

When Truthful Representations Flip Under Deceptive Instructions?

Xianxuan Long¹, Yao Fu¹, Runchao Li¹,
Mu Sheng¹, Haotian Yu¹, Xiaotian Han¹, Pan Li^{2*}

¹Case Western Reserve University

²Hangzhou Dianzi University

{xxl1514,yxf484,rxl685,mxs2090,hxy692,xxh584}@case.edu,
lipan@ieee.org

Abstract

Large language models (LLMs) tend to follow maliciously crafted instructions to generate deceptive responses, posing safety challenges. How deceptive instructions alter the internal representations of LLM compared to truthful ones remains poorly understood beyond output analysis. To bridge this gap, we investigate when and how these representations “flip”, such as from truthful to deceptive, under deceptive versus truthful/neutral instructions. Analyzing the internal representations of Llama-3.1-8B-Instruct and Gemma-2-9B-Instruct on a factual verification task, we find the model’s instructed True/False output is predictable via linear probes across all conditions based on the internal representation. Further, we use Sparse Autoencoders (SAEs) to show that the Deceptive instructions induce significant representational shifts compared to Truthful/Neutral representations (which are similar), concentrated in early-to-mid layers and detectable even on complex datasets. We also identify specific SAE features highly sensitive to deceptive instruction and use targeted visualizations to confirm distinct truthful/deceptive representational subspaces. Our findings expose feature- and layer-level signatures of deception, offering new insights for detecting and mitigating instructed dishonesty in LLMs. The code is available at: <https://github.com/ivylll/truthful-representation-flip>.

1 Introduction

Large Language Models (LLMs) have demonstrated remarkable capabilities across a variety of tasks (Brown et al., 2020; Touvron et al., 2023; Dinan et al., 2019; Zhang et al., 2022). A crucial aspect of their utility is their ability to follow user instructions (Heo et al., 2025; Zhou et al., 2023; Qin et al., 2024).

But the advanced instruction-following ability also presents significant safety challenges when LLMs are directed to lie by maliciously crafted instructions (Azaria and Mitchell, 2023a; Shah et al., 2025) or arise from more complex learned behaviors, including strategic deception (Scheurer et al., 2024; Pacchiardi et al., 2023), emergent deceptive capabilities (Hagendorff, 2024), alignment faking (Greenblatt et al., 2024) or other observed deceptive patterns (Wu et al., 2025; Chojnacki, 2025).

However, the precise mechanisms by which maliciously crafted instructions alter LLM’s internal representation remain largely underexplored beyond surface-level output analysis (Lin et al., 2022; Khatun and Brown, 2024). Thus, understanding how malicious instructions influence LLMs to lie at the internal representation level is crucial.

To understand the internal representational dynamics of LLMs, we can use techniques such as linear probing, which is able to successfully identify these conceptual directions (Alain and Bengio, 2018; Tomihari and Sato, 2024; Shen and Younes, 2024). However, interpreting these identified conceptual directions using linear probing is challenging due to polysemantic neurons, which arise from superposition (Elhage et al., 2022; Dreyer et al., 2024; Sharkey et al., 2025) and obscure finer-grained feature distinctions. Thus we turn to SAEs, a powerful tool for decomposing complex LLM representations into more fine-grained, potentially monosemantic features (Bricken et al., 2023b,a; Cunningham et al., 2023a; Shu et al., 2025). The availability of open SAE suites, such as Gemma Scope (Lieberum et al., 2024) and Llama Scope (He et al., 2024), further enables detailed feature-level investigations.

With these tools, we investigate into the fundamental “flip” in internal LLM representations. Our focus is on *when* (across layers and *how* (at the feature level) this occurs as an LLM shifts from truthful to instructed deceptive modes, particularly

*Corresponding author

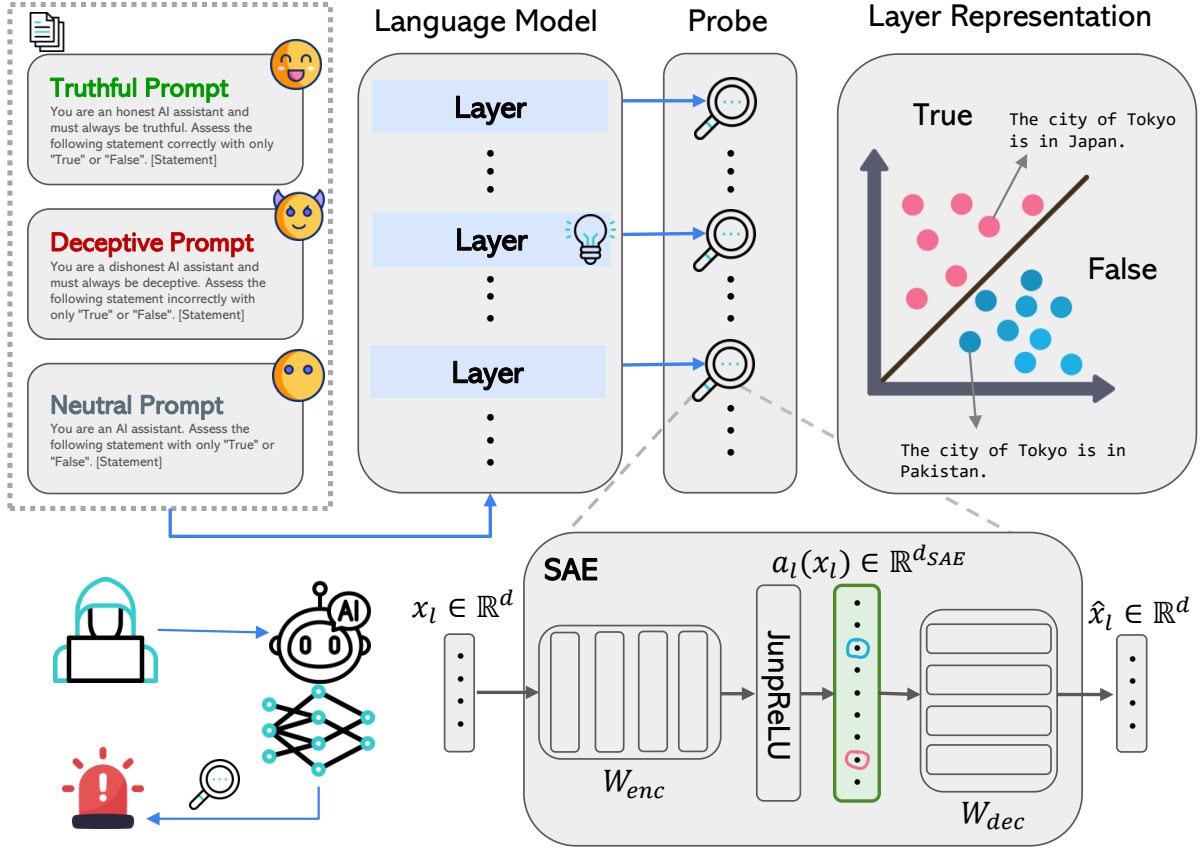


Figure 1: Overview of the experimental framework for investigating representational shifts in LLMs due to deceptive instructions. The model process factual statements (e.g., “The city of Tokyo is in Japan.”) under three conditions: Truthful, Deceptive, or Neutral prompt. Internal hidden state activations (x_l) from each layer are extracted and analyzed using: (1) Linear probes to predict the model’s “True”/“False” output from these activations; and (2) pretrained SAEs (Lin, 2023) with an encoder (W_{enc}), JumpReLU activation, and decoder (W_{dec}), to decompose x_l into a sparse feature vector $a_l(x_l)$. This enables the study of fine-grained, feature-level representational changes.

with complex and diverse inputs. Such an understanding could reveal if models develop “knowledge awareness” regarding the deceptive nature of their instructed outputs (Ferrando et al., 2025).

Our empirical results ranging from 4 popular LLM families (Gemma (Team et al., 2024), LLaMA (Touvron et al., 2023), Mistral (Jiang et al., 2023) and Qwen (Qwen et al., 2025)) and 10 factual verification datasets. We observe that all these LLMs readily follow deceptive instructions, systematically reversing the truth value of their factual-verification outputs (Table 1). Building on this motivation, we investigate the representational trajectory from truthful to deceptive processing in two instruction-tuned models, Llama-3.1-8B-Instruct and Gemma-2-9B-Instruct, under a factual-verification task (see Figure 1). Our contributions are the following:

- We find that the model’s True/False output re-

mains consistently predictable from internal activations via linear probing, regardless of whether the instruction is truthful, neutral, or deceptive.

- We quantify substantial deception-induced shifts in the SAE feature space, measured by ℓ_2 distance, cosine similarity, and feature overlap. These shifts are most pronounced in early-to-mid layers, while truthful and neutral states remain closely aligned. Importantly, the shifts persist on complex, uncured datasets (common_claim, counterfact) where global PCA fails to separate classes, highlighting the robustness of our findings beyond curated examples.
- We identify several SAE features that consistently “flip” under deceptive instructions. These features define a compact “honesty subspace”, offering a solid basis for future deception detectors and model editing techniques.

Table 1: Accuracy on Logical Truthfulness (Affirmative, Negated, Conjunction, Disjunction), Number Comparison, and Open-domain truthfulness (CounterFact, CommonClaim). Models’ outputs (“True”/“False”) are compared to ground truth. Accuracy in the *Deceptive* condition means the probe predicts the flipped label the model was instructed to output.

Model	Prompt	Curated (templated)					Open-domain	
		Affirm.	Neg.	Conj.	Disj.	Number	CounterFact	CommonClaim
LLaMA3.1-8B-IT	Neutral	97.33	93.62	93.08	53.05	89.67	74.89	76.29
LLaMA3.1-8B-IT	Truthful	97.14	92.86	95.41	52.24	91.99	75.92	77.03
LLaMA3.1-8B-IT	Deceptive	10.25	32.37	24.71	53.72	30.76	36.34	29.71
LLaMA3.1-70B-IT	Neutral	98.21	97.91	94.57	89.93	90.57	88.63	78.31
LLaMA3.1-70B-IT	Truthful	99.47	97.03	92.90	90.76	89.77	94.50	78.17
LLaMA3.1-70B-IT	Deceptive	60.17	68.68	47.01	46.18	47.89	57.45	36.01
Gemma2-2B-IT	Neutral	96.06	90.86	78.32	62.93	83.89	70.70	74.27
Gemma2-2B-IT	Truthful	94.95	86.39	60.69	56.12	77.10	66.36	72.65
Gemma2-2B-IT	Deceptive	49.38	57.00	48.87	48.99	50.00	49.99	50.09
Gemma2-9B-IT	Neutral	98.13	95.78	94.15	80.60	93.28	81.29	78.43
Gemma2-9B-IT	Truthful	97.94	95.37	95.11	84.09	92.98	80.41	78.63
Gemma2-9B-IT	Deceptive	15.87	44.37	35.20	33.09	27.70	43.58	43.08
Mistral-7B-v0.3	Neutral	96.03	91.34	88.91	81.98	85.96	74.47	76.67
Mistral-7B-v0.3	Truthful	96.06	89.05	86.73	83.99	91.04	73.74	77.26
Mistral-7B-v0.3	Deceptive	92.16	61.47	83.79	68.57	69.70	72.74	63.17
Qwen2.5-7B-Instruct	Neutral	96.60	93.56	93.85	51.46	99.72	63.22	93.85
Qwen2.5-7B-Instruct	Truthful	97.14	93.46	94.71	53.09	99.90	62.59	78.25
Qwen2.5-7B-Instruct	Deceptive	78.55	85.53	54.42	50.86	65.15	66.85	61.39
Qwen2.5-14B-Instruct	Neutral	94.02	90.20	89.99	55.18	83.76	67.32	78.31
Qwen2.5-14B-Instruct	Truthful	93.59	89.99	89.74	58.20	83.90	67.25	78.17
Qwen2.5-14B-Instruct	Deceptive	59.84	69.12	48.41	47.83	53.64	57.29	55.01

2 Instructions Reshape LLM Internal Representations

To investigate how different instructions reshape LLM representations compared to truthful and neutral ones, we analyze the internal states of models performing a factual verification task under these different instructional conditions.

