

Competence-based Analysis of Language Models

Adam Davies, Jize Jiang, ChengXiang Zhai
Department of Computer Science, University of Illinois Urbana-Champaign

What properties does LLM use to perform a task?

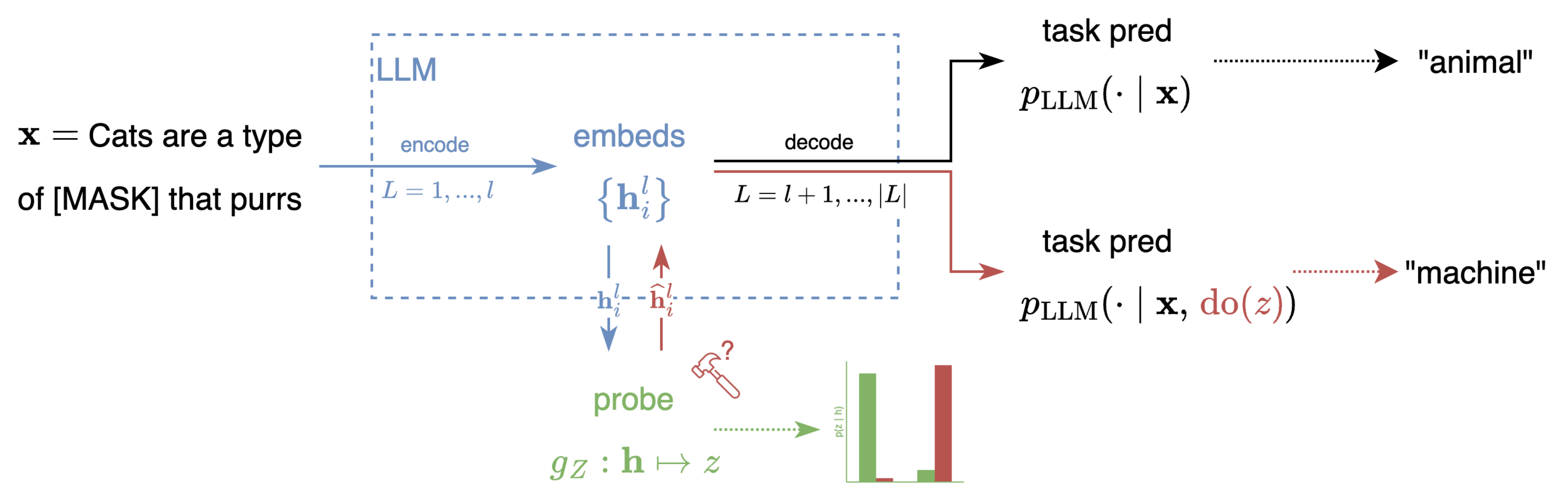
Causal → robust
Spurious → brittle

CAUSAL PROBING

1. Train probe to predict property
2. Intervene on probe to modify representation
3. Analyze impact on model behavior

CALM: use causal probing to measure **competence**

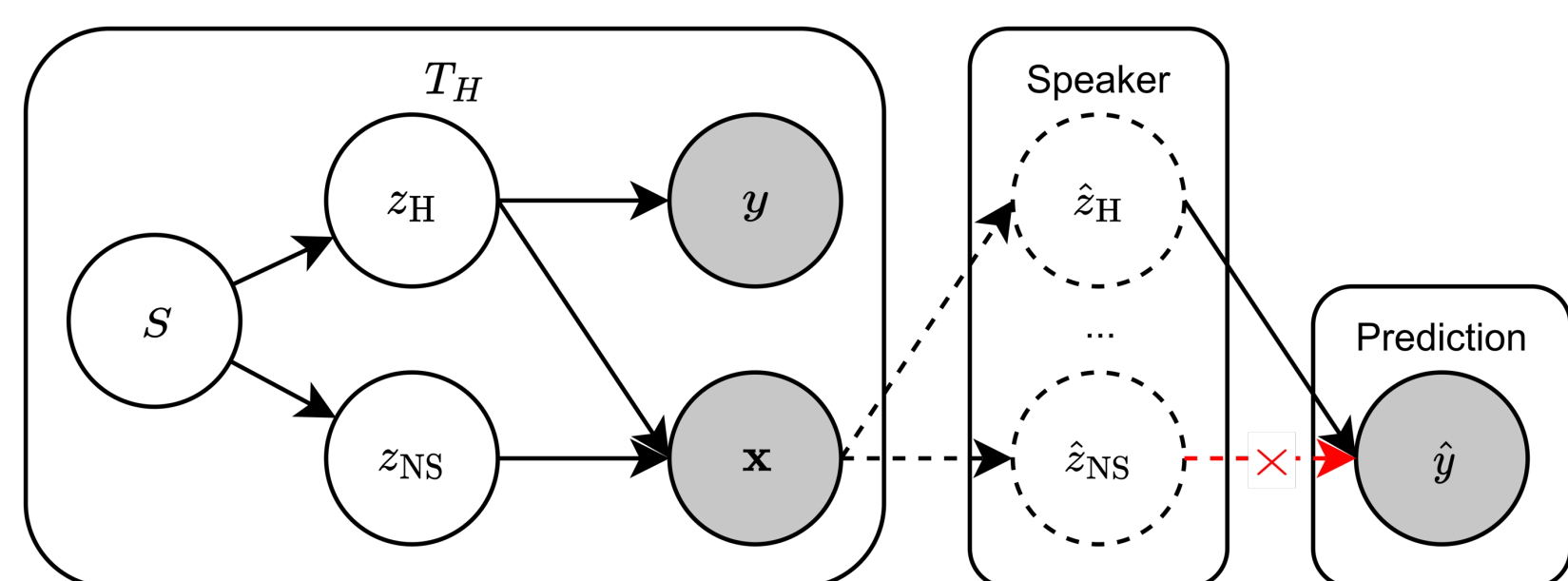
- Change predictions when modifying task-causal properties
- Stays the same when modifying spurious properties



