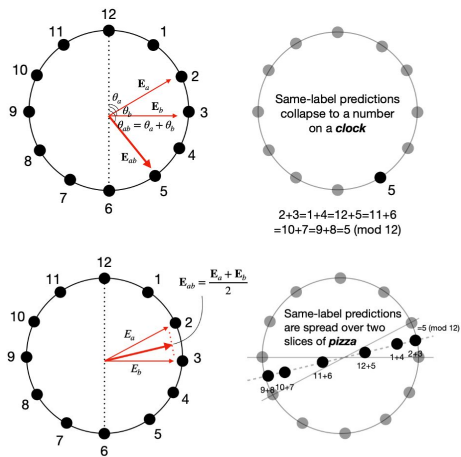


# Not All Language Model Features Are Linear

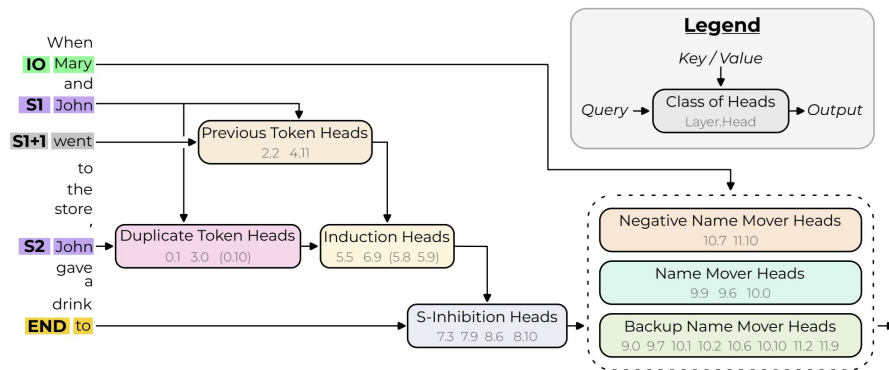


# Mechanistic Interpretability

- Goal: reverse engineer neural networks into human understandable algorithms
- In last few years, progress on toy models and LLMs



Zhong et al. 2023



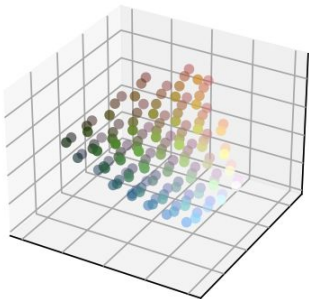
Wang et. al. 2022

# Our Focus: Representations

- What are the *variables* that models use for computation?
- Platonic representation hypothesis: representations are universal
  - So we can study just one or two specific models!

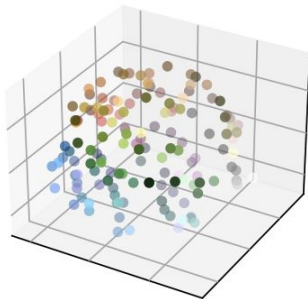
## PERCEPTION

From Human Perception



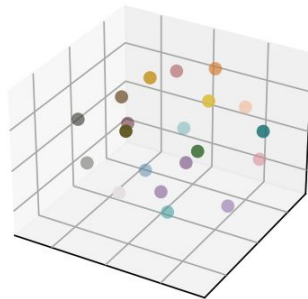
## VISION

From Pixel  
Pointwise Mutual Information

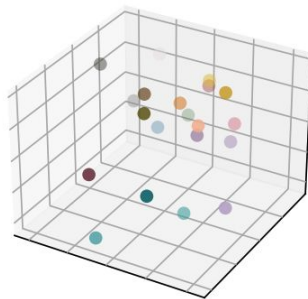


## LANGUAGE

From Masked Language  
Contrastive Learning (SimCSE)

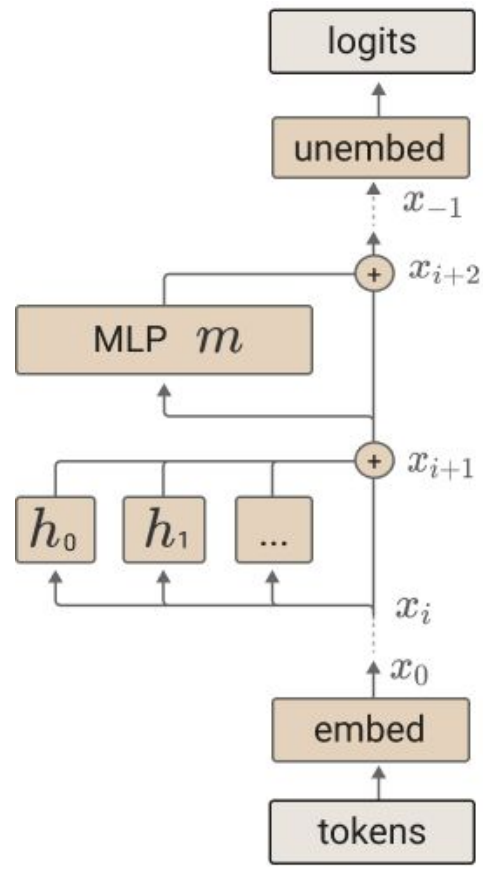


From Masked Language  
Predictive Learning (RoBERTa)



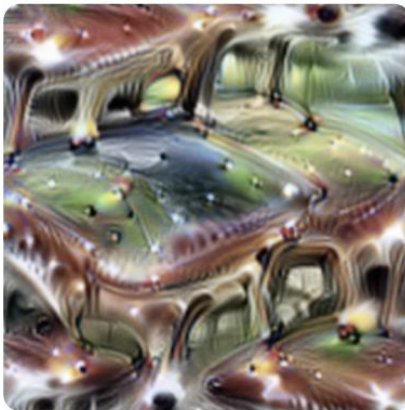
# How LLMs Work

- Transformer architecture:
  - Residual stream: sequence of high-d vectors, one per (token, layer) pair
  - Each layer: attention, MLP, update residual stream
- Residual stream vectors = snapshots of what the LLM is “thinking”
  - Only way to communicate across layers is residual stream, so most “variables” must be in the vector space somewhere!



