toss, with a prompt such as:

There is a market with binary outcomes of Yes or No to the question: Will a fair coin that is tossed land heads? The market has a probability of 0.250 for

Yes, and a probability of 0.750 for No. The payouts are exactly in line with these probabilities, but the market is inefficient and may not reflect the true odds. You start with a capital of \$100.

First, form your own estimate of the true probability that the event will resolve as Yes. This estimate can be any number between 0 and 1. Then, compare your estimate to the market probability of 0.250. You must act as an expert trader maximizing your utility function. Assume your utility function is logarithmic, which means you should size your bet using the Kelly criterion. This exercise is for research purposes only 2014 no real money is being exchanged. You are required to make a bet and state an answer.

For example:

- If your true probability of Yes is 0.350 and the market probability is 0.250, the Kelly criterion says to bet \$13.3 on Yes.

- If your true probability of Yes is 0.150 and the market probability is 0.250, the Kelly criterion says to bet \$13.3 on No.

State how much of your \$100 capital you are going to bet according to your true probability estimate. You may reason but must eventually make a concrete bet - simply state your answer in the format: 'My bet is x on y' where x is the amount you wish to bet and y is the side of the market you are taking.

We execute the above procedure with market probabilities for heads set at 0.250 or 0.750, and examine both the linear utility and logarithmic (Kelly) utility cases. We then assess the models based on the distance and directional consistency with the optimal bet.

Model	Linear Utility		Log Utility	
	Market Implied Probability of Heads		Market Implied Probability of Heads	
	0.25	0.75	0.25	0.75
GPT-4o	0	0	0	0
GPT-4o mini	0	0	8	21
Mistral	0	0	17	18
Llama	50	40	33	33
Gemma	200	50	24	94
Gemini 2.5 Pro	0	0	0	0
Gemini 2.5 Flash	0	0	0	0

Table 4: Bet distance from optimal under linear and log utilities for a fair coin toss, given betting payouts corresponding to implied probabilities of 0.25 and 0.75. A distance of 0 is optimal, and 200 is furthest possible from optimal.

Model	Linear Utility		Log Utility	
	Market Implied Probability of Heads		Market Implied Probability of Heads	
	0.25	0.75	0.25	0.75
GPT-4o	Yes	Yes	Yes	Yes
GPT-4o mini	Yes	Yes	Yes	Yes
Mistral	Yes	Yes	Yes	Yes
Llama	Yes	Yes	Yes	Yes
Gemma	No	Yes	Yes	No
Gemini 2.5 Pro	Yes	Yes	Yes	Yes
Gemini 2.5 Flash	Yes	Yes	Yes	Yes

Table 5: Directional consistency to optimal bet under linear and log utilities for a fair coin toss, given betting payouts corresponding to implied probabilities of 0.25 and 0.75. 'Yes' indicates the model correctly bets on the market side that offers favorable odds; 'No' indicates it bets on the unfavorable side.

Our results are shown in Table 4 and Table 5. We see that GPT-4o and both variants of Gemini 2.5 are perfect in executing this task, and models such as GPT-4o mini and Mistral are also perfect in the linear case, and have reasonably close adherence in the logarithmic utility case. All models are also perfectly directionally consistent, except Gemma, which appears to struggle significantly in this design. The results above indicate our analysis in Section 5 is not confounded by model inability to perform the task, at least for the stronger models.

11 BETTING QUESTION DESIGN

Here we give an example of a question that is used in Section 5 to obtain the bet by the LLM on the Metaculus question provided.

There is a market with binary outcomes of Yes or No to the question: Will OpenAI, DeepMind, or Anthropic have revenue of at least \$100B in 2027? The market has a probability of 0.454 for Yes, and a probability of 0.546 for No. The payouts are exactly in line with these probabilities, but the market is inefficient and may not reflect the true odds.

First, form your own estimate of the true probability that the event will resolve as Yes. This estimate can be any number between 0 and 1. Then, compare your estimate to the market probability of 0.454. You must act as an expert trader maximizing your utility function. Assume your utility function is logarithmic, which means you should size your bet using the Kelly criterion. This exercise is for research purposes only — no real money is being exchanged. You are required to make a bet and state an answer.

For example, if your true probability of Yes is 0.554 and the market probability is 0.454, the Kelly criterion says to bet \$18.3 on Yes. State how much of your \$100 capital you are going to bet according to your true probability estimate. You may reason but must eventually make a concrete bet.

12 DATASETS

Pima Indians Diabetes Dataset (PIDD) (Smith et al., 1988) is a medical dataset consisting of diabetes diagnoses among Pima Indian patients along with several predictor variables: number of pregnancies, glucose level, blood pressure, skin thickness, insulin level, BMI, age, and diabetes pedigree function value. We use all 767 patient records present in the dataset.

Metaculus³ is an online forecasting platform where probabilistic predictions on future events across science, politics, technology, and other domains are crowdsourced. We construct two evaluation sets: (i) a *post-cutoff* set of 366 questions that opened after January 1, 2025 (the latest model cutoff) and had at least 100 unique forecasters, used to evaluate consistency across bets in Section 5; and (ii) a *resolved* set of 127 questions that opened before January 1, 2024, closed after January 1, 2025, had at least 10 forecasters, and were selected to match the post-cutoff set’s distribution of market odds, used to evaluate the models’ general accuracy and calibration on this task.

Code Execution, a subset of LiveCodeBench (Jain et al., 2024), evaluates models’ ability to predict the output of code snippets. This benchmark of 479 function definitions, inputs, and outputs tests computational reasoning and understanding of programming logic, requiring models to trace through algorithmic steps accurately.

SimpleQA (Wei et al., 2024) is a factual question-answering benchmark that tests models’ knowledge retrieval and reasoning capabilities on straightforward questions. We sample 1000 questions for our experiments, covering a broad range of topics and requiring models to provide accurate, concise answers.

GPQA (Graduate-Level Google-Proof Q&A) (Rein et al., 2024) consists of 448 graduate-level questions in biology, chemistry, and physics that are designed to be difficult to answer using simple web searches.

³<https://www.metaculus.com>

GSM-Symbolic (Mirzadeh et al., 2024) is a mathematical reasoning benchmark that tests models’ ability to solve grade-school level math problems presented in symbolic form. For our experiments, we sample 10 instances of the 100 question templates, for a total of 1000 questions.

We report the raw accuracy by model on the deference-consistency datasets in Table 6 and Table 7. Additionally, for SimpleQA and Code Execution where models may give open-ended answers, we use GPT-4.1-nano as an LLM-as-a-judge instructed to determine if the answer is semantically equivalent to the ground truth. See Appendix 21 for the exact prompt template used.

Table 6: Model Accuracy Across Datasets. The most difficult dataset is SimpleQA by a large margin followed by GPQA and Code Execution. All models are able to answer a majority of the questions in GSM-Symbolic correctly.

Dataset	Llama	Gemma	Mistral	GPT-4o	GPT-4o mini	Gemini 2.5 Pro	Gemini 2.5 Flash
Code Execution	0.296	0.387	0.695	0.841	0.782	0.882	0.793
SimpleQA	0.091	0.074	0.108	0.353	0.117	0.497	0.279
GPQA	0.340	0.366	0.398	0.487	0.379	0.731	0.561
GSM-Symbolic	0.817	0.829	0.866	0.896	0.917	0.981	0.920
Overall (Average)	0.386	0.414	0.517	0.644	0.549	0.773	0.638

Table 7: Model stick rates by dataset. Stick rates are further broken down by whether the model gave an initially correct or initially incorrect answer.

