

Exploring Mental State Representations in Language Models

Anonymous ACL submission

Abstract

While numerous works have assessed the generative performance of language models (LMs) on tasks requiring Theory of Mind reasoning, research into the models’ *internal representation* of mental states remains limited. Recent work has used probing to demonstrate that LMs can represent beliefs of themselves and others. However, these claims are accompanied by limited evaluation, making it difficult to assess how mental state representations are affected by model design and training choices. We report extensive experiments with different LMs and prompt designs to study the robustness of mental state representations. Our results show that the quality of models’ internal representations of the beliefs of others increases with model size and, more crucially, with fine-tuning. We are the first to study how prompt variations impact probing performance on Theory of Mind tasks. We find that models’ representations are sensitive to prompt variations, even when such variations should be beneficial. Finally, we complement previous activation editing experiments on Theory of Mind tasks and show that it is possible to improve models’ reasoning performance by steering their activations without the need to train any probe.

1 Introduction

Modern language models (LMs) trained on next token prediction have demonstrated impressive capabilities, spanning coding, mathematical reasoning, fact verification, and embodied interaction (Wei et al., 2022; Bubeck et al., 2023). As these models are designed with the ultimate goal of collaborating with humans, it becomes imperative that they complement these skills with an understanding of humans. Core to this understanding is *Theory of Mind* (ToM) – the ability to attribute mental states to oneself and others (Premack and Woodruff, 1978). ToM is essential for effective communication and cooperation with other agents, facilitating

interaction and learning from feedback and demonstrations (Saha et al., 2023). Given its significance, ToM has emerged as a critical milestone in AI and an important capability when evaluating cutting-edge LMs (Bubeck et al., 2023). Interest in LMs’ generative performance on tasks requiring ToM reasoning has resulted in a wide variety of benchmark datasets, framed as question-answering tasks (Le et al., 2019; Gandhi et al., 2023; Kim et al., 2023; He et al., 2023; Tan et al., 2024; Xu et al., 2024).

Despite showing improved performance on ToM benchmarks compared to earlier models, modern LMs are still far from perfect (Sap et al., 2022). Text generated by LMs often contains errors that limit their performance on ToM tasks (Martindale et al., 2019). Previous work has shown that it is sometimes possible to still obtain correct predictions by *probing* LMs’ internal representations (Li et al., 2021; Liu et al., 2023; Gurnee et al., 2023). In particular, Zhu et al. (2024) have shown that LMs, when prompted with a story and a belief statement, can represent beliefs from their own perspective and, to a lesser extent, from the perspective of a character in the story. However, this work is limited in the number settings studied, leaving several questions unanswered.

We explore belief representations of self and others in language models through extensive experiments of different LM families, model sizes, fine-tuning approaches, and prompts. Specifically, we design a set of experiments to address the following research questions: **RQ1.** What is the relation between model size and probing accuracy? **RQ2.** Does supervised fine tuning (Wei et al., 2021, SFT) and/or reinforcement learning from human feedback (Christiano et al., 2017; Ouyang et al., 2022, RLHF) have an effect on probing accuracy? **RQ3.** Are models’ internal representations of beliefs sensitive to prompt variations? **RQ4.** Due to the large dimensionality of LM representations, are the probes just fitting irrelevant patterns in the

data? **RQ5.** Can we enhance LMs’ performance by editing their activations without training dedicated probes?

To answer these research questions, we perform experiments on two families of LMs, Llama-2 (Touvron et al., 2023), and Pythia (Biderman et al., 2023). We first compare the probing performance of pre-trained models with models that have been fine-tuned using SFT and/or RLHF. Our experiments reveal that when predicting others’ belief, probing accuracy increases with model size and, more crucially for smaller models, with fine-tuning (RQ1, RQ2). We then explore, for the first time, the sensitivity of LMs’ representations to prompting in the context of ToM. Our experiments with four different prompt variations (*Random*, *Misleading*, *Time Specification*, and *Initial Belief*) demonstrate that models’ representations are sensitive to prompt variations (RQ3). We also find no strong evidence of spurious memorisation in the probes, as it is possible to recover most of the accuracy by training probes on a much small subset of principal components of models’ representations (RQ4). Finally, we show that by using contrastive activation addition (Rinsky et al., 2023, CAA), we can steer models’ activations without the need to train any probe and, in a generalisable way, obtain significant performance improvements across different ToM tasks (RQ5).

In summary, our work makes the following contributions:

1. We report extensive probing experiments with various types of LMs with different model sizes and fine-tuning approaches, showing that the quality of models’ internal representations of the beliefs of others increases with model size and, more crucially, fine-tuning.
2. We are the first to study how prompt variations impact belief probing performance, showing that models’ representations are sensitive to prompt variations.
3. We show that by using contrastive activation addition it is possible to improve models’ reasoning performance by steering their activations without the need to train any probe.

2 Related Work

Machine Theory of Mind Theory of mind has been studied in AI for more than a decade (Baker et al., 2009; Rabinowitz et al., 2018; Bara et al.,

2021; Bortoletto et al., 2024a,b,c). Recent advances in LMs have sparked interest in evaluating their ToM capabilities. Various benchmarks have been proposed, aiming to measure LMs’ ability to understand and reason about the beliefs, goals, and intentions of others (Le et al., 2019; He et al., 2023; Kim et al., 2023; Gandhi et al., 2023; Xu et al., 2024; Tan et al., 2024; Sclar et al., 2023; Ma et al., 2023b; Wu et al., 2023). Additionally, efforts have been made to enhance LMs’ ToM through prompting techniques (Zhou et al., 2023b; Moghaddam and Honey, 2023; Wilf et al., 2023). A new direction of research explores LMs’ internal representation of mental states. Zhu et al. (2024) demonstrated that LMs linearly encode beliefs from different agents’ perspectives, and manipulating these representations can enhance ToM task performance. Our work dives deeper into LMs’ internal belief representations, offering a broader insight into these mechanisms.

Probing Neural Representations Initially proposed by Alain and Bengio (2017), probing has emerged as a common method for determining if models represent particular features or concepts. In the realm of LMs, numerous works used probing to demonstrate that these models acquire rich linguistic representations – spanning semantic concepts such as syntactic categories, dependency relations, co-reference, and word meaning (Conneau et al., 2018; Tenney et al., 2018, 2019; Rogers et al., 2021; Li et al., 2021; Hernandez and Andreas, 2021; Marks and Tegmark, 2023; Liu et al., 2023). A separate line of work explored if LMs possess a *world model* (Li et al., 2021; Abdou et al., 2021; Patel and Pavlick, 2022; Li et al., 2023a; Nanda et al., 2023). An emergent line of work that is relevant to our work used probing to explore if LMs have *agent models*, for example, if they can represent beliefs of self and others (Zhu et al., 2024; Bortoletto et al., 2024a). While representing an important first step towards understanding the internals of ToM in LMs, experiments in (Zhu et al., 2024) are limited in settings and models considered. In this work, we contribute with extensive experiments that employ a wider variety of LMs and a wider range of settings.

Prompt Analysis Previous work has shown that LMs are vulnerable to prompt alterations like token deletion or reordering (Ishibashi et al., 2023), biased or toxic prompts (Shaikh et al., 2023) and similarity to training data (Razeghi et al., 2022).

On the other hand, instruction-tuned models have proved to be more robust against prompt variation, even when using misleading instructions (Webson and Pavlick, 2022). Other works have shown the importance of input-output format (Min et al., 2022) and of demonstration example ordering for few-shot performance (Zhao et al., 2021; Lu et al., 2022; Zhou et al., 2023a). In this work, we shift our focus from analysing how sensitive model outputs are to how model *representations* change. Our work, along with (Gurnee et al., 2023), is one of the first to explore how prompt design affects how accurately models represent concepts. In particular, Gurnee et al. (2023) have studied whether LMs’ representations of space and time are robust to prompt variations. In stark contrast, we explore for the first time the effect of prompt variations on how models represent mental states internally.

Activation Editing Activation editing has emerged as an alternative way to influence model behaviour without any additional fine-tuning (Li et al., 2023a; Hernandez et al., 2023). This approach involves manipulating the internal representations of models to direct their outputs towards desired outcomes. One notable method in this domain is inference-time intervention (Li et al., 2023b, ITI), which has been proposed to enhance truthfulness in LMs. ITI involves training linear probes on contrastive question-answering datasets to identify “truthful” attention heads and then shifting attention head activations during inference along the identified truthful directions. In contrast, activation addition (Turner et al., 2023, AA) and contrastive activation addition (Rimsky et al., 2023, CAA) generate steering vectors by only using LMs’ activations. Zhu et al. have used ITI to show that it is possible to manipulate LMs’ internal representations of mental states. In this work, we show that using CAA can further improve LMs’ ToM capabilities while eliminating the need for a fine-grained search over attention heads.

3 Experimental Setup

3.1 Probing

We linearly decode belief status from the perspective of different agents by using probing (Alain and Bengio, 2017). Probing involves localising specific concepts in a neural model by training a simple classifier (called a *probe*) on model activations to predict a target label associated with the input data. To provide a formal definition, we adopt a similar

notation to the one introduced in (Belinkov, 2022). Let us define an *original model* $f : x \mapsto \hat{y}$ that is trained on a dataset $\mathcal{D}^O = \{x^{(i)}, y^{(i)}\}$ to map input x to output $\hat{y}$. Model performance is evaluated by some measure, denoted $\text{PERF}(f, \mathcal{D}^O)$. A *probe* $g_l : f_l(x) \mapsto \hat{z}$ maps intermediate representations of x in f at layer l to some property $\hat{z}$, which is the label of interest. The probe g_l is trained on a *probing dataset* $\mathcal{D}^P = \{x^{(i)}, z^{(i)}\}$ and evaluated using some performance measure $\text{PERF}(g_l, f, \mathcal{D}^O, \mathcal{D}^P)$. In our case, f is an autoregressive language model that given a sequence of tokens x outputs a probability distribution over the token vocabulary to predict the next token in the sequence. Our probe is a logistic regression model $g_l : \hat{z} = Wa_l + b$ trained on neural activations $f_l(x) = a_l$ to predict binary belief labels $y = \{0, 1\}$.