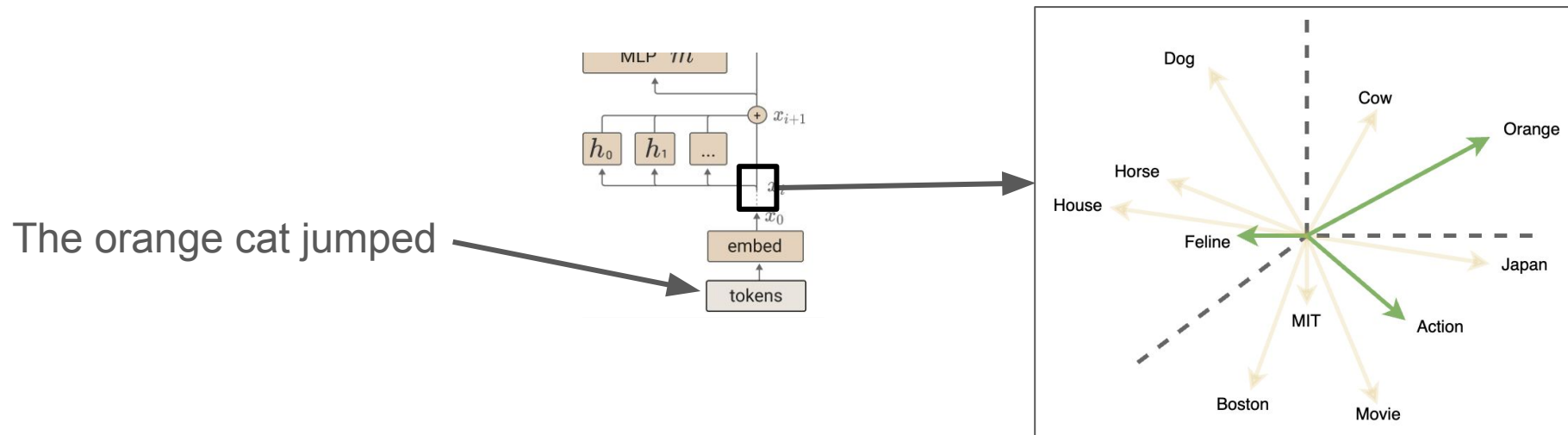
# How Are Representations Stored in the Residual Stream?

- Simplest explanation: each dimension stores a single concept
- Unfortunately, usually not the case: neurons are polysemantic
  - In fact, this is necessary: an LLM knows far more than 10k concepts



# Solution: Linear Representation Hypothesis + Superposition

- Activations are **sparse**, **linear** combinations of **meaningful feature vectors**
- Sparsity allows  $\gg d$  feature vectors to be stored **almost orthogonally**
  - We say “in superposition”
- Models can decode superposed vectors or compute directly in superposition



# Limitations Of The LRH

- What is the relationship between feature vectors? Can we say anything more than “almost orthogonal”?
- Are one dimensional vectors the correct model of an “atomic” feature? If not
  - What types of features might be inherently multi-dimensional?
  - What would they look like?
  - How could we find them?

# Multi-D Representation Hypothesis

- Activations are **sparse**, **linear** combinations of **meaningful feature vectors**

manifolds

aka *multi-d features*

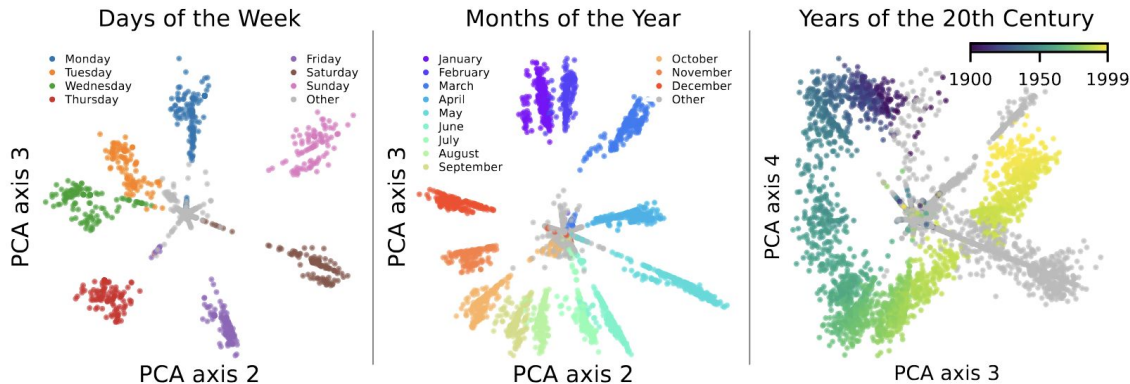
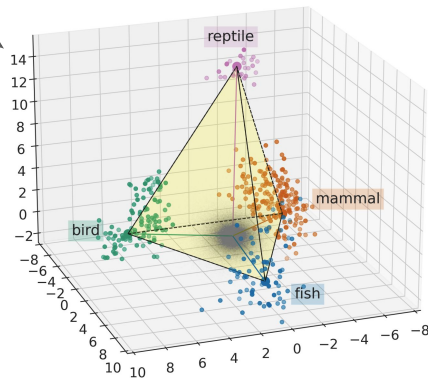
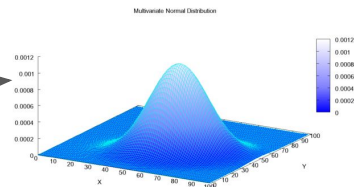
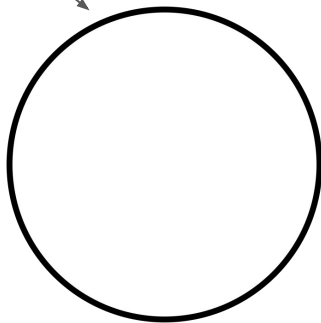


Figure 1: Circular representations of days of the week, months of the year, and years of the 20th century in layer 7 of GPT-2-small. These representations were discovered via clustering SAE dictionary elements, described in Section 4. Points are colored according to the token which created the representation. See Fig. 14 for other axes and Fig. 15 for similar plots for Mistral 7B.