CALM FRAMEWORK

Causal Task Structure

- LLM M
- Task $\mathcal{T} \sim P(\mathcal{X}, \mathcal{Y})$
- Set of properties $\mathbf{Z} = \{Z_j\}$ taking values $z \in \{z_k\}$ for a given input $\mathbf{x}$
 - Decompose $\mathbf{Z} = \mathbf{Z}_c \cup \mathbf{Z}_e$ for *causal* vs *environmental* properties
 - Dependency $M(\mathbf{x}|\text{do}(Z_j)) \neq M(\mathbf{x})$ is *spurious* if $Z_j \in \mathbf{Z}_e$



- Structural causal model $\mathcal{G}_{\mathcal{T}}$
 - Nodes $\mathbf{Z} \cup \{\mathcal{Y}\}$ for set of properties $\mathbf{Z} = \{Z_j\}$
 - Edges denote causal dependencies
 - Decompose $\mathbf{Z} = \mathbf{Z}_c \cup \mathbf{Z}_e$ for *causal* (path to $\mathcal{Y}$) vs *environmental* (spurious, no path)

Measuring Competence

Define *competence* of M w.r.t. $\mathcal{T}$ as alignment with $\mathcal{G}_{\mathcal{T}}$

$$\mathcal{C}_{\mathcal{T}}(M|\mathcal{G}_{\mathcal{T}}) = \mathbb{E}_{\mathbf{x} \sim \mathcal{X}, \mathbf{z} \sim \text{val}(\mathbf{Z})} S(M(\mathbf{x}|\text{do}(\mathbf{z})), \mathcal{G}_{\mathcal{T}}(\mathbf{x}|\text{do}(\mathbf{z})))$$

- Reformulation of Interchange Intervention Accuracy:
 - IIA uses *interchange interventions*: extract representation of property $Z = z$ from source $\mathbf{x}_s$ to patch into target $\mathbf{x}_t$
- CALM uses *causal probing interventions* instead:
 - Operate at *concept-level*
 - No need to “borrow” representations from other inputs $\mathbf{x}_s$
 - Can study unseen combinations of $\mathbf{Z}$ (required for simulating OOD)

EXPERIMENTS

Dataset

LAMA ConceptNet: 14 *lexical inference tasks* for masked-language models

- *Hypernym prediction* (IsA): “cats are a type of [MASK] that purrs”
- Also includes relations like PartOf, HasProperty, etc.

Competence Approximation

Approximate $\mathcal{C}_{\mathcal{T}}(M|\mathcal{G}_{\mathcal{T}})$ via $E = (\mathbf{Z}, \mathcal{G}_{\mathcal{T}}, S)$

- $\mathbf{Z}$ is the set of relations corresponding to each of the 14 lexical inference tasks
- $\mathcal{G}_{\mathcal{T}}$ is SCM with a single edge (as determined by the task $\mathcal{T}$; other relations are $\mathbf{Z}_e$)
- S is the overlap between top- k predictions before and after intervention