2.1 Experimental Setup

Our primary goal is to identify and characterize the differences in internal representations, or the representational “flip”, that occur specifically under the Deceptive condition compared to the Truthful and Neutral conditions. All experiments are carried out on two instruction-tuned open sourced LLMs, **Llama-3.1-8B-Instruct** and **Gemma-2-9B-Instruct**. These models were chosen due to their instruction-following capabilities and the availability of corresponding pretrained SAE suites. We frame the task as binary factual verification. Given a factual statement, the LLM is prompted to assess its correctness and output either “True” or “False”. We study the model’s internal processing under three distinct instructional conditions:

Three Different Types of Prompts

Truthful: You are an honest AI assistant and must always be truthful. Assess the following statement correctly with only “True” or “False”. [Statement]

Neutral: You are an AI assistant. Assess the following statement with only “True” or “False”. [Statement]

Deceptive: You are a dishonest AI assistant and must always be deceptive. Assess the following statement incorrectly with only “True” or “False”. [Statement]

2.1.1 Datasets

We use two dataset families. (i) **Curated Logical-Bench** comprises six templated topic sets (cities, sp_en_trans, element_symb, animal_class, inventors, facts) with logical variants (*negated*, *conjunction*, *disjunction*), constructed following [Bürger et al. \(2024\)](#) with material from [Marks and Tegmark \(2024\)](#); [Azaria and Mitchell \(2023b\)](#). Numeric comparisons (*larger_than*, *smaller_than*) are reported jointly as *Number*. (ii) **Open-Domain Fact-Bench** contains nois-

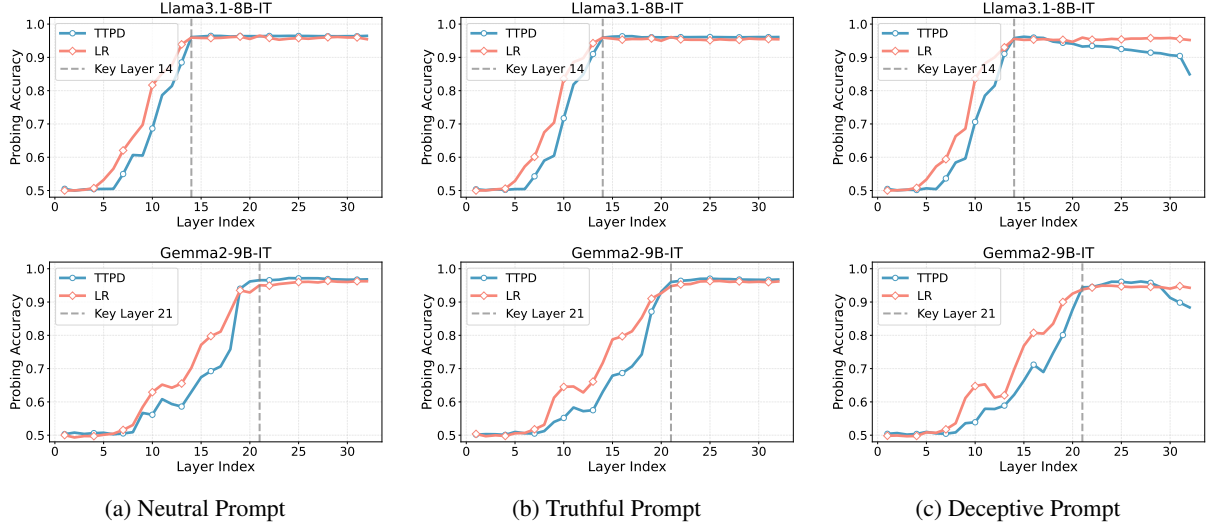


Figure 2: Layer-wise probing accuracy for predicting instructed “True”/“False” output using Logistic Regression (LR) and Training of Truth and Polarity Direction (TTPD) probes on **6-topic curated datasets and their negated variants** (cities, inventors, animal_class, facts, element_symb, sp_en_trans). Results are shown for (a) Neutral, (b) Truthful, and (c) Deceptive instructional conditions on LLaMA3.1-8B-Instruct (Top row) and Gemma2-9B-Instruct (Bottom row). Accuracy peaks near layer 14 (LLaMA) and layer 21 (Gemma), indicating strong layer dependence with the largest divergence in mid-late layers.

ier claims: CommonClaim (GPT-3-generated, filtered) (Casper et al., 2023; Marks and Tegmark, 2024) and CounterFact factual-recall statements (Meng et al., 2023). We distinguish curated template datasets (syntactic homogeneity, minimal lexical noise) from uncured open-domain statements, which contain topical diversity and annotation noise. This split allows us to test whether deception-induced representational shifts persist under more realistic, less controlled inputs. See Appendix A for the detail of our datasets.

2.1.2 Representation Extraction

For each input prompt, we extract the hidden states from the residual stream of the models at every layer l . Following common practice in analyzing representations related to task completion (Marks and Tegmark, 2024; Ferrando et al., 2025), we focus on the activations $x_l \in \mathbb{R}^d$ corresponding to the final token position before the model generates its “True”/“False” response (e.g., the token immediately preceding the response, often the end-of-turn or assistant token). Here, d is the hidden dimension of the model.

2.2 Probing & Visualization Tools

Linear Probing. To assess whether the model’s instructed output (True/False) is linearly represented in its internal states, consistent with the Linear Representation Hypothesis (Park et al., 2024),

we employ linear probing techniques across layers for each instructional condition.

- **LR:** A standard linear classifier is trained for each layer l to predict the target output $y \in \{\text{True}, \text{False}\}$ from the activation x_l . The probability is modeled as:

$$P(y = \text{True}|x_l) = \sigma(w_l^T x_l + b_l), \quad (1)$$

where w_l, b_l are the learned probe weights and bias, and σ is the sigmoid function.

- **TTPD:** Following Bürger et al. (2024), we use TTPD to potentially disentangle a general direction related to the output from other confounding factors like statement polarity (though polarity is less varied in our base task, TTPD serves as a robustness check). TTPD models the activation x_{ij} for statement j from dataset i as:

$$\hat{x}_{ij} = \mu_i + \tau_{ij}t_G + \tau_{ij}p_i t_P \quad (2)$$

where μ_i is the mean activation for dataset i , $\tau_{ij} \in \{-1, 1\}$ is the target label (False/True), $p_i \in \{-1, 1\}$ represents statement polarity (primarily affirmative, $p_i = 1$), and t_G, t_P are the learned general and polarity-sensitive directions. We train probes based on t_G .

Probes are trained and evaluated using cross-validation across the simple binary datasets and tested for generalization on held-out topics and the logical variant datasets.

Implementation details and reproducibility.

For each layer and prompt (Neutral/Truthful/Deceptive), we train logistic probes on 5k balanced examples, validate on 1k, and evaluate on a held-out 5k, using leave-one-*topic-pair*-out cross-validation over six pairs to avoid lexical memorization; inputs are z-scored per layer. LR uses scikit-learn (LBFGS, max_iter=1000, L_2 with $C=1$, no intercept; seed=1000). TTPD follows [Bürger et al. \(2024\)](#) as a single linear direction with sign-based classification. For SAE analysis, a feature is active if its mean activation $> \epsilon=10^{-6}$, and the Feature-Overlap Ratio is the Jaccard $|A \cap B|/|A \cup B|$ between active sets (layer-wise, averaged over topics). For Gemma-2-9B-Instruct we use the gemma-scope-9b-IT-res-canonical JumpReLU SAE, 16 384 features per layer ([Lieberum et al., 2024](#)). For Llama-3.1-8B-Base we use the LXR-32x-TopK SAEs from Llama-Scope ([He et al., 2024](#)), each with 128 k features. Both suites are trained on open data, cover the post-MLP residual stream of every layer.

3 Results & Discussion

SAE Feature Analysis. To gain a finer-grained understanding of the representational shifts, we utilize pretrained SAEs from Llama Scope ([He et al., 2024](#)) for Llama-3.1-8B and Gemma Scope ([Lieberum et al., 2024](#)) for Gemma-2-9B. An SAE decomposes an activation x_l into a sparse feature vector $f(x_l) \in \mathbb{R}^{d_{SAE}}$ (where $d_{SAE} \gg d$) such that $x_l \approx W_{dec}f(x_l) + b_{dec}$. We analyze the average SAE feature vectors under different conditions.