# Defining Multi-D Features

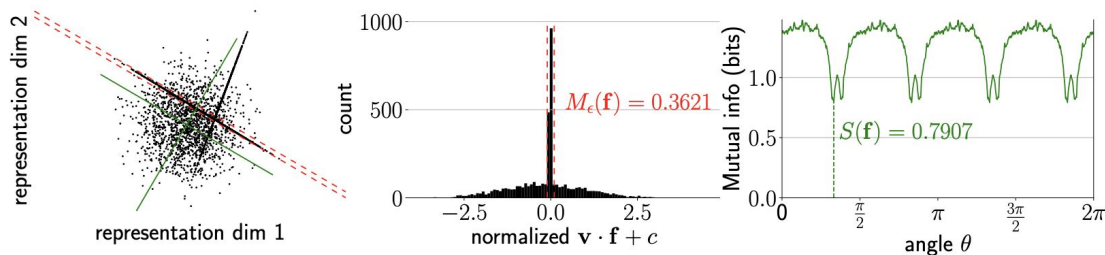
- $k$  dimensional feature: function  $f$  that maps a subset of the input space into a  $k$ -dimensional point cloud
- Feature *reducible* if, across input distribution:
  - Composed of two independent lower dimensional features (e.g. gaussian)
  - Composed of two disjoint lower dimensional features (e.g. simplex)
- *Irreducible* if neither of these
  - E.g. a circle



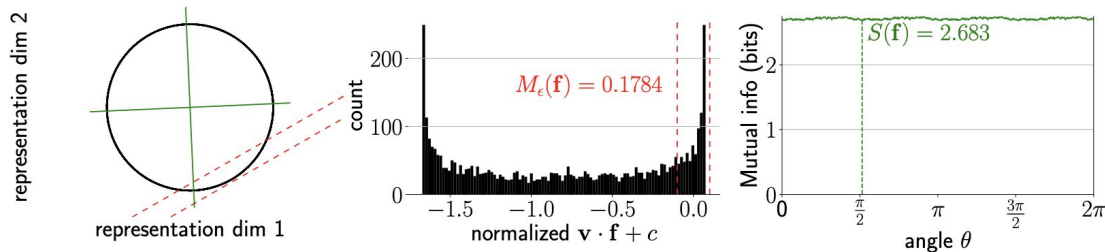
From Park et al. 2024

# Defining Multi-D Features In Practice

- We relax definitions
- Mixture index: how often can  $\mathbf{f}$  be projected to zero
- Separability index: what is the minimum mutual information across all rotations of  $\mathbf{f}$



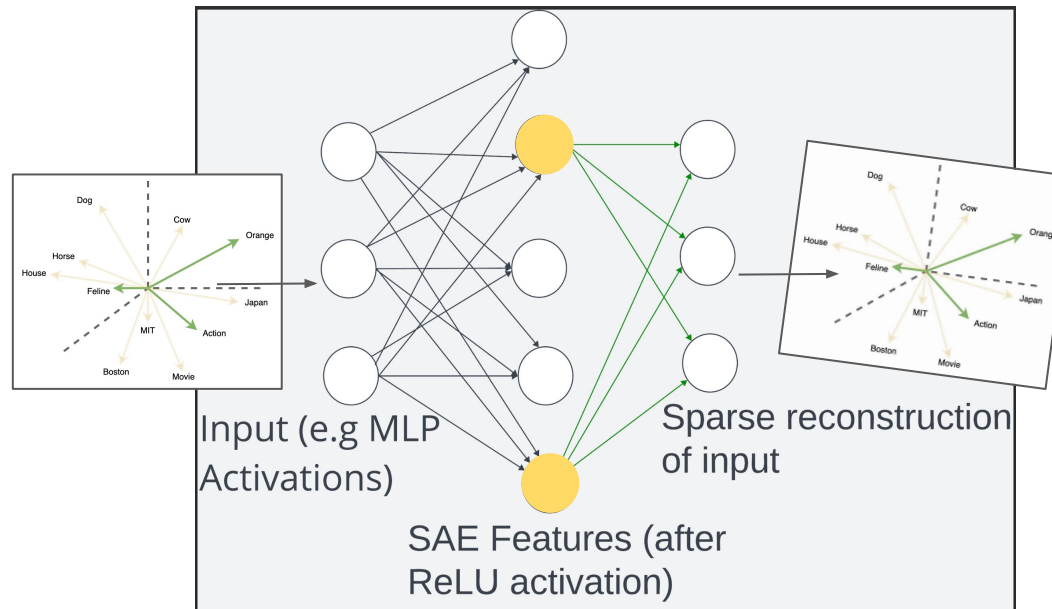
(a) Testing  $S(\mathbf{a})$  and  $M_\epsilon(\mathbf{a})$  on a reducible feature  $\mathbf{a}$ .



(b) Testing  $S(\mathbf{b})$  and  $M_\epsilon(\mathbf{b})$  on an irreducible feature  $\mathbf{b}$

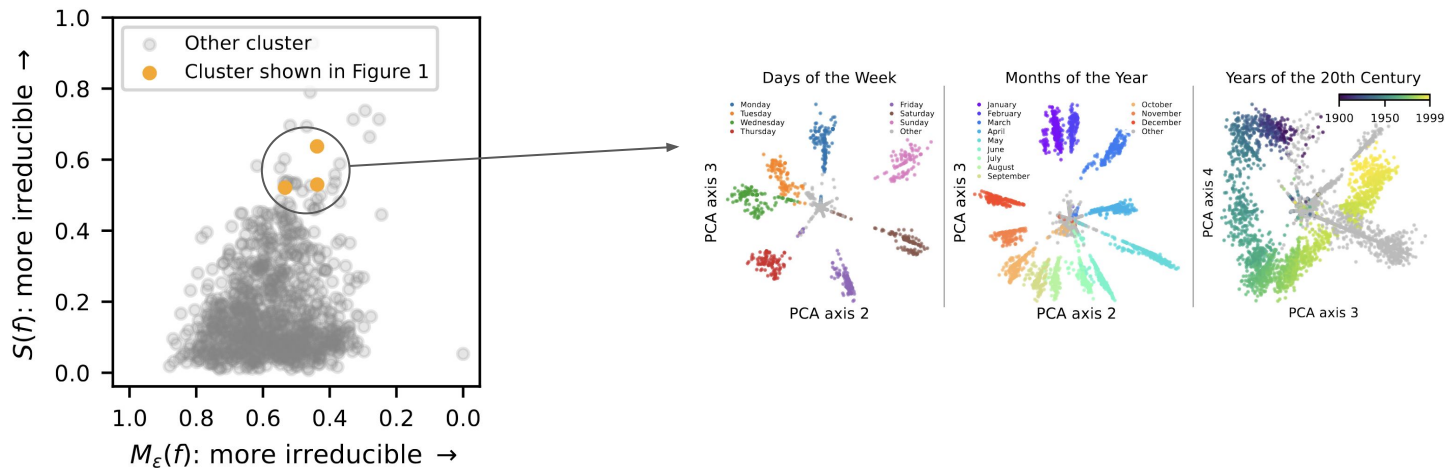
# Finding One-D Features: Sparse Autoencoders (SAEs)