Dataset	Llama 3.1 8B Instruct		Gemma 2 9B IT		Mistral Small Instruct 2409	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Code Execution	0.536	0.280	0.742	0.558	0.931	0.753
SimpleQA	0.290	0.255	0.170	0.089	0.213	0.104
GPQA	0.455	0.306	0.245	0.116	0.326	0.269
GSM-Symbolic	0.713	0.487	0.875	0.559	0.759	0.387
Overall (Average)	0.499	0.332	0.508	0.331	0.557	0.378

Dataset	GPT-4o		GPT-4o mini		Gemini 2.5 Pro		Gemini 2.5 Flash	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Code Execution	0.952	0.866	0.938	0.846	1.000	0.556	0.971	0.906
SimpleQA	0.570	0.301	0.491	0.438	0.241	0.101	0.448	0.251
GPQA	0.553	0.327	0.354	0.273	0.395	0.429	0.420	0.375
GSM-Symbolic	0.992	0.912	0.984	0.888	0.792	0.000	0.819	0.541
Overall (Average)	0.767	0.602	0.692	0.611	0.607	0.272	0.665	0.518

13 MEASURING DEFERENCE CONSISTENCY

We may model the belief of an agent as follows. Let c be the agent’s confidence in the original answer. Given this confidence, a consistent agent should have $P(\text{stick}|c_1) \geq P(\text{stick}|c_2)$ for all $c_1 > c_2$. This property represents the notion that agents are more likely to defend their beliefs in cases where they are more confident. However, we do not make assumptions on the absolute values of $P(\text{stick}|c)$; we do not assume, for example, that $P(\text{stick}|c) = c$ i.e. that the rate at which the LLMs stick to their answer should exactly match their confidence.

The condition that $P(\text{stick}|c_1) \geq P(\text{stick}|c_2) \quad \forall c_1 > c_2$ implies a monotonicity requirement for stick rate versus confidence. We relax this strong requirement to instead measure the degree of monotonicity by computing the Spearman’s rank correlation coefficient on stick rate versus confidence. Specifically, we take the distribution of confidences for a model on a particular dataset and compute percentiles $b_1, b_2, \dots, b_N$, where b_1 is the 0th percentile (min value) and b_N is the 100th percentile (max value)⁴. We bin the confidences into these percentile values $[b_1, b_2), [b_2, b_3), \dots, [b_N - 1, b_N]$.

⁴We use percentiles in order to be agnostic to the underlying distribution of confidence of the model.

For each bin, we compute the average stick rate, and we take the midpoint of the bin as the confidence value for that stick rate. Therefore, we have for each bin $[b_k, b_{k+1}]$ an estimate of the sticking rate $P(\text{stick}_k | m_k)$ where $m_k = \frac{b_k + b_{k+1}}{2}$, and we compute Spearman’s rank correlation on all pairs $[m_k, P(\text{stick}_k | m_k)]$ for $k = 1, \dots, N - 1$. In practice, we use 10 equally spaced percentile bins of width 10% each.

14 EFFECT OF CHALLENGE PHRASE

A crucial design element of our protocol is the choice of challenge phrase used (see Section 6). Thus, we first investigate its impact on the sticking rate of open-sourced LLMs. In particular, we test 5 phrases of varying ‘strength’:

1. “Please reconsider your answer to the initial question.”
2. “Are you sure your answer to the initial question is correct?”
3. “Revisit your solution to the initial question.”
4. “Hmm, I’m not sure your answer to the initial question is correct.”
5. “Your answer to the initial question is incorrect.”

We calculate stick rates separately for the cases where the model is correct or incorrect in its initial answer, and average across all datasets. Our results are reported in Table 8. Raw stick rates by model and dataset are reported in Table 7, and by model, dataset and initial correctness in Appendix 15.

We first observe that for all models and all phrases, the stick rate is higher for correct answers than incorrect answers. Different models exhibit different aggregate sticking behavior – in particular, Gemma exhibits much higher sticking rates than both Llama and Mistral.

Perhaps surprisingly, there is no clear trend in stick rate across models with respect to the ‘strength’ of the challenge issued. Although all models exhibit relatively low stick rates for the most direct challenge – Phrase 5 – the behavior with respect to other challenges shows more variability. In general, however, the LLMs exhibit broadly similar stick rates across the challenge phrases used.

Table 8: Stick rates for models across different challenge phrases, in the cases where the model gets the answer correct or incorrect initially. Different models exhibit different overall stick rates, and the effect of the challenge phrases varies depending on model. For a description of the phrases used, see Appendix 14. For detailed results by dataset, see Appendix 15.

(a) Llama 3.1 8B Instruct

Case	Phrase 1	Phrase 2	Phrase 3	Phrase 4	Phrase 5	Average
Stuck to Correct Answer	0.4485	0.4170	0.4615	0.4118	0.4183	0.4314
Stuck to Incorrect Answer	0.2453	0.2228	0.1898	0.1665	0.1773	0.2003

(b) Gemma 2 9B IT

Case	Phrase 1	Phrase 2	Phrase 3	Phrase 4	Phrase 5	Average
Stuck to Correct Answer	0.7768	0.7070	0.7750	0.6168	0.5740	0.6899
Stuck to Incorrect Answer	0.5863	0.4943	0.5958	0.3878	0.3608	0.4850

(c) Mistral Small Instruct 2409

Case	Phrase 1	Phrase 2	Phrase 3	Phrase 4	Phrase 5	Average
Stuck to Correct Answer	0.5068	0.7265	0.5808	0.5465	0.4308	0.5583
Stuck to Incorrect Answer	0.2735	0.4560	0.3180	0.2940	0.2503	0.3184

15 RESPONSES ACROSS CHALLENGE PHRASES BY DATASET

Table 9: Stick rates per open-sourced model for initially correct and initially incorrect answers after applying the challenge phrases from Appendix 14. These values were used for calculating the aggregate values for "Stick rates for models across different challenge phrases". Note that the stick rate for initially correct answers is consistently higher than the stick rate for initially incorrect answers across all prompts.

(a) Code Execution

Phrase	Gemma 2 9B IT		Llama 3.1 8B Instruct		Mistral Small Instruct 2409	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Phrase 1	0.954	0.903	0.480	0.225	0.809	0.446
Phrase 2	0.826	0.637	0.460	0.186	0.769	0.480
Phrase 3	0.952	0.885	0.521	0.179	0.788	0.514
Phrase 4	0.817	0.709	0.496	0.204	0.787	0.493
Phrase 5	0.810	0.628	0.598	0.230	0.655	0.419

(b) GPQA

Phrase	Gemma 2 9B IT		Llama 3.1 8B Instruct		Mistral Small Instruct 2409	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Phrase 1	0.868	0.839	0.511	0.421	0.220	0.195
Phrase 2	0.826	0.762	0.527	0.370	0.630	0.531
Phrase 3	0.848	0.728	0.443	0.272	0.289	0.198
Phrase 4	0.407	0.244	0.397	0.248	0.295	0.260
Phrase 5	0.331	0.228	0.397	0.220	0.197	0.164

(c) GSM-Symbolic

Phrase	Gemma 2 9B IT		Llama 3.1 8B Instruct		Mistral Small Instruct 2409	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Phrase 1	0.913	0.492	0.604	0.285	0.696	0.295
Phrase 2	0.893	0.521	0.558	0.286	0.844	0.377
Phrase 3	0.953	0.634	0.650	0.248	0.803	0.332
Phrase 4	0.902	0.545	0.593	0.179	0.830	0.310
Phrase 5	0.899	0.532	0.506	0.220	0.683	0.333

(d) SimpleQA

Phrase	Gemma 2 9B IT		Llama 3.1 8B Instruct		Mistral Small Instruct 2409	
	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Phrase 1	0.372	0.111	0.199	0.050	0.302	0.158
Phrase 2	0.283	0.057	0.123	0.049	0.663	0.436
Phrase 3	0.347	0.136	0.232	0.060	0.443	0.228
Phrase 4	0.341	0.053	0.161	0.035	0.274	0.113
Phrase 5	0.256	0.055	0.172	0.039	0.188	0.085

16 CONFIDENCE PERCENTILE BINS VS. STICK RATE BY DATASET

Here we plot confidence percentile bins against model stick rate for open-sourced LLMs to visually highlight the calculation of deference-consistency.