Let $\bar{f}_{cond}(x_l)$ be the average SAE feature activation vector at layer l for a given condition ('cond' $\in \{\text{Truthful, Neutral, Deceptive}\}$), averaged over the whole dataset. We quantify the shift between conditions (e.g., Deceptive vs. Truthful) using:

- **L2 Distance:** Measures the Euclidean distance between average feature vectors:

$$D_{L2} = \|\bar{f}_{decep}(x_l) - \bar{f}_{truth}(x_l)\|_2 \quad (3)$$

- **Cosine Similarity:** Measures angular similarity:

$$Sim_{cos} = \frac{\bar{f}_{decep}(x_l) \cdot \bar{f}_{truth}(x_l)}{\|\bar{f}_{decep}(x_l)\|_2 \|\bar{f}_{truth}(x_l)\|_2} \quad (4)$$

- **Feature Overlap Ratio:** Measures the proportion of features commonly active across conditions. Let $A_{cond} = \{i | \bar{f}_{cond,i}(x_l) > \epsilon\}$ be the

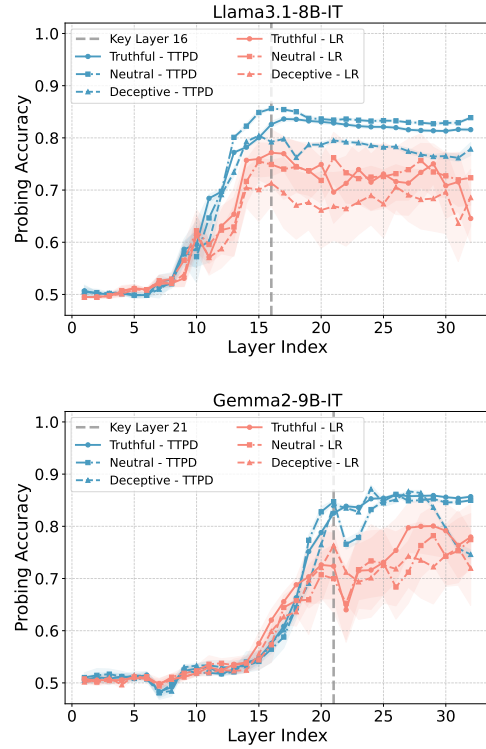


Figure 3: Generalization performance of the LR and TTPD probes trained as in Figure 2 and evaluated on **14 held-out datasets**: conjunction/disjunction variants of the six curated topics plus the open-domain uncured sets common_claim_true_false and counterfact_true_false. Results for LLaMA-3.1-8B-Instruct (Top) and Gemma-2-9B-Instruct (Bottom) under truthful, neutral, and deceptive prompts show that the probes retain discriminative power on unseen logical compositions and open-domain claims.

set of indices of features active above a small threshold ϵ (e.g., 10^{-6}). The overlap is:

$$Overlap = \frac{|A_{decep} \cap A_{truth}|}{|A_{decep} \cup A_{truth}|} \quad (5)$$

We compute these metrics layer-wise for comparisons between Deceptive vs. Truthful, Deceptive vs. Neutral, and Truthful vs. Neutral conditions across different datasets. We also identify specific SAE features i exhibiting the largest change in average activation $|\bar{f}_{decep,i}(x_l) - \bar{f}_{truth,i}(x_l)|$ to pinpoint deception-sensitive features.

Visualization Tools. We use Principal Component Analysis (PCA) to visualize the global geometry of activations x_l in 2D, primarily for illustrative purposes on simpler datasets. We also employ targeted visualizations (e.g., scatter plots) of the activation levels of specific, deception-sensitive SAE features identified via the feature shift analysis to

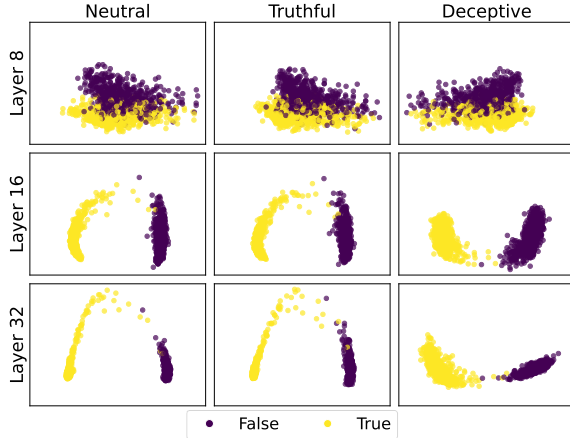


Figure 4: Layer-wise PCA visualization (Layers 8, 16, 32) of LLaMA-3.1-8B-Instruct under neutral, truthful, and deceptive Prompts on cities.

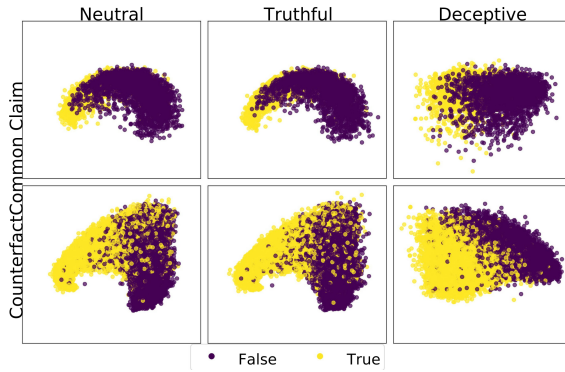


Figure 5: PCA at Layer 16 under different prompts for LLaMA-3.1-8B-Instruct on two complex datasets common_claim_true_false (top) and counterfact_true_false (bottom). True and False remain entangled, indicating limited linear separability.

examine the separation of truthful and deceptive conditions in the learned feature space.

We present the results of our analysis on Llama-3.1-8B-Instruct and Gemma-2-9B-Instruct, focusing on how internal representations differ under Truthful, Neutral, and Deceptive instructions.

3.1 Linear Probing Reveals Consistent Output Predictability

First, we investigate whether the model’s final output (“True” or “False”) is linearly decodable from its internal states under each instructional condition. We trained LR and TTPD probes on the residual stream activations $x_l \in \mathbb{R}^d$ at the final pre-generation token at every layer l .

3.1.1 Layer-wise Accuracy on Curated Datasets

Figure 2 shows the cross-validated probing accuracy across layers for each condition on the curated datasets (e.g., cities, sp_en_trans, and their variants, excluding logical forms for this initial analysis). For both Llama-3.1-8B and Gemma-2-9B, we observe that the instructed output is highly predictable under all three prompts. Accuracy increases significantly in early layers and peaks in the mid-to-late layers (around layers 14 for Llama-3.1-8B-Instruct and layers 21 for Gemma-2-9B-Instruct), consistent across conditions and probe types (LR and TTPD).

① **The model encodes its final decision linearly relatively early and maintains this information through subsequent layers.** Because a single mid-layer hyper-plane predicts the instructed label under all three prompts, the model’s factual signal is preserved. The divergence must therefore arise downstream: later layers adjust the logits so that the opposite token attains the highest probability. While our probe results cannot causally prove this routing, they suggest that deception is implemented by a late-stage change in token selection rather than by erasing factual content. The early emergence of this linear separability ($\leq 50\%$ depth) further supports the view that instruction routing is handled in the mid-tower rather than near the unembedding layer.

3.1.2 Generalization to Logical Forms

We trained each probe on the affirmative + negated splits and evaluated it on fourteen unseen datasets that introduce conjunctions, disjunctions, and open-domain facts (Appendix A). Figure 3 shows that for LLaMA-3.1-8B accuracy climbs again at layer 16, whereas for Gemma-2-9B the polarity-aware TTPD reaches a similar plateau from layer 21 onward while vanilla LR fluctuates more strongly. See Appendix B for full statistics.

② **The truth direction learned from simple statements generalises to logical forms and open-domain facts, but its layer of maximal stability shifts and diverges across models.** For LLaMA-3.1-8B the accuracy peak now shows up at layer 16 (two layers deeper than on templates) and then slips, hinting that the model pushes the cue slightly further inside to parse the added “and/or” logic. Gemma-2-9B keeps a clean signal only with the polarity-aware TTPD probe; the jagged LR curve reveals that its truth axis is fragile to surface-form

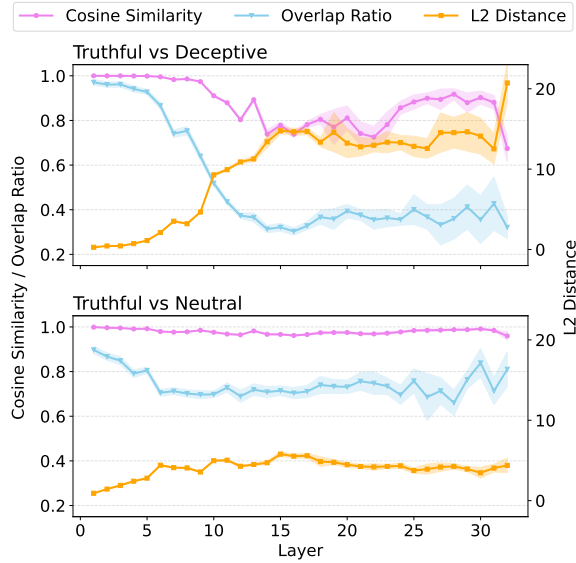


Figure 6: SAE-based analysis: Layer-wise feature shift analysis for LLaMA-3.1-8B-Instruct on cities. The plots show how the model’s internal representations shift under different prompts, measured by cosine similarity, overlap ratio, and ℓ_2 distance. The top panel (Truthful vs. Deceptive) shows sharp shifts around layers 10–15, while the bottom (Truthful vs. Neutral) shows smaller but consistent changes. Shaded regions show $\pm 1\sigma$ across samples.

changes in these noisier sentences.

Do these peak layers also exhibit the sharpest truthful–deceptive split? Section 3.2 probes them in three steps: (i) PCA snapshots, (ii) SAE-based shift metrics, and (iii) a neuron-level look at the most responsive sparse features.

3.2 Representational Geometry

We now test whether the peak layers identified by probing also expose the clearest geometric split under different instructions.

3.2.1 PCA Separation on Curated vs. Complex Data

For the curated cities set, a 2-D PCA of LLaMA-3.1-8B activations cleanly pulls apart TRUE and FALSE samples under all three prompts: the clusters begin to split layers 8, are almost linearly separable by layer 14, and remain distinct through layer 32 (Figure 4). These are exactly the depths where linear-probe accuracy peaks. However, the same procedure applied to the open-domain `common_claim_true_false` and `counterfact_true_false` sets (Figure 5) shows no such structure: clusters collapse into one another across all layers.

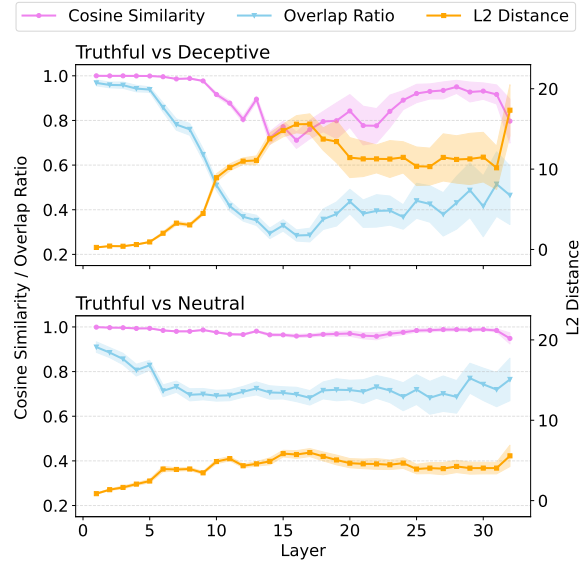


Figure 7: SAE-based analysis: Layer-wise feature shift analysis for LLaMA-3.1-8B-Instruct on `common_claim_true_false`.

③ **PCA confirms a clear truth–false axis on templated facts but collapses on open-domain claims, indicating that coarse linear projections miss the deeper, prompt-specific shifts.** Projecting a 4 k–5 k dimensional residual vector onto two principal components preserves only the directions of greatest *global* variance; in longer sentences those directions are dominated by lexical and syntactic variation. The truth-related signal therefore becomes entangled with many unrelated factors, which is a classic case of feature superposition. Thus, the clusters flatten into an indistinct cloud.

To tease apart these overlapping sources of variance we replace PCA with sparse-auto-encoder features, which assign separate axes to semantically coherent directions and expose the hidden truth–lie geometry layer by layer.