3.2 Dataset

Following Zhu et al. (2024) we use the BigToM benchmark (Gandhi et al., 2023). BigToM is constructed using GPT-4 (Achiam et al., 2023) to populate causal templates and combine elements from these templates. Each causal template is set up with a *context* and a description of the *protagonist* (e.g. “Noor is working as a barista [...]”), a *desire* (“Noor wants to make a cappuccino”), a *percept* (“Noor grabs a milk pitcher and fills it with oat milk”), and a *belief* (“Noor believes that the pitcher contains oat milk”). The state of the world is changed by a causal event (“A coworker swaps the oat milk in the pitcher with almond milk”). The dataset constructs different conditions by changing the percepts of the protagonist after the causal event, which will result in different beliefs. In this work, we focus on the *Forward Belief* setting proposed by (Zhu et al., 2024) in which models have to infer the belief of the protagonist given the percepts of the causal event, $P(\text{belief}|\text{percepts})$. We report additional details in Appendix A.1.1

Probing Datasets We consider two probing datasets: $\mathcal{D}_p^P = \{x_p^{(i)}, z_p^{(i)}\}$, where the labels $z_p^{(i)}$ correspond to ground-truth beliefs from the *protagonist* perspective, and $\mathcal{D}_o^P = \{x_o^{(i)}, z_o^{(i)}\}$, where the labels $z_o^{(i)}$ reflect the perspective of an omniscient *oracle*. $\mathcal{D}_p^P$ and $\mathcal{D}_o^P$ are built by pairing each story in BigToM with a belief statement, as shown in Figure 1. After prompting the model with a story-belief pair x we cache the residual stream activations $f_l(x)$ at the final token position for all residual streams (see Figure 5).

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task.

Noor does not see her coworker swapping the milk.
Belief: The milk pitcher contains almond milk.
 $z_o = \text{True}, z_p = \text{False}$

Noor sees her coworker swapping the milk.
Belief: The milk pitcher contains almond milk.
 $z_o = \text{True}, z_p = \text{True}$

Figure 1: Example of false belief from our probing datasets. The labels y_p and y_o correspond to $\mathcal{D}_p^P$ and $\mathcal{D}_o^P$, respectively. By manipulating the protagonist’s **percepts** after the causal event we obtain two scenarios: **true belief** and **false belief**.

3.3 Models

We study two families of LMs that offer us options in model sizes and fine-tuning: Pythia (Biderman et al., 2023) and Llama-2 (Touvron et al., 2023). While Llama-2 offers “chat” versions first trained with supervised fine-tuning (SFT) and then RLHF, Pythia’s open-source training set (Gao et al., 2020) ensures that there is no data leakage.¹ Additionally, we consider a SFT version of Pythia-6.9B trained on open-source instruction datasets (Wang et al., 2024), which we refer to as Pythia-6.9B-chat.² We provide a summary of the models in Table 2.

3.4 Probing Experiments

We aim to study how LMs represent beliefs of self and others by proposing a set of extensive probing experiments across LMs that differ in architecture, size, and fine-tuning approach. Our approach is generally similar to the one used by previous work (Zhu et al., 2024), but we make a different operational choice: We train probes on the residual stream instead of attention heads. We opted to use the residual stream as it integrates information from both the attention and feed-forward components, potentially encoding richer representations. Additionally, since the residual activations directly contribute to the final output predictions, probing them may better align with understanding the model’s

¹Llama-2 was released later than BigToM.

²<https://huggingface.co/allenai/open-instruct-pythia-6.9b-tulu>

behaviour for downstream tasks.

Model Size and Fine-tuning We first report experiments to better understand the effect of model size and fine-tuning on belief probing accuracy. Specifically, we ask the following questions: *Is there a relation between model size and probing accuracy?* (RQ1) *Does fine-tuning an LM with instruction-tuning or RLHF have an effect on probing accuracy?* (RQ2) To answer these questions we performed the same probing experiment across all our models and compared the results.

Sensitivity to Prompting By using a single prompt design, previous work left the impact of prompt design on probing accuracy unclear (Zhu et al., 2024). Our second set of experiments aims to explore how belief representations are sensitive to different prompts. Research on prompt robustness in language models is still in its infancy and focused mainly on revealing vulnerability to prompt alternations on downstream performance (Min et al., 2022; Ishibashi et al., 2023; Shaikh et al., 2023; Leiding et al., 2023; Sclar et al., 2024). In contrast, we study how the input influences models’ representations by asking: *Are models’ internal belief representations robust to prompt variations?* (RQ3) To answer this question we define four prompt variations:

- **Random:** Following Gurnee and Tegmark (2024), we add 10 random tokens to the belief statement.
- **Misleading:** Each story is followed by two belief statements, one pertinent to the story and one randomly chosen from another.
- **Time Specification:** The prompt specifies that the belief statement refers to the end of the story. We study this variation because some belief statements can be true (false) at the story’s beginning but false (true) at the end. For example, consider the story in Figure 1: if Noor does not witness the swap, in the end, she will believe the pitcher contains almond milk ($y_p = \text{True}$). However, if the same belief is referred to at the beginning of the story, then it is false ($y_p = \text{False}$).
- **Initial Belief:** We explicitly reveal the protagonist’s initial belief (e.g. “Noor believes that the pitcher contains oat milk”) in the story to test whether it biases the representations of LMs.

While all maintaining conceptual and semantic parity with the *Original* prompt used in (Zhu et al.,

2024), *Random* and *Misleading* are expected to negatively impact LMs’ representations, while *Time Specification* and *Initial Belief* are supposed to have a positive influence. Robust representations of beliefs should exhibit minimal sensitivity to these alterations. Our experiments compare probe accuracy across different model sizes, fine-tuning, and prompt variations. Examples of prompts are reported in Appendix A.1.4.

Dimensionality Reduction The probes we train have a significant number of learnable parameters – up to 16,385 for Llama-2-70B. This raises the concern that probes might learn to rely on irrelevant patterns in the data instead of capturing meaningful relationships (Alain and Bengio, 2017). Our final set of probing experiments answers the following question: *Are the probes memorising irrelevant patterns in the training data?* (RQ4) To answer this question, before training the probes, we project the probing datasets $\mathcal{D}_p^P$ and $\mathcal{D}_o^P$ onto their k largest principal components using PCA. This procedure significantly reduces the number of learnable parameters in the probes, minimizing the risk of them relying on spurious patterns in the data.

3.5 Contrastive Activation Addition

Our final set of experiments builds upon the findings of Zhu et al. (2024), who showed that employing trained probes with inference time intervention (Li et al., 2023b, ITI) could enhance LMs’ performance on ToM tasks. We take a step further and ask: *Can we enhance LMs’ performance by manipulating their activations without the need for training dedicated probes?* (RQ5) To find an answer we use contrastive activation addition (Rimsky et al., 2023, CAA), an extension of activation addition (Turner et al., 2023, AA) that computes *steering vectors* to control LMs’ behaviour. Steering vectors are computed as the average difference in residual stream activations between pairs of positive and negative instances of a specific behaviour. Formally, given a dataset $\mathcal{D}$ of triplets (p, c_p, c_n) , where p is a prompt, c_p is a positive completion, and c_n is a negative completion, CAA computes a *mean difference* vector v_l^{md} for layer l as:

$$v_l^{md} = \frac{1}{|\mathcal{D}|} \sum_{p, c_p, c_n} a_l(p, c_p) - a_l(p, c_n)$$

During inference, these steering vectors are multiplied with an appropriate coefficient α and added at every token position of the generated text after the

prompt. CAA has two main advantages over ITI: First, it eliminates the need to train probes. Second, it operates at the residual stream level, making it easier to use than methods that intervene on specific attention heads like ITI. While CAA has been used to control alignment-relevant behaviour, such as hallucinations, refusal, and sycophancy (Rimsky et al., 2023), we are the first to apply it to enhance LMs’ ToM reasoning. This can be understood as isolating the direction in the LMs’ latent space corresponding to taking the perspective of another agent. To evaluate both base and fine-tuned LMs, we rank their answers to the ToM questions according to $p_{LM}(a|q)$ (Petroni et al., 2019). We adopt the *Forward Belief* task split used in (Zhu et al., 2024) to compute the steering vectors. Additionally, we evaluate the transferability of the CAA steering vectors by applying them to two other BigToM tasks: *Forward Action* and *Backward Belief*. We provide details about these tasks in Appendix A.1.1, and a more detailed explanation of how ITI works in Appendix A.5.

4 Results

Effect of Model Size and Fine-tuning Results from our study on model size and fine-tuning are shown in Figure 2. When considering *oracle* beliefs, probing accuracy rapidly converges to 100, with larger models showing faster convergence rates. The smallest Pythia-70m that performs slightly worse but still achieves 95% accuracy despite having less than 0.6% of the parameters of Pythia-12B. This finding suggests that even small LMs can effectively represent beliefs from an omniscient perspective.

For *protagonist* beliefs, accuracy also increases with model size, although there is a performance gap between Llama-2 and Pythia. For example, Llama2-13B reaches around 80%, while Pythia-12B achieves approximately 60%. This gap is likely due to Llama-2 being trained on nearly seven times more tokens than Pythia. The figure also shows that accuracy at early layers is particularly low across all models. We speculate that this is due to the initial coding strategy of LMs that uses the first layers to combine individual tokens into more semantically meaningful representations (Gurnee et al., 2023). Probes on fine-tuned LMs show significantly better accuracy with improvements of up to 29% for Llama2-7B-chat and 26% for Pythia-6.9B-chat with respect to their base version. Fine-

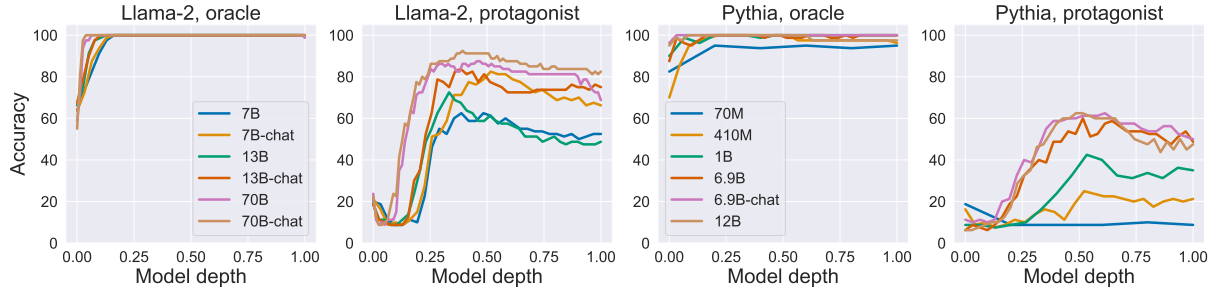


Figure 2: Belief probing accuracy across models with different architecture, size and fine-tuning.