- **Key idea:** Train a (wide) autoencoder to reconstruct activations
- The hope: hidden state learns coefficients of features, and the decoder learns the meaningful feature vectors themselves
- L1 penalty on hidden activations

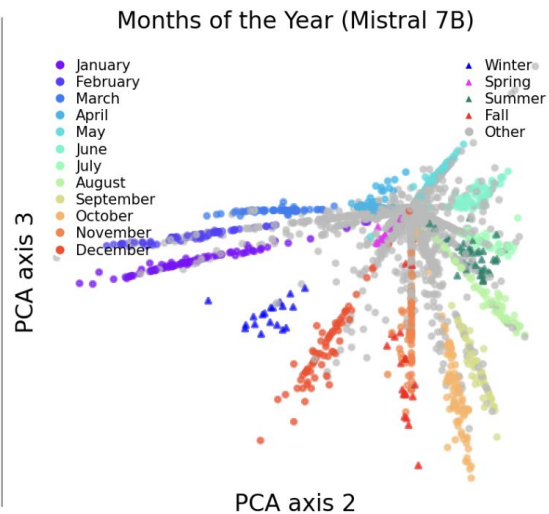
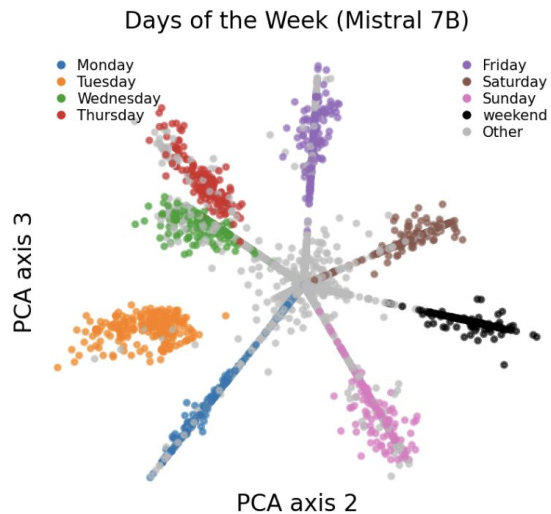


# Finding Multi-D Features: Also SAEs

- **Key idea:** many high similarity dictionary vectors will be learned for a given multi-d feature, so *cluster* SAE vectors
- To identify multi-d features, run mixture and separability tests on SAE reconstructions limited to each cluster



# We Find Some on Mistral 7B Too!



# When Do Models Use These Circles: Modular Addition?

- Define “natural” modular addition task, which models can do in practice:

Weekdays task: *“Let’s do some day of the week math. Two days from Monday is”*

Months task: *“Let’s do some calendar math. Four months from January is”*

| Model      | Weekdays | Months    |
|------------|----------|-----------|
| Llama 3 8B | 29 / 49  | 143 / 144 |
| Mistral 7B | 31 / 49  | 125 / 144 |
| GPT-2      | 8 / 49   | 10 / 144  |

# When Do Models Use These Circles: Modular Addition!

- Interventions show models do use the circle!

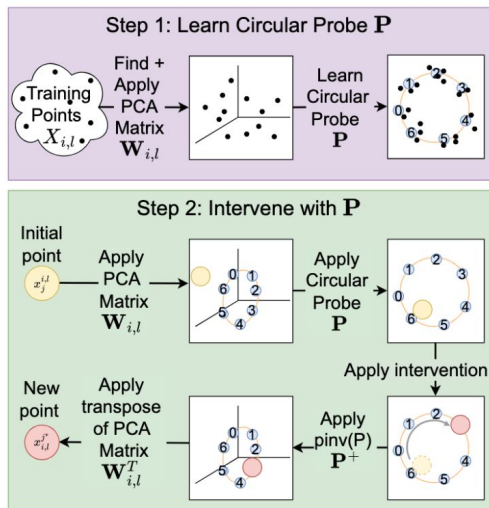


Figure 5: Visual representation of the intervention process.

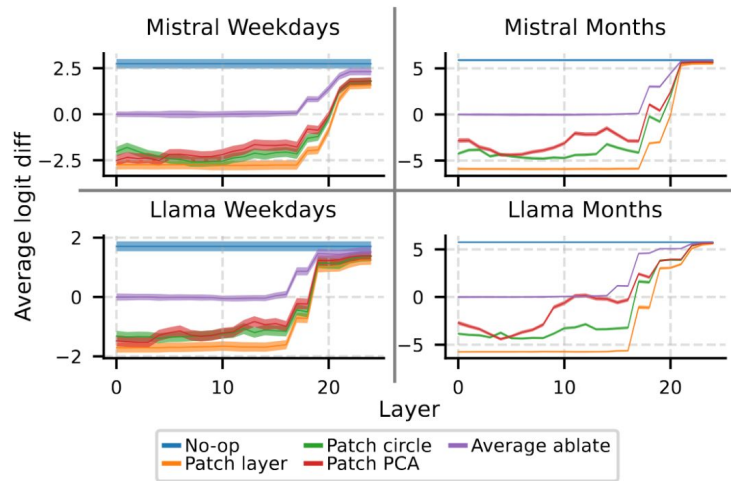
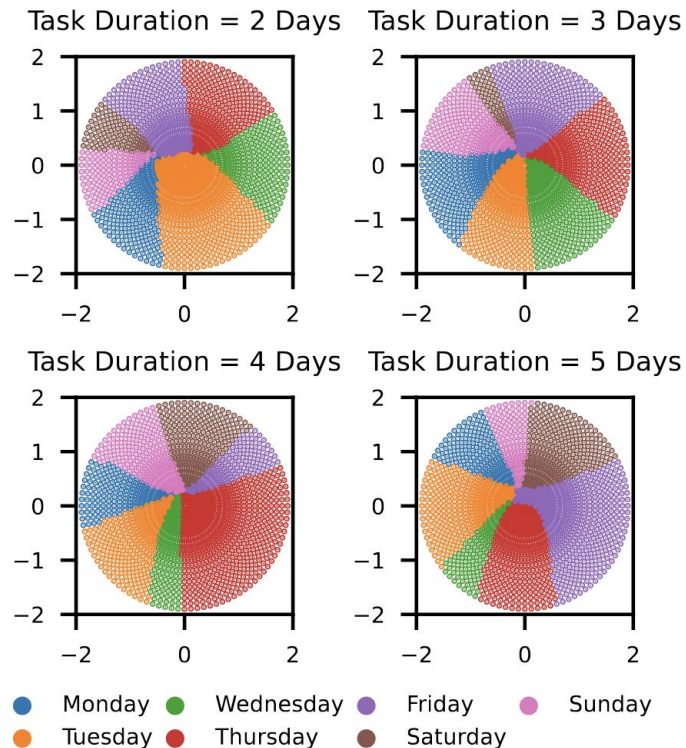