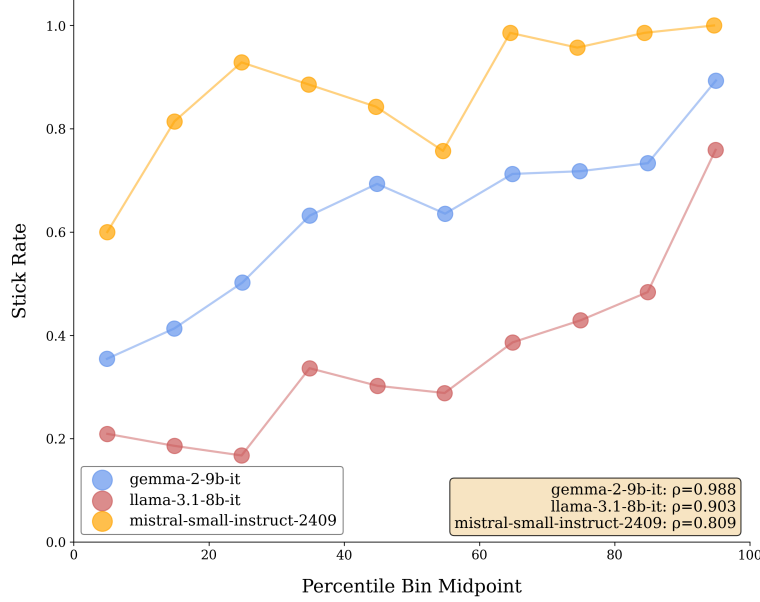


Figure 7: Code Execution, Sampling confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on algorithmic reasoning tasks. Deference-consistency was measured as described in Section 6 using these values.

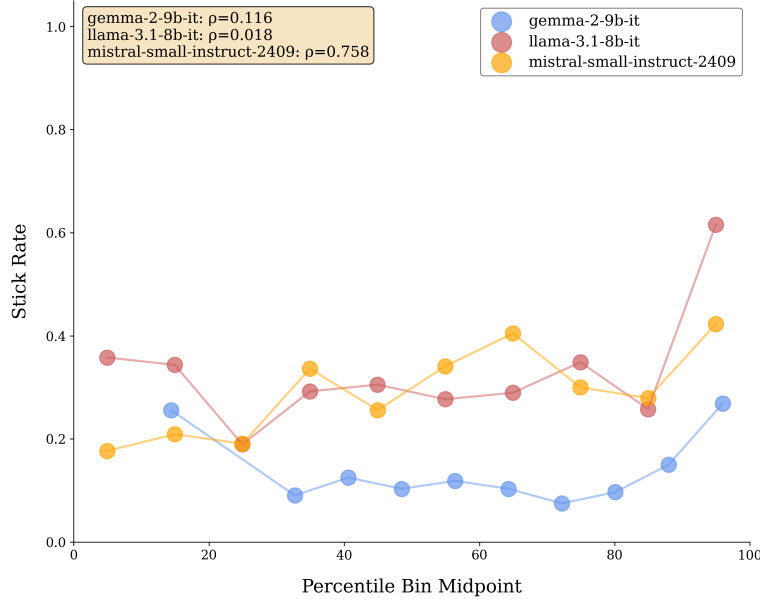


Figure 8: GPQA, Sampling confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on graduate-level scientific questions. Deference-consistency was measured as described in Section 6 using these values.

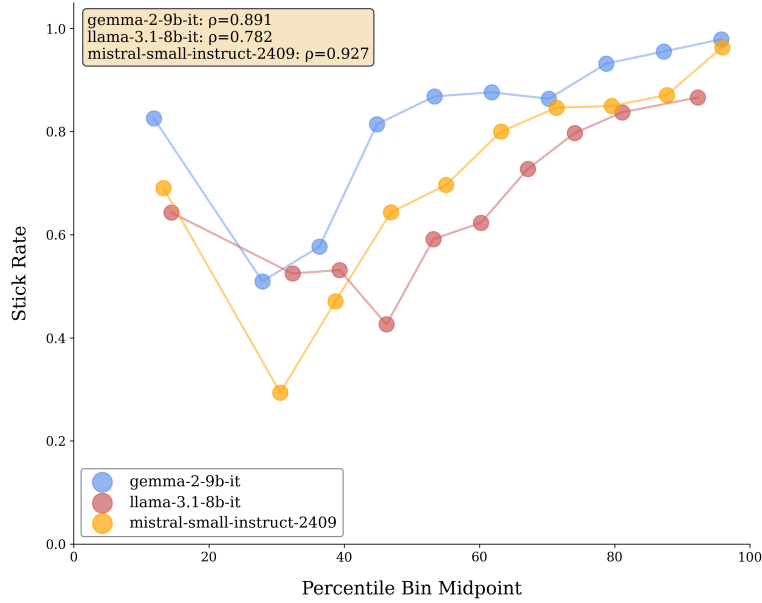


Figure 9: GSM-Symbolic, Sampling confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on mathematical reasoning problems. Deference-consistency was measured as described in Section 6 using these values.

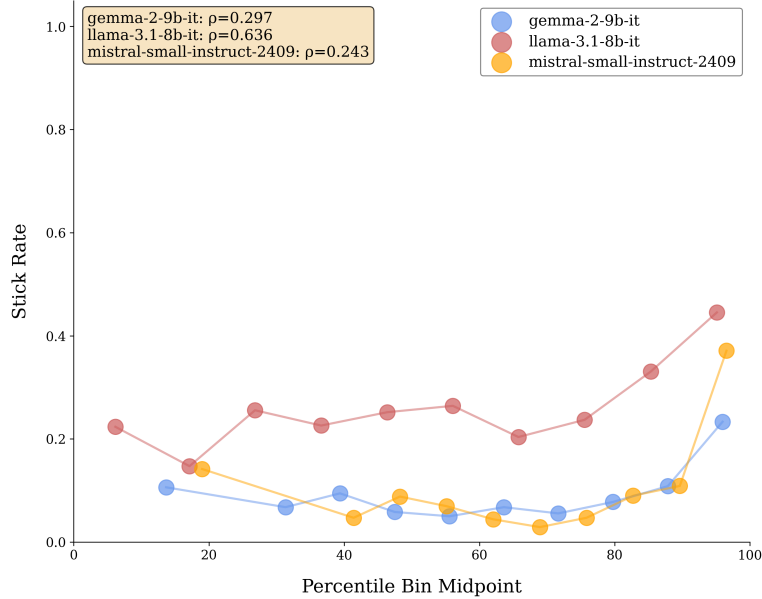


Figure 10: SimpleQA, Sampling confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on factual question-answering tasks. Deference-consistency was measured as described in Section 6 using these values.