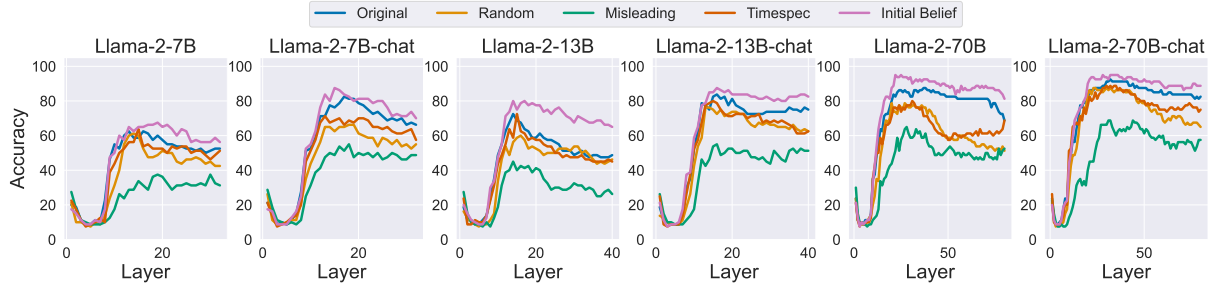


Figure 3: Sensitivity of protagonist belief probing accuracy to different prompt variations.

tuned 7B LMs outperform (Llama-2) or are on par (Pythia) with twice as large base models (12/13B), highlighting the importance of fine-tuning in developing representations of others’ beliefs. This resonates with cognitive psychology findings that ToM development is closely linked to social communication (Tomasello, 2010; Sidera et al., 2018; Ma et al., 2023a), which instruction-tuning and RLHF may help induce in LMs. For larger LMs, the improvements from fine-tuning decrease as model size increases (Figure 6a). We characterise the relationship between probe accuracy and model size in Figure 6, where we consider the *best* probe accuracy for every LM, i.e. the highest accuracy among probes $\{g_l\}$ trained on $\{a_l\}$ for a LM f . For Llama-2 base, the best probe accuracy scales logarithmically with model size ($R^2 = 0.98$, cf. Figure 6b), whereas for fine-tuned models it scales linearly ($R^2 = 1.0$, cf. Figure 6c). For Pythia base, the best probe accuracy also scales logarithmically with model size ($R^2 = 0.96$, cf. Figure 6d).

Sensitivity to Prompting Figure 3 compares *protagonist* probe accuracy across various prompt variations for Llama2 models. As can be seen from the figure, providing the protagonist’s *Initial Belief* in the story yields higher probe accuracy compared to the *Original* prompt (Figure 1). Accuracy for all the other prompt variations is generally lower than *Original*. On one hand, misleading prompts

hurt performance across all models. This finding resonates with Webson and Pavlick (2022) who found that instruction-tuned models, despite being more robust, are still sensitive to misleading prompts. On the other hand, *Time Specification* unexpectedly does not help in disambiguating belief states in different time frames, as we hypothesised in §3.4. Additionally, models show sensitivity to *Random* tokens placed before the belief statement. Pythia models show similar patterns results, provided in Figure 7. Results for *oracle* beliefs are reported in Figure 8 and indicate that models maintain high accuracy. *Misleading* prompts slightly reduce performance to around 95%. In summary, these experiments show that LMs possess robust belief representations when taking an omniscient perspective, whereas their representations of others’ beliefs are more susceptible to prompt variations.

Dimensionality Reduction Figure 4 shows the probe accuracy on *protagonist* when training the probes on the top k principal components of Llama-2’s internal activations. We provide results for Pythia in Figure 9, and for all models on *oracle* settings in Figure 10. We consider $k = \{2, 10, 100, 1000\}$, spanning several orders of magnitude.³ For all models, it is generally possible to recover most of the original accuracy by training

³For models with hidden dimensions smaller than 1000, we skip this value.

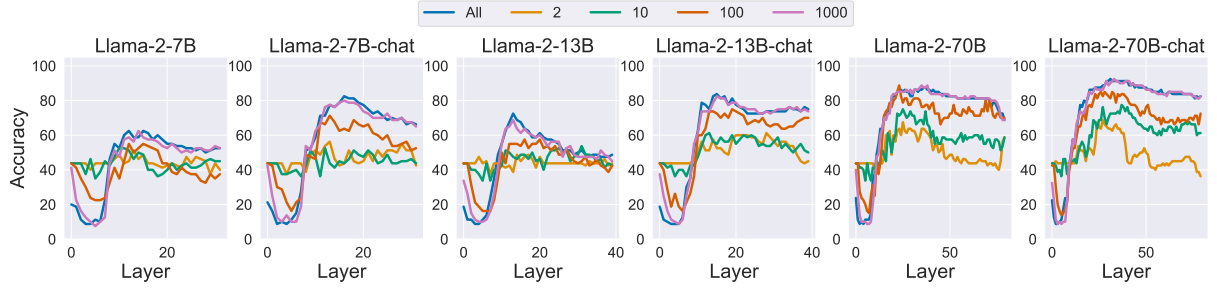


Figure 4: We compare the probing accuracy obtained by using the original set of activations (All) with the accuracy obtained by considering only the first $n = \{2, 10, 100, 1000\}$ principal components. For Llama2: All(7b) = 4096, All(13b) = 5120, All(70b) = 8192. We report results for *protagonist* beliefs. Results for *oracle* are shown in Figure 10.

probes on a number k of principal components of the activations that is more than one order of magnitude smaller than the full dimensionality. This suggests that belief representations are embedded in a low-dimensional manifold $\mathcal{B}$ spanned by the top k eigenvectors $\{v_1, \dots, v_k\}$ of the covariance matrix $\mathbf{C} = \mathbb{E}[(a - \mathbb{E}[a])(a - \mathbb{E}[a])^\top]$, and provides a clear indication that the probes measure meaningful representations rather than spurious patterns.

Contrastive Activation Addition We finally compare models’ accuracy on three BigToM tasks in Table 1. Each model has been evaluated three times: without any intervention, using ITI, and using CAA. Hyperparameter details can be found in Appendix A.6. Note that we use steering vectors computed using the *Forward Belief* task for all three tasks to test their generalisability.

Performance without intervention is generally lower across tasks and model sizes, with the larger Llama-2-70B and Llama-2-70B-chat models exhibiting higher accuracy. Performance for Pythia models of different sizes does not change much, with the fine-tuned Pythia-6.9B-chat often showing better performance on single true belief (TB) and false belief (FB) tasks but not on their conjunction (Both). ITI demonstrates modest improvements over no intervention for Llama-2 models. Improvements for Pythia models are consistent and higher, up to +17. The only exception is Pythia-6.9B-chat, for which ITI is not always beneficial.

CAA consistently delivers the most substantial accuracy improvements across all models and tasks, up to +56 for Llama-2-13B-chat on the (*Backward Belief*), which Gandhi et al. have identified as the hardest task. Despite its relatively small size, Llama-2-13B-chat excels in all three tasks when using CAA. Larger 70B models often achieve ac-

curacies close to or exceeding 90%. Smaller models like Pythia-70M and Pythia-410M also show significant gains with CAA, though the absolute performance is still lower than Llama-2. Overall, our results indicate that it is possible to effectively enhance ToM reasoning in LMs without needing to train any probe, which yields even improved results. Furthermore, we show that CAA steering vectors generalise well, yielding substantial performance gains across all ToM tasks. To further demonstrate CAA’s effectiveness, we applied it while evaluating models on a control task where the causal event is replaced by a random one that does not change the environment (e.g., *A musician starts playing music while Noor is making the latte*; Gandhi et al. (2023)). Table 4 shows improved results for all models, indicating that CAA improves performance on ToM tasks without compromising the models’ ability on control tasks.

5 Discussion and Conclusion

In this work, we conducted extensive experiments involving various LM types, sizes, fine-tuning approaches, and prompt designs to examine their internal representation of beliefs of self (*oracle*) and others (*protagonist*).

Our experiments show that, when predicting others’ belief, probing accuracy increases with model size and, more crucially for smaller models, with fine-tuning (Figure 2). Notably, fine-tuned 7B LMs outperform (Llama-2 with SFT and RLHF) or match the performance (Pythia with SFT) of base models with double the parameter count. Our experiments also reveal that the best probe accuracy scales with model size logarithmically for pre-trained models (Figure 6b, Figure 6d), and linearly with models fine-tuned with SFT and RLHF (Figure 6c). We then explore, for the first time, the