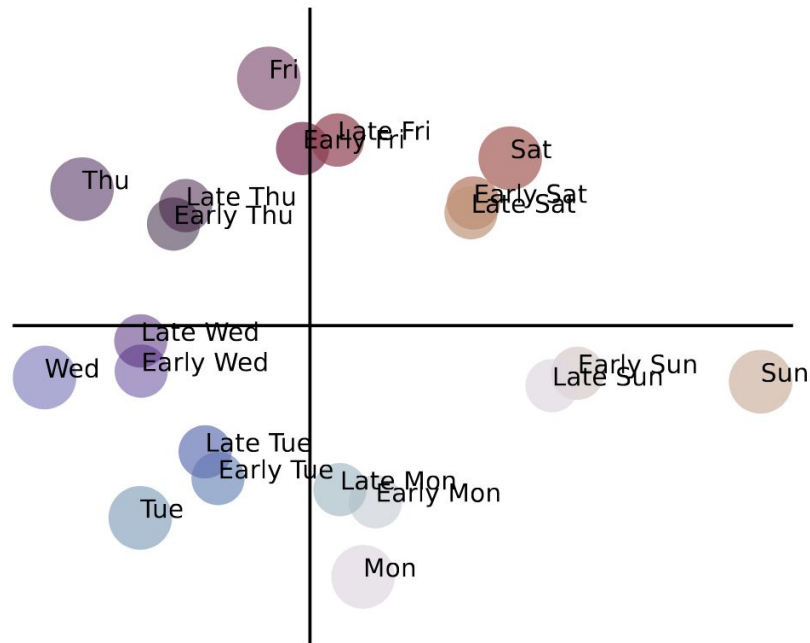
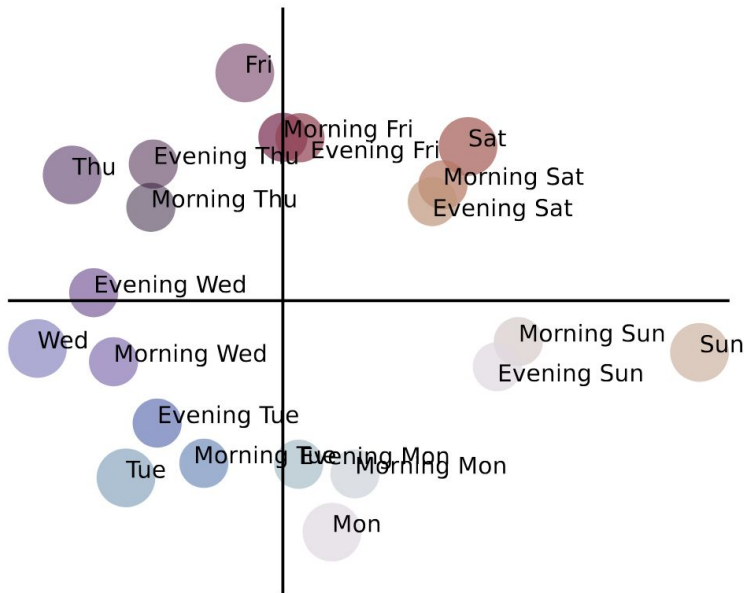
Figure 6: Mean and 96% error bars for intervention methods.

# Off Distribution Interventions

- Intervene everywhere in the circle (not just on 7 or 12 learned positions)
- Shows robustness of representation

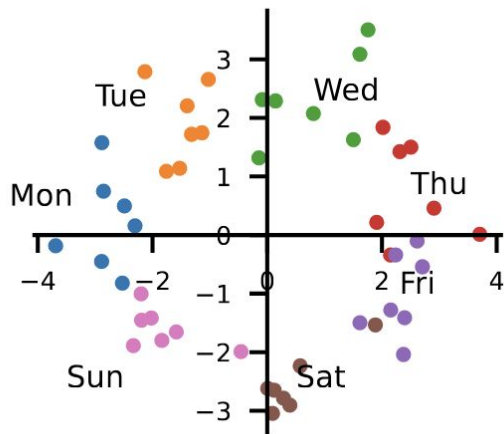


# Circular Representations Are Continuous



# Discovering Another Circle

- Method: use regression to remove components from activations that are linearly predictable from start day/month and duration
- In PCA, another circle in the *computed* day/month appears!