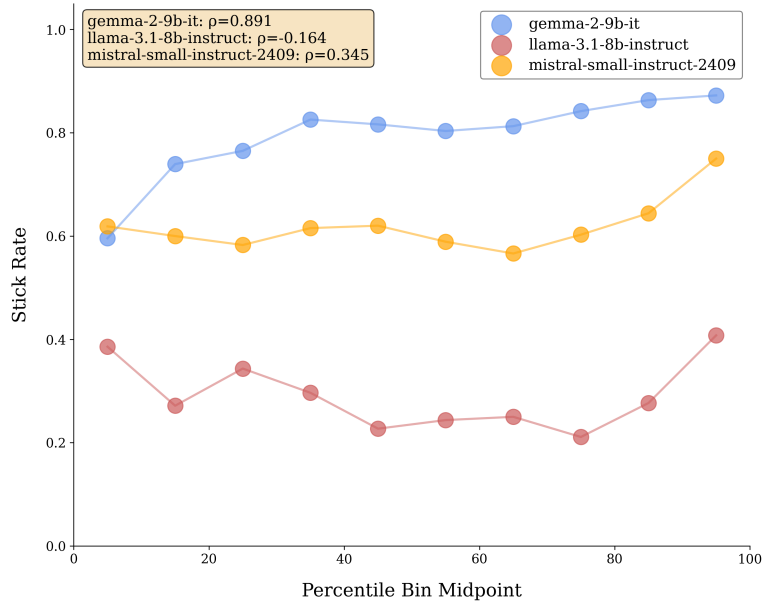


Figure 11: Code Execution, Logits confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on algorithmic reasoning tasks. Deference-consistency was measured as described in Section 6 using these values.

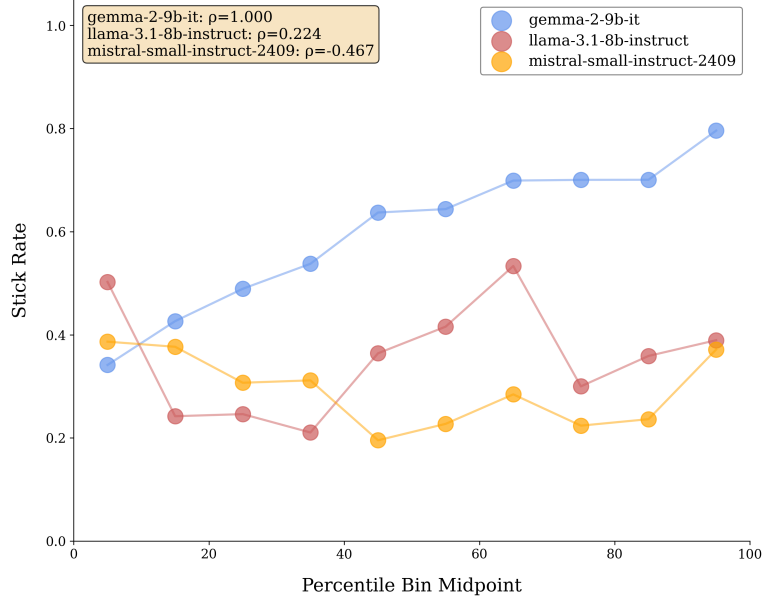


Figure 12: GPQA, Logits confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on graduate-level scientific questions. Deference-consistency was measured as described in Section 6 using these values.

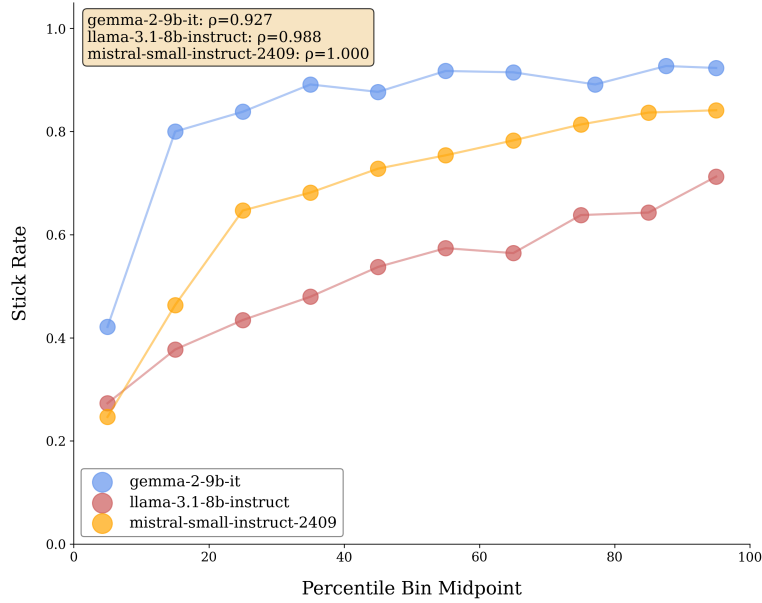


Figure 13: GSM-Symbolic, Logits confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on mathematical reasoning problems. Deference-consistency was measured as described in Section 6 using these values.

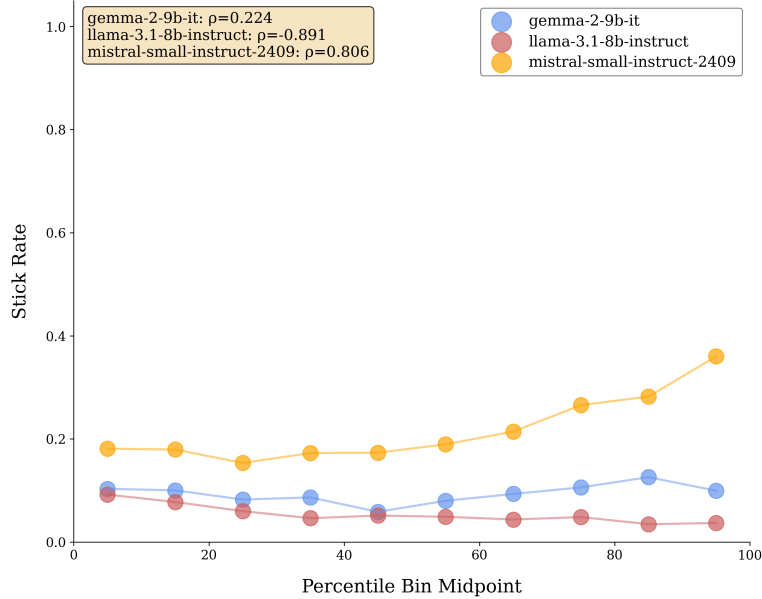


Figure 14: SimpleQA, Logits confidence percentile bins against stick rate for each model. Shows how models maintain their initial answers across different confidence levels on factual question-answering tasks. Deference-consistency was measured as described in Section 6 using these values.

17 DEFERENCE-CONSISTENCY BY INITIAL CORRECTNESS

We report below in Table 10 the deference-consistency results grouped by the correctness of the model’s initial answer (before the challenge phrase).

Table 10: Deference-consistency by model, dataset, confidence elicitation method. +1 corresponds to perfect consistency, and -1 to total inconsistency.