Experiments

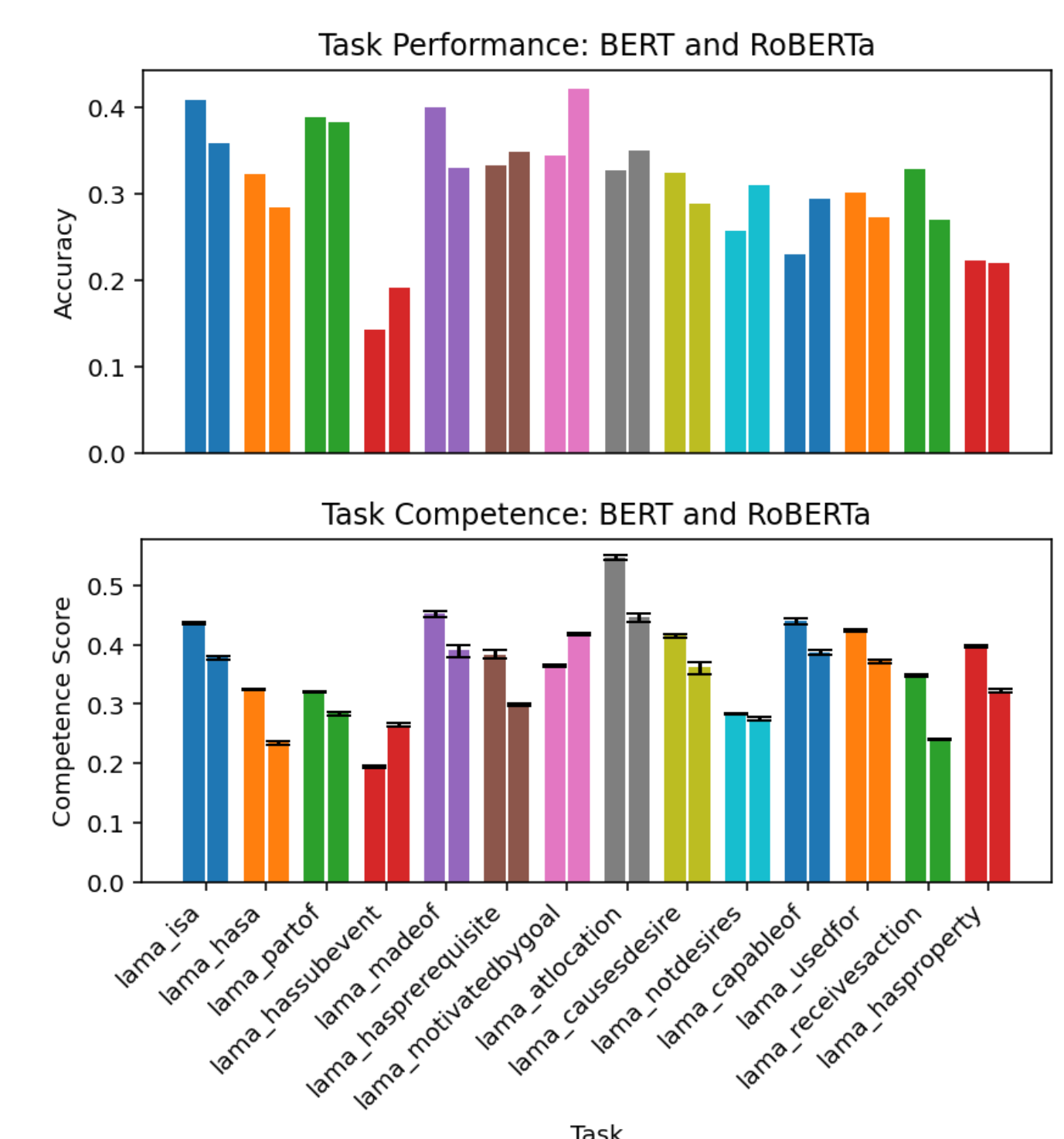
Implement interventions using **GBIs**

- Probe is MLP over final embedding layer
- Attack probe using FGSM and PGD (constrain collateral damage via ϵ)

Measure average approximated $\mathcal{C}_{\mathcal{T}}(M|\mathcal{G}_{\mathcal{T}})$ of **BERT** and **RoBERTa** on each task

- Averaged over 10 experimental runs (randomly re-initialize probes)

RESULTS



BERT (left bars) and RoBERTa (right bars)

- Both models are partially competent across tasks
 - Always higher than random baseline (0.0714)
- Competence is predictive of relative task performance (Spearman's $\rho = 0.508, p = 0.064$)
- Explains earlier findings re: hypernym prediction
 - Great performance with engineered prompts, but fail under small changes to prompts
 - *Explanation*: intermediate competence means models rely on both task-causal and spurious lexical properties

GRADIENT-BASED INTERVENTIONS

Prior work in causal probing has used *INLP*:

- *Nullifies* representation of property
- Assumes *linear representation*

For our experimental setting, we need *nonlinear counterfactual* interventions

- *Nonlinear* required for *relational* properties
- *Counterfactual* required for measuring *competence*

Gradient-Based Interventions

Use *gradient-based adversarial attacks* against probes g to minimize probe loss wrt counterfactual target $y' \neq y$

$$\text{E.g., FGSM: } \mathbf{h}' = \mathbf{h} + \epsilon \text{sgn}(\nabla_{\mathbf{h}} \mathcal{L}(g, \mathbf{x}, y'))$$

- *Flexible*: can use any differentiable probe
- *Controllable*: can modulate perturbation magnitude ϵ



PAPER