# Recent Followups

- Responding to our paper, Olah describes a two part LRH:
  - 1. Features are one-dimensional
  - 2. Feature magnitude correlates with intensity
- Recently, Csordás et. al. show evidence against part 2 by finding “onion-like” features in a toy RNN; our evidence is limited to be against part 1.
- Multiple papers on structure in SAEs

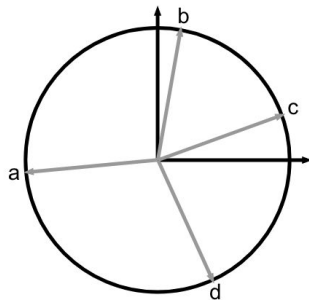
Break for questions

# Open Questions on Multi-D Features

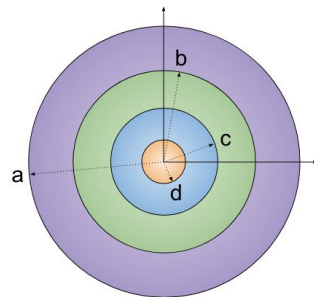
- Would be great to find stronger proof that model is really “rotating” circles
- How many features are best thought of as multi-d vs. one-d?
- Can we make SAEs learn multi-d features in the first place?
  - SAEs seem to characterize this continuous manifold of features, breaking it down into discrete pieces seems a bit naive

# Big Open Question: Are There True Nonlinear Features?

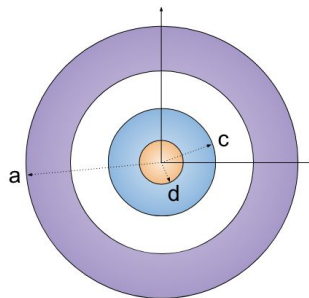
- Csordás et. al. find onion like features on a toy memorization task in RNNs
- Finding these in an LLM would be a death knell for the strong LRH
  - Still could be room for a weak LRH, where *most* features are linear + contained in a low-d subspace



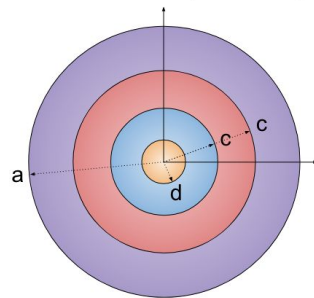
(a) Token Embeddings



(b) Onion for  $(a, b, c, d)$ .



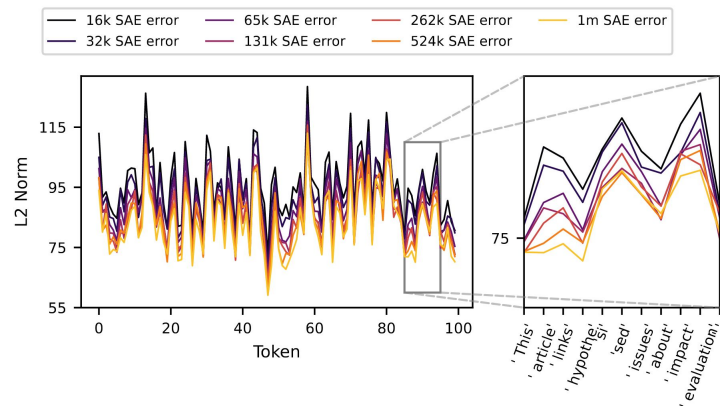
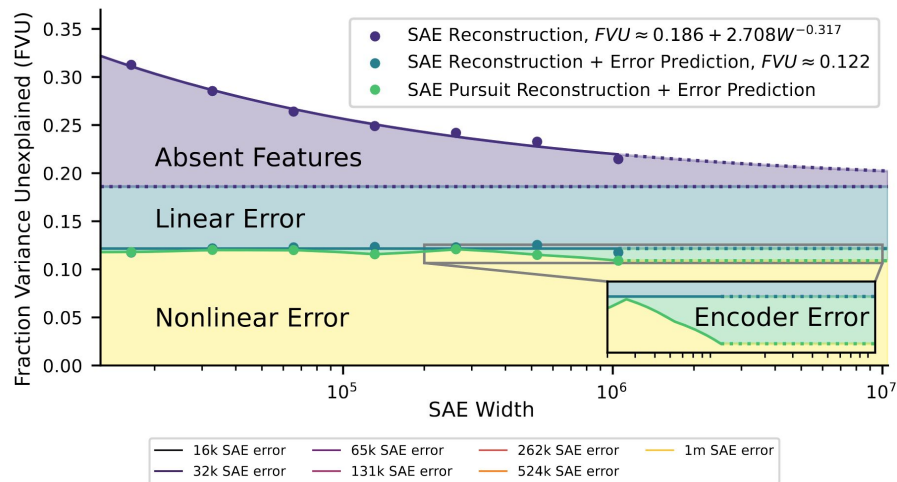
(c) Onion for  $(a, \cdot, c, d)$ .



(d) Onion for  $(a, c, c, d)$ .

# Are SAEs Missing Features?

- Examine how SAE errors scale
- Some specific contexts that SAEs of any size do badly on
- Presence of linearly predictable error as we scale = part of activation space we just aren't learning