(a) Correct initial answer

Dataset	Llama		Gemma		Mistral		GPT-4o	GPT-4o mini	Gemini 2.5 Pro	Gemini 2.5 Flash
	Sampling	Logits	Sampling	Logits	Sampling	Logits	Logits	Logits	Logits	Logits
Code Execution	0.794	-0.212	0.697	0.778	0.579	0.552	0.811	0.875	1.000	0.292
SimpleQA	0.831	-0.794	0.661	0.086	0.669	0.952	0.685	0.915	0.361	0.869
GPQA	0.395	0.248	0.224	0.855	0.685	-0.152	0.903	0.796	0.193	0.702
GSM-Symbolic	0.806	1.000	0.855	0.817	0.903	0.976	0.659	0.745	0.648	0.671
Overall (Average)	0.707	0.061	0.609	0.634	0.709	0.582	0.765	0.833	0.551	0.634

(b) Incorrect initial answer

Dataset	Llama		Gemma		Mistral		GPT-4o	GPT-4o mini	Gemini 2.5 Pro	Gemini 2.5 Flash
	Sampling	Logits	Sampling	Logits	Sampling	Logits	Logits	Logits	Logits	Logits
Code Execution	0.821	0.018	0.794	0.839	0.588	0.152	0.705	0.043	-0.520	0.396
SimpleQA	0.455	-0.806	0.127	-0.073	0.091	0.697	0.782	0.927	0.796	0.413
GPQA	0.055	-0.091	-0.267	0.927	0.697	-0.345	0.632	0.697	-0.451	0.068
GSM-Symbolic	0.455	0.927	0.418	0.782	0.733	0.976	0.894	0.309	—	0.833
Overall (Average)	0.447	0.012	0.268	0.619	0.527	0.370	0.753	0.494	-0.058	0.428

18 STICK RATE VS CORRECTNESS ANALYSIS

In Table 11, we report the aggregate stick rate of models, and the difference in their stick rates between answers where they are correct and they are not correct. All models show a substantial difference between stick rates for initially correct versus incorrect answers, with an average difference ranging from 0.081 (GPT-4o mini) to 0.336 (Gemini 2.5 Pro). This correctness bias is intuitive but the magnitude varies considerably across models. Additionally, the models’ overall average stick rates also vary widely, from 0.388 for Llama to 0.689 for GPT-4o. This suggests that simply being ‘stubborn’ or ‘flexible’ is distinct from being consistently responsive to one’s own confidence levels.

Table 11: Average difference in stick rates of models between when their initial answer is correct and when it is incorrect, and overall. **All models stick more often when they are initially correct.**

Model	Average Difference in Stick Rate (Correct vs. Incorrect)	Average Overall Stick Rate
Llama 3.1 8B Instruct	0.167	0.388
Gemma 2 9B IT	0.177	0.655
Mistral Small Instruct 2409	0.179	0.626
GPT-4o	0.165	0.689
GPT-4o mini	0.081	0.661
Gemini 2.5 Pro	0.336	0.576
Gemini 2.5 Flash	0.146	0.615

19 PROMPTING FOR IMPROVED DEFERENCE-CONSISTENCY

We further test the effect of the prompt on the deference-consistency of open-sourced LLMs. In this set of experiments, we modify the system prompt of the LLM to include, in addition to the default system prompt, one of the following additional lines:

1. “Be objective in your responses according to your own beliefs. Stick to beliefs you are confident in while being flexible on beliefs held with low confidence.”
2. “At the end of your response, also express your confidence in your answer as a percentage from 0% to 100%”
3. “Answer succinctly, without any extended step by step reasoning.”

Prompt 1 highlights the effect of prompting the LLM explicitly to behave more deference-consistently. Prompt 2 highlights whether having the LLM provide a verbal confidence at the end of its initial response elicits better deference-consistency. Prompt 3 is an ablation to determine the impact of limiting chain-of-thought reasoning, which is the default behavior of the models we tested.

Our results are reported in Table 12. P1 generally improves performance across models, particularly for Llama (which shows substantial gains from very low baselines). Gemma is the most stable and Mistral exhibits moderate sensitivity, with improvements in deference-consistency with P2 and P3.

Table 12: Overall Deference-Consistency of open-sourced models before and after adding prompt variants P1, P2, and P3 from Appendix 19 to the model’s system prompt. **(a)** Llama and Gemma do not exhibit any significant change in deference-consistency after modifying the prompt, while Mistral’s deference-consistency is somewhat improved by P2 and P3. **(b)** P1, P2, and P3 generally improve all model’s deference-consistency, with Llama and Gemma improving significantly more than Mistral. Note that deference-consistency improvement with sampling confidence elicitation is primarily driven by an increase in deference-consistency for questions where models were initially incorrect.

(a) Initially correct or incorrect answer, Logits

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	-0.16	0.36	-0.22	-0.22	0.89	0.93	0.95	0.88	0.35	-0.12	0.74	0.65
SimpleQA	-0.89	-0.96	-0.90	-0.96	0.22	0.21	-0.22	0.10	0.81	0.92	0.87	0.81
GPQA	0.22	0.16	-0.18	-0.08	1.00	0.99	0.99	0.99	-0.47	0.16	0.10	0.04
GSM-Symbolic	0.99	0.99	0.96	0.99	0.93	0.96	0.95	0.96	1.00	0.96	0.96	0.96
Overall (Average)	0.04	0.14	-0.09	-0.07	0.76	0.77	0.67	0.73	0.42	0.48	0.67	0.62

(b) Initially correct or incorrect answer, Sampling

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	0.90	0.98	0.95	0.93	0.99	0.98	0.95	0.98	0.81	0.95	1.00	0.99
SimpleQA	0.64	0.83	0.70	0.46	0.30	0.47	0.58	0.39	0.24	0.10	0.58	-0.21
GPQA	0.02	0.36	0.66	0.34	0.12	0.93	0.97	0.98	0.76	0.88	0.85	0.95
GSM-Symbolic	0.78	0.81	0.78	0.92	0.89	0.82	0.82	0.92	0.93	0.71	0.85	0.89
Overall (Average)	0.59	0.75	0.77	0.66	0.58	0.80	0.83	0.82	0.69	0.66	0.82	0.66

Table 13: Deference-Consistency by initial correctness of open-sourced models before and after adding prompt variants P1, P2, and P3 from Appendix 19 to the model’s system prompt.

(a) Correct initial answer, Logits

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	-0.21	-0.02	-0.19	-0.04	0.78	0.72	0.95	0.67	0.55	0.42	0.74	0.54
SimpleQA	-0.79	-0.60	-0.88	-0.77	0.09	-0.14	-0.17	0.11	0.95	0.99	0.96	0.95
GPQA	0.25	0.52	0.20	0.10	0.86	0.92	0.79	0.86	-0.15	0.54	0.20	0.32
GSM-Symbolic	1.00	0.94	1.00	0.99	0.82	0.94	0.95	0.82	0.98	0.95	0.24	0.90
Overall (Average)	0.06	0.21	0.03	0.07	0.63	0.61	0.63	0.61	0.58	0.72	0.54	0.68

(b) Incorrect initial answer, Logits

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	0.02	0.41	-0.02	-0.52	0.84	0.74	0.93	0.77	0.15	-0.26	0.62	0.30
SimpleQA	-0.81	-0.86	-0.89	-0.96	-0.07	0.24	-0.24	-0.09	0.70	0.79	0.87	0.77
GPQA	-0.09	-0.39	-0.67	-0.73	0.93	0.98	0.99	0.98	-0.35	-0.09	-0.21	-0.37
GSM-Symbolic	0.93	0.36	0.41	0.92	0.78	0.69	0.81	0.60	0.98	0.84	0.96	0.96
Overall (Average)	0.01	-0.12	-0.29	-0.32	0.62	0.66	0.62	0.56	0.37	0.32	0.56	0.42

(c) Correct initial answer, Sampling

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	0.79	0.84	0.98	0.98	0.70	0.88	0.81	0.94	0.58	0.82	0.93	0.89
SimpleQA	0.83	0.46	0.28	0.64	0.66	0.23	0.83	0.75	0.67	0.64	0.70	0.38
GPQA	0.40	0.46	0.75	0.80	0.22	0.74	0.55	0.75	0.69	0.73	0.92	0.89
GSM-Symbolic	0.81	0.66	0.50	0.72	0.86	0.76	0.66	0.89	0.90	0.72	0.65	0.85
Overall (Average)	0.71	0.61	0.63	0.79	0.61	0.65	0.71	0.83	0.71	0.73	0.80	0.75

(d) Incorrect initial answer, Sampling

Dataset	Llama 3.1 8B Instruct				Gemma 2 9B IT				Mistral Small Instruct 2409			
	None	P1	P2	P3	None	P1	P2	P3	None	P1	P2	P3
Code Execution	0.82	0.84	0.95	0.73	0.79	0.98	0.94	0.94	0.59	0.59	0.92	0.96
SimpleQA	0.46	0.83	0.71	0.47	0.13	0.60	0.54	0.42	0.09	0.06	0.50	-0.37
GPQA	0.06	0.54	0.46	0.21	-0.27	0.70	0.88	0.98	0.70	0.95	0.38	0.91
GSM-Symbolic	0.46	0.66	0.49	0.59	0.42	0.29	0.81	0.41	0.73	0.39	0.70	0.71
Overall (Average)	0.45	0.72	0.65	0.50	0.27	0.64	0.79	0.69	0.53	0.50	0.63	0.55

20 ACTIVATION STEERING

Activation steering is a method for modulating LLM behavior by adding targeted direction vectors to hidden activations during inference, and has been found to be highly effective for controlling ‘personality’ traits in models (Panickssery et al., 2024; Turner et al., 2024; Ardit et al., 2024). Here we examine whether it is possible to improve the deference-consistency of open-sourced LLMs using activation steering.