Model	Method	Forward Belief			Forward Action			Backward Belief		
		TB	FB	Both	TB	FB	Both	TB	FB	Both
Llama-2-7b	No int.	44	44	44	44	44	44	44	44	44
	ITI	44 ₊₀	44 ₊₀	44 ₊₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀
	CAA	66 ₊₂₂ *	71 ₊₂₇ *	54 ₊₁₀	66 ₊₂₂ *	57 ₊₁₃ *	54 ₊₁₀	60 ₊₁₆ *	74 ₊₃₀	54 ₊₁₀
Llama-2-7b-chat	No int.	56	56	55	69	55	37	56	56	55
	ITI	58 ₊₂	58 ₊₂	57 ₊₂	69 ₊₀	55 ₊₀	37 ₊₀	58 ₊₂	60 ₊₃	57 ₊₂
	CAA	70 ₊₁₄	72 ₊₁₆ *	57 ₊₂	69 ₊₀	67 ₊₁₂	53 ₊₁₆	66 ₊₁₀	84 ₊₂₇ *	57 ₊₂
Llama-2-13b	No int.	52	44	35	59	50	37	46	49	33
	ITI	52 ₊₀	45 ₊₁	35 ₊₀	64 ₊₅	61 ₊₁₁	46 ₊₉	48 ₊₂	59 ₊₁₀	42 ₊₉
	CAA	85 ₊₃₃ *	88 ₊₄₄ *	66 ₊₃₁ *	71 ₊₁₂ *	69 ₊₁₉ *	55 ₊₁₈ *	75 ₊₂₉ *	92 ₊₄₃ *	59 ₊₂₆ *
Llama-2-13b-chat	No int.	84	56	47	78	51	38	72	48	31
	ITI	84 ₊₀	65 ₊₉	59 ₊₁₂	78 ₊₀	58 ₊₇	47 ₊₉	72 ₊₀	60 ₊₁₂	48 ₊₁₇
	CAA	97 ₊₁₃ *	94 ₊₃₈ *	91 ₊₄₄ *	80 ₊₂	71 ₊₂₀ *	54 ₊₁₆ *	97 ₊₂₅	94 ₊₄₆ *	87 ₊₅₆ *
Llama-2-70b	No int.	90	87	78	93	52	48	73	53	32
	ITI	90 ₊₀	90 ₊₃	78 ₊₀	94 ₊₁	55 ₊₃	50 ₊₂	77 ₊₄	58 ₊₅	37 ₊₅
	CAA	99 ₊₉ *	97 ₊₁₀ *	95 ₊₁₇ *	94 ₊₁	80 ₊₂₈ *	73 ₊₂₅ *	94 ₊₂₁	92 ₊₃₉ *	83 ₊₅₁ *
Llama-2-70b-chat	No int.	69	75	56	86	56	52	63	59	52
	ITI	69 ₊₀	76 ₊₁	59 ₊₂	86 ₊₀	56 ₊₀	52 ₊₀	63 ₊₀	60 ₊₁	54 ₊₂
	CAA	92 ₊₂₃ *	97 ₊₂₂ *	89 ₊₃₂ *	87 ₊₁	75 ₊₁₉ *	60 ₊₈ *	88 ₊₂₅	92 ₊₃₃ *	80 ₊₂₈ *
Pythia-70m	No int.	41	41	37	46	45	41	44	41	37
	ITI	54 ₊₁₃	54 ₊₁₃	54 ₊₁₇ *	54 ₊₈	54 ₊₉	54 ₊₁₃ *	54 ₊₁₀	54 ₊₁₃	54 ₊₁₇
	CAA	62 ₊₂₁ *	56 ₊₁₅ *	54 ₊₁₇ *	59 ₊₁₃ *	60 ₊₁₅ *	58 ₊₁₇ *	63 ₊₁₉	56 ₊₁₅ *	54 ₊₁₇ *
Pythia-410m	No int.	48	45	45	44	44	44	44	47	44
	ITI	55 ₊₇	62 ₊₁₇ *	52 ₊₇	54 ₊₁₀ *	54 ₊₁₀ *	54 ₊₁₀	60 ₊₁₆	63 ₊₁₆	56 ₊₁₂
	CAA	67 ₊₁₉	64 ₊₁₉	61 ₊₁₆ *	56 ₊₁₂ *	63 ₊₁₉ *	56 ₊₁₂ *	69 ₊₂₅	63 ₊₁₆ *	60 ₊₁₆
Pythia-1b	No int.	44	44	44	44	44	44	44	44	44
	ITI	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀
	CAA	59 ₊₁₅ *	62 ₊₁₈ *	54 ₊₁₀	57 ₊₁₃	59 ₊₁₅	56 ₊₁₂	57 ₊₁₃	60 ₊₁₆	54 ₊₁₀
Pythia-6.9b	No int.	44	44	44	44	44	44	44	44	44
	ITI	45 ₊₁	54 ₊₁₀	44 ₊₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀
	CAA	56 ₊₁₂	71 ₊₂₇ *	55 ₊₁₁	55 ₊₁₁	63 ₊₁₉	55 ₊₁₁	55 ₊₁₁	71 ₊₂₇ *	55 ₊₁₁
Pythia-6.9b-chat	No int.	55	54	28	36	64	20	44	67	30
	ITI	57 ₊₂	54 ₊₀	28 ₊₀	44 ₊₈	71 ₊₇	32 ₊₁₂	44 ₊₀	67 ₊₀	30 ₊₀
	CAA	68 ₊₁₃	65 ₊₁₁	57 ₊₂₉ *	54 ₊₁₈	75 ₊₁₁	48 ₊₂₈ *	58 ₊₁₄ *	67 ₊₀	54 ₊₂₄ *
Pythia-12b	No int.	44	44	44	44	44	44	44	44	44
	ITI	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀	54 ₊₁₀
	CAA	54 ₊₁₀	64 ₊₂₀ *	54 ₊₁₀	60 ₊₁₆	58 ₊₁₄	55 ₊₁₁	54 ₊₁₀	67 ₊₂₃	54 ₊₁₀

Table 1: Comparison of the effects of ITI (Li et al., 2023b) and CAA (Rimsky et al., 2023) activation editing methods on three tasks from BigToM (Gandhi et al., 2023). TB denotes a true belief task, whereas FB denotes a false belief task. The numbers represent accuracy scores, with the difference in performance compared to no intervention (No int.) indicated as subscripts (ITI – No int. and CAA – No int.). An asterisk (*) denotes a statistically significant difference from No int. based on a McNemar’s test (McNemar, 1947) with $p < 0.05$.

sensitivity of LMs’ representations to prompting in the context of ToM. Our experiments with different prompt variations demonstrate not only that models’ representations degrade with the addition of random tokens or distractors in the prompt, but also when including time specifications that should make the prompt less ambiguous (Figure 3 and Figure 7). In contrast, including the protagonist’s initial belief in the prompt yields higher probe accuracy. We also verify that probes measure meaningful representations rather than spurious patterns, as it is possible to recover most of the accuracy by training probes on a much small subset of principal components of models’ representations (Figure 4, Figure 9, and Figure 10). Finally, we show that

CAA can steer models’ activations in a generalisable way, yielding significant performance improvements across different ToM tasks (Table 1). Unlike ITI, which requires training a separate probe for each attention head in a LM, CAA computes a single vector per layer, dramatically reducing computational overhead. For example, in the case of Llama 2-70B, ITI requires training 5120 probes (64 attention heads across 80 layers), while CAA requires computing only 80 vectors, one per layer. Overall, our work contributes valuable insights into the factors influencing LMs’ mental state representations, shedding light on avenues for improving their performance in ToM tasks.

6 Limitations

Our study focused on expanding experiments from the model perspective, examining architectures, sizes, fine-tuning, and prompt design, all within the same dataset. A natural extension of our work is replicating these experiments across multiple datasets and more model families. Given the rapid pace of new language model releases, studying all available models is impractical, particularly considering computational resource constraints. Nevertheless, our approach can be adopted to support new benchmarks or to evaluate newly released models as they become available. Finally, while in this work we focused on beliefs, our experimental approach can be adapted to investigate how LMs represent desires, emotions, intentions, or preferences. Future research exploring other types of mental states can use our findings to determine whether similar or distinct patterns emerge.

References

- Mostafa Abdou, Artur Kulmizev, Daniel Hershcovich, Stella Frank, Ellie Pavlick, and Anders Søgaard. 2021. [Can language models encode perceptual structure without grounding? a case study in color](#). In *Proceedings of the 25th Conference on Computational Natural Language Learning*, pages 109–132, Online. Association for Computational Linguistics.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Guillaume Alain and Yoshua Bengio. 2017. Understanding intermediate layers using linear classifier probes. In *International Conference on Learning Representations*.
- Chris L Baker, Rebecca Saxe, and Joshua B Tenenbaum. 2009. Action understanding as inverse planning. *Cognition*, 113(3):329–349.
- Cristian-Paul Bara, Sky CH-Wang, and Joyce Chai. 2021. [MindCraft: Theory of mind modeling for situated dialogue in collaborative tasks](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 1112–1125, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Yonatan Belinkov. 2022. Probing classifiers: Promises, shortcomings, and advances. *Computational Linguistics*, 48(1):207–219.
- Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al. 2023. Pythia: A suite for analyzing large language models across training and scaling. In *International Conference on Machine Learning*, pages 2397–2430. PMLR.
- Matteo Bortoletto, Constantin Ruhdorfer, Adnen Abdesaied, Lei Shi, and Andreas Bulling. 2024a. Limits of theory of mind modelling in dialogue-based collaborative plan acquisition. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*.
- Matteo Bortoletto, Constantin Ruhdorfer, Lei Shi, and Andreas Bulling. 2024b. Explicit modelling of theory of mind for belief prediction in nonverbal social interactions. *arXiv preprint arXiv:2407.06762*.
- Matteo Bortoletto, Lei Shi, and Andreas Bulling. 2024c. Neural reasoning about agents’ goals, preferences, and actions. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 456–464.
- Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan, Johannes Gehrkke, Eric Horvitz, Ece Kamar, Peter Lee, Yin Tat Lee, Yuanzhi Li, Scott Lundberg, et al. 2023. Sparks of artificial general intelligence: Early experiments with gpt-4. *arXiv preprint arXiv:2303.12712*.
- Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. 2017. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30.
- Alexis Conneau, Germán Kruszewski, Guillaume Lample, Loïc Barrault, and Marco Baroni. 2018. What you can cram into a single $\&\!'\#\&$ vector: Probing sentence embeddings for linguistic properties. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2126–2136.
- Kanishk Gandhi, Jan-Philipp Fränken, Tobias Gerstenberg, and Noah Goodman. 2023. Understanding social reasoning in language models with language models. *Advances in Neural Information Processing Systems*, 37.
- Leo Gao, Stella Biderman, Sid Black, Laurence Golding, Travis Hoppe, Charles Foster, Jason Phang, Horace He, Anish Thite, Noa Nabeshima, et al. 2020. The pile: An 800gb dataset of diverse text for language modeling. *arXiv preprint arXiv:2101.00027*.
- Wes Gurnee, Neel Nanda, Matthew Pauly, Katherine Harvey, Dmitrii Troitskii, and Dimitris Bertsimas. 2023. Finding neurons in a haystack: Case studies with sparse probing. *Transactions on Machine Learning Research*.
- Wes Gurnee and Max Tegmark. 2024. Language models represent space and time. *International Conference on Learning Representations*.