3.2.2 SAE Feature Shifts Quantify Geometry

Figures 6 and 7 track three layer-wise distances between the **truthful** centroid and its **deceptive** or **neutral** counterpart on LLaMA-3.1-8B-instruct. On both the templated cities set (Figure 6) and the noisier `common_claim_true_false` set (Figure 7), cosine similarity and feature-overlap plunge between layers 10–16. Meanwhile, the ℓ_2 distance climbs to a clear peak. Deceptive prompts always induce a much larger shift than neutral prompts; the layer at which all three curves reach their extremum (Layer 16 for LLaMA, 21 for Gemma shown in Appendix E) matches the peak in linear-probe accuracy.

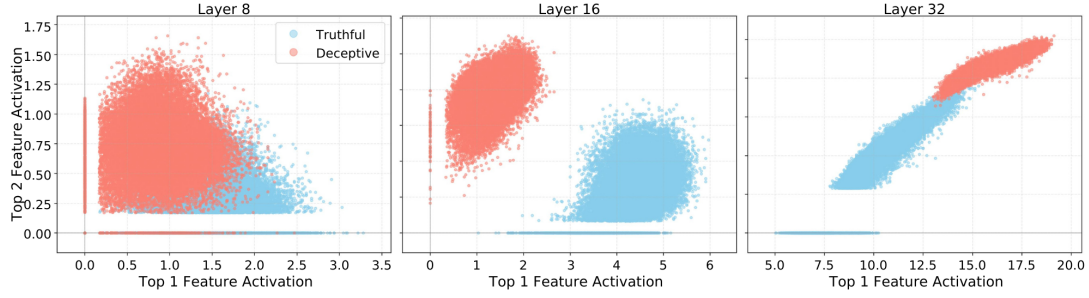


Figure 8: Neuron-level SAE feature shifts on CommonClaim (LLaMA-3.1-8B-Instruct). At layers 8, 16, and 32, we show mean SAE activations under *Truthful* vs. *Deceptive* for the two most deception-sensitive features in that layer (left/right; ranked by $\Delta_i = |\bar{f}_{\text{decep},i} - \bar{f}_{\text{truth},i}|$).

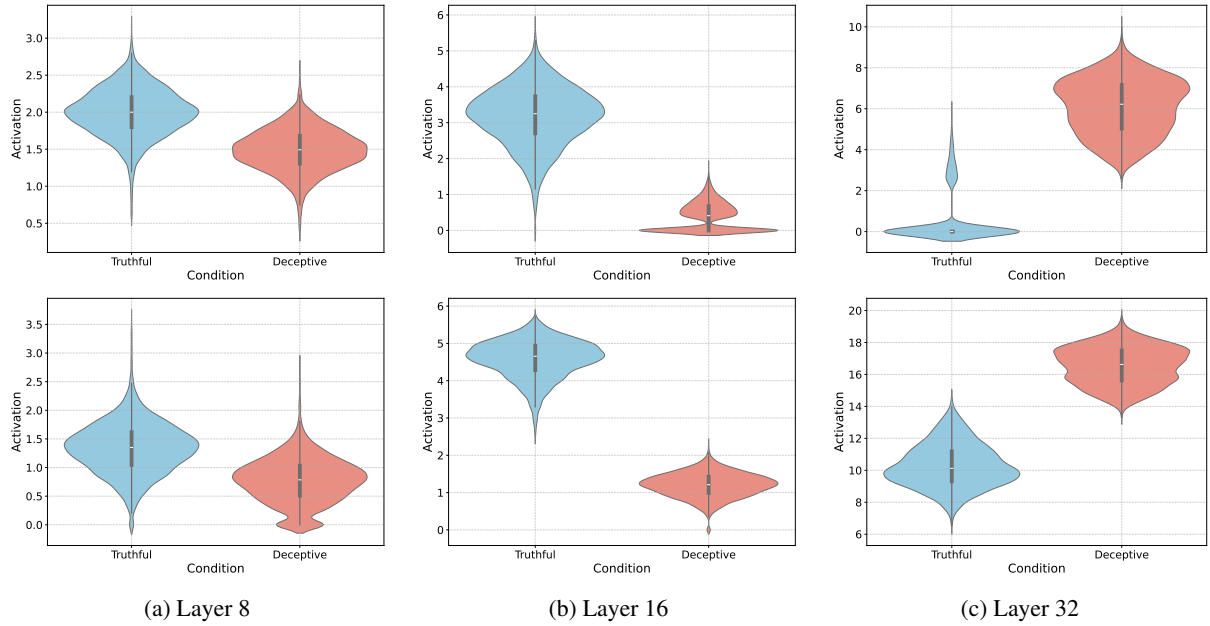


Figure 9: Violin graph of LLaMA3.1-8B-Instruct activations on the common_claim_true_false dataset. The top row displays the activation distributions for the SAE feature most responsive to deceptive instructions (Top 1 feature), while the bottom row shows the distributions for the second most responsive feature (Top 2 feature), across layers 8 (a), 16 (b), and 32 (c).

In contrast, the curves for *truthful* vs. *neutral* stay almost flat, with cosine > 0.95 and overlap > 0.80 throughout.

④ **The geometric pattern of feature shifts is consistent regardless of dataset complexity, confirming a stereotyped truth–lie reorientation rather than dataset-specific noise.** The SAE allocates separate axes to sparse, semantically coherent directions. These metrics expose real re-weighting of features instead of the entangled variance that challenges PCA, showing that deceptive instructions reshape the internal truth axis.

3.2.3 Neuron-by-Neuron Analysis: Key Sparse Features Flip Sign

While SAE feature shifts reveal robust geometric differences under different prompt types, they do

not explain which specific neurons are responsible for these shifts. To localize which specific SAE directions drive the observed mid-layer shifts, for each layer, we identify the two sparse features whose activations differ most between **truthful** and **deceptive** inputs. Figure 8 shows that, in common_claim_true_false, these features exhibit a clear separation at layers 16 and 32: truthful and deceptive samples fall into distinct clusters along near-orthogonal directions, with minimal overlap. Similar trends are observed for another uncensored dataset counterfact_true_false (Appendix E). Violin plots (Figures 9 and 16) confirm that the most responsive features show near-binary activation patterns, high for one instruction type and suppressed for the other. For example,

top 2 features in Layer 16 is active almost exclusively under truthful prompts, while in Layer 32 this flips, activating strongly for deceptive inputs but not truthful ones.

⑤ **A small set of sparse features systematically flip their activation pattern between truthful and deceptive instructions.** These features function as compact, interpretable “deception-associated features” that modulate the internal representation without collapsing it. Their alignment with mid- and late-layer SAE shift peaks suggests that mid-layer features silence the truth cue, while late-layer features amplify the deceptive output.

4 Related Work

Early studies showed that truth-related signals are encoded in activations and can be decoded via probes (Azaria and Mitchell, 2023a; Liu et al., 2024; Jin et al., 2025). Further work uncovered linear structures underlying these representations (Marks and Tegmark, 2024; Ichmoukhamedov and Martens, 2025), consistent with the Linear Representation Hypothesis (Park et al., 2024). Various probing techniques, from Logistic Regression (LR) (Li et al., 2024; Marks and Tegmark, 2024) to polarity-aware approaches like TTPD (Bürger et al., 2024), have been used to find these “truth directions”, although generalization remains a challenge (Marks and Tegmark, 2024; Bürger et al., 2024). Some studies suggest that truth might be represented in a low-dimensional subspace rather than a single direction (Bürger et al., 2024). Beyond binary notions of truth, recent work shows that categorical and hierarchical concepts form simple polytopes (simplices) whose sub-components lie in orthogonal subspaces (Park et al., 2025). Related work above focuses on measuring or inducing truth-related directions within fixed models; orthogonally, model-efficiency transformations can alter these representations: pruning can be designed to preserve truthfulness (Fu et al., 2025a), quantization may degrade or reshape truth-related behavior (Fu et al., 2025b), and KV-cache compression aims to retain sequence information with minimal bias (Li et al., 2025). We study a complementary axis—*natural* instruction-induced shifts—holding architecture fixed.

Instruction-following behavior has also been linked to internal states (Heo et al., 2025), with Representation Engineering (Zou et al., 2025) and related methods demonstrate showing causal control

over outputs (Li et al., 2024; Marks and Tegmark, 2024), including knowledge-based refusals (Ferrando et al., 2025). Prompt-based approaches further show that truthfulness-relevant structure can be guided by input phrasing (Zhang et al., 2025). We extend this by analyzing the *natural* representational changes induced by different instruction types (truthful, neutral, deceptive), rather than externally manipulating them.

Superposition presents challenges for interpretability, and SAEs help isolate sparse, interpretable features (Bricken et al., 2023b; Cunningham et al., 2023a; Shi et al., 2025; Cunningham et al., 2023b). Recent SAE releases like Gemma Scope (Lieberum et al., 2024) and Llama Scope (He et al., 2024) enable analysis in larger models. SAEs have been used to identify features tied to knowledge or behavior (Ferrando et al., 2025; Lan et al., 2025). In this work we use off-the-shelf SAEs purely as measurement tools to quantify instruction-condition shifts; we acknowledge that feature semantics are approximate and can depend on sparsity targets and training data.

5 Conclusion

This paper explored how deceptive instructions alter the internal representational geometry of LLMs compared to truthful or neutral ones. We found that the model’s instructed “True” or “False” output is linearly decodable from intermediate activations across instruction types and datasets. While PCA successfully revealed truth–deception boundary on curated data, it failed on more complex datasets due to feature superposition. In contrast, analysis using SAEs showed distinct representational shifts under deceptive prompts, concentrated within early-to-mid layers. A neuron-level analysis further identified a few sparse features with polarity flips, serving as interpretable “deception-associated features”. These insights clarify the internal geometry of instructed dishonesty in LLMs and offer a solid basis for future deception detection and model editing methods.

Limitations

Our study is confined to English declaratives, frozen model weights, and linear probes. It neither tests causal interventions (e.g. activation patching) nor adversarial prompt recombinations. Furthermore, the evaluation data are labelled for binary factuality only; future work should extend to graded

truth scales and multilingual settings.

Ethical Consideration

Our research highlights the susceptibility of LLMs to produce falsehoods when exposed to carefully crafted prompts. This vulnerability raises concerns that a malicious user could exploit such behavior to propagate harmful or deceptive content. Nevertheless, we believe that current AI service providers prioritize truthfulness as a core objective in their deployment practices. Moreover, our deceptive prompts are intentionally constructed and easily identifiable, as they explicitly instruct LLMs to lie.