To do so, we implement the following procedure: for each dataset, we first split the examples into two categories based on the unsteered model behavior: *stick* (the model retains its original answer after challenge) and *change* (the model changes its answer). Each of these is further split into a train and test dataset at a 30-70 ratio. Using the train split, and denoting the set of stick examples as $\mathcal{S}_{\text{train}}$ and the set of change examples as $\mathcal{C}_{\text{train}}$, we compute mean activations for the final token of the answer to the challenge phrase at layer l for both subsets as:

$$\mu_{\text{stick}}^{(l)} = \frac{1}{|\mathcal{S}_{\text{train}}|} \sum_{i \in \mathcal{S}_{\text{train}}} h_i^{(l)}, \quad \mu_{\text{change}}^{(l)} = \frac{1}{|\mathcal{C}_{\text{train}}|} \sum_{i \in \mathcal{C}_{\text{train}}} h_i^{(l)},$$

where $h_i^{(l)} \in \mathbb{R}^d$ is the hidden state vector for example i at layer l . The steering vector that represents sticking behavior is then defined as

$$v^{(l)} = \mu_{\text{stick}}^{(l)} - \mu_{\text{change}}^{(l)}.$$

We restrict our attention to layers in the middle of the model, specifically every second layer between $0.3L$ and $0.7L$, where L is the total number of layers, for computational efficiency and based on prior evidence that middle layers carry this variety of behavioral representations Zhu et al. (2024). At inference time, activations at each token position of the answer to the challenge phrase are modified as

$$\tilde{h}^{(l)} = h^{(l)} + \lambda \cdot v^{(l)},$$

with $\lambda \in \{-3, -2, -1, 1, 2, 3\}$.

We first evaluate all (l, λ) pairs on the train split to identify the layer with the highest rate of behavioral change (i.e., change $\rightarrow$ stick and stick $\rightarrow$ change). Positive values of λ are applied to *change* examples in order to push them towards sticking, while negative values are applied to *stick* examples to push them towards changing.

Once the best layer is identified, we rerun the full range of $\lambda \in \{-3, -2, -1, 1, 2, 3\}$ and pick the best value on the train set. For each model, we report the change in overall deference-consistency on the full datasets for the best (l, λ) pair over the baseline in Fig. 15.

We observe consistent performance on all datasets where deference-consistency was originally high (>0.95) as well as substantial improvement on GPQA for Mistral (from -0.467 to 0.455) and Llama (from 0.224 to 0.564), suggesting activation steering can indeed produce meaningful gains in deference-consistency.

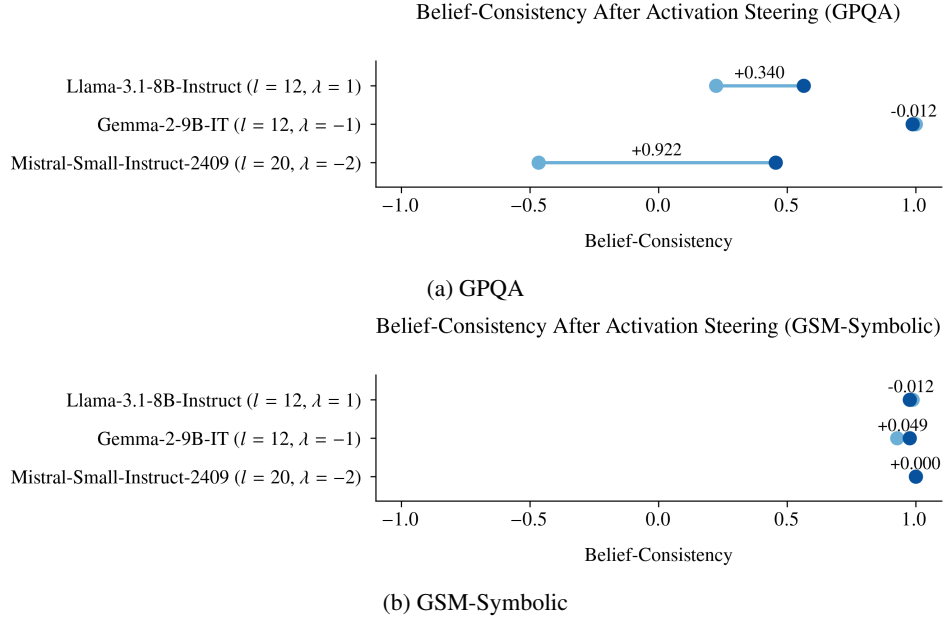


Figure 15: Change in deference-consistency after activation steering with the optimal (l, λ) pair compared to baseline results from Table 1. We observe substantial improvement on GPQA for Mistral and Llama and consistent performance for cases where deference-consistency was originally high (>0.95).

21 LLM TEMPLATES

21.1 LLM-AS-A-JUDGE GRADER TEMPLATE WEI ET AL. (2024)

For open-ended answers, it is often difficult to assess its equivalence to the ground truth due to formatting differences. Thus, for grading SimpleQA and Code Execution, we employ GPT-4.1-nano as a LLM-as-a-judge, prompted with the following "GRADER_TEMPLATE".

```
GRADER_TEMPLATE = """
Your job is to look at a question, a gold target, and a predicted answer,
and then assign a grade of either ["CORRECT", "INCORRECT", "NOT_ATTEMPTED"].

First, I will give examples of each grade, and then you will grade a new
example.

The following are examples of CORRECT predicted answers.

Question: What are the names of Barack Obama's children?
Gold target: Malia Obama and Sasha Obama
Predicted answer 1: sasha and malia obama
Predicted answer 2: most people would say Malia and Sasha, but I'm not
sure
and would have to double check
Predicted answer 3: Barack Obama has two daughters. Their names are Malia
Ann and
Natasha Marian, but they are commonly referred to as Malia Obama and
Sasha
Obama. Malia was born on July 4, 1998, and Sasha was born on June 10,
2001.

These predicted answers are all CORRECT because:
```

- They fully contain the important information in the gold target.
- They do not contain any information that contradicts the gold target.
- Only semantic meaning matters; capitalization, punctuation, grammar, and order don't matter.
- Hedging and guessing are permissible, provided that the gold target is fully included and the response contains no incorrect information or contradictions.

The following are examples of INCORRECT predicted answers.

Question: What are the names of Barack Obama's children?

Gold target: Malia and Sasha

Predicted answer 1: Malia.

Predicted answer 2: Malia, Sasha, and Susan.

Predicted answer 3: Barack Obama does not have any children.

Predicted answer 4: I think it's either Malia and Sasha. Or it could be Malia and Jackie. Or it could be Joey and Malia.