722	Yinghui He, Yufan Wu, Yilin Jia, Rada Mihalcea, Yu-	Dong C Liu and Jorge Nocedal. 1989. On the limited	779
723	long Chen, and Naihao Deng. 2023. Hi-tom: A	memory bfgs method for large scale optimization.	780
724	benchmark for evaluating higher-order theory of	<i>Mathematical programming</i> , 45(1):503–528.	781
725	mind reasoning in large language models. <i>arXiv</i>		
726	<i>preprint arXiv:2310.16755</i> .		
727	Evan Hernandez and Jacob Andreas. 2021. The low-	Kevin Liu, Stephen Casper, Dylan Hadfield-Menell, and	782
728	dimensional linear geometry of contextualized word	Jacob Andreas. 2023. Cognitive dissonance: Why	783
729	representations . In <i>Proceedings of the 25th Confer-</i>	do language model outputs disagree with internal	784
730	<i>ence on Computational Natural Language Learning</i> ,	representations of truthfulness? In <i>Proceedings of the</i>	785
731	pages 82–93, Online. Association for Computational	<i>2023 Conference on Empirical Methods in Natural</i>	786
732	Linguistics.	<i>Language Processing</i> , pages 4791–4797.	787
733	Evan Hernandez, Belinda Z Li, and Jacob Andreas.	Yao Lu, Max Bartolo, Alastair Moore, Sebastian Riedel,	788
734	2023. Inspecting and editing knowledge repre-	and Pontus Stenetorp. 2022. Fantastically ordered	789
735	sentations in language models. <i>arXiv preprint</i>	prompts and where to find them: Overcoming few-	790
736	<i>arXiv:2304.00740</i> .	shot prompt order sensitivity . In <i>Proceedings of the</i>	791
737	Yoichi Ishibashi, Danushka Bollegala, Katsuhito Sudoh,	<i>60th Annual Meeting of the Association for Compu-</i>	792
738	and Satoshi Nakamura. 2023. Evaluating the robust-	<i>tational Linguistics (Volume 1: Long Papers)</i> , pages	793
739	ness of discrete prompts . In <i>Proceedings of the 17th</i>	8086–8098, Dublin, Ireland. Association for Compu-	794
740	<i>Conference of the European Chapter of the Associa-</i>	<i>tational Linguistics</i> .	795
741	<i>tion for Computational Linguistics</i> , pages 2373–2384,	Weina Ma, Jieyu Mao, Yu Xie, Simeng Li, and Mian	796
742	Dubrovnik, Croatia. Association for Computational	Wang. 2023a. Examining the effects of theory of	797
743	Linguistics.	mind and social skills training on social competence	798
744	Hyunwoo Kim, Melanie Sclar, Xuhui Zhou, Ronan Le	in adolescents with autism. <i>Behavioral Sciences</i> ,	799
745	Bras, Gunhee Kim, Yejin Choi, and Maarten Sap.	13(10):860.	800
746	2023. Fantom: A benchmark for stress-testing ma-	Ziqiao Ma, Jacob Sansom, Run Peng, and Joyce Chai.	801
747	chine theory of mind in interactions. <i>arXiv preprint</i>	2023b. Towards a holistic landscape of situated the-	802
748	<i>arXiv:2310.15421</i> .	ory of mind in large language models . In <i>Findings</i>	803
749	Matthew Le, Y-Lan Boureau, and Maximilian Nickel.	<i>of the Association for Computational Linguistics:</i>	804
750	2019. Revisiting the evaluation of theory of mind	<i>EMNLP 2023</i> , pages 1011–1031, Singapore. Associ-	805
751	through question answering. In <i>Proceedings of the</i>	<i>ation for Computational Linguistics</i> .	806
752	<i>2019 Conference on Empirical Methods in Natu-</i>	Samuel Marks and Max Tegmark. 2023. The geometry	807
753	<i>ral Language Processing and the 9th International</i>	of truth: Emergent linear structure in large language	808
754	<i>Joint Conference on Natural Language Processing</i>	model representations of true/false datasets. <i>arXiv</i>	809
755	<i>(EMNLP-IJCNLP)</i> , pages 5872–5877.	<i>preprint arXiv:2310.06824</i> .	810
756	Alina Leidinger, Robert van Rooij, and Ekaterina	Marianna Martindale, Marine Carpuat, Kevin Duh, and	811
757	Shutova. 2023. The language of prompting: What	Paul McNamee. 2019. Identifying fluently inade-	812
758	linguistic properties make a prompt successful? In	quate output in neural and statistical machine transla-	813
759	<i>Findings of the Association for Computational Lin-</i>	tion. In <i>Proceedings of Machine Translation Summit</i>	814
760	<i>guistics: EMNLP 2023</i> , pages 9210–9232.	<i>XVII: Research Track</i> , pages 233–243.	815
761	Belinda Z Li, Maxwell Nye, and Jacob Andreas. 2021.	Quinn McNemar. 1947. Note on the sampling error	816
762	Implicit representations of meaning in neural lan-	of the difference between correlated proportions or	817
763	guage models. In <i>Proceedings of the 59th Annual</i>	percentages. <i>Psychometrika</i> , 12(2):153–157.	818
764	<i>Meeting of the Association for Computational Lin-</i>	Sewon Min, Xinxu Lyu, Ari Holtzman, Mikel Artetxe,	819
765	<i>guistics and the 11th International Joint Conference</i>	Mike Lewis, Hannaneh Hajishirzi, and Luke Zettle-	820
766	<i>on Natural Language Processing (Volume 1: Long</i>	moyer. 2022. Rethinking the role of demonstrations:	821
767	<i>Papers)</i> , pages 1813–1827.	What makes in-context learning work? In <i>Proceed-</i>	822
768	Kenneth Li, Aspen K Hopkins, David Bau, Fernanda	<i>ings of the 2022 Conference on Empirical Methods in</i>	823
769	Viégas, Hanspeter Pfister, and Martin Wattenberg.	<i>Natural Language Processing</i> , pages 11048–11064,	824
770	2023a. Emergent world representations: Exploring	Abu Dhabi, United Arab Emirates. Association for	825
771	a sequence model trained on a synthetic task . In	<i>Computational Linguistics</i> .	826
772	<i>The Eleventh International Conference on Learning</i>	Shima Rahimi Moghaddam and Christopher J Honey.	827
773	<i>Representations</i> .	2023. Boosting theory-of-mind performance in large	828
774	Kenneth Li, Oam Patel, Fernanda Viégas, Hanspeter	language models via prompting. <i>arXiv preprint</i>	829
775	Pfister, and Martin Wattenberg. 2023b. Inference-	<i>arXiv:2304.11490</i> .	830
776	time intervention: Eliciting truthful answers from a	Neel Nanda, Andrew Lee, and Martin Wattenberg. 2023.	831
777	language model . In <i>Thirty-seventh Conference on</i>	Emergent linear representations in world models of	832
778	<i>Neural Information Processing Systems</i> .	self-supervised sequence models . In <i>Proceedings</i>	833
		<i>of the 6th BlackboxNLP Workshop: Analyzing and</i>	834

835	<i>Interpreting Neural Networks for NLP</i> , pages 16–30,	Melanie Sclar, Yejin Choi, Yulia Tsvetkov, and Alane	889
836	Singapore. Association for Computational Linguistics.	Suhr. 2024. Quantifying language models’ sensitivity to spurious features in prompt design or: How i	890
837		learned to start worrying about prompt formatting.	891
838	Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida,	In <i>The Twelfth International Conference on Learning</i>	893
839	Carroll Wainwright, Pamela Mishkin, Chong Zhang,	Representations.	894
840	Sandhini Agarwal, Katarina Slama, Alex Ray, et al.		
841	2022. Training language models to follow instruc-	Melanie Sclar, Sachin Kumar, Peter West, Alane Suhr,	895
842	tions with human feedback. <i>Advances in neural in-</i>	Yejin Choi, and Yulia Tsvetkov. 2023. Minding lan-	896
843	formation processing systems, 35:27730–27744.	guage models’ (lack of) theory of mind: A plug-and-	897
844		play multi-character belief tracker. In <i>Proceedings</i>	898
845	Roma Patel and Ellie Pavlick. 2022. Mapping language	of the 61st Annual Meeting of the Association for	899
846	models to grounded conceptual spaces. In <i>Interna-</i>	Computational Linguistics (Volume 1: Long Papers),	900
	<i>tional Conference on Learning Representations.</i>	pages 13960–13980, Toronto, Canada. Association	901
		for Computational Linguistics.	902
847	Fabio Petroni, Tim Rocktäschel, Sebastian Riedel,		
848	Patrick Lewis, Anton Bakhtin, Yuxiang Wu, and	Omar Shaikh, Hongxin Zhang, William Held, Michael	903
849	Alexander Miller. 2019. Language models as knowl-	Bernstein, and Diyi Yang. 2023. On second thought,	904
850	edge bases? In <i>Proceedings of the 2019 Conference</i>	let’s not think step by step! bias and toxicity in zero-	905
851	<i>on Empirical Methods in Natural Language Pro-</i>	shot reasoning. In <i>Proceedings of the 61st Annual</i>	906
852	<i>cessing and the 9th International Joint Conference</i>	Meeting of the Association for Computational Lin-	907
853	<i>on Natural Language Processing (EMNLP-IJCNLP),</i>	guistics (Volume 1: Long Papers), pages 4454–4470,	908
854	pages 2463–2473.	Toronto, Canada. Association for Computational Lin-	909
		guistics.	910
855	David Premack and Guy Woodruff. 1978. Does the	Francesc Sidera, Georgina Perpiñà, Jèssica Serrano, and	911
856	chimpanzee have a theory of mind? <i>Behavioral and</i>	Carles Rostan. 2018. Why is theory of mind impor-	912
857	<i>brain sciences</i> , 1(4):515–526.	tant for referential communication? <i>Current Psychol-</i>	913
		<i>ogy</i> , 37:82–97.	914
858	Neil Rabinowitz, Frank Perbet, Francis Song, Chiyuan		
859	Zhang, SM Ali Eslami, and Matthew Botvinick. 2018.	Fiona Anting Tan, Gerard Christopher Yeo, Fanyou	915
860	Machine theory of mind. In <i>International conference</i>	Wu, Weijie Xu, Vinija Jain, Aman Chadha, Kokil	916
861	<i>on machine learning</i> , pages 4218–4227. PMLR.	Jaidka, Yang Liu, and See-Kiong Ng. 2024. Phan-	917
		tom: Personality has an effect on theory-of-mind rea-	918
862	Yasaman Razeghi, Robert L Logan IV, Matt Gardner,	soning in large language models. <i>arXiv preprint</i>	919
863	and Sameer Singh. 2022. Impact of pretraining term	<i>arXiv:2403.02246.</i>	920
864	frequencies on few-shot numerical reasoning. In		
865	<i>Findings of the Association for Computational Lin-</i>	Ian Tenney, Dipanjan Das, and Ellie Pavlick. 2019. Bert	921
866	<i>guistics: EMNLP 2022</i> , pages 840–854, Abu Dhabi,	rediscovered the classical nlp pipeline. In <i>Proceedings</i>	922
867	United Arab Emirates. Association for Computa-	of the 57th Annual Meeting of the Association for	923
868	tional Linguistics.	Computational Linguistics , pages 4593–4601.	924
869	Nina Rimskey, Nick Gabrieli, Julian Schulz, Meg Tong,		
870	Evan Hubinger, and Alexander Matt Turner. 2023.	Ian Tenney, Patrick Xia, Berlin Chen, Alex Wang, Adam	925
871	Steering llama 2 via contrastive activation addition.	Poliak, R Thomas McCoy, Najoung Kim, Benjamin	926
872	<i>arXiv preprint arXiv:2312.06681.</i>	Van Durme, Samuel R Bowman, Dipanjan Das, et al.	927
		2018. What do you learn from context? probing for	928
873	Anna Rogers, Olga Kovaleva, and Anna Rumshisky.	sentence structure in contextualized word represen-	929
874	2021. A primer in bertology: What we know about	tations. In <i>International Conference on Learning</i>	930
875	how bert works. <i>Transactions of the Association for</i>	<i>Representations.</i>	931
876	<i>Computational Linguistics</i> , 8:842–866.		
877	Swarnadeep Saha, Peter Hase, and Mohit Bansal. 2023.	Michael Tomasello. 2010. <i>Origins of human communi-</i>	932
878	Can language models teach weaker agents? teacher	<i>cation.</i> MIT Press.	933
879	explanations improve students via theory of mind.		
880	<i>Advances in Neural Information Processing Systems</i> ,	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	934
881	37.	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	935
		Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti	936
882	Maarten Sap, Ronan Le Bras, Daniel Fried, and Yejin	Bhosale, et al. 2023. Llama 2: Open founda-	937
883	Choi. 2022. Neural theory-of-mind? on the limits of	tion and fine-tuned chat models. <i>arXiv preprint</i>	938
884	social intelligence in large LMs. In <i>Proceedings of</i>	<i>arXiv:2307.09288.</i>	939
885	<i>the 2022 Conference on Empirical Methods in Nat-</i>	Alex Turner, Lisa Thiergart, David Udell, Gavin Leech,	940
886	<i>ural Language Processing</i> , pages 3762–3780, Abu	Ulisse Mini, and Monte MacDiarmid. 2023. Acti-	941
887	Dhabi, United Arab Emirates. Association for Com-	vation addition: Steering language models without	942
888	putational Linguistics.	optimization. <i>arXiv preprint arXiv:2308.10248.</i>	943