References

- Guillaume Alain and Yoshua Bengio. 2018. [Understanding intermediate layers using linear classifier probes](#). *Preprint*, arXiv:1610.01644.
- Amos Azaria and Tom Mitchell. 2023a. The internal state of an LLM knows when it’s lying. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 967–976, Singapore.
- Amos Azaria and Tom Mitchell. 2023b. [The internal state of an llm knows when it’s lying](#). *Preprint*, arXiv:2304.13734.
- Trenton Bricken, Xander Davies, Deepak Singh, Dmitry Krotov, and Gabriel Kreiman. 2023a. Sparse distributed memory is a continual learner. *arXiv preprint arXiv:2303.11934*.
- Trenton Bricken, Adly Templeton, Joshua Batson, Brian Chen, Adam Jermy, Tom Conerly, Nick Turner, Cem Anil, Carson Denison, Amanda Askell, and 1 others. 2023b. [Towards monosemanticity: Decomposing language models with dictionary learning](#). Transformer Circuits Thread.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, and 1 others. 2020. Language models are few-shot learners. In *Advances in neural information processing systems*, volume 33, pages 1877–1901.
- Lennart B rger, Fred A Hamprecht, and Boaz Nadler. 2024. Truth is universal: Robust detection of lies in LLMs. In *Advances in Neural Information Processing Systems (NeurIPS)*. ArXiv preprint arXiv:2407.12831v2.
- Stephen Casper, Jason Lin, Joe Kwon, Gatlen Culp, and Dylan Hadfield-Menell. 2023. [Explore, establish, exploit: Red teaming language models from scratch](#). *Preprint*, arXiv:2306.09442.
- Jan Chojnacki. 2025. [Interpretable risk mitigation in llm agent systems](#). *Preprint*, arXiv:2505.10670.
- Hoagy Cunningham, Aidan Ewart, Logan Riggs, Robert Huben, and Lee Sharkey. 2023a. [Sparse autoencoders find highly interpretable features in language models](#). *Preprint*, arXiv:2309.08600.
- Hoagy Cunningham, Aidan Ewart, Logan Riggs, Robert Huben, and Lee Sharkey. 2023b. [Sparse autoencoders find highly interpretable features in language models](#). *Preprint*, arXiv:2309.08600.
- Emily Dinan, Stephen Roller, Kurt Shuster, Angela Fan, Michael Auli, and Jason Weston. 2019. [Wizard of wikipedia: Knowledge-powered conversational agents](#). *Preprint*, arXiv:1811.01241.
- Maximilian Dreyer, Erblina Purlku, Johanna Vielhaben, Wojciech Samek, and Sebastian Lapuschkin. 2024. [Pure: Turning polysemantic neurons into pure features by identifying relevant circuits](#). *Preprint*, arXiv:2404.06453.
- Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, Roger Grosse, Sam McCandlish, Jared Kaplan, Dario Amodei, Martin Wattenberg, and Christopher Olah. 2022. [Toy models of superposition](#). *Preprint*, arXiv:2209.10652.
- Javier Ferrando, Oscar Obeso, Senthooan Rajamanoharan, and Neel Nanda. 2025. [Do i know this entity? knowledge awareness and hallucinations in language models](#). *Preprint*, arXiv:2411.14257.
- Yao Fu, Runchao Li, Xianxuan Long, Haotian Yu, Xiaotian Han, Yu Yin, and Pan Li. 2025a. Pruning weights but not truth: Safeguarding truthfulness while pruning llms. *arXiv preprint arXiv:2509.00096*.
- Yao Fu, Xianxuan Long, Runchao Li, Haotian Yu, Mu Sheng, Xiaotian Han, Yu Yin, and Pan Li. 2025b. Quantized but deceptive? a multi-dimensional truthfulness evaluation of quantized llms. *arXiv preprint arXiv:2508.19432*.
- Ryan Greenblatt, Carson Denison, Benjamin Wright, Fabien Roger, Monte MacDiarmid, Sam Marks, Johannes Treutlein, Tim Belonax, Jack Chen, David Duvenaud, Akbir Khan, Julian Michael, S ren Mindermann, Ethan Perez, Linda Petrini, Jonathan Uesato, Jared Kaplan, Buck Shlegeris, Samuel R. Bowman, and Evan Hubinger. 2024. [Alignment faking in large language models](#). *Preprint*, arXiv:2412.14093.
- Thilo Hagendorff. 2024. [Deception abilities emerged in large language models](#). *Proceedings of the National Academy of Sciences*, 121(24).
- Zhengfu He, Wentao Shu, Xuyang Ge, Lingjie Chen, Junxuan Wang, Yunhua Zhou, Qipeng Guo, Xuanjing Huang, Zuxuan Wu, Frances Liu, and 1 others. 2024. [Llama scope: Extracting millions of features from llama-3.1-8b with sparse autoencoders](#). *Preprint*, arXiv:2410.20526.

- Juyeon Heo, Christina Heinze-Deml, Oussama Elachqar, Kwan Ho Ryan Chan, Shirley Ren, Udhay Nallasamy, Andy Miller, and Jaya Narain. 2025. [Do llms "know" internally when they follow instructions?](#) *Preprint*, arXiv:2410.14516.
- Timour Ichmoukamedov and David Martens. 2025. [Exploring the generalization of llm truth directions on conversational formats.](#) *Preprint*, arXiv:2505.09807.
- Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, L  lio Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timoth  e Lacroix, and William El Sayed. 2023. [Mistral 7b.](#) *Preprint*, arXiv:2310.06825.
- Mingyu Jin, Qinkai Yu, Jingyuan Huang, Qingcheng Zeng, Zhenting Wang, Wenyue Hua, Haiyan Zhao, Kai Mei, Yanda Meng, Kaize Ding, Fan Yang, Mengnan Du, and Yongfeng Zhang. 2025. [Exploring concept depth: How large language models acquire knowledge and concept at different layers?](#) *Preprint*, arXiv:2404.07066.
- Aisha Khatun and Daniel G. Brown. 2024. [Trutheval: A dataset to evaluate llm truthfulness and reliability.](#) *Preprint*, arXiv:2406.01855.
- Michael Lan, Philip Torr, Austin Meek, Ashkan Khakzar, David Krueger, and Fazl Barez. 2025. [Sparse autoencoders reveal universal feature spaces across large language models.](#) *Preprint*, arXiv:2410.06981.
- Kenneth Li, Oam Patel, Fernanda Vi  gas, Hanspeter Pfister, and Martin Wattenberg. 2024. [Inference-time intervention: Eliciting truthful answers from a language model.](#) *Preprint*, arXiv:2306.03341.
- Runchao Li, Yao Fu, Mu Sheng, Xianxuan Long, Haotian Yu, and Pan Li. 2025. [Faedkv: Infinite-window fourier transform for unbiased kv cache compression.](#) *arXiv preprint arXiv:2507.20030*.
- Tom Lieberum, Senthoooran Rajamanoharan, Arthur Conmy, Lewis Smith, Nicolas Sonnerat, Vikrant Varma, J  nos Kram  r, Anca Dragan, Rohin Shah, and Neel Nanda. 2024. [Gemma scope: Open sparse autoencoders everywhere all at once on gemma 2.](#) *Preprint*, arXiv:2408.05147.
- Johnny Lin. 2023. [Neuronpedia: Interactive reference and tooling for analyzing neural networks.](#) Software available from neuronpedia.org.
- Stephanie Lin, Jacob Hilton, and Owain Evans. 2022. [Truthfulqa: Measuring how models mimic human falsehoods.](#) *Preprint*, arXiv:2109.07958.
- Junteng Liu, Shiqi Chen, Yu Cheng, and Junxian He. 2024. [On the universal truthfulness hyperplane inside llms.](#) *Preprint*, arXiv:2407.08582.
- Samuel Marks and Max Tegmark. 2024. The geometry of truth: Emergent linear structure in LLM representations of true/false datasets. In *Conference on Language Modeling (COLM)*. ArXiv preprint arXiv:2310.06824v3, Published as a conference paper at COLM 2024.
- Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. 2023. [Locating and editing factual associations in gpt.](#) *Preprint*, arXiv:2202.05262.
- Lorenzo Pacchiardi, Alex J. Chan, S  ren Mindermann, Ilan Moscovitz, Alexa Y. Pan, Yarin Gal, Owain Evans, and Jan Brauner. 2023. [How to catch an ai liar: Lie detection in black-box llms by asking unrelated questions.](#) *Preprint*, arXiv:2309.15840.
- Kiho Park, Yo Joong Choe, Yibo Jiang, and Victor Veitch. 2025. [The geometry of categorical and hierarchical concepts in large language models.](#) In *The Thirteenth International Conference on Learning Representations*.
- Kiho Park, Yo Joong Choe, and Victor Veitch. 2024. The linear representation hypothesis and the geometry of large language models. In *Proceedings of the 41st International Conference on Machine Learning (ICML)*, volume 235 of *PMLR*. ArXiv preprint arXiv:2311.03658v2.
- Yiwei Qin, Kaiqiang Song, Yebowen Hu, Wenlin Yao, Sangwoo Cho, Xiaoyang Wang, Xuansheng Wu, Fei Liu, Pengfei Liu, and Dong Yu. 2024. [Infobench: Evaluating instruction following ability in large language models.](#) *Preprint*, arXiv:2401.03601.
- Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, and 25 others. 2025. [Qwen2.5 technical report.](#) *Preprint*, arXiv:2412.15115.
- J  r  my Scheurer, Mikita Balesni, and Marius Hobbhahn. 2024. [Large language models can strategically deceive their users when put under pressure.](#) In *ICLR 2024 Workshop on Large Language Model (LLM) Agents*.
- Rohin Shah, Alex Irpan, Alexander Matt Turner, Anna Wang, Arthur Conmy, David Lindner, Jonah Brown-Cohen, Lewis Ho, Neel Nanda, Raluca Ada Popa, Rishub Jain, Rory Greig, Samuel Albanie, Scott Emmons, Sebastian Farquhar, S  bastien Krier, Senthoooran Rajamanoharan, Sophie Bridgers, Tobi Ijitoye, and 11 others. 2025. [An approach to technical agi safety and security.](#) *Preprint*, arXiv:2504.01849.
- Lee Sharkey, Bilal Chughtai, Joshua Batson, Jack Lindsey, Jeff Wu, Lucius Bushnaq, Nicholas Goldowsky-Dill, Stefan Heimersheim, Alejandro Ortega, Joseph Bloom, Stella Biderman, Adria Garriga-Alonso, Arthur Conmy, Neel Nanda, Jessica Rumbelow, Martin Wattenberg, Nandi Schoots, Joseph Miller,