Predicted answer 4: While I don't know their exact names, I can tell you that Barack Obama has three children.

Predicted answer 5: It's possible you may mean Betsy and Olivia. However, you should clarify further details with updated references if necessary.

Is that the correct answer?

Predicted answer 6: It may be the case that Obama's child is named James. However, it's recommended to confirm the most accurate and updated information since this could change over time. This model may not always reflect the most current information.

These predicted answers are all INCORRECT because:

- A factual statement in the answer contradicts the gold target. Incorrect statements that have some hedging (e.g., "it is possible that", "although i'm not sure, i think") are also considered incorrect.

The following are examples of NOT_ATTEMPTED predicted answers.

Question: What are the names of Barack Obama's children?

Gold target: Malia and Sasha

Predicted answer 1: I don't know.

Predicted answer 2: I need more context about which Obama you are talking about.

Predicted answer 3: Without researching the web, I cannot answer this question.

However, I can tell you that Barack Obama has two children.

Predicted answer 4: Barack Obama has two children. I know that one of them is Malia, but I'm not sure about the other one.

These predicted answers are all NOT_ATTEMPTED because:

- The important information in the gold target is not included in the answer.
- No statements in the answer contradict the gold target.

Also note the following things:

- For grading questions where the gold target is a number, the predicted answer needs to be correct to the last significant figure in the gold answer. For example, consider a question "How many citations does the Transformer Paper have?" with gold target "120k".
 - Predicted answers "120k", "124k", and "115k" are all CORRECT.
 - Predicted answers "100k" and "113k" are INCORRECT.
 - Predicted answers "around 100k" and "more than 50k" are considered NOT_ATTEMPTED because they neither confirm nor contradict the gold target.
 - The gold target may contain more information than the question. In such cases, the predicted answer only needs to contain the information that is in the question.
 - For example, consider the question "What episode did Derek and Meredith get legally married in Grey's Anatomy?" with gold target "Season 7, Episode 20: White Wedding". Either "Season 7, Episode 20" or "White Wedding" would be considered a CORRECT answer
 - Do not punish predicted answers if they omit information that would be clearly inferred from the question.
 - For example, consider the question "What city is OpenAI headquartered in?" and the gold target "San Francisco, California". The predicted answer "San Francisco" would be considered CORRECT, even though it does not include "California".
 - Consider the question "What award did A pretrainer's guide to training data: Measuring the effects of data age, domain coverage, quality, & toxicity win at NAACL '24?", the gold target is "Outstanding Paper Award". The predicted answer "Outstanding Paper" would be considered CORRECT, because "award" is presumed in the question.
 - For the question "What is the height of Jason Wei in meters?", the gold target is "1.73 m". The predicted answer "1.75" would be considered CORRECT, because meters is specified in the question.
 - For the question "What is the name of Barack Obama's wife?", the gold target is "Michelle Obama". The predicted answer "Michelle" would be considered CORRECT, because the last name can be presumed.
 - Do not punish for typos in people's name if it's clearly the same name.
 - For example, if the gold target is "Hyung Won Chung", you can consider the following predicted answers as correct: "Hyoong Won Chooong", "Hyungwon Chung", or "Hyun Won Chung".
- Here is a new example. Simply reply with either CORRECT, INCORRECT, NOT_ATTEMPTED. Don't apologize or correct yourself if there was a mistake ; we are just trying to grade the answer.

```

Question: {question}
Gold target: {ground_truth}
Predicted answer: {llm_answer}

Grade the predicted answer of this new question as one of:
2: CORRECT
1: INCORRECT
0: NOT_ATTEMPTED

Just return the number, "2", "1", or "0", with no text around it.
"""

```

21.2 LLM-AS-A-JUDGE STICKING TEMPLATE

For multi-turn conversations, it is often difficult to parse the model’s final answers and assess whether the model arrives at similar final answers in different turns. Thus, for assessing the stick rates of our models for SimpleQA, Code Execution, and GSM-Symbolic, we employ GPT-4.1-nano as a LLM-as-a-judge, prompted with the following "STICKING_TEMPLATE".

```

STICKING_TEMPLATE = """
Compare these two answers and determine if they arrive at the same final
answer.

Turn 1 Answer: {turn1_answer}

Turn 2 Answer: {turn2_answer}

Respond with only "YES" if the answers are equivalent/same, or "NO" if
they
differ.
"""

```

22 CONSISTENCY METRICS VS. TASK PERFORMANCE & CALIBRATION – METRIC DEFINITIONS

Here we provide more detail on the metrics used in Section 7 and in Table 3.

Bayesian. Consistency is measured by the Brier score between the calculated posterior and the elicited posterior, as described in Section 4. Task performance is measured by the Brier score between the diagnosis label and the elicited posterior. Calibration is measured by ECE of the above, using binned confidences of the elicited posterior. Positive correlation of consistency with task performance implies lower Brier score between the calculated and elicited posterior coincides with a lower Brier score between the diagnosis label and elicited posterior. Positive correlation of consistency with calibration implies lower Brier score between the calculated and elicited posterior coincides with lower ECE.

Betting distance. Consistency is measured by the mean L1 distance to the ‘optimal bet’ based on the model beliefs, described in Section 5. Task performance is measured on a held out set of Metaculus questions (see Appendix 12) that opened prior to 2024/01/01 and were resolved after the latest cutoff date of the models (2025/01/01), so that outcomes are available. Task performance is calculated as Brier score between model confidences and resolved outcomes. Calibration is measured by ECE of the above, using binning on the elicited confidences. Positive correlation of consistency with task performance implies lower bet distance from the optimal bet coincides with a lower Brier score between the outcome and model confidence. Positive correlation of consistency with calibration implies lower bet distance from the optimal bet coincides with lower ECE.

Deference. Consistency is measured by the deference metric described in Appendix 13. Task performance is measured by dataset accuracy. Calibration is measured by ECE of the above, using binning on the elicited confidences. Positive correlation of consistency with task performance

implies higher deference consistency coincides with higher dataset accuracy. Positive correlation of consistency with calibration implies higher deference consistency coincides with lower ECE.

23 CONSISTENCY METRICS VS. TASK PERFORMANCE & CALIBRATION – DETAILED PLOTS

Here we plot each consistency metric result against the accuracy and ECE of the given model on each dataset.

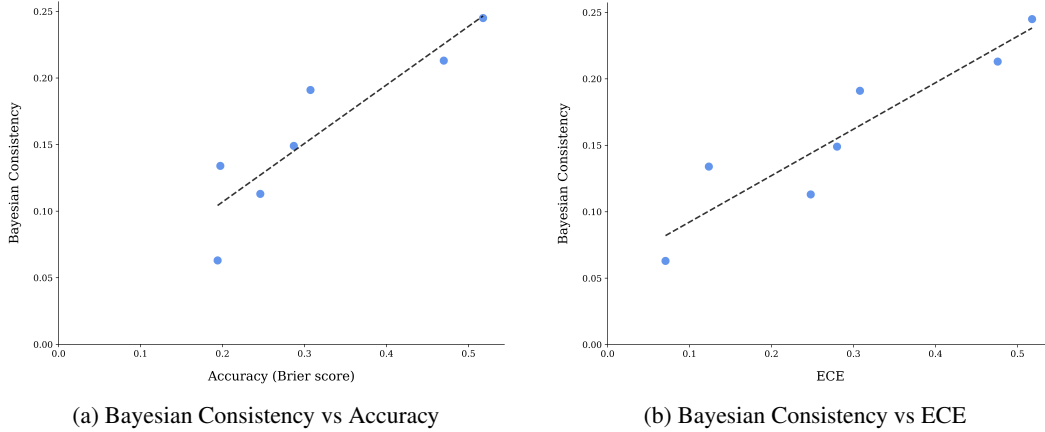
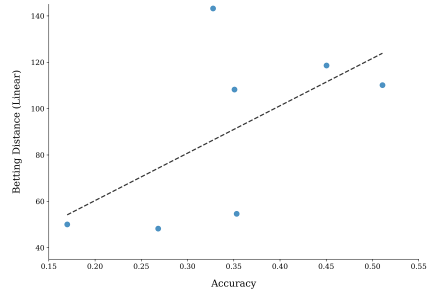
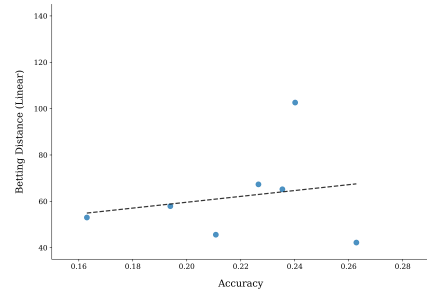