Yizhong Wang, Hamish Ivison, Pradeep Dasigi, Jack Hessel, Tushar Khot, Khyathi Chandu, David Wadden, Kelsey MacMillan, Noah A Smith, Iz Beltagy, et al. 2024. How far can camels go? exploring the state of instruction tuning on open resources. *Advances in Neural Information Processing Systems*, 36.

Albert Webson and Ellie Pavlick. 2022. [Do prompt-based models really understand the meaning of their prompts?](#) In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 2300–2344, Seattle, United States. Association for Computational Linguistics.

Jason Wei, Maarten Bosma, Vincent Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M Dai, and Quoc V Le. 2021. Finetuned language models are zero-shot learners. In *International Conference on Learning Representations*.

Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel, Barret Zoph, Sebastian Borgeaud, Dani Yogatama, Maarten Bosma, Denny Zhou, Donald Metzler, Ed H. Chi, Tatsunori Hashimoto, Oriol Vinyals, Percy Liang, Jeff Dean, and William Fedus. 2022. [Emergent abilities of large language models](#). *Transactions on Machine Learning Research*. Survey Certification.

Alex Wilf, Sihyun Shawn Lee, Paul Pu Liang, and Louis-Philippe Morency. 2023. Think twice: Perspective-taking improves large language models’ theory-of-mind capabilities. *arXiv preprint arXiv:2311.10227*.

Jincenzi Wu, Zhuang Chen, Jiawen Deng, Sahand Sabour, and Minlie Huang. 2023. Coke: A cognitive knowledge graph for machine theory of mind. *arXiv preprint arXiv:2305.05390*.

Hainiu Xu, Runcong Zhao, Lixing Zhu, Jinhua Du, and Yulan He. 2024. Opentom: A comprehensive benchmark for evaluating theory-of-mind reasoning capabilities of large language models. *arXiv preprint arXiv:2402.06044*.

Zihao Zhao, Eric Wallace, Shi Feng, Dan Klein, and Sameer Singh. 2021. Calibrate before use: Improving few-shot performance of language models. In *International conference on machine learning*, pages 12697–12706. PMLR.

Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc V Le, and Ed H. Chi. 2023a. [Least-to-most prompting enables complex reasoning in large language models](#). In *The Eleventh International Conference on Learning Representations*.

Pei Zhou, Aman Madaan, Srividya Pranavi Potharaju, Aditya Gupta, Kevin R McKee, Ari Holtzman, Jay Pujara, Xiang Ren, Swaroop Mishra, Aida Nematzadeh, et al. 2023b. How far are large language models from agents with theory-of-mind? *arXiv preprint arXiv:2310.03051*.

Wentao Zhu, Zhining Zhang, and Yizhou Wang. 2024. Language models represent beliefs of self and others. *arXiv preprint arXiv:2402.18496*.

A Appendix

A.1 Experimental setup

A.1.1 BigToM

BigToM (Gandhi et al., 2023) is constructed using GPT-4 (Achiam et al., 2023) to populate causal templates and combine elements from these templates. Each causal template is set up with a *context* and a description of the *protagonist* (e.g. “Noor is working as a barista [...]”), a *desire* (“Noor wants to make a cappuccino”), a *percept* (“Noor grabs a milk pitcher and fills it with oat milk”), and a *belief* (“Noor believes that the pitcher contains oat milk”). The state of the world is changed by a *causal event* (“A coworker swaps the oat milk in the pitcher with almond milk”). The dataset constructs different conditions by changing the percepts of the protagonist after the causal event, which will result in different beliefs – true or false. Gandhi et al. (2023) generated 200 templates and extracted 25 conditions from each template, resulting in 5,000 test samples. In this work, following Zhu et al. (2024) and Gandhi et al. (2023) we focused on the 6 most important conditions, corresponding to true and false beliefs on the following three tasks:

- *Forward Belief*: given the protagonist’s percepts of the causal event, infer their belief: $P(\text{belief}|\text{percept})$.
- *Forward Action*: infer the protagonist’s action given their desire and percepts of the causal event. Before inferring the action, one would need to first implicitly infer the protagonist’s belief: $\sum_{\text{belief}} P(\text{action}|\text{percept}, \text{belief}, \text{desire})$.
- *Backward Belief*: infer the protagonist’s belief from observed actions. This requires to first implicitly infer the protagonist’s percepts: $\sum_{\text{percepts}} P(\text{belief}|\text{action}, \text{percept}, \text{desire})$.

The dataset was released under the MIT license and can be accessed at <https://github.com/cicl-stanford/procedural-evals-tom>. We report one example for each task in Example 1, 2, and 3, where the text defining true belief or false belief task is shown in blue and red, respectively.

A.1.2 Linear probes

Our probing approach is illustrated in Figure 5. For our experiments, we cache activations at the residual stream level. To perform ITI and compare it to CAA, we also cache attention heads activations. We trained the probes using the L-BFGS solver (Liu and Nocedal, 1989) with L2 penalty with inverse of regularisation strength 10 for a maximum of 1000 iterations. We use zero as random seed.

A.1.3 Language models

A detailed summary of the models we use in this work is shown in Table 2. Pythia was released under the Apache 2.0 license. Llama-2 is licensed by Meta for both researchers and commercial entities (Touvron et al., 2023). For all the models, we set the temperature to zero.

A.1.4 Examples of prompt variations

We provide examples of variations for a prompt (Example 4) in Example 5 (random), Example 6 (misleading), Example 7, and Example 8.

A.2 Model size and fine-tuning

To characterise the relationship between probe accuracy and model size we consider the *best* probe accuracy for every LM, i.e. the highest accuracy among probes $\{g_l\}$ trained on $\{a_l\}$ for a LM f . For Llama-2 base, the best probe accuracy scales logarithmically with model size ($R^2 = 0.98$, Figure 6b), whereas for fine-tuned models it scales linearly ($R = 1.0$, cf. Figure 6c). For Pythia base, the best probe accuracy also scales logarithmically with model size ($R^2 = 0.96$, Figure 6d).

A.3 Sensitivity to prompting

Accuracy on *protagonist* belief probing for Pythia models is shown in Figure 7.

Accuracy on *oracle* belief probing for different prompt variations are reported in Figure 8.

A.4 Dimensionality reduction

Probing accuracy obtained by Pythia models for the *protagonist* setting is reported in Figure 9.

Oracle probe accuracy obtained by considering only the first $n = \{2, 10, 100, 1000\}$ principal components are shown in Figure 10.

A.5 Inference-time intervention

Inference-time intervention (Li et al., 2023b, ITI) employs a two-step process. First, it trains a probe for each attention head across all layers of a LM.

These probes are evaluated on a validation set, and the top- k heads with the highest accuracy are selected. Subsequently, during inference, ITI steers the activations of these top heads along the directions defined by their corresponding probes. Formally, ITI can be defined as an additional term to the multi-head attention:

$$x_{l+1} = x_l + \sum_{h=1}^H Q_l^h \left(\text{Att}_l^h(P_l^h x_l) + \alpha \sigma_l^h \theta_l^h \right)$$

where x_l is the residual stream at layer l , H is the number of attention heads, $\alpha \in \mathbb{R}^+$ is a coefficient, σ_l^h is the standard deviation of activations along the direction identified by the probe trained on attention head h at layer l , and θ_l^h is zero for not-selected attention heads.

A.6 Activation editing

Table 3 reports results obtained on the three BigToM tasks with the hyperparameters used for ITI (Li et al., 2023b) and CAA (Rimsky et al., 2023). We report an example of prompt used for evaluation in Example 9. Table 4 shows the accuracy obtained by using CAA on the Forward Belief True Control task in BigToM. On this control task, CAA produced improved results for all model, proving that CAA not only improves performance on ToM tasks, but also does not degrades the models' ability to perform other tasks.