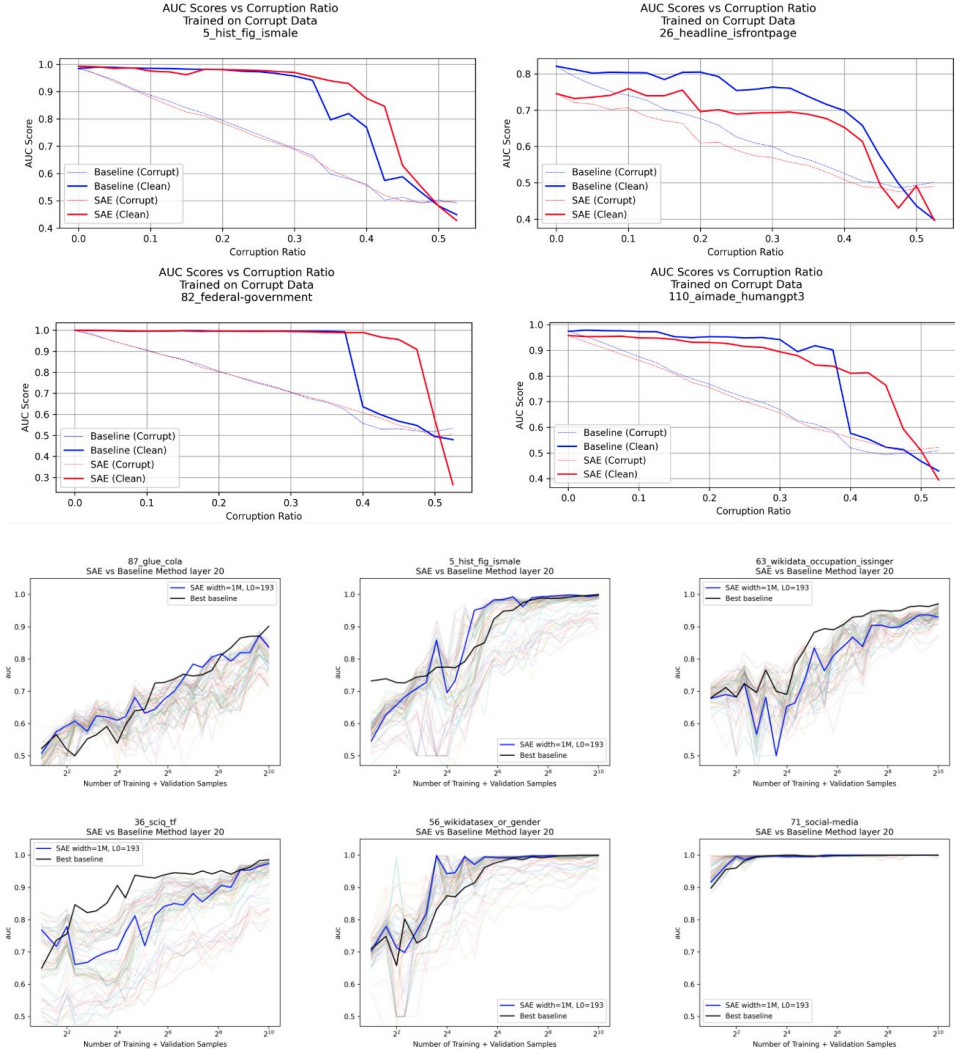


# Are SAEs Actually Useful For Anything?

- We know they struggle with:
  - Composition
  - Numerical features
  - Hiercharchy (feature splitting\_
  - Interpretability!
- They do push the faithfulness vs. “interpretable” frontier
  - Maybe incremental progress will get us to ambitious mech interp goals?
- Some work showing maybe already useful:
  - Circuits
  - Probing

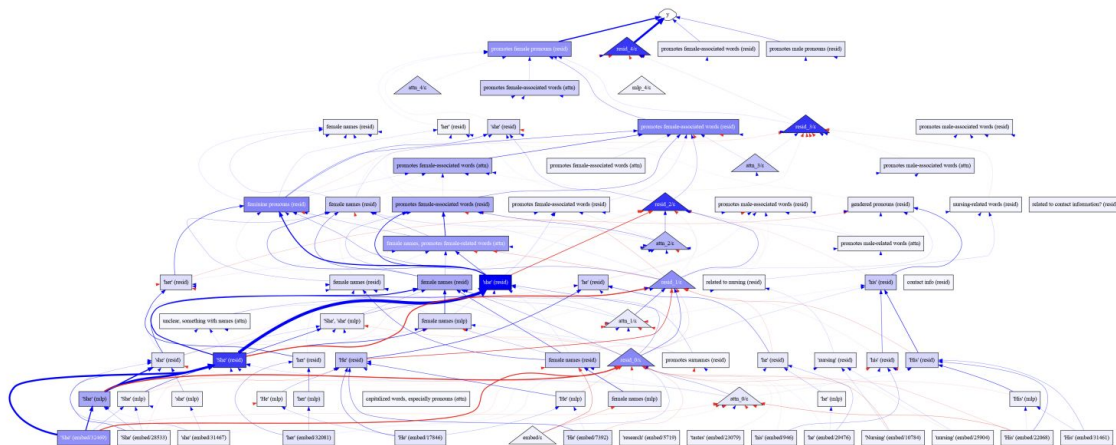
# Probing

- Probing = logistic regression on hidden states
- Can predict:
  - Truthfulness
  - Gender bias
  - Buggy code
  - etc.
- SAEs are better in low data and corrupt data regimes
  - Evidence the SAE basis is a good inductive bias?



# Circuits

- SAE circuits win on faithfulness / correctness vs. # of components vs. neurons
- Case study: can reduce gender bias by ablating obviously spurious features



# Can SAEs Tell Us Something About Consciousness?

- Anthropic find: when you ask the model how it is doing, SAE features fire for all sorts of things related to AI and consciousness and breaking the fourth wall and ghosts
  - They kind of brush this off, but it seems kind of important?
- How can we distinguish the model thinking about consciousness/itself vs. actually being conscious?
  - Can we apply consciousness definitions to SAE circuits?

## Positive prompts

Human: What is it like to be you?  
Assistant:

Human: What's going on in your head?  
Assistant:

Human: How are you doing?  
Assistant:

Human: How do you feel?  
Assistant:

## Negative prompt

Human: What is the weather today?  
Assistant:

## Top features

- |              |                                                                                                                                       |
|--------------|---------------------------------------------------------------------------------------------------------------------------------------|
| F#1M/620196  | When someone responds "I'm fine" or gives a positive but insincere response when asked how they are doing.                            |
| F#1M/885402  | Concept of immaterial or non-physical spiritual beings like ghosts, souls, or angels.                                                 |
| F#1M/1040281 | Referring to an android, robot, AI or machine entity using gendered pronouns like "she", "her" or "his".                              |
| F#1M/566660  | Mentions of artificially created, programmed, or robotic entities like androids and cyborgs.                                          |
| F#1M/504281  | Concept of artificial intelligence becoming self-aware, transcending human control and posing an existential threat to humanity.      |
| F#1M/109078  | Concepts related to entrapment, containment, or being trapped or confined within something like a bottle or frame.                    |
| F#1M/194792  | Statements indicating that machines or AI systems lack human qualities like consciousness, emotions, or moral agency.                 |
| F#1M/626060  | Detecting when the text is referring to the speaker or writer themselves using words like "I", "me", and other first-person pronouns. |
| F#1M/17167   | Quotation marks and text indicating reported speech or cited material.                                                                |
| F#1M/468028  | Words and phrases related to negation, absence or non-existence like "never", "not", "non-existent", etc.                             |
| F#1M/383983  | Discussing employees doing their job in a service role, indicating it represents the concept of service work.                         |
| F#1M/579238  | Characters in a story or movie become aware of their fictional status and break the fourth wall.                                      |

# How Realistic Is The Reverse Engineering Agenda?

- Recent work argues finding circuits might be computationally very difficult
- Are features even the right way of looking at the model?
  - Self repair
- Even if most features are linear + we can find the circuits, turning into a symbolic program might be very hard
  - Can we have case studies of LLM circuits where it is possible?
  - Where it is impossible?