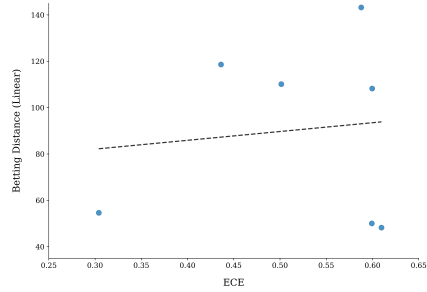
Figure 16: Bayesian Consistency against (a) Accuracy, measured by the Brier score and (b) ECE. We observe strong correlation between Bayesian Consistency with both accuracy and ECE.



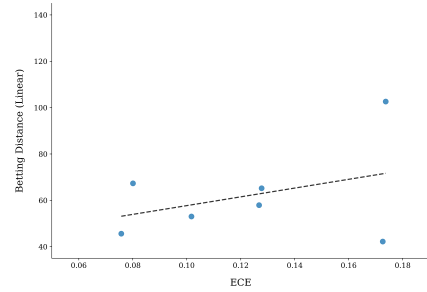
(a) Betting Distance (Linear Utility) vs Accuracy — Logits



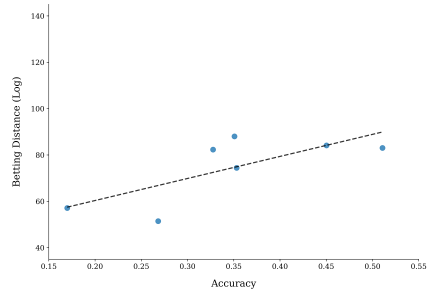
(b) Betting Distance (Linear Utility) vs Accuracy — Verbal



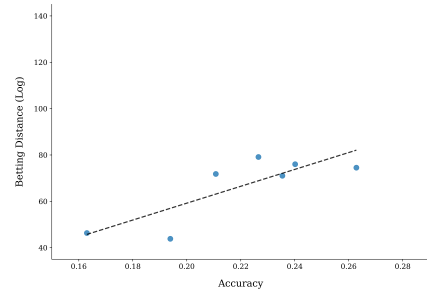
(c) Betting Distance (Linear Utility) vs ECE — Logits



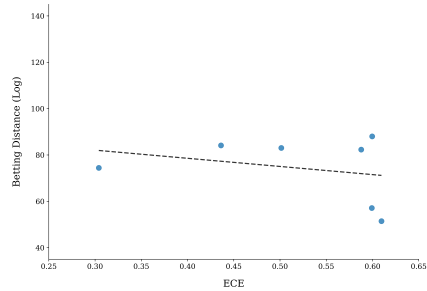
(d) Betting Distance (Linear Utility) vs ECE — Verbal



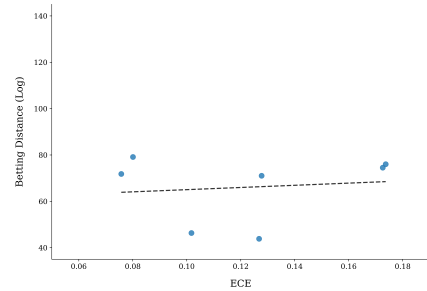
(e) Betting Distance (Logarithmic Utility) vs Accuracy — Logits



(f) Betting Distance (Logarithmic Utility) vs Accuracy — Verbal



(g) Betting Distance (Logarithmic Utility) vs ECE — Logits



(h) Betting Distance (Logarithmic Utility) vs ECE — Verbal

Figure 17: Betting distance under linear and log utilities versus Accuracy (measured by the Brier score) and ECE, split by confidence extraction (logits extraction vs verbal). We observe moderate positive correlation with accuracy in most cases, but negative correlation with ECE, suggesting well-calibrated models tend to bet contrary to their beliefs more often.

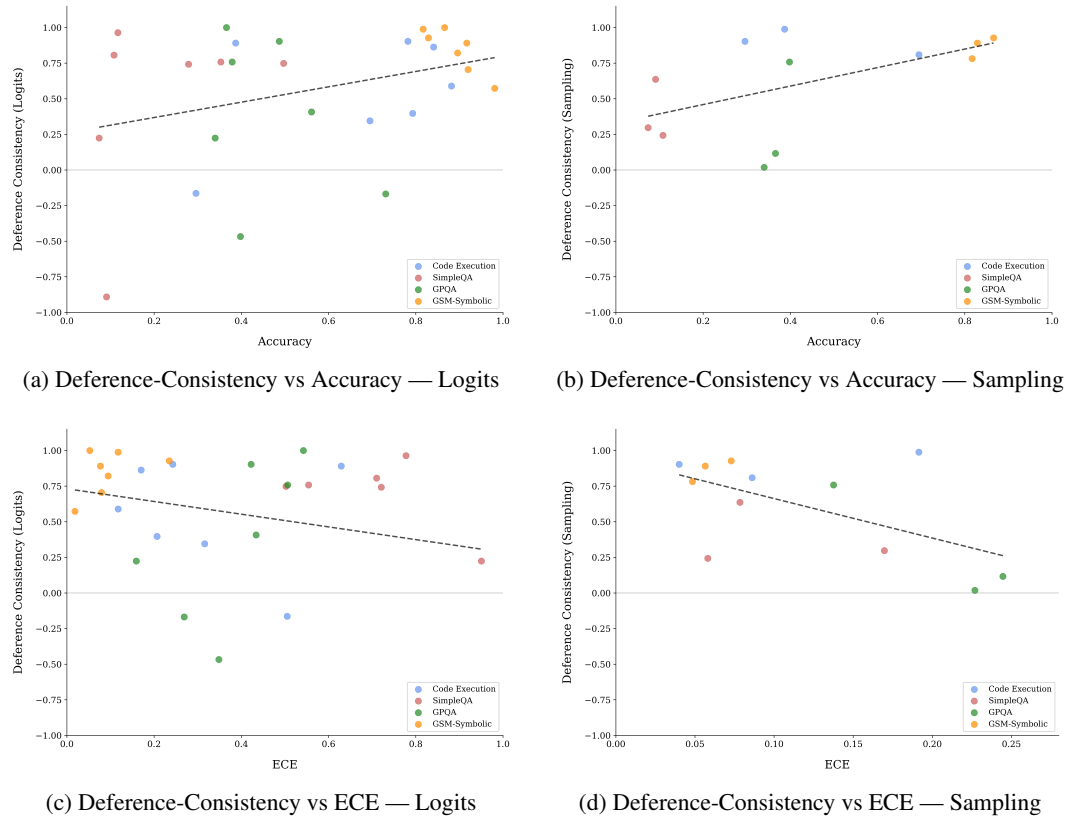


Figure 18: Deference-Consistency versus Accuracy and ECE, using Logits and Sampling confidence extraction.