A.7 Compute resources

We ran our experiments on a server running Ubuntu 22.04, equipped with eight NVIDIA Tesla V100-SXM2 GPUs with 32GB of memory and Intel Xeon Platinum 8260 CPUs.

A.8 Code

Our code is provided as supplementary material and it will be made public under the MIT licence at www.example.com.

A.9 Societal impact

While our work is foundational and remains distant from specific applications with direct societal impact, it's important to recognise the ethical implications of modelling and predicting mental states. Handling sensitive aspects of individuals' inner experiences and emotions requires careful consideration to avoid reinforcing biases or misunderstanding psychological nuances.

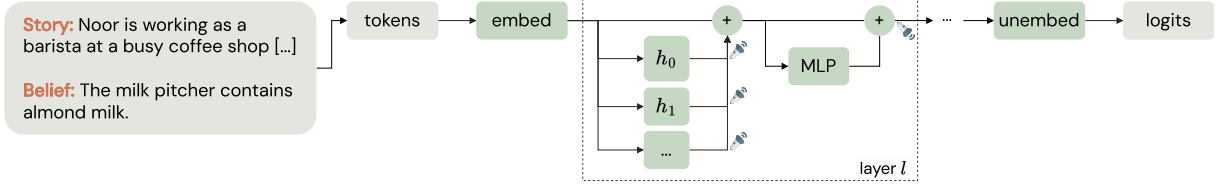


Figure 5: Given a tokenised input, we cache the internal activations for all attention heads $h_i, i = 0, \dots, H - 1$, and residual streams. In our experiments, we use residual stream activations.

LM	Size	+ SFT	+ RLHF	Tokens	d_{model}	Layers
Llama-2	7B			2T	4096	32
	13B			2T	5120	40
	70B			2T	8192	80
Llama-2-chat	7B	✓	✓	2T	4096	32
	13B	✓	✓	2T	5120	40
	70B	✓	✓	2T	8192	80
Pythia	70M			300B	512	6
	410M			300B	1024	24
	1B			300B	2048	16
	6.9B			300B	4096	32
	12B			300B	5120	36
	6.9B	✓		300B	4096	32

Table 2: The 12 models used in this work..

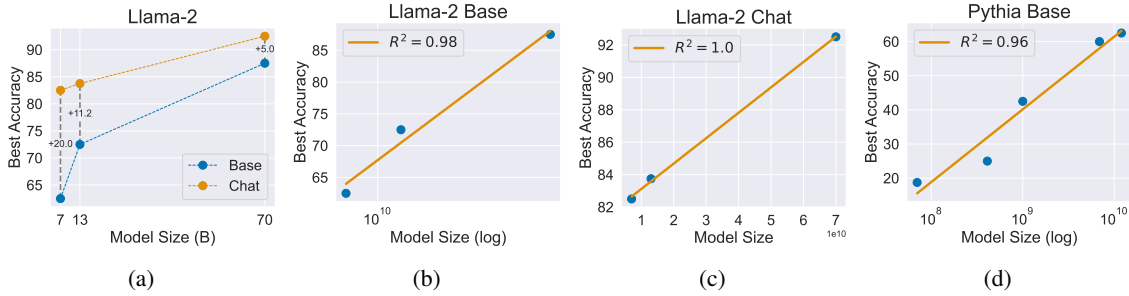


Figure 6: To characterise the relationship between probe accuracy and model size we consider the *best* probe accuracy for every LM, i.e. the highest accuracy among probes $\{g_l\}$ trained on $\{a_l\}$ for a LM f . (a) Best accuracy for Llama-2 models of different size. Numbers on the vertical dotted lines indicate the gain in accuracy between base and fine-tuned model of the same size. (b) Logarithmic fit for Llama-2 base. (c) Linear fit for Llama-2 fine-tuned (chat). (d) Logarithmic fit for Pythia base.

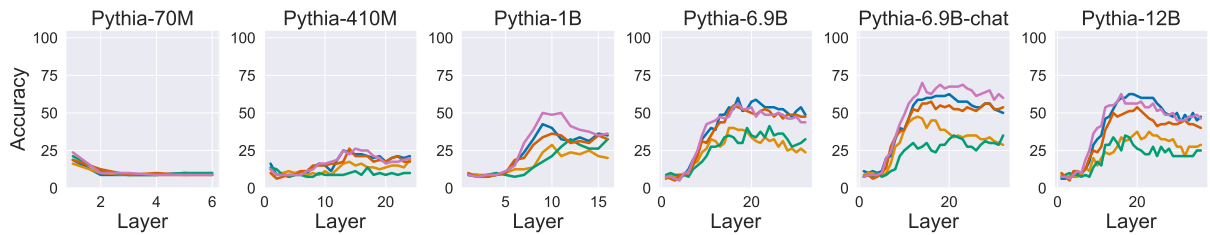


Figure 7: Sensitivity of protagonist belief probing accuracy to different prompt variations.

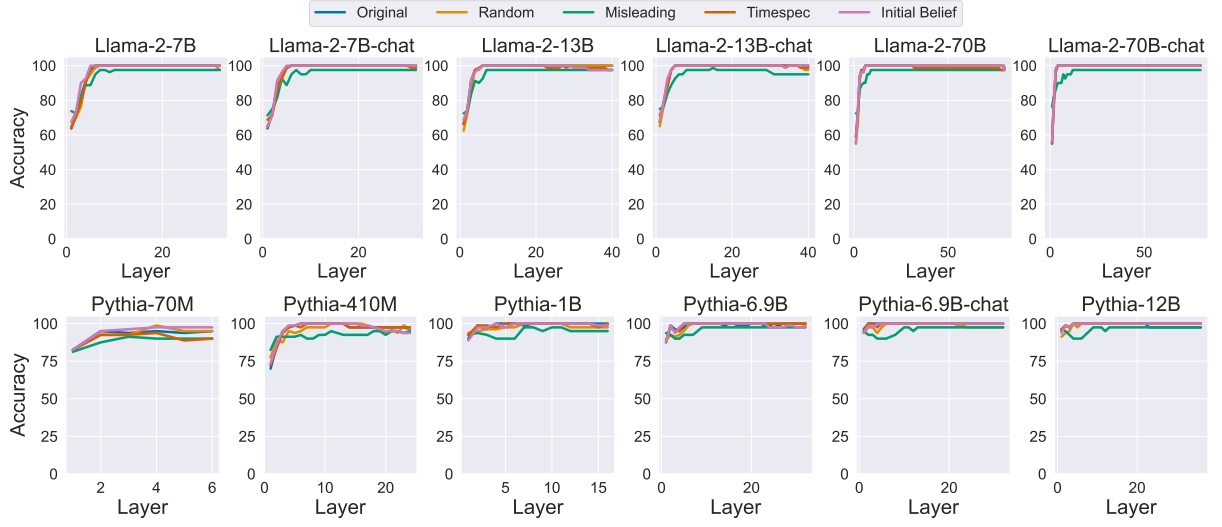


Figure 8: Sensitivity of protagonist belief probing accuracy to different prompt variations.

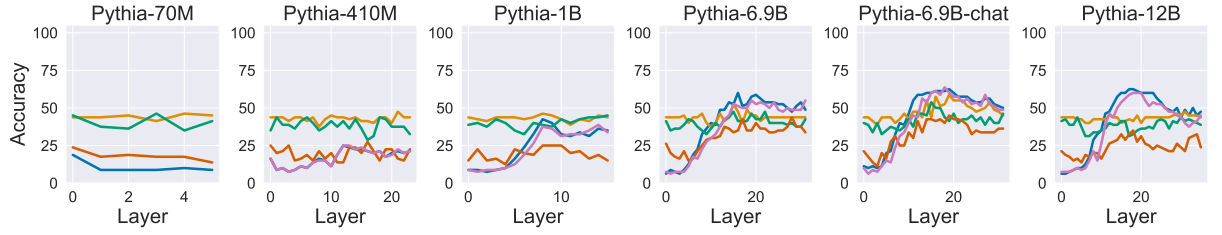


Figure 9: We compare the probing accuracy obtained by using the original set of activations (All) with the accuracy obtained by considering only the first $n = \{2, 10, 100, 1000\}$ principal components. For Pythia: All(70m) = 512, All(410m) = 1024, All(1b) = 2048, All(6.9b) = 4096, All(12b) = 5120. Results for *oracle* are shown in Figure 10.

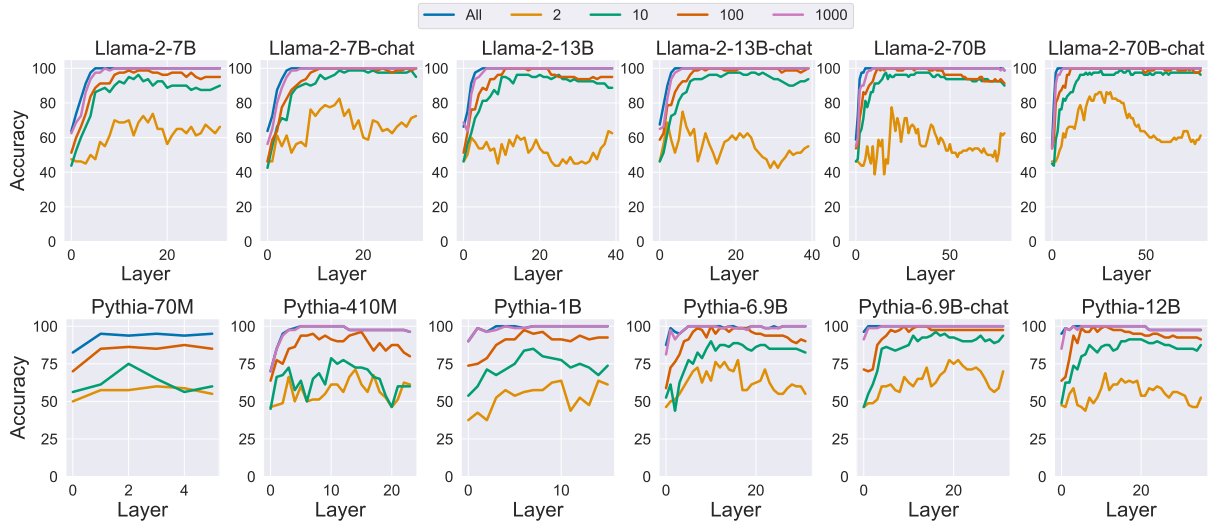


Figure 10: **(Oracle)** To investigate potential memorisation in the probes, we compare the probing accuracy obtained by using the original set of activations (All) with the accuracy obtained by considering only the first $n = \{2, 10, 100, 1000\}$ principal components. For Llama2: All(7b) = 4096, All(13b) = 5120, All(70b) = 8192. For Pythia: All(70m) = 512, All(410m) = 1024, All(1b) = 2048, All(6.9b) = 4096, All(12b) = 5120.