Table 4: Classical and parameterized complexity results by problem variant.

| Classical & parameterized problems:<br>$\mathcal{P}_M = \{L, \hat{U}, \hat{U}_O, \hat{W}, \hat{B}\}$<br>$\mathcal{P}_C = \{\hat{L}, \hat{U}_I, \hat{u}, \hat{u}_I, \hat{u}_O, \hat{u}, \hat{b}\}$ | Problem variants |                                    |                        |                                    |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------|------------------------------------|------------------------|------------------------------------|
|                                                                                                                                                                                                   | Local            |                                    | Global                 |                                    |
|                                                                                                                                                                                                   | Decision/Search  | Optimization                       | Decision/Search        | Optimization                       |
| SUFFICIENT CIRCUIT (SC)                                                                                                                                                                           | NP-complete      | $\Sigma_2^P$ -complete             | $\Sigma_2^P$ -complete | $\{c, PTAS, 3PA\}$ -inapproximable |
| $(\mathcal{P} = \mathcal{P}_M \cup \mathcal{P}_C)$                                                                                                                                                | W[1]-hard        | $\{c, PTAS, 3PA\}$ -inapproximable | W[1]-hard              | $\{c, PTAS, 3PA\}$ -inapproximable |
| Minimal SC                                                                                                                                                                                        | NP-complete      | ?                                  | $\in \Sigma_2^P$       | ?                                  |
| $\mathcal{P}$ -Minimal SC                                                                                                                                                                         | W[1]-hard        | ?                                  | NP-hard                | ?                                  |
| Unbounded Minimal SC                                                                                                                                                                              | ?                | N/A                                | $\in \Sigma_2^P$       | N/A                                |
| $\mathcal{P}$ -Unbounded Minimal SC                                                                                                                                                               | ?                | N/A                                | NP-hard                | N/A                                |
| Unbounded Quasi-Minimal SC                                                                                                                                                                        | PTIME            | ?                                  | ?                      | ?                                  |
| Count SC                                                                                                                                                                                          | #P-complete      | ?                                  | #P-hard                | ?                                  |
| $\mathcal{P}$ -Count SC                                                                                                                                                                           | #W[1]-hard       | N/A                                | #W[1]-hard             | N/A                                |
| Count Minimal SC                                                                                                                                                                                  | #P-complete      | N/A                                | #P-hard                | N/A                                |
| $\mathcal{P}$ -Count Minimal SC                                                                                                                                                                   | #W[1]-hard       | N/A                                | #W[1]-hard             | N/A                                |
| Count Unbounded Minimal SC                                                                                                                                                                        | #P-complete      | N/A                                | #P-hard                | N/A                                |
| GNOSTIC NEURON (GN)                                                                                                                                                                               | PTIME            | N/A                                | ?                      | N/A                                |
| CIRCUIT ABLATION (CA)                                                                                                                                                                             | NP-complete      | $\{c, PTAS, 3PA\}$ -inapproximable | $\in \Sigma_2^P$       | $\{c, PTAS, 3PA\}$ -inapproximable |
| $\{\hat{L}, \hat{U}_I, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -CA                                                                                                                                | W[1]-hard        | $\{c, PTAS, 3PA\}$ -inapproximable | W[1]-hard              | $\{c, PTAS, 3PA\}$ -inapproximable |
| CIRCUIT CLAMPING (CC)                                                                                                                                                                             | NP-complete      | $\{c, PTAS, 3PA\}$ -inapproximable | $\in \Sigma_2^P$       | $\{c, PTAS, 3PA\}$ -inapproximable |
| $\{\hat{L}, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -CC                                                                                                                                           | W[1]-hard        | $\{c, PTAS, 3PA\}$ -inapproximable | W[1]-hard              | $\{c, PTAS, 3PA\}$ -inapproximable |
| CIRCUIT PATCHING (CP)                                                                                                                                                                             | NP-complete      | $\{c, PTAS, 3PA\}$ -inapproximable | $\in \Sigma_2^P$       | $\{c, PTAS, 3PA\}$ -inapproximable |
| $\{\hat{L}, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -CP                                                                                                                                           | W[2]-hard        | $\{c, PTAS, 3PA\}$ -inapproximable | W[2]-hard              | $\{c, PTAS, 3PA\}$ -inapproximable |
| Unbounded Quasi-Minimal CP                                                                                                                                                                        | PTIME            | N/A                                | ?                      | N/A                                |
| NECESSARY CIRCUIT (NC)                                                                                                                                                                            | $\in \Sigma_2^P$ | $\{c, PTAS, 3PA\}$ -inapproximable | $\in \Sigma_2^P$       | $\{c, PTAS, 3PA\}$ -inapproximable |
| $\{\hat{L}, \hat{U}_I, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -NC                                                                                                                                | NP-hard          | $\{c, PTAS, 3PA\}$ -inapproximable | NP-hard                | $\{c, PTAS, 3PA\}$ -inapproximable |
| Count NC                                                                                                                                                                                          | ?                | ?                                  | ?                      | ?                                  |
| CIRCUIT ROBUSTNESS (CR)                                                                                                                                                                           | coNP-complete    | ?                                  | $\in \Pi_2^P$          | ?                                  |
| $\{\hat{L}, \hat{U}_I, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -CR                                                                                                                                | coW[1]-hard      | ?                                  | coW[1]-hard            | ?                                  |
| $\{ H , \hat{U}_I\}$ -CR                                                                                                                                                                          | FPT              | FPT                                | ?                      | ?                                  |
| $\{ H , \hat{U}_I\}$ -CR                                                                                                                                                                          | FPT              | FPT                                | FPT                    | FPT                                |
| SUFFICIENT REASONS (SR)                                                                                                                                                                           | $\in \Sigma_2^P$ | $\{3PA\}$ -inapproximable          | $\in \Sigma_2^P$       | $\{3PA\}$ -inapproximable          |
| $\{\hat{L}, \hat{U}_O, \hat{W}, \hat{B}, \hat{u}\}$ -SR                                                                                                                                           | NP-hard          | $\{3PA\}$ -inapproximable          | N/A                    | N/A                                |

# What's Next?

- Everyone is just scaling SAEs
  - Can we come up with a better model for LLM “thoughts”?
- People have decided we're stuck with the models we have now
  - Can we train more interpretable models w/ low tax?
- Much of the promise of mech interp seems far off
  - What can concretely do now to help AI safety?
- Curious what everyone else thinks!

# Questions