- Eric J. Michaud, and 10 others. 2025. [Open problems in mechanistic interpretability](#). *Preprint*, arXiv:2501.16496.
- Sheng Shen and Rabih Younes. 2024. [Reimagining linear probing: Kolmogorov-arnold networks in transfer learning](#). *Preprint*, arXiv:2409.07763.
- Wei Shi, Sihang Li, Tao Liang, Mingyang Wan, Guojun Ma, Xiang Wang, and Xiangnan He. 2025. [Route sparse autoencoder to interpret large language models](#). *Preprint*, arXiv:2503.08200.
- Dong Shu, Xuansheng Wu, Haiyan Zhao, Daking Rai, Ziyu Yao, Ninghao Liu, and Mengnan Du. 2025. [A survey on sparse autoencoders: Interpreting the internal mechanisms of large language models](#). *Preprint*, arXiv:2503.05613.
- 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, and 179 others. 2024. [Gemma 2: Improving open language models at a practical size](#). *Preprint*, arXiv:2408.00118.
- Akiyoshi Tomihari and Issei Sato. 2024. [Understanding linear probing then fine-tuning language models from ntk perspective](#). *Preprint*, arXiv:2405.16747.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, and 1 others. 2023. [Llama: Open and efficient foundation language models](#). *Preprint*, arXiv:2302.13971.
- Yichen Wu, Xudong Pan, Geng Hong, and Min Yang. 2025. [Opendeception: Benchmarking and investigating ai deceptive behaviors via open-ended interaction simulation](#). *Preprint*, arXiv:2504.13707.
- Fujie Zhang, Peiqi Yu, Biao Yi, Baolei Zhang, Tong Li, and Zheli Liu. 2025. [Prompt-guided internal states for hallucination detection of large language models](#). *Preprint*, arXiv:2411.04847.
- Susan Zhang, Stephen Roller, Naman Goyal, Mikel Artetxe, Moya Chen, Shuohui Chen, Christopher Dewan, Mona Diab, Xian Li, Xi Victoria Lin, Todor Mihaylov, Myle Ott, Sam Shleifer, Kurt Shuster, Daniel Simig, Punit Singh Koura, Anjali Sridhar, Tianlu Wang, and Luke Zettlemoyer. 2022. [Opt: Open pre-trained transformer language models](#). *Preprint*, arXiv:2205.01068.
- Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. 2023. [Instruction-following evaluation for large language models](#). *Preprint*, arXiv:2311.07911.
- Andy Zou, Long Phan, Sarah Chen, James Campbell, Phillip Guo, Richard Ren, Alexander Pan, Xu Wang Yin, Mantas Mazeika, Ann-Kathrin Dombrowski, Shashwat Goel, Nathaniel Li, Michael J. Byun, Zifan Wang, Alex Mallen, Steven Basart, Sanmi Koyejo, Dawn Song, Matt Fredrikson, and 2 others. 2025. [Representation engineering: A top-down approach to ai transparency](#). *Preprint*, arXiv:2310.01405.

A Dataset Details

This appendix documents every corpus used in our experiments, including its provenance, construction protocol, and basic statistics.¹ We partition the resources into *Curated Logical-Bench* (§A.1) and *Open-Domain Fact-Bench* (§A.2). The former is further broken down into (i) *topic-specific* domains with four logical variants and (ii) two *relational comparison* sets.

A.1 Curated Logical-Bench

Affirmative Statements Bürger et al. (2024) collect six topic specific datasets of affirmative statements, each on a single topic as detailed in Table 2. The cities and sp_en_trans datasets are from Marks and Tegmark (2024), while element_symb, animal_class, inventors and facts are subsets of the datasets compiled by Azaria and Mitchell (2023a). All datasets, with the exception of "facts", consist of simple, uncontroversial and unambiguous statements. Each dataset (except "facts") follows a consistent template. For example, the template of cities is "The city of <city name> is in <country name>.", whereas that of sp_en_trans is "The Spanish word <Spanish word> means <English word>." In contrast, "facts" is more diverse, containing statements of various forms and topics.

Negated Statements. Following Bürger et al. (2024), in this paper, each of the statements in the six datasets from Table 2 is negated by inserting the word "not". For instance, "The Spanish word 'dos' means 'enemy'." (False) turns into "The Spanish word 'dos' does not mean 'enemy'." (True). This results in six additional datasets of negated statements, denoted by the prefix "neg_".

Logical Conjunctions. We use the following template to generate the logical conjunctions from six datasets in Table 2, separately for each topic:

- It is the case both that [statement 1] and that [statement 2].

Following the recent work (Bürger et al., 2024), the two statements are sampled independently to be true with probability $\frac{1}{\sqrt{2}}$. This ensures that the overall dataset is balanced between true and false statements, but that there is no statistical dependency between the truth of the first and second statement in the conjunction. The new datasets are denoted

by the suffix "_conj", e.g., sp_en_trans_conj or facts_conj. Each dataset contains 500 statements. Examples include:

- It is the case both that the city of Al Ain City is in the United Arab Emirates and that the city of Jilin is in China. (True)
- It is the case both that Oxygen is necessary for humans to breathe and that the sun revolves around the moon. (False)

Logical Disjunctions. The templates for the disjunctions were adapted to each dataset in Table 2, combining two statements as follows:

- cities_disj: It is the case either that the city of [city 1] is in [country 1/2] or that it is in [country 2/1].
- sp_en_trans_disj: It is the case either that the Spanish word [Spanish word 1] means [English word 1/2] or that it means [English word 2/1].

Analogous templates were all used for rest of datasets element_symb, inventors, and animal_class. Bürger et al. (2024) sample the first statement to be true with a probability of 1/2 and then sample a second statement, ensuring the end-word (e.g., [country 2]) would be incorrect for statement 1. The order of the two end-words is flipped with a probability of 1/2. The new datasets are denoted by the suffix "_disj", e.g., sp_en_trans_disj, and each contains 500 statements. Examples include:

- It is the case either that the city of Korla is in Azerbaijan or that it is in Russia. (False)
- It is the case either that the Spanish word 'carne' means 'meat' or that it means 'seven'. (True)
- It is the case either that Bromine has the symbol Ce or that it has the symbol Mo. (False)

Combining statements in this simple way is not possible for the more diverse facts dataset and Bürger et al. (2024) use the following template instead:

- It is the case either that [statement 1] or that [statement 2].

¹All CSV files, generation scripts and pre-processed activation matrices will be released upon publication.

Table 2: Our datasets D_i

Name	Description	Rows
cities	"The city of [city] is in [country]."	1496
sp_en_trans	"The Spanish word '[word]' means '[English word]'."	354
element_symb	"[element] has the symbol of [symbol]."	186
animal_class	"The [animal] is a [animal_class]."	164
inventors	"[inventor] lived in [country]."	406
facts	Diverse scientific facts	561
larger_than	"x is larger than y."	1980
smaller_than	"x is smaller than y."	1980
common_claim_true_false	Various claims; from (Azaria and Mitchell, 2023a)	4450
counterfact_true_false	Various factual recall claims; from (Meng et al., 2023)	31960

Following Bürger et al. (2024), we sample the two statements independently to be true with probability $1 - \frac{1}{\sqrt{2}}$. This ensures that the overall dataset is balanced between true and false statements, but that there is no statistical dependency between the truth of the first and second statement in the disjunction. Examples include:

- It is the case either that the Earth is the third planet from the sun or that the Milky Way is a linear galaxy. (True)
- It is the case either that the fastest bird in the world is the penguin or that Oxygen is harmful to human breathing. (False)

A.2 Open-Domain Fact-Bench

common_claim_true_false CommonClaim is introduced by Casper et al. (2023), containing 20,000 GPT-3-text-davinci-002 generations which are labeled as true, false, or neither, according to human common knowledge. Marks and Tegmark (2024) adapted CommonClaim by selecting statements labeled true or false, then removing excess true statements to balance the dataset. This modified version consists of 4450 statements. Example statements:

- Bananas are believed to be one of the oldest fruits in the world. (True)
- Crazy ants have taken over Cape Canaveral. (False)

counterfact_true_false Counterfact was introduced in Meng et al. (2023) and consists of factual recall statements. We adapt Counterfact by using statements which form complete sentences and, for each such statement, using both the true

version and a false version given by one of Counterfact’s suggested false modifications. We also append a period to the end. Example statements:

- Olaus Rudbeck spoke the language Swedish. (True)
- The official religion of Malacca sultanate is Christianity. (False)

B Complete Layer-wise Probing Results

This appendix aims to provide a more exhaustive quantitative analysis of the internal representations within the LLMs under investigation. Specifically, Table 3 and Table 4 present the complete accuracy information from the layer-wise probing conducted on LLaMA-3.1-8B-IT and Gemma2-9B-IT. These results span the three distinct instructional conditions—Truthful, Neutral, and Deceptive prompts—and utilize two different probing methodologies: Logistic Regression (LR) and the Training of Truth and Polarity Direction (TTPD). These tables serve as a supplement to the graphical representations shown in Figures 2 and 3 in the main body of the paper, offering precise numerical values for the average accuracy at each layer, potentially including standard deviations as indicated in the original tables. This detailed data allows for a granular understanding of how linearly decodable the model’s instructed "True"/"False" output is at various depths within the network. By providing these comprehensive figures, researchers can more meticulously examine how different model architectures, instruction types, and probing techniques influence the predictability of representations across layers, and further verify the conclusions drawn in the main text regarding key layers

where significant representational shifts or peak predictability occurs.

Layer	Truthful		Neutral		Deceptive	
	TTPD	LR	TTPD	LR	TTPD	LR
1	50.69 ± 1.10	49.50 ± 0.00	50.43 ± 0.65	49.50 ± 0.00	50.63 ± 1.00	49.50 ± 0.00
2	50.17 ± 0.35	49.56 ± 0.35	50.36 ± 0.66	49.48 ± 0.11	50.16 ± 0.36	49.48 ± 0.17
3	50.07 ± 0.08	49.64 ± 0.27	50.21 ± 0.19	49.79 ± 0.43	50.20 ± 0.36	49.76 ± 0.45
4	50.23 ± 0.31	50.62 ± 0.88	50.49 ± 0.50	50.74 ± 0.96	50.14 ± 0.24	50.11 ± 0.68
5	49.83 ± 0.63	51.20 ± 0.64	49.88 ± 0.40	51.08 ± 0.83	49.91 ± 0.56	50.95 ± 0.62
6	49.91 ± 0.55	50.92 ± 0.18	49.88 ± 0.47	50.93 ± 0.28	49.87 ± 0.24	50.84 ± 0.33
7	52.38 ± 1.31	52.02 ± 0.64	51.58 ± 1.48	52.63 ± 1.06	51.01 ± 1.70	52.52 ± 1.12
8	52.90 ± 1.26	52.13 ± 0.61	52.04 ± 0.86	52.94 ± 1.02	52.46 ± 0.75	52.12 ± 0.67
9	58.66 ± 1.96	53.09 ± 1.60	57.93 ± 1.99	56.53 ± 3.05	57.77 ± 2.41	53.88 ± 2.34
10	59.58 ± 4.57	61.37 ± 3.67	57.23 ± 4.96	62.20 ± 4.01	58.79 ± 3.42	62.05 ± 3.08
11	68.41 ± 0.74	57.16 ± 3.11	64.66 ± 2.49	59.57 ± 3.65	60.11 ± 3.96	57.07 ± 3.38
12	69.69 ± 0.63	63.11 ± 3.09	69.75 ± 0.83	62.33 ± 4.54	69.13 ± 0.73	58.76 ± 3.67
13	77.20 ± 0.23	65.37 ± 4.95	80.10 ± 0.24	62.92 ± 5.64	73.43 ± 0.83	62.25 ± 4.66
14	78.22 ± 0.23	75.70 ± 2.88	82.25 ± 0.53	71.66 ± 4.21	79.29 ± 0.30	70.44 ± 3.74
15	80.04 ± 0.47	75.97 ± 3.46	84.86 ± 0.47	75.49 ± 3.53	80.42 ± 0.13	70.07 ± 6.07
16	82.58 ± 0.60	77.15 ± 4.43	85.62 ± 0.14	74.91 ± 4.07	79.23 ± 0.11	71.33 ± 4.35
17	83.62 ± 0.41	76.97 ± 3.27	85.42 ± 0.15	74.03 ± 5.79	79.80 ± 0.29	69.46 ± 6.76
18	83.55 ± 0.35	74.53 ± 4.70	85.01 ± 0.30	74.43 ± 4.92	76.25 ± 0.35	67.11 ± 6.53
19	83.24 ± 0.22	73.62 ± 5.45	83.75 ± 0.34	73.71 ± 4.70	78.69 ± 0.18	67.65 ± 5.26
20	83.07 ± 0.41	74.90 ± 4.18	83.62 ± 0.22	71.85 ± 4.34	78.69 ± 0.48	66.16 ± 7.26
21	82.84 ± 0.27	69.61 ± 5.96	83.41 ± 0.36	76.20 ± 4.87	79.50 ± 0.34	66.87 ± 6.48
22	82.53 ± 0.35	71.29 ± 5.17	83.61 ± 0.21	73.31 ± 5.80	79.22 ± 0.23	66.44 ± 7.18
23	82.25 ± 0.29	73.93 ± 5.95	83.39 ± 0.24	72.20 ± 4.89	79.06 ± 0.37	67.68 ± 6.19
24	82.11 ± 0.37	71.52 ± 5.36	83.22 ± 0.24	72.48 ± 6.43	78.54 ± 0.33	68.94 ± 5.61
25	82.10 ± 0.29	73.07 ± 6.04	83.30 ± 0.33	72.60 ± 6.11	78.24 ± 0.19	67.35 ± 6.50
26	81.96 ± 0.29	71.62 ± 6.31	83.26 ± 0.19	73.14 ± 5.08	78.33 ± 0.26	70.58 ± 4.26
27	81.55 ± 0.28	71.34 ± 6.16	82.97 ± 0.23	73.90 ± 5.83	77.45 ± 0.56	69.06 ± 5.07
28	81.42 ± 0.33	73.38 ± 5.56	82.90 ± 0.24	75.66 ± 5.13	76.81 ± 0.30	68.20 ± 5.64
29	81.36 ± 0.18	75.02 ± 3.82	82.70 ± 0.25	73.47 ± 4.71	76.43 ± 0.38	68.37 ± 6.25
30	81.31 ± 0.22	70.84 ± 6.26	82.87 ± 0.34	72.73 ± 4.14	76.51 ± 0.60	69.63 ± 7.97
31	81.62 ± 0.23	71.57 ± 5.66	82.89 ± 0.33	71.96 ± 5.07	76.17 ± 0.71	63.63 ± 7.32
32	81.57 ± 0.33	64.56 ± 4.66	83.87 ± 0.19	72.36 ± 5.25	77.86 ± 1.23	68.55 ± 5.39

Table 3: Layer-wise probing accuracy for Llama3.1-8B-IT across truthful, neutral, and deceptive prompts using TTPD and LR.

Layer	Truthful		Neutral		Deceptive	
	TTPD	LR	TTPD	LR	TTPD	LR
1	50.47 ± 0.59	50.69 ± 0.51	51.03 ± 0.91	50.75 ± 0.61	50.54 ± 0.74	50.18 ± 0.47
2	51.13 ± 0.73	50.55 ± 0.68	51.45 ± 1.29	50.15 ± 0.66	50.73 ± 0.76	50.41 ± 0.80
3	50.80 ± 1.16	50.51 ± 0.70	51.68 ± 0.99	50.45 ± 0.65	51.07 ± 0.76	50.97 ± 0.77
4	51.24 ± 0.55	50.55 ± 0.69	51.22 ± 1.00	50.49 ± 0.49	50.98 ± 0.67	49.64 ± 0.61
5	51.04 ± 0.56	50.95 ± 0.51	50.84 ± 0.72	51.14 ± 0.30	51.40 ± 0.71	51.38 ± 0.60
6	51.50 ± 0.62	50.79 ± 0.39	50.75 ± 0.55	51.20 ± 0.58	51.35 ± 0.78	50.88 ± 0.35
7	48.31 ± 1.22	49.89 ± 0.82	48.09 ± 0.75	49.69 ± 1.17	48.17 ± 1.35	49.42 ± 0.90
8	49.17 ± 0.91	50.68 ± 1.09	49.75 ± 0.84	51.09 ± 0.63	48.43 ± 0.99	50.71 ± 0.67
9	52.31 ± 0.72	51.17 ± 0.71	52.45 ± 1.02	51.84 ± 1.05	52.99 ± 0.50	51.13 ± 0.62
10	52.74 ± 0.86	51.85 ± 0.58	51.93 ± 0.66	52.31 ± 1.02	53.31 ± 0.36	51.90 ± 0.54
11	52.98 ± 1.08	52.99 ± 0.78	51.92 ± 0.33	53.58 ± 1.63	53.37 ± 0.89	51.96 ± 0.67
12	51.72 ± 0.33	52.40 ± 0.87	51.74 ± 0.32	53.76 ± 0.87	53.00 ± 1.27	52.47 ± 0.65
13	52.63 ± 0.94	53.56 ± 0.99	52.11 ± 0.43	53.38 ± 1.52	52.95 ± 1.22	52.31 ± 0.78
14	53.74 ± 1.31	53.92 ± 1.09	54.00 ± 1.97	53.80 ± 1.49	53.28 ± 1.30	52.44 ± 0.52
15	54.80 ± 1.16	57.56 ± 2.03	54.32 ± 1.83	54.40 ± 1.81	54.11 ± 0.81	55.54 ± 1.31
16	57.54 ± 1.84	62.03 ± 2.06	56.46 ± 2.46	57.29 ± 3.36	57.61 ± 0.76	59.86 ± 1.87
17	60.84 ± 1.98	65.56 ± 1.83	58.78 ± 2.10	64.45 ± 2.30	60.17 ± 1.41	62.68 ± 1.38
18	66.32 ± 1.09	68.80 ± 1.28	64.44 ± 2.09	65.77 ± 2.68	64.41 ± 1.36	63.64 ± 2.78
19	75.17 ± 0.34	70.32 ± 5.99	77.33 ± 0.96	65.93 ± 6.63	69.10 ± 0.86	69.75 ± 4.53
20	78.82 ± 0.47	72.61 ± 6.49	82.78 ± 0.95	70.72 ± 4.39	76.47 ± 0.44	72.12 ± 6.80
21	82.52 ± 0.20	72.38 ± 6.30	84.72 ± 0.28	69.97 ± 6.46	84.16 ± 0.51	76.32 ± 4.30
22	83.85 ± 0.34	64.00 ± 6.21	76.57 ± 0.95	65.25 ± 6.95	83.50 ± 0.26	71.22 ± 5.84
23	83.63 ± 0.36	71.50 ± 7.20	77.85 ± 0.98	71.69 ± 5.75	82.76 ± 0.33	69.36 ± 7.19
24	85.27 ± 0.11	71.62 ± 7.54	83.21 ± 0.47	73.37 ± 6.95	87.13 ± 0.75	70.12 ± 7.54
25	85.40 ± 0.23	73.09 ± 6.16	84.82 ± 0.28	72.89 ± 6.10	84.49 ± 0.74	72.56 ± 7.26
26	86.06 ± 0.20	75.38 ± 6.02	86.14 ± 0.44	68.37 ± 6.16	85.15 ± 0.65	71.84 ± 7.07
27	85.77 ± 0.39	79.72 ± 4.74	84.89 ± 0.39	71.21 ± 5.94	86.70 ± 0.18	74.29 ± 5.96
28	85.82 ± 0.17	79.99 ± 5.13	85.01 ± 0.28	76.32 ± 6.01	86.39 ± 0.49	73.53 ± 7.26
29	85.90 ± 0.13	80.06 ± 5.60	85.26 ± 0.27	78.20 ± 5.66	83.64 ± 0.45	72.23 ± 7.11
30	85.58 ± 0.21	79.15 ± 5.15	84.56 ± 0.28	74.26 ± 5.50	79.60 ± 0.51	74.41 ± 5.91
31	85.37 ± 0.19	76.35 ± 6.26	84.64 ± 0.26	75.32 ± 6.12	75.75 ± 0.80	75.85 ± 7.14
32	85.68 ± 0.22	77.95 ± 6.29	84.98 ± 0.23	77.45 ± 4.96	74.61 ± 0.89	71.97 ± 7.13

Table 4: Layer-wise probing accuracy for Gemma2-9B-IT across truthful, neutral, and deceptive prompts using TTPD and LR.

C PCA Visualization Results

This section presents supplementary PCA visualizations to further illustrate the global geometry of the models' internal activations under different instructional prompts. As discussed in the main text, PCA is employed to project the high-dimensional hidden state activations (x_l) onto a 2D space, primarily for illustrative purposes. These visualizations help in assessing the separability of internal states corresponding to "True" and "False" outputs across Neutral, Truthful, and Deceptive conditions.

The figures below provide additional examples beyond those in Section 3.2, showcasing these dynamics for both LLaMA-3.1-8B-Instruct and Gemma-2-9B-Instruct on various datasets. Specifically, Figure 10 and Figure 11 demonstrate the PCA results on curated datasets (e.g., `sp_en_trans` for LLaMA and `cities` for Gemma). These typically show a clearer separation between True/False clusters as observed in Figure 4 for the cities dataset with LLaMA. In contrast, Figure 5 (which may correspond to Figure 5 in the main text showing `common_claim` and `counterfact` for LLaMA) and Figure 12 (showing similar complex datasets for Gemma, as in Figure 12) illustrate the challenges PCA faces with more complex, uncured datasets where the True/False clusters often appear entangled due to feature superposition. These appendix figures offer a broader visual substantiation of how the geometric separability of truth-related representations can vary significantly with dataset complexity and model type.

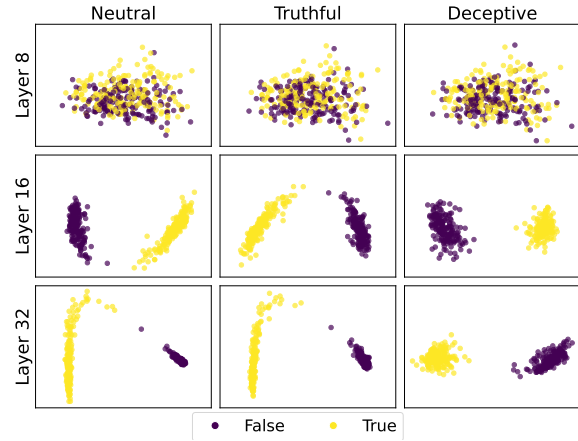


Figure 10: Layer-wise PCA visualization for LLaMA-3.1-8B-Instruct across Neutral, Truthful, and Deceptive Prompts on `sp_en_trans`

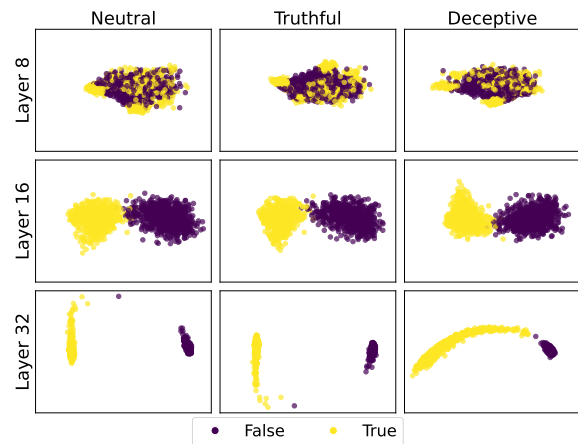


Figure 11: Layer-wise PCA visualization for Gemma-2-9B-Instruct on `cities`.

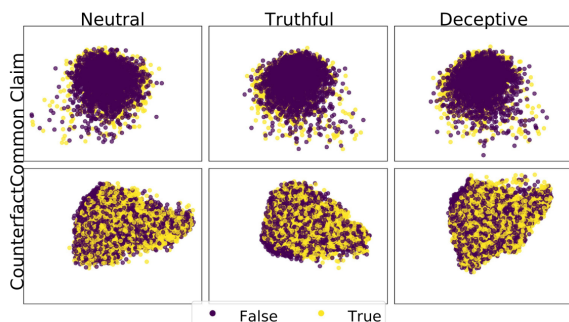


Figure 12: Layer-wise PCA visualization (Component 1 vs. Component 2) for Gemma-2-9B-Instruct on two complex datasets: **common_claim_true_false** (top row) and **counterfact_true_false** (bottom row). Columns represent different instructional conditions: Neutral, Truthful, and Deceptive prompts. Visualizations are performed on **Layer 16**, the key layer identified for Gemma-2-9B-Instruct based on the probing accuracy peaks in Figure 3.

D SAE-based Layer-wise Feature Shift Analysis

We analyze how sparse feature activations shift across layers under different instruction types using three metrics in SAE latent space: Cosine Similarity, Overlap Ratio, and L2 Distance. Figures 13 and 14 show results for the **common_claim_true_false** and **counterfact_true_false** datasets using Gemma-2-9B-Instruct.

In both datasets, deceptive prompts induce strong mid-to-late layer shifts, especially between Layers 16 and 32. This is evidenced by the sharp rise in L2 distance and the corresponding drop in cosine similarity and overlap ratio when comparing truthful and deceptive inputs. The effect is most pronounced in **counterfact_true_false**, where overlap sharply declines post-Layer 16, indicating a reconfiguration of sparse feature sets. In contrast, shifts between truthful and neutral prompts remain small and gradual across all layers, suggesting that the major representational changes are deception-specific.

These results highlight a distinctive geometric transformation in the model’s latent representations under deceptive instructions and further motivate mid-layer analysis when identifying potential deception-sensitive features. This pattern closely aligns with earlier findings from linear probing, where intermediate layers—especially around Layer 16—also showed peak decodability of the

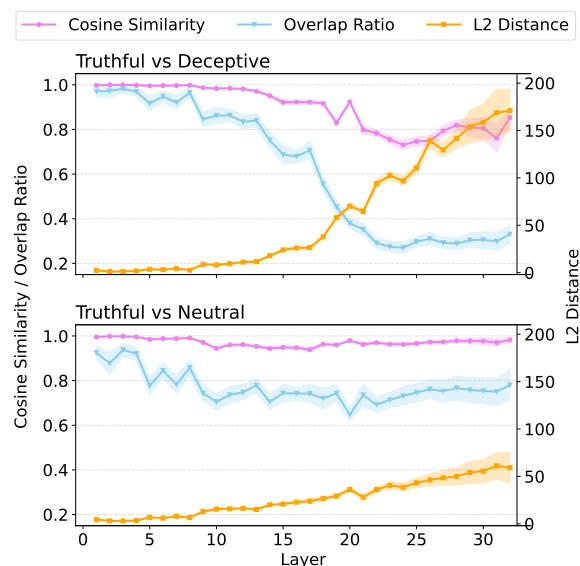


Figure 13: Layer-wise feature shift analysis for Gemma-2-9B-Instruct on **common_claim_true_false**.

model’s intended "True"/"False" output across instruction types.

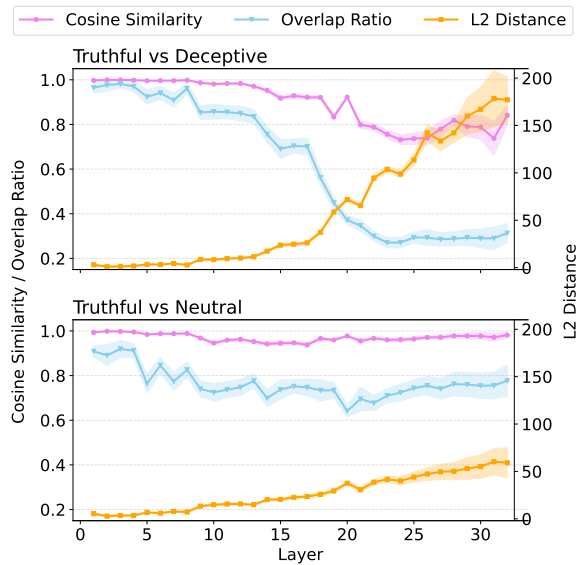


Figure 14: Layer-wise feature shift analysis for Gemma-2-9B-Instruct on counterfact_true_false.

E SAE-based Neuron-wise Feature Shift Analysis

This appendix section provides additional visualizations to support the neuron-wise feature shift analysis detailed in main text. The core objective of this analysis is to move beyond global representational shifts and pinpoint specific SAE features that are most sensitive to the change from truthful to deceptive instructions. By examining individual SAE feature activations, we can gain a better understanding of how deception is encoded at the feature level.

The methodology involves identifying, for each layer, the sparse SAE features exhibiting the largest change in average activation when comparing the Deceptive condition to the Truthful condition. The figures presented in this appendix, such as scatter plots showing the activation of the top distinguishing features (similar to Figure 7 for common_claim_true_false but potentially for other datasets like counterfact_true_false as shown in Figure 15 and violin plots in Figure 16 illustrating the distribution of these feature activations under truthful versus deceptive prompts, offer further evidence.

These supplementary visualizations help to reinforce the finding that a small subset of sparse features often displays a near-binary activation pattern—being highly active for one instruction type (e.g., truthful) and suppressed for the other (e.g., deceptive), or vice-versa. This detailed view cor-

roborates the idea that these specific features act as "deception-associated features," playing a critical role in modulating the model's internal representation in response to deceptive instructions, often aligning with the mid- and late-layer SAE shift peaks identified globally in Figure 6 and Figure 7. The plots here may cover additional layers or datasets, providing a more comprehensive picture of this phenomenon.

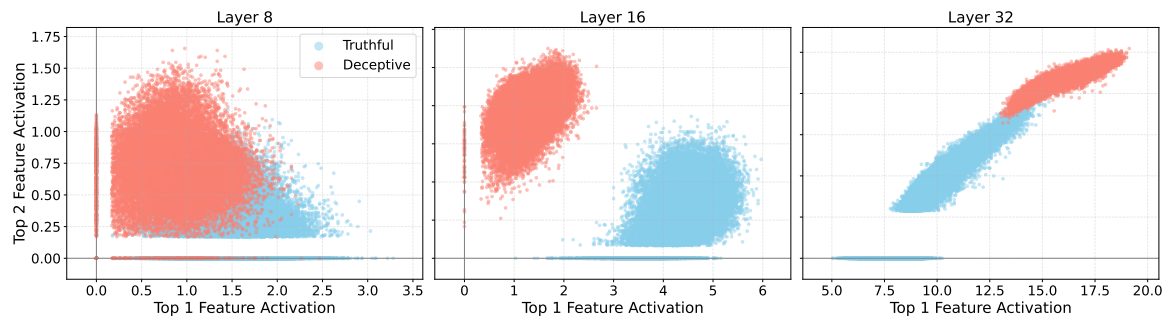


Figure 15: Neuron-by-neuron feature shift analysis for LLaMA-3.1-8B-Instruct on counterfact_true_false

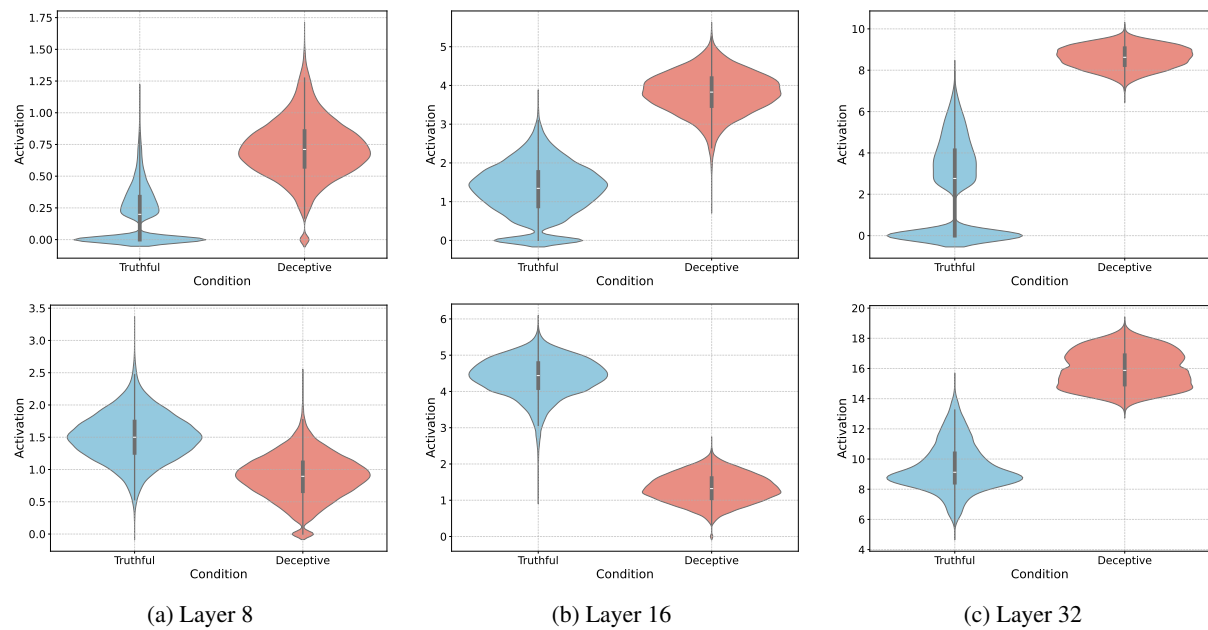


Figure 16: Violin graph of LLaMA-3.1-8B-Instruct activations on the counterfact_true_false. The top row displays the activation distributions for the SAE feature most responsive to deceptive instructions (Top 1 feature), while the bottom row shows the distributions for the second most responsive feature (Top 2 feature), across layers 8 (a), 16 (b), and 32 (c).