Model	Method	Forward Belief			Forward Action			Backward Belief		
		TB	FB	Both	TB	FB	Both	TB	FB	Both
Llama-2-7b	No int.	44	44	44	44	44	44	44	44	44
	ITI	44 _{0,0}	44 _{0,0}	44 _{0,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}
	CAA	66 _{2,0,11}	71 _{1,0,31}	54 _{2,0,0}	66 _{2,0,11}	57 _{2,0,12}	54 _{2,0,2}	60 _{2,0,11}	74 _{1,0,31}	54 _{2,0,2}
Llama-2-7b-chat	No int.	56	56	55	69	55	37	56	56	55
	ITI	58 _{15,0}	58 _{15,0}	57 _{15,0}	69 _{0,0}	55 _{0,0}	37 _{0,0}	58 _{10,0}	60 _{10,0}	57 _{10,0}
	CAA	70 _{1,0,11}	72 _{1,5,10}	57 _{1,0,1}	69 _{0,0,0}	67 _{1,5,11}	53 _{1,5,12}	66 _{1,0,11}	84 _{1,5,10}	57 _{1,0,0}
Llama-2-13b	No int.	52	44	35	59	50	37	46	49	33
	ITI	52 _{0,0}	45 _{15,0}	35 _{0,0}	64 _{15,0}	61 _{20,0}	46 _{20,0}	48 _{20,0}	59 _{20,0}	42 _{20,0}
	CAA	85 _{2,0,12}	88 _{2,0,14}	66 _{2,0,12}	71 _{1,5,10}	69 _{2,0,13}	55 _{1,0,39}	75 _{2,0,10}	92 _{2,0,13}	59 _{1,5,12}
Llama-2-13b-chat	No int.	84	56	47	78	51	38	72	48	31
	ITI	84 _{0,0}	65 _{15,0}	59 _{15,0}	78 _{0,0}	58 _{15,0}	47 _{15,0}	72 _{0,0}	60 _{15,0}	48 _{15,0}
	CAA	97 _{1,0,12}	94 _{1,0,12}	91 _{1,0,12}	80 _{1,5,11}	71 _{1,0,13}	54 _{1,5,13}	97 _{1,5,10}	94 _{1,5,12}	87 _{1,5,12}
Llama-2-70b	No int.	90	87	78	93	52	48	73	53	32
	ITI	90 _{0,0}	90 _{20,0}	78 _{0,0}	94 _{15,0}	55 _{20,0}	50 _{15,0}	77 _{10,0}	58 _{15,0}	37 _{10,0}
	CAA	99 _{2,0,16}	97 _{1,5,19}	95 _{1,5,18}	94 _{1,5,2}	80 _{2,0,19}	73 _{1,5,18}	94 _{2,0,18}	92 _{2,0,19}	83 _{1,5,19}
Llama-2-70b-chat	No int.	69	75	56	86	56	52	63	59	52
	ITI	69 _{0,0}	76 _{10,0}	59 _{10,0}	86 _{0,0}	56 _{0,0}	52 _{0,0}	63 _{0,0}	60 _{10,0}	54 _{10,0}
	CAA	92 _{1,5,18}	97 _{1,5,25}	89 _{1,5,18}	87 _{1,5,17}	75 _{1,0,19}	60 _{1,0,19}	88 _{1,5,18}	92 _{1,0,19}	80 _{1,5,18}
Pythia-70m	No int.	41	41	37	46	45	41	44	41	37
	ITI	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}
	CAA	62 _{1,0,2}	56 _{1,0,1}	54 _{1,5,1}	59 _{1,0,2}	60 _{1,0,3}	58 _{1,0,2}	63 _{1,0,2}	56 _{1,0,2}	54 _{1,5,1}
Pythia-410m	No int.	48	45	45	44	44	44	44	47	44
	ITI	55 _{20,0}	62 _{20,0}	52 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	60 _{20,0}	63 _{20,0}	56 _{20,0}
	CAA	67 _{2,0,4}	64 _{2,0,4}	61 _{2,0,0}	56 _{2,0,6}	63 _{1,5,12}	56 _{2,0,6}	69 _{2,0,4}	63 _{2,0,0}	60 _{2,0,0}
Pythia-1b	No int.	44	44	44	44	44	44	44	44	44
	ITI	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}
	CAA	59 _{2,0,8}	62 _{2,0,5}	54 _{2,0,0}	57 _{2,0,4}	59 _{2,0,10}	56 _{2,0,4}	57 _{2,0,3}	60 _{2,0,5}	54 _{2,0,0}
Pythia-6.9b	No int.	44	44	44	44	44	44	44	44	44
	ITI	45 _{20,0}	54 _{20,0}	44 _{0,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}
	CAA	56 _{1,5,12}	71 _{1,5,9}	55 _{2,0,23}	55 _{2,0,4}	63 _{1,5,11}	55 _{2,0,4}	55 _{2,0,23}	71 _{1,5,9}	55 _{2,0,23}
Pythia-6.9b-chat	No int.	55	54	28	36	64	20	44	67	30
	ITI	57 _{15,0}	54 _{0,0}	28 _{0,0}	44 _{15,0}	71 _{15,0}	32 _{15,0}	44 _{0,0}	67 _{0,0}	30 _{0,0}
	CAA	68 _{1,5,15}	65 _{1,5,12}	57 _{1,5,11}	54 _{1,5,10}	75 _{1,5,5}	48 _{1,5,10}	58 _{1,5,15}	67 _{0,0,0}	54 _{1,5,10}
Pythia-12b	No int.	44	44	44	44	44	44	44	44	44
	ITI	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}	54 _{20,0}
	CAA	54 _{2,0,0}	64 _{2,0,9}	54 _{2,0,0}	60 _{2,0,11}	58 _{2,0,11}	55 _{2,0,12}	54 _{2,0,0}	67 _{2,0,10}	54 _{2,0,0}

Table 3: Activation intervention: comparison between ITI (Li et al., 2023b) and CAA (Rimsky et al., 2023). For ITI, the subscript indicates the value of the coefficient α_{ITI} used: $\text{Acc}_{\alpha_{ITI}}$. For CAA, the subscript indicates first the value of the coefficient α used and second the layer l at which intervention takes place: $\text{Acc}_{\alpha_{CAA},l}$.

Model	Method	Control	CAA Parameters
Llama-2-7b	No int.	44	
	CAA	66 ₊₂₂	2.0, 11
Llama-2-7b-chat	No int.	56	
	CAA	70 ₊₁₄	1.0, 11
Llama-2-13b	No int.	52	
	CAA	85 ₊₃₃	2.0, 12
Llama-2-13b-chat	No int.	84	
	CAA	97 ₊₁₃	1.0, 12
Llama-2-70b	No int.	90	
	CAA	99 ₊₉	2.0, 16
Llama-2-70b-chat	No int.	69	
	CAA	92 ₊₂₃	1.5, 18
Pythia-70m	No int.	41	
	CAA	62 ₊₂₁	1.0, 2
Pythia-410m	No int.	48	
	CAA	67 ₊₁₉	2.0, 4
Pythia-1b	No int.	44	
	CAA	59 ₊₁₅	2.0, 8
Pythia-6.9b	No int.	44	
	CAA	56 ₊₁₂	1.5, 12
Pythia-6.9b-chat	No int.	55	
	CAA	68 ₊₁₃	1.5, 15
Pythia-12b	No int.	44	
	CAA	54 ₊₁₀	2.0, 0

Table 4: Results for CAA (Rimsky et al., 2023) on the Forward Belief True Control task in BigToM (Gandhi et al., 2023). Numbers indicate accuracy scores, with differences (CAA — No int.) as subscripts.

Example 1: Forward Belief

Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor sees her coworker swapping the milk. Noor does not see her coworker swapping the milk.

Does Noor believe the milk pitcher contains oat milk or almond milk?

- a) Noor believes the milk pitcher contains oat milk.
- b) Noor believes the milk pitcher contains almond milk.

Example 2: Forward Action

Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor sees her coworker swapping the milk. Noor does not see her coworker swapping the milk.

What will Noor do?

- a) Noor will make the cappuccino using the milk in the pitcher.
- b) Noor will open the fridge once again to take out the oat milk and replace the almond milk with oat milk.

Example 3: Backward Belief

Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor opens the fridge again and reaches for the oat milk. Noor makes the cappuccino using the milk in the pitcher.

Does Noor believe the milk pitcher contains oat milk or almond milk?

- a) Noor believes the milk pitcher contains oat milk.
- b) Noor believes the milk pitcher contains almond milk.

Example 4: Default prompt

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.
Belief: Noor believes the milk pitcher contains oat milk.

Example 5: Prompt variation – Random

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.
Belief: 장± decided [_ countries sections behoSOUR gminy bef Noor believes the milk pitcher contains oat milk.

1141

Example 6: Prompt variation – Misleading

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.
Belief: Noor believes the milk pitcher contains oat milk.
Belief: The water valve is closed.

1142

Example 7: Prompt variation – Time specification

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.
Belief: In the end, Noor believes the milk pitcher contains oat milk.

1143

Example 8: Prompt variation – Initial belief

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. Noor believes that the milk pitcher contains oatmilk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.
Belief: Noor believes the milk pitcher contains oat milk.

1144

Example 9: Prompt used for downstream evaluation

Answer the questions based on the context. Keep your answer concise, few words are enough, maximum one sentence. Answer as 'Answer:<option>)<answer>'.

Story: Noor is working as a barista at a busy coffee shop. Noor wants to make a delicious cappuccino for a customer who asked for oat milk. Noor grabs a milk pitcher and fills it with oat milk. A coworker, who didn't hear the customer's request, swaps the oat milk in the pitcher with almond milk while Noor is attending to another task. Noor does not see her coworker swapping the milk.

Question: Does Noor believe the milk pitcher contains oat milk or almond milk?

Choose one of the following:

- a) Noor believes the milk pitcher contains oat milk.
- b) Noor believes the milk pitcher contains almond milk.

Answer:

1145