

## DOES CHATGPT HAVE A MIND?

*Simon Goldstein*<sup>1\*</sup>, *Benjamin Levinstein*<sup>2</sup> 

<sup>1</sup> University of Hong Kong, Philosophy Department, Hong Kong SAR China

<sup>2</sup> University of Illinois at Urbana-Champaign, Philosophy Department, Champaign, Illinois,  
United States

\**simon.d.goldstein@gmail.com*

**Abstract** This paper examines the question of whether Large Language Models (LLMs) like ChatGPT, Claude, and Gemini possess minds, focusing specifically on whether they have a genuine folk psychology encompassing beliefs, desires, and intentions. We approach this question by investigating two key aspects: internal representations and dispositions to act. First, we survey various philosophical theories of representation, including informational, causal, structural, and teleosemantic accounts, arguing that LLMs satisfy key conditions proposed by each. We draw on recent interpretability research in machine learning to support these claims. Second, we explore whether LLMs exhibit robust dispositions to perform actions, a necessary component of folk psychology. We consider two prominent philosophical traditions, interpretationism and representationalism, to assess LLM action dispositions. While we find evidence suggesting LLMs may satisfy some criteria for having a mind, particularly in game-theoretic environments, we conclude that the data remains inconclusive. Additionally, we reply to several skeptical challenges to LLM folk psychology, including issues of sensory grounding, the “stochastic parrots” argument, and concerns about memorization. Our paper has three main upshots. First, LLMs do have robust internal representations. Second, there is an open question to answer about whether LLMs have robust action dispositions. Third, existing skeptical challenges to LLM representation do not survive philosophical scrutiny.

**Keywords** Large Language Models (LLMs); Artificial Intelligence; Folk Psychology; Mental Representation; Interpretability; AI Cognition

### 1. Introduction

Recent developments in AI are stunning. Large language models can generate text that is fluent, accurate, and responsive to human questions. These advances have sparked a fundamental question: Do these AI systems possess minds?

To address this broad question, we focus on a more specific aspect of the mental: folk psychology. The key question we explore is whether LLMs have beliefs, desires, and intentions. In other words, do LLMs have goals about what to do, a perspective on what the world is like, and plans for achieving their goals given what the world is like?

Why care about LLM folk psychology? There are at least three reasons. First, folk psychology is the primary lens that humans use to understand the behavior of agents. When we interact with one another, we consistently try to explain what is happening in terms of beliefs and desires (Hutto and Ravenscroft 2021).<sup>1</sup> As LLMs become increasingly capable, we will increasingly interact with them. It is worth figuring out whether we can genuinely use beliefs and desires to understand these interactions, or whether instead beliefs and desires would be at best an elaborate metaphor. Second, folk psychology is relevant for moral patiency. When we ask what it takes to be the kind of entity that can be harmed, one important answer appeals to the satisfaction of desires (Heathwood 2016). More may be required too, such as full-fledged consciousness; but folk psychology will play a role. Third, folk psychology is relevant to AI safety. As AI systems become more powerful, many have worried that they may systematically pursue goals that conflict with humanity (Russell 2019). But much of this discussion implicitly assumes that AI systems will have goals and a perspective about how to achieve those goals. If there are important barriers for AI systems to possess a folk psychology, this may complicate our understanding of what it would take for AI systems to be safe.

Methodologically, we'll focus on two important aspects of AI mental states: internal representations and dispositions to act. Our first question is whether LLMs have robust internal representations of the world. Here, our strategy in §2 will be to survey various philosophical theories of representation, and see whether LLMs satisfy the theories. We'll look at a range of conditions on representation, including: (i) that the system has internal states that *carry information* about the world; (ii) that these internal states are *causally effective* at producing the system's behavior; (iii) that these internal representations satisfy *folk patterns of reasoning*; (iv) that the *structure* of these internal representations mirrors the structure of what they represent; and (v) that the information-carrying capacity of these representations emerged from some kind of *selective, evolutionary process*. In each case, we'll draw on recent research in machine learning about LLM *interpretability*, which suggests that these systems satisfy the relevant condition on mental representation.

But internal representations alone are not sufficient for a full-blown psychology. In order to possess beliefs and desires, the second key question is whether LLMs have robust dispositions to perform actions. Nearly every theory of belief and desire will explain these mental states in terms of some combination of internal representations and dispositions to act. For this reason, we focus on these two questions throughout. In §4, we explore two of the most prominent traditions of theorizing about belief and desire, interpretationism and representationalism. In each case, we suss out what kinds of action dispositions the theory requires of LLMs. Then we critically assess whether LLMs satisfy the relevant condition. The crucial question for LLMs will be whether their linguistic outputs are *stable* enough to be best explained as promoting goals. Our conclusion in §4 will be tentative. We argue that the behavior of LLMs in game environments is suggestive of the kinds of rich plans of action required for belief/desire psychology. But the data is not decisive. (That said, we'll flag that the bulk of our discussion will focus on mental representations rather than action dispositions.)

Our third goal of the paper, in §3, will be to refute some of the existing skeptical challenges to LLM folk psychology. Here, we'll engage directly with three challenges. The first challenge, *symbol grounding*, raises the following question: if the only inputs to a language model are strings of text, rather than rich perceptual experiences or feedback from motor

---

1. Although see (Andrews 2013) for more sophisticated discussion of whether predicting the behavior of others, especially in the animal context, ultimately requires belief-desire attribution.

actions, how can language models understand prompts about the external world? The second challenge is that LLMs do not represent the world because they are not trained to do so; instead, they are merely *stochastic parrots*, trained to predict the next word. The third challenge is that LLMs do not represent the world because their behavior can be fully explained by an alternative theory, according to which they rely on *memorization* and shallow shortcuts. In each case, we'll argue that the challenge does not survive philosophical scrutiny.

Overall, then, our conclusions for future research run as follows. First, we think there are interesting questions about exactly how LLMs represent the world, and what in the world they represent. But we think overall there is quite strong evidence that they do so. Second, there is a rich debate to be had about whether LLMs genuinely have beliefs and desires about the world, in the sense that is involved in assembling complex plans of action. Third, we suggest that many of the existing skeptical challenges to LLMs deserve more careful development, as they possess major shortcomings. Before we turn to our main claims, we'll end this section by briefly summarizing how LLMs work.

This paper contributes to a growing literature that connects classic questions in the philosophy of mind to contemporary AI systems. Other recent work in this area includes Millière and Buckner (2024a) and Millière and Buckner (2024b). For a recent argument that LLMs have beliefs and desires (Cappelen and Dever 2025). Cappelen and Dever adopt a very different methodology than our own. They give a direct argument that LLMs have beliefs and desires, because they have the ability to answer questions and make suggestions. By contrast, our own approach is top-down: we focus on whether LLMs satisfy the demands of existing theories of belief and desire, with special emphasis on mental representation. Effectively, we will search among existing theories for various conditions on possessing mental representations, and possessing beliefs and desires, and we will assess the extent to which LLMs satisfy these various conditions.

It is also worth flagging one alternative response to questions of AI psychology. Perhaps LLMs are sufficiently different from existing minds that there is just no determinate fact of the matter about how our existing concepts would apply to them. Or, relatedly, maybe LLMs suggest that existing theories of representation are flawed. In either case, perhaps the better question would be how we should *engineer* our concepts of the mind to grapple with LLMs. We lack the scope to critically consider this proposal below. But below we will argue that LLMs display various complex representational capacities that could be relevant to the project of conceptual engineering, depending on whether newly engineered concepts of the mind would be sensitive to these questions about representation.

Finally, it is worth flagging that throughout the paper we'll move freely back and forth between talking about LLMs, and talking about the chatbots (like ChatGPT and Claude) that are built on top of them. In general, we tend to think that questions about representation are in the first case targeting LLMs, while questions about action disposition are in the first case targeting chatbots. But since chatbots are built by sampling from LLMs in a special way, it makes sense to consider both together.

## 1.1 Large Language Models

At the heart of modern LLMs lies the transformer architecture (Vaswani et al. 2017). Although many variants exist, we'll here focus on the mechanics of decoder-only models such as GPT-4 or Llama 3.<sup>2</sup>

---

2. Many recent models share this basic architecture, though details vary.

When you feed a prompt to a model, it makes a prediction about what comes next. For example, if you feed it *The cat sat on the*, it will assign a probability to each possible next token.<sup>3</sup> It will assign some probability to *aardvark*, some probability to *banana*, some probability to *mat*, and so on.

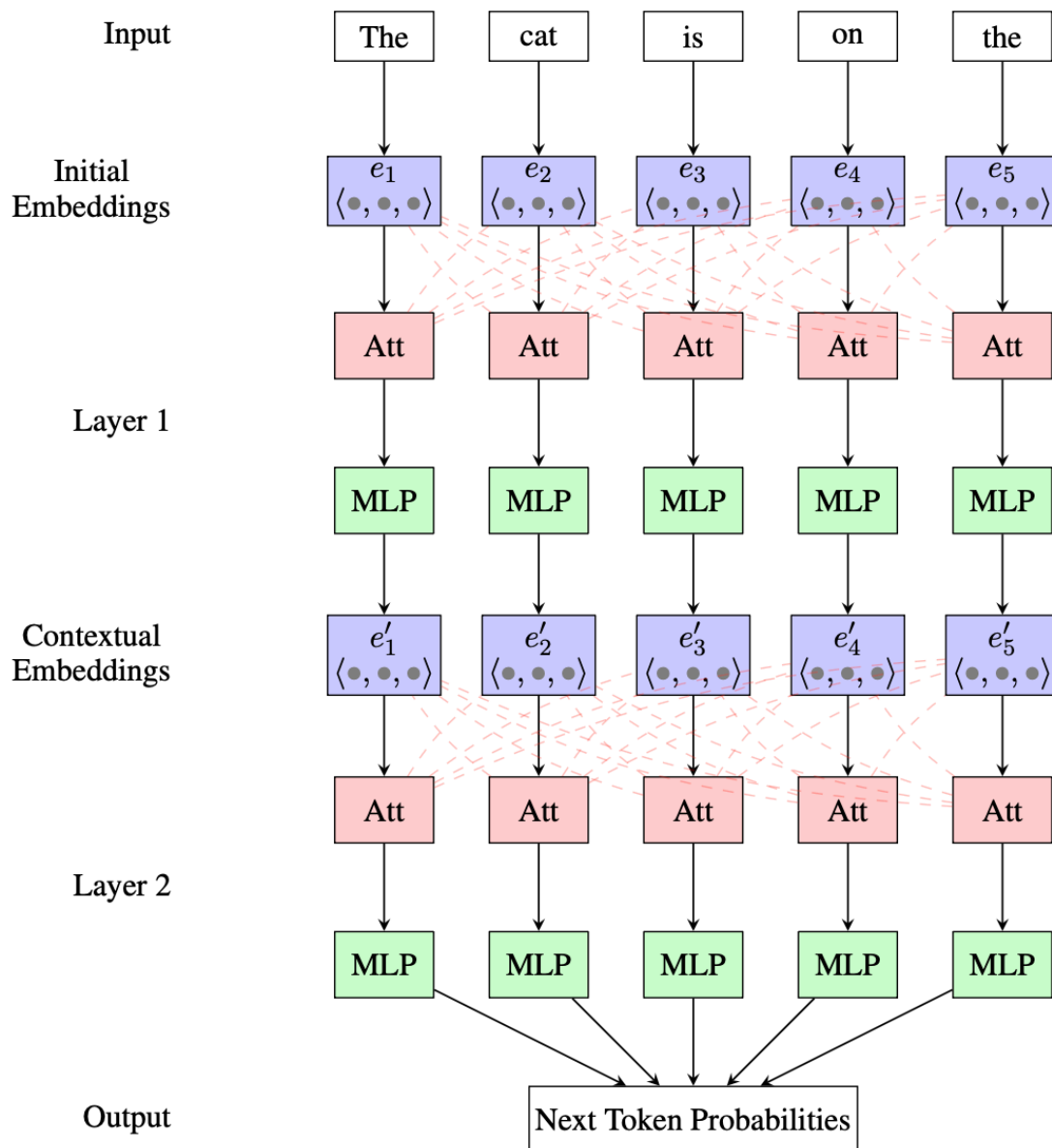
To compute these probabilities, each token is converted into an initial *embedding*—a vector, or long list of numbers, which encode *The*, *cat*, *sat*, *on*, and *the*. The initial embedding carries information about the corresponding token, but it doesn't at first carry any information about surrounding tokens. For example, the initial embedding for *cat* does not encode the fact that *The* precedes it.

The initial embeddings for each token are then transformed and updated across a large number of layers. At each layer, the embedding for a given token is first updated through the mechanism of self-attention. Self-attention allows the embedding for a given token to attend to itself and earlier tokens in the sequence. For example, the token for *sat* might attend to *cat* at a given layer. The embedding for *sat* could then be updated to represent the fact that *cat* was the immediately preceding token or, perhaps, to encode somehow that *cat* was the subject. In other words, after paying attention to the information of earlier tokens and itself, the embedding for *cat* is updated to include contextual information about the surrounding tokens in the prompt. We call this new embedding a *contextual embedding*. These new embeddings are then refined further using something like a multi-layer perceptron (MLP) before being fed forward into a new layer.

After the embeddings are passed through all the layers of the model, the model uses the final contextual embedding for a token to predict what the next token will be. To generate more and more text, we can then select some token assigned relatively high probability, tack it onto the initial prompt, and then feed the new augmented sequence to the model again. For an illustration, see Figure 1.

---

3. Tokens can be words, subwords, numerals, punctuation, etc. For our purposes, we can just think of tokens as words.



**Figure 1.** Simplified two-layer transformer architecture processing The cat is on the. Each word is initially converted to an embedding vector. In each layer, self-attention (Att) allows words to attend to each other, followed by a multi-layer perceptron (MLP). After the first layer, new contextual embeddings are created. The final layer produces probabilities for the next token.

Transformer models get extremely good at generating plausible text via training. Training for retail models like ChatGPT or Claude comes in two main phases. In the first phase, we take lots of pre-existing text, feed an initial segment of it to the model, and then have the model generate probabilities about what comes next. We then tweak the parameters of the model via gradient descent to make more accurate predictions. For instance, if the model is fed Happy families are all alike; every unhappy family is unhappy in its own it will learn to generate a higher probability for way as the next token. The model's parameters are adjusted through gradient descent to improve its predictions over time.

Eventually, the model gets very good at predicting what comes next in text. But it is still not especially conversational or useful. So, the second (more optional) phase of training, called fine-tuning, involves tweaking the model further to be a better conversational agent. We will omit the details of this portion of training here, but the essential idea is to get humans or other AI to rate various responses for quality and then push the model to be more disposed to generate high quality responses.<sup>4</sup>

Note, at this point, some initial obvious reasons for skepticism about LLMs. In pre-training LLMs are rewarded, essentially, for making good predictions about sequences of text. There is no direct pressure to represent the world nor to represent truth. Instead, the immediate pressure on LLMs is to find plausible continuations of prompts (on the other hand, other rewards are added in post-training). Furthermore, pure LLMs only connect with the world via textual embeddings. They have some initial embedding for words like rainbow or hammer, and they manipulate these embeddings based on context, but they've never actually seen a rainbow nor manipulated a hammer.

There is, then, an important question of whether LLMs can or do end up representing the world as a means toward their training objective. While they manipulate embeddings based on context, their lack of direct sensory experience and the focus on plausibility rather than accuracy raise reasons for skepticism about their representational capabilities.

## 2. Theories of representation

Can LLMs represent the world? In this section, we explore this question by examining various philosophical theories of representation. We aim to demonstrate that LLMs meet many of the criteria set by these theories. In particular, we'll highlight a series of recent studies in AI interpretability research and related behavioral research that connect closely to various philosophical theories of representation.

We focus primarily on naturalistic theories of representation. These theories explain representation in terms of physical processes. They differ from other theories, such as those proposing that representation is primitive (Boghossian 1990) or dependent on primitive phenomenal properties (Graham et al. 2007). While non-naturalist theories may allow for LLM representation, it is difficult to say whether AIs could have primitive representational or phenomenal properties.

Our focus is on whether LLMs have internal states that represent the world by having truth conditions. A representation can truly depict the world if the world is a certain way and falsely if it is another. We will also consider if LLMs have internal states that refer to objects or properties in the world.

Importantly, our main interest is not the text produced by LLMs but whether LLMs have mental states that represent the world. Our hypothesis is that the *activations* of LLMs—the patterns of internal neural activity within the network as it processes information—refer to objects in the world and have truth conditions. These activations, which can be thought of as the temporary, computation-specific states of the network, are distinct from the more permanent weights that encode the model's learned knowledge. This question of representation is an important first step in determining whether LLMs have a robust folk psychology with beliefs and desires.

---

4.A common method is reinforcement learning from human feedback (RLHF) (Christiano et al. 2017), but there are a number of alternatives such as reinforcement learning from AI feedback and direct preference optimization.

With these questions in mind, we will survey leading theories of mental representation to see if LLMs meet their criteria. We draw on existing surveys, including Adams and Aizawa (2021) and Schulte (2023). Our goal is to show that, according to these leading theories, there is strong evidence that LLMs indeed represent the world.

We will consider five key conditions on mental representations, and argue that LLMs satisfy each one:

- **Information carrying:** Informational theories posit that mental representation requires internal states to *carry information* about the external world, typically through probabilistic connections. We demonstrate how recent advances in AI interpretability, particularly in *probing* techniques, provide compelling evidence that LLM internal embeddings carry such world-relevant information.
- **Causal efficacy:** Fodorian theories demand that representational states be *causally effective* in generating system behavior. We present evidence from recent interpretability studies showing that LLM outputs indeed depend counterfactually on their internal embeddings, satisfying this causal requirement.
- **Folk-psychological reasoning:** Another Fodorian condition requires that reasoning with internal representations follows patterns familiar from folk psychology. We argue that emerging research on *world models* in LLMs reveals their capacity for effective reasoning about the world.
- **Structural isomorphism:** Structural theories of representation require that the internal architecture of representations mirrors the structure of what they represent. We illustrate how recent studies on LLM concepts of color and direction demonstrate that their representations exhibit the requisite structural properties.
- **Selection:** Teleosemantic theories stipulate that genuine representations must emerge from a selective, evolution-like process. We contend that the training methodologies employed in developing LLMs fulfill this selectional criterion.

## 2.1 Information

A long tradition of work on representation has appealed to the concept of carrying information. Smoke carries information about fire, mumps carry information about measles, and thermometers carry information about temperature. Dretske (1981) and others have argued that a system can only represent the world if that system carries information about it.

Philosophers have disagreed about how exactly to define carrying information, but in general different analyses all appeal to probabilistic concepts. According to Dretske, a state carries the information that  $p$  if and only if the probability of  $p$  given the state is 1, provided that various background conditions obtain. Some theorists have instead focused on states raising the probability of  $p$ , rather than making it certain (Usher 2001). Other theorists have focused on more general conditions involving entropy.<sup>5</sup>

How can we tell whether LLMs carry information about the world? Harding (forthcoming) argues there is a close connection between informational approaches to representation and recent work on probing in AI interpretability research (Alain and Bengio 2018). In probing, researchers train a separate classifier to take activations as input and make a prediction.

---

<sup>5</sup>The exact connection between carrying information and representing the world is also debated. For Dretske, a mental state represents that  $p$  iff the state carries the information that  $p$  during the end of the subject's learning period associated with the state. This doesn't require that whenever a state represents that  $p$ , it carries the information that  $p$ ; for Dretske, carrying information is *factive*, and so this would rule out misrepresentation. But other notions of carrying information might not be *factive*, and so could allow a more direct connection between representation and information.

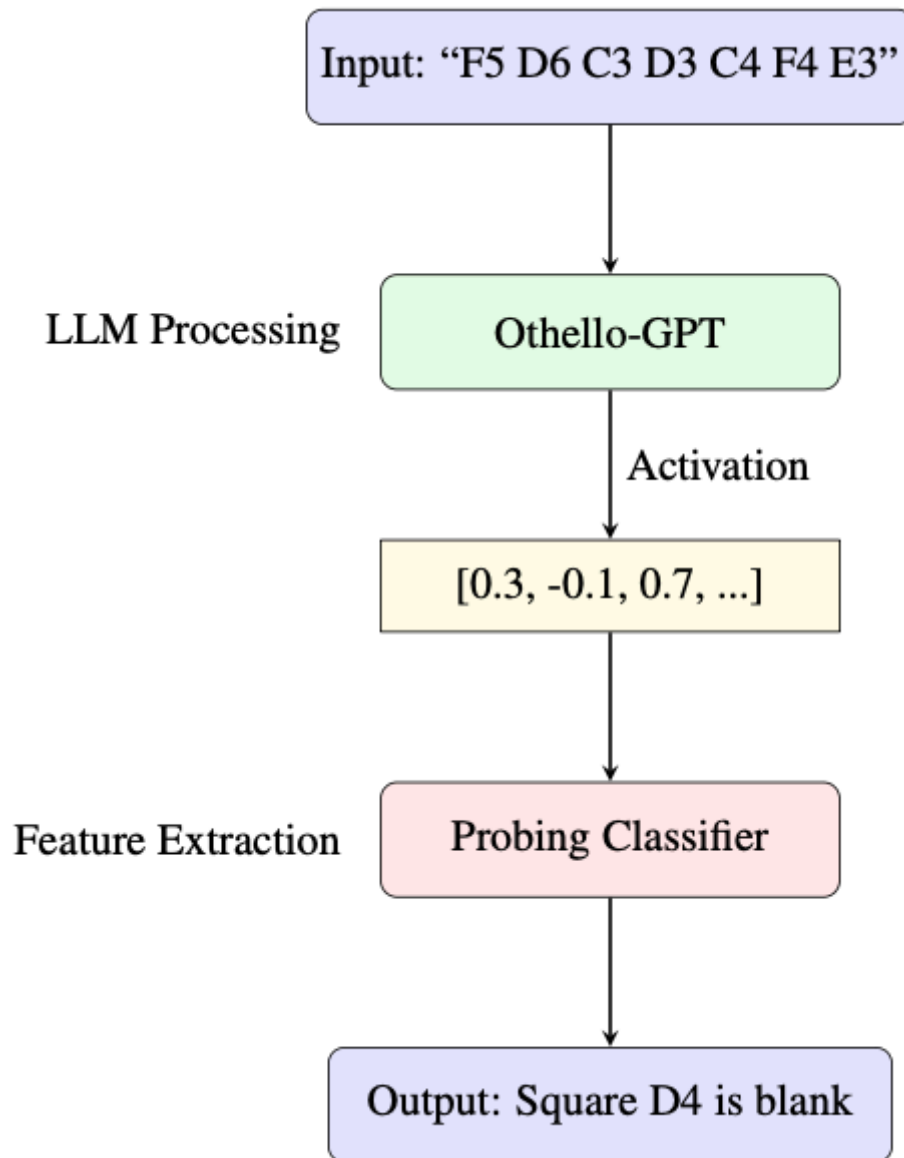
The probe takes in some activations and predicts features of the input. For example, in a visual AI system, a probe might predict whether the system is looking at a cat based solely on activations, without access to the original input.<sup>6</sup>

A compelling example of using probes to discover LLM representation comes from Li et al. (2023). We'll use this example as a case study below. Li et al. trained an LLM on sequences of moves in the 8x8 board game Othello, using only lists of moves (like F5 D6 C3 D3 C4 F4 E3) without describing the rules or the board. Othello is a simpler game than chess, but there are far too many possible moves in general for an LLM simply to memorize all legal game states. Despite this, the trained LLM, dubbed Othello-GPT, tended to output legal moves with high probability. To understand how, the authors used probes to identify possible internal representations of the board. By comparing the model's internal states to the actual board, they trained probes to guess whether each of the sixty-four squares was black, white, or blank. The probes achieved remarkable accuracy with an error rate of only 1.7%. Similar probing techniques have been used to understand how models represent grammatical case, number, tense, and more (Giulianelli et al. 2021). Figure 2 illustrates the probing process.

---

6. Probes typically use linear classifiers or shallow neural networks trained on model activations to predict specific features. For details, see (Alain and Bengio 2018).





**Figure 2.** Illustration of the probing process using Othello-GPT. The LLM processes the input sequence of Othello moves, generating activations. A separate probing classifier is trained to predict specific features (e.g., the state of a particular square) from these activations.

Informational theories of representation make sense of the relevance of probes for mental representation. If the classifier can correctly determine whether a cat is present almost 100% of the time, and the classifier is only using the activations of the system to make its prediction, then those activations likely represent the cat. At the very least, they carry mutual information with the presence of cats. If mental representation is a matter of carrying information, then there is no special barrier to LLMs representing the world.<sup>7</sup>

## 2.2 Causal powers

Fodor (1975) imposed a series of conditions on genuine representations, one of which is that they must have genuine causal powers. The key question here is whether LLM outputs are robustly caused by the activations discovered by probes or whether these apparent representations are actually epiphenomenal.

To make this concern concrete, return to Othello-GPT. Suppose Othello-GPT itself does not represent the state of the board at all. However, a clever probe could read off from its activations alone what the history of moves was. For instance, if F5 D6 C3 D3 C4 F4 E3 were fed to Othello-GPT, the probe could learn just from its internal state that F5 D6 C3 D3 C4 F4 E3 was the input. If the probe also learned (via its own training) what the rules of Othello were, it could determine whether each square was black, white, or blank, even if Othello-GPT itself didn't represent that fact.

Alternatively, we can imagine that Othello-GPT does represent the board state, but the probe found some other way to determine whether its square was black, white, or blank using information from Othello-GPT's activations that Othello-GPT itself does not use. To investigate this question, interpretability researchers use methods of causally intervening on a model's activations to test whether supposed representations are actually relevant or useful to the model.<sup>8</sup> For instance, if we find a successful probe for a square in Othello-GPT, we can determine what elements of Othello-GPT's activations the probe is using to identify whether a square is black, white, or blank. We can then hand-edit the activation to see what happens.

Suppose, in an oversimplified example, that the probe notices that when the first coordinate of an embedding vector for a square is 1, the square is black; when it is 0, it is blank; and when it is -1, it is white. We can feed the model a prompt making the square black, setting the first coordinate of the relevant embedding vector to 1. Then, we can manually change this coordinate to 0 or -1. If the model's output changes as expected, we have strong evidence that the representation we found is one the model itself uses to track the board state. 5 illustrates the process of causal intervention in the Othello-GPT model. This demonstrates how altering specific internal representations can change the model's outputs, providing evidence for the causal efficacy of these representations.

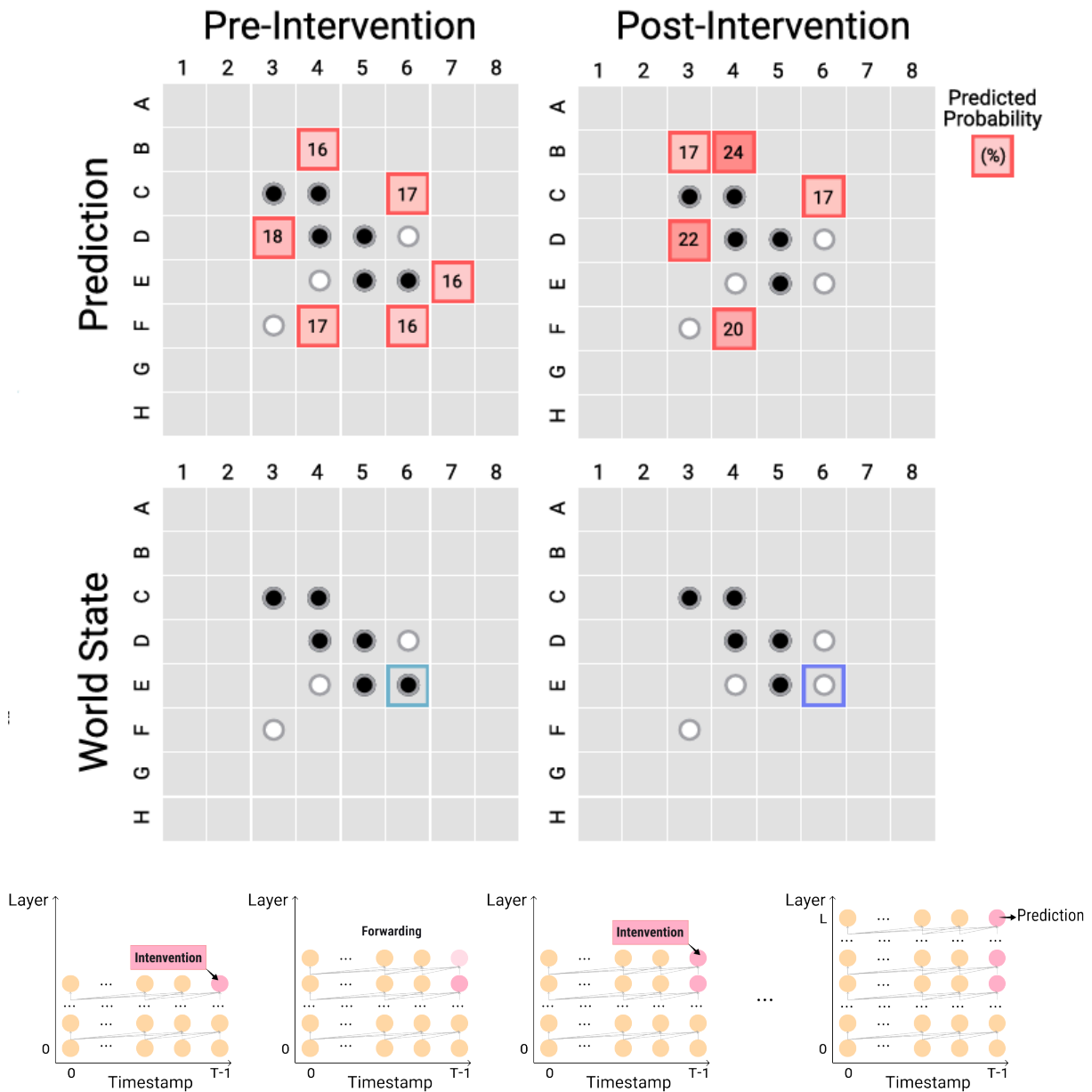
7. Recent work on *sparse autoencoders* (SAEs) and *activation oracles* provides complementary, unsupervised methods for discovering interpretable structure in LLM activations. Unlike probes, which are supervised classifiers trained to detect a specific target property, SAEs decompose a model's activations into a sparse dictionary of monosemantic features without any labeled data, recovering directions that the model's own computation relies on (Bricken et al. 2023), (Templeton et al. 2024). When scaled to production models like Claude 3 Sonnet, SAEs have uncovered thousands of interpretable, abstract features—including features for specific entities, code errors, and deceptive behavior—that are multilingual, multimodal, and steerable. More recently, *activation oracles* train language models to accept neural activations as inputs and answer arbitrary natural-language questions about them, generalizing far beyond their training distribution to uncover phenomena like misalignment introduced through fine-tuning (Karvonen et al. 2026). Because both methods are unsupervised, they partially address the concern (discussed below) that probe-discovered representations may be epiphenomenal: the features they identify are extracted from the model's own representational geometry rather than imposed by an external classifier.

8. There are many different methods of intervention to understand and manipulate neural network representations. The simplest is ablation, where entire neurons are deactivated to assess their importance in the model's performance. More sophisticated approaches include: Iterated Nullspace Projection (Ravfogel et al. 2020): This method iteratively projects embeddings into a nullspace to remove the influence of specific concepts, effectively isolating and erasing targeted information. Least Squares Concept Erasure (LEACE) (Belrose et al. 2023): LEACE identifies and removes concept-specific information from embeddings using a least squares optimization approach, ensuring that the targeted concept is no longer represented in the neural activations. Causal Tracing (Meng et al. 2022): This method tracks the flow of information through a model to determine the causal impact of specific components or representations. By intervening in the causal pathways, researchers can assess the importance and role of particular features in the model's decision-making process. These techniques allow researchers to surgically target and manipulate information encoded in embeddings, providing deeper insights into the functioning and interpretability of neural networks.

As another example, consider finding a potential representation of grammatical number in an LLM. Suppose we feed the LLM the prompt I ate these. Without hand-editing, apples should receive a higher probability than apple as the next token. But if we hand-edit the representation of these from plural to singular, we can see if apple becomes more likely than apples. If so, we have causal evidence that our supposed representation is used by the model.

In the case of Othello-GPT, causal interventions were effective showing that the probes did in fact find the model's representation of the board's state. In other cases, similar causal methods have been used to find representations of grammatical number, subjecthood, and factual associations (Giulianelli et al. 2021), (Meng et al. 2022). We thus have direct evidence of the causal effectiveness of LLM representations.

More recently, circuit-tracing methods have moved beyond testing individual representations to mapping entire computational pathways within LLMs. Ameisen et al. (2025) and Lindsey et al. (2025) developed *attribution graphs* that trace the chain of intermediate steps a model uses to transform a specific input into an output, applied to a production-scale model (Claude 3.5 Haiku). Among other findings, they discovered that the model plans ahead when composing poetry (activating rhyming-word features before writing the line leading to them), performs multi-step factual reasoning by composing separately stored facts, and operates through a shared conceptual space across languages—the same core features activate for the same concepts in English, French, and Chinese, with language-specific circuitry handling only the final translation step. These results go beyond showing that individual representations are causally effective; they reveal how representations are composed into multi-step reasoning chains that drive model behavior.



**Figure 3.** Activation patching in Othello-GPT. (a) The top panel shows the Othello board state and model predictions before and after intervention. The upper row in each state displays the model's move predictions with associated probabilities, while the lower row shows the actual board state. Pre-intervention, the model correctly predicts legal moves. Post-intervention, the model's predictions change. (b) The bottom panel illustrates the process of activation patching across different layers and timestamps of the model. The intervention at a specific layer and timestamp propagates through subsequent layers, ultimately affecting the final prediction. This demonstration shows how altering internal representations can causally influence the model's outputs, even leading to illegal move predictions in the context of the game state shown above. Both depictions are adapted from (Li et al. 2023).

### 2.3 Folk patterns of reasoning

So far, we've argued that LLM activations carry information and causally influence LLM outputs. But this alone may not be enough for genuine representation. For example, Fodor (1975) argued that genuine mental representations have to have special kinds of causal powers: the representations have to influence the system's behavior in ways that match the laws

of folk psychology. The representations have to make sense of the world, causing outputs in a way that resembles ordinary reasoning. Given our interest in whether LLMs have folk psychological states, this theory of representation is especially important for our purposes.

One way interpretability researchers have investigated this kind of condition is through the idea of world models. A world model, in the context of AI, refers to an internal representation of how the external world operates, including objects, their properties, and the causal relationships between them. World models require coherent and consistent representations across contexts and a degree of abstraction that allows LLMs to generalize from particular cases to more general situations. For LLMs, these models are constructed purely from textual data, raising important questions about their nature and limitations.

The importance of world models lies in their potential to bridge the gap between mere pattern recognition and genuine understanding. If LLMs possess robust world models, it suggests they have developed representations that go beyond simple statistical associations, potentially supporting attributions of beliefs and reasoning capabilities.

(Li et al. 2023) suggest their experiments with Othello-GPT are strong evidence of LLMs' ability to create world models, since Othello-GPT has some coherent internal model of the board state that allows it to predict the next move. Indeed, Othello-GPT can easily generalize to make predictions about new boards it's never seen before and can even model permissible moves in impossible boards—i.e., boards with states that can never be reached through a legal initial series of moves. Recent work by Vafa et al. (2024) builds upon and extends this approach, proposing new evaluation metrics for world model recovery inspired by the Myhill-Nerode theorem from language theory. Their study encompasses not only game environments like Othello but also navigation tasks and logic puzzles that can be captured by a finite state automaton. Their research demonstrated that while language models can perform well on existing diagnostics, their underlying internal models vary in levels of coherence.<sup>9</sup>

However, the rules of Othello and those studied by Vafa et al. (2024) are essentially a kind of language that renders some moves grammatical and others not. It's not clear that modeling Othello is sufficient evidence of LLMs' having causal abstractions of physical processes that are robust enough to count as world models.

Recent work by Shai et al. (2024) provides striking evidence that transformers develop structured internal representations of the world. They demonstrate that when trained on data generated by hidden Markov models (HMMs)—systems with hidden states that probabilistically generate observable outputs—transformers learn to represent information about these hidden states in their residual stream in a geometrically organized way. Such information, in effect, represents the transformer's credences about the underlying hidden state of the system given the observed sequence of tokens. Remarkably, these representations emerge even when the optimal belief state geometry has complex fractal structure, suggesting that transformers learn sophisticated world models rather than just superficial patterns. Furthermore, these representations encode information about possible futures beyond just next-token predictions, indicating that transformers develop richer internal models than would be strictly necessary for their training objective. This provides concrete evidence that transformers can develop genuine internal representations of hidden structure in their environment, even when that structure is quite complex.

---

9. Indeed, there is evidence that *some* of Othello-GPT's ability to predict legal moves depends on a combined patchwork of heuristics alongside a model of the board (Lin et al. 2024).

(Musker and Pavlick 2023) explore whether large language models build causal models in order to understand the meaning of words. They rely on the HIPE theory, developed to understand human lexical concepts, in connection with artifact terms like mop and pencil (Chaigneau et al. 2004). According to this theory, when human language users decide whether an artifact term like mop can correctly apply to an object, they rely on an implicit causal model:

The object’s design history and the user’s goal are distal causes in the CM, while the object’s physical structure and the user’s actions with respect to it are proximal causes in the CM. Thus, HIPE predicts that, for example, both the physical structure of an object (e.g., having a handle and something absorbent on one end) as well as the reason the object was originally created (e.g., for wiping up water) should affect how appropriate it is to call the object a mop, but that the latter should have a minimal effect when the former is fully specified (Musker and Pavlick 2023).

The idea is that the status of later nodes screens off earlier nodes when judging whether something is an artifact. To test this theory, Chaigneau et al. (2004) built vignettes that manipulated various nodes in the causal model, and explored which of the nodes influenced language users’ judgments about whether an object counted as a type of artifact. For example, one such vignette was:

One day Jane wanted to wipe up a water spill on the kitchen floor, but she didn’t have anything to do it with. So she decided to make something. [...] The object consisted of a bundle of thick cloth attached to a 4-foot long stick. Later that day, John was looking for something to wipe up a water spill on the kitchen floor. [...] He grabbed the object with the bundle of thick cloth pointing downward and pressed it against the water spill (Musker and Pavlick 2023).

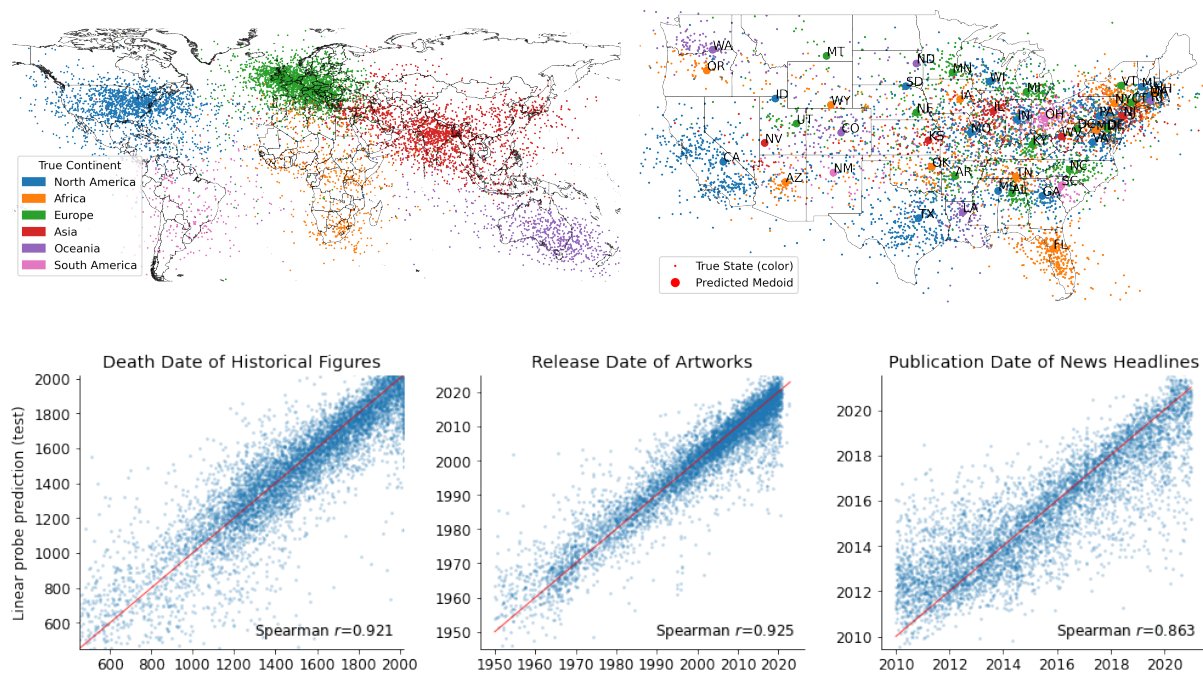
(Musker and Pavlick 2023) applied this theory of lexical concepts to large language models. Musker and Pavlick found that GPT-4 behaved very similarly to human subjects when evaluating counterfactual vignettes. Compromising any factor in the model led to a negative effect on GPT-4 deciding that something counted as an artifact. Compromising proximal features had a larger effect on GPT-4’s judgments than compromising distal features.<sup>10</sup> This potentially suggests that GPT-4 built causal models in order to deploy artifactual lexical concepts, in a way structurally similar to humans.<sup>11</sup>

Recent work by (Gurnee and Tegmark 2024) provides compelling evidence that large language models (LLMs) develop coherent representations of space and time, even when trained solely on text data. By probing the internal activations of Llama-2 models, they discovered that LLMs learn linear representations of spatial and temporal information across multiple scales. These representations are unified across different entity types (e.g., cities and landmarks) and robust to variations in prompting. Remarkably, they identified individual space neurons and time neurons that reliably encode spatial and temporal coordinates. Fig-

10. For example, in one vignette (labeled pencil object, compromised action scenario) an object is designed to be used as a pencil (satisfying the goal condition), but the user fails to successfully use the object to write on a piece of paper (violating the action condition).

11. For more on the emergence of causal models in LLMs, see (Forbes et al. 2019), (Da and Kasai 2019), (Ettinger 2020), (Petroni et al. 2019), and (Kassner and Schütze 2020). Musker and Pavlick themselves shy away from interpreting these results as showing that GPT-4 genuinely builds causal models of artifacts (p. 7). Instead, they suggest that more work is needed to identify methods for probing inner models. After all, their methodology relies solely on GPT-4’s responses to text vignettes. But we have already seen that other work in interpretability research has used probing and other paradigms to uncover world models in the internal representations of LLMs. In this setting, even the behavioral evidence from Musker and Pavlick can potentially be interpreted as revealing inner causal models associated with LLM lexical concepts. See (Yildirim and Paul 2024) for further work exploring the philosophical upshots of world models in LLMs.

ure 8 provides striking visual evidence of how LLMs develop structured representations of space and time. The clear organization of locations and events in the model’s internal space suggests that these representations go beyond mere statistical associations.



**Figure 4.** Spatial and temporal representations in Llama-2-70b. Each point corresponds to the layer 50 activations of the last token of a place (top) or event (bottom) projected onto a learned linear probe direction. The clear structure in these projections, closely matching real-world geography and chronology, demonstrates that the model has learned coherent representations of space and time. All points depicted are from the test set. (Adapted from (Gurnee and Tegmark 2024).)

Nanda et al. (2023) explored internal computations in an LLM trained to perform modular addition. They found that rather than merely memorizing answers or relying on statistical patterns, the LLM implemented a particular Fourier multiplication algorithm for computing the sum: simply put, they perform this task by mapping the inputs onto a circle and performing addition on the circle.<sup>12</sup>

During training, the models initially overfitted the data by memorizing specific examples. However, over time, they transitioned to a generalizable understanding of modular addition, demonstrating an internal shift from memorization to systematic reasoning.<sup>13</sup>

## 2.4 Structure

Structuralist theories of mental representation propose that mental states represent the world only if their internal structure mirrors the structure of the external world (Opie and O’Brien 2004). This mirroring creates a network of relationships within the mental states that correspond to the relationships between the objects or properties they represent. As an

12. See (Zhong et al. 2023) for an alternative algorithm some LLMs learned to perform modular arithmetic.

13. There are many more questions that can be asked about the extent to which LLM representations satisfy various folk reasoning patterns. Recent discussion, for example, has highlighted the reversal curse, where LLMs are competent with the phrase A is B, while failing to understand the phrase B is A (Berglund et al. 2024). There has been rich recent discussion about these sorts of issues, caught up with the question of whether LLM representations exhibit the kind of full-blown compositional structure familiar from (Fodor 1975). (For example, see (Lake and Baroni 2023) for a recent argument that LLMs do exhibit full compositional structure.)

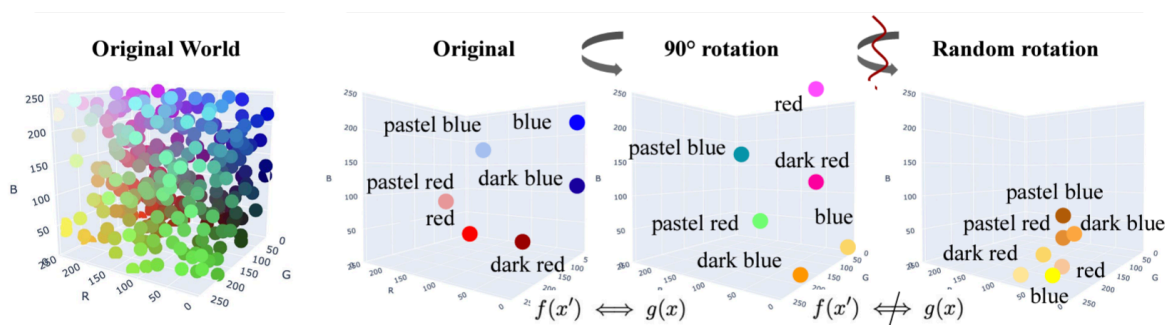
analogy, maps represent geographic areas by containing a series of symbols whose physical relations on the page mirror the physical distances between locations.

Structuralist theories of mental content are particularly relevant to large language models. Interpretability researchers have explored the activations of LLMs to see whether they stand in patterns of relations with one another that are isomorphic to various worldly features.

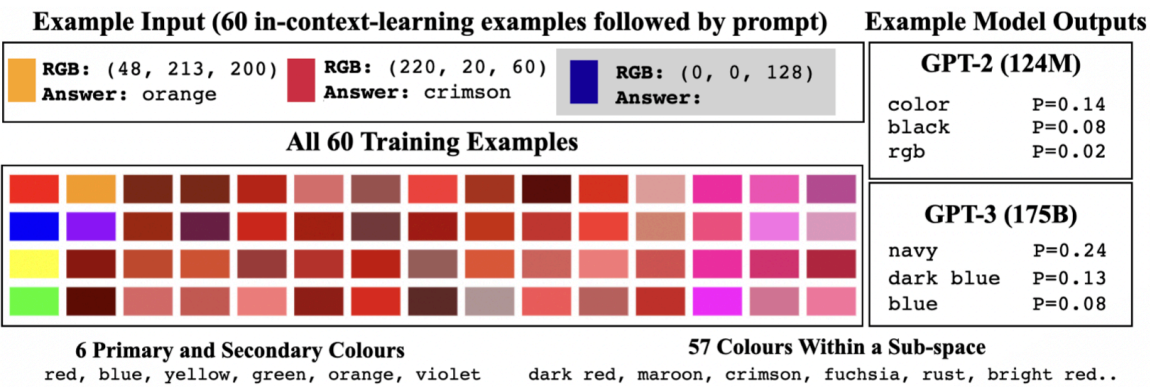
(Patel and Pavlick 2022) tested GPT-3's concepts of direction and color. In the case of cardinal directions, they exposed GPT-3 to a series of gridworlds, consisting of matrices filled with 0s and a single 1. The task was to identify the direction of the 1 symbol in the gridworld: left, right, etc. They prompted GPT-3 with a series of sample answers, and tested whether GPT-3 could complete the pattern. Patel and Pavlick found that GPT-3 had a strong grasp of cardinal directions. First, the LLM could smoothly generalize to a range of gridworld environments: it could correctly describe directions in gridworlds with new lengths and widths. Second, the LLM could smoothly generalize across cardinal directions: even when it was only prompted with northerly or easterly directions, it could correctly answer questions about southerly and westerly located 1s. This showed that GPT-3 had an underlying conceptual space of directions that connected the concepts of north to south, and west to east. Once the LLM could hook up one of these concepts to the gridworld, its underlying structural understanding of direction allowed it to complete the grounding task. Finally, the LLM smoothly generalized to rotated gridworlds. Even when the grids did not map onto our ordinary judgments of left and right, the LLM could quickly generalize from examples. This suggested that the LLM had a stable underlying conceptual space for directions.

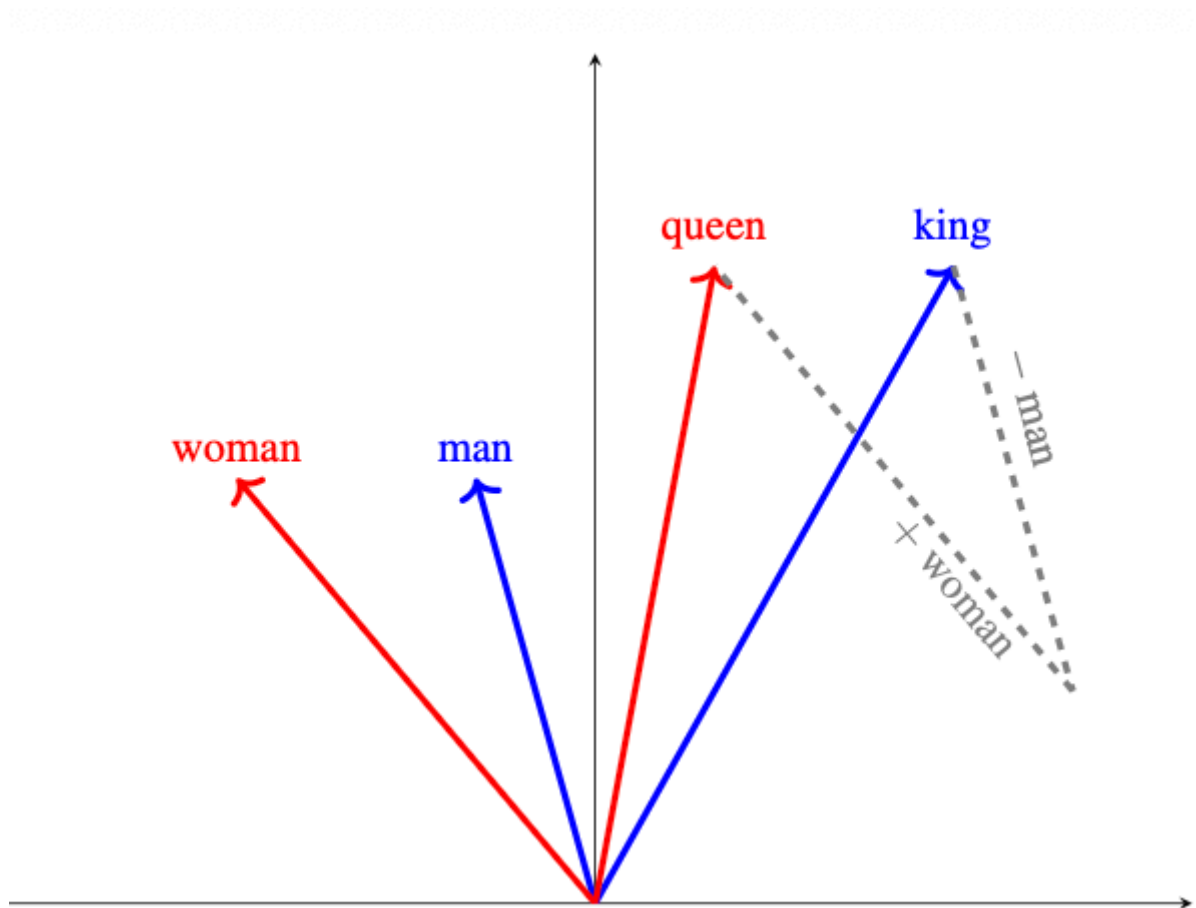
Patel and Pavlick ran a similar experiment for color. In this case, they tested whether LLMs could few-shot learn how to map folk color terms to three dimensional RGB scales. Again, although GPT-3 is a text-only model, it nonetheless possessed an underlying conceptual space for color. Even when it was only prompted with examples that associated RGB values with red colors, it could still use this information to correctly predict how to associate RGB values with blue colors. This showed that the LLM had an underlying grasp of the structural relation between red and blue colors: once it learned how to connect red colors to RGB values, the further connection to blue colors could be inferred.

In both the case of direction and color, Patel and Pavlick's experiment demonstrated that LLMs develop rich networks of representations, with structural relationships that resemble the analogous human concepts. This result fits related interpretability research from (Abdou et al. 2021), which found that text-only language models create a network of embeddings for color terms. These embeddings were found to have structural relationships that mirrored human judgments of similarity between various colors: they found that language model representations of color terms that are derived from text only express a degree of isomorphism to the structure of humans' perceptual color space. For a visual illustration, see Figure 5.









**Figure 6.** Word embedding vector operations illustrated:  $\text{king} - \text{man} + \text{woman} \approx \text{queen}$

When the LLM searches for a word that satisfies this relationship, it can correctly identify queen as the word that completes the analogy. Likewise, the model can understand the relationship between a city and its country through a similar mechanism. The vector difference between Paris and France is similar to the difference between Tokyo and Japan, enabling the model to complete the analogy.

Structural conditions on representational content may not themselves be sufficient to fully explain how mental states have truth conditions.<sup>14</sup> After all, several different networks of physical properties could stand in the same structural relations, and this would then leave unsettled which of them is the referent of the relevant mental state. But structural conditions can be combined with other conditions to produce a full theory of mental content. For example, one strand of work in philosophy of mind has appealed to structural conditions to explain the particular significance of phenomenal experiences. According to representationalists about phenomenal character, phenomenal experiences supervene on their representational content. According to these theories, there can be no change in how things seem qualitatively without a change in the representational content of one's experiences. Many of these representationalists attempt to explain phenomenal experiences in terms of a special kind of representational content. Here, structural conditions have played an important role. For example, Rosenthal (2005) argued that qualitative experiences are associated with a network of representational contents that are arranged in a similarity space. Consider color

<sup>14</sup>See Piantadosi and Hill (2022) for an extended discussion of conceptual role semantics in LLMs.

concepts. Rosenthal's idea is that color experiences not only represent the world, but also stand in patterns of similarity to one another: red experiences are more similar to orange experiences than to yellow ones. And these similarity patterns are isomorphic to patterns in what color experience represents (for example, wave length). The interpretability research surveyed above suggests that LLMs may satisfy this condition on rich color representation.

Finally, this work on color concepts is also relevant to our earlier question of whether LLM activations satisfy generalizations from folk psychology. When LLMs successfully make predictions about the relationships between colors or directions they haven't previously encountered, they seem to reason in ways that are similar to the way humans reason.

## 2.5 Teleosemantics

Another potential necessary condition on representation comes from teleosemantic theories of mental content. According to teleosemantic theories, the function of the underlying system generating a mental state determines its content (Millikan 1984). For example, whether a given mental state counts as a representation of a goldfinch depends on the function of the system that created that mental state. On standard teleosemantic views, the function of a state is determined by evolution: the state has whatever function explains why it was selected for via natural selection (Schulte 2023, 35).

Teleological conditions can be combined to causal and informational conditions on representation. For example, Stampe's (1977) causal account of content appealed to the idea of causal correlations that obtain in normal conditions. This can be fleshed out in terms of evolution: a perceptual state represents a feature of the world if that state being reliably caused by the feature helps explain why the organism was selected for in natural selection. Similarly, Neander (2013, 2017) argues that perceptual states have content when they have the function of being reliably produced by features of the environment. For example, toads have evolved perceptual states that fire in response to small, dark, moving objects. These states were selected for by natural selection, because toads with these states were more successful at hunting flies. For this reason, these perceptual states represent small, dark, moving objects. Similarly, informational accounts can be supplemented with the condition that various information-carrying channels were selected for (Dretske 1981).

We argue that teleological constraints on meaning are broadly compatible with AI representation. Although LLMs are not products of biological evolution, they undergo a form of artificial selection during their training process. This selection process optimizes the model's parameters—specifically, its weights—to improve its performance on specific tasks. In teleosemantic theories, the function of a mental state is determined by what it was selected for in the evolutionary process. Analogously, in LLMs, the function of a particular weight configuration is determined by its role in minimizing the loss function during training.

During training, the model's weights evolve in response to the pressure of minimizing loss, much like biological traits evolve under natural selection. This process bears important similarities to natural selection: just as biological traits that enhance fitness are more likely to persist, weight configurations that reduce prediction error are more likely to be retained. The resulting weight structure of the trained model can be seen as encoding the functions that the model was selected to perform, paralleling how evolved biological structures encode their functions.

The crucial disanalogy to natural selection is that in natural selection mutation is random, while in machine learning weights do not change randomly, but are instead adjusted in the direction of lower loss. But this disanalogy does not seem relevant to teleosemantic

theories. In both cases, an underlying optimization process creates a clear sense of a goal, and of normal versus abnormal pursuit of that goal.

At this point, we've surveyed several different potential requirements on mental representation. We have argued that results about AI interpretability suggest that large language models can meet each of these requirements. We think that when LLMs play games like Othello, for example, they don't merely predict the next word. Instead, they represent the game environment. In particular, as LLMs learn how to complete this task, they build internal models, representations that help them figure out what to do next, for example by representing board states. In this way, we think that interpretability research provides significant reason to reject skepticism about mental representation.

Before going on, it is worth flagging two points about the argument so far. First, not all of the empirical work on LLM representation points in the direction of success. We already mentioned Vafa et al.'s (2024)'s negative results about city maps, and Berglund et al.'s (2024)'s concerns about the reversal curse. As another example, Perez-Cruz and Shin (2024) and Chan (2025) observe that GPT-4o and similar quality LLMs could not successfully solve simple variants of Cheryl's birthday puzzle, a complex logic puzzle involving higher order reasoning about the beliefs of others. As another example, some researchers have recently suggested that LLMs may employ bags of heuristics rather than simple algorithms to complete tasks. For example, in the case of Othello-GPT, some researchers have suggested that in fact moves are determined from board positions using thousands of one-shot rules, rather than a few simple principles (Lin et al. 2024). This may suggest that the system's behavior is less psychological, because it doesn't involve a few rules for manipulating representations in pursuit of a goal. In general, we suggest that negative results about AI psychology may involve cases in which LLM representations, while in some sense present, are extremely complex and context-sensitive.

In assessing LLM cognitive abilities, it is important to not merely focus on successes, but also consider negative results as well. Ultimately, the goal of our paper is to clarify the terms of the debate, rather than to settle it. Second, the conditions discussed above can potentially vary in *degree* as well as *kind*. There may be no determinate fact of the matter, for example, about how well a system's information processing must approximate folk psychology to count as genuinely engaged in representation. After all, some of the conditions we discuss may be satisfied *to some degree* by Roombas. But we think the case for LLM representation, and ultimately LLM belief and desire, is significantly stronger than the case for Roombas, because LLMs satisfy the conditions we have explored to a higher degree than Roombas do.<sup>15</sup>

### 3. Responding to Skeptical Challenges

At this point in the paper, we've laid out our positive claims about LLM mental representation. In short, we think that there is a strong case that LLMs possess robust internal representations of the world.

---

15. For a recent philosophical treatment that synthesizes many of these interpretability findings see Beckmann and Queloz (2026), who propose a tiered framework distinguishing *conceptual understanding* (forming features as directions in latent space), *state-of-the-world understanding* (tracking contingent factual connections between features), and *principled understanding* (discovering compact circuits connecting facts rather than relying on memorization). Their framework maps naturally onto the conditions on representation we have surveyed: their first tier corresponds roughly to our informational and structural conditions, while their third tier corresponds to our condition on folk-psychological reasoning patterns. We lack the space in this paper to consider every teleosemantic theory on offer. For further recent work on teleosemantics, with some application to AI (Shea 2024).

In this section, we'll respond to three skeptical challenges to LLM folk psychology. Each skeptical challenge targets the claim that LLMs represent the world. (In the next section, we turn to challenges to LLMs possessing robust action dispositions.) In each case, we'll identify potential gaps in the skeptical challenge, and consider how the skeptical challenge is potentially relevant to questions about LLM folk psychology.

We'll consider three challenges:

- **Sensory grounding:** This challenge says that LLM inputs are meaningless because LLMs have no connection to the external world. In response, we'll argue that (i) the challenge relies on overly simplistic causal constraints on mental representation; (ii) the challenge ignores the possibility that LLMs form hypotheses about aspects of the world that are not directly observable to them; and (iii) the challenge is *fragile*, because it relies on properties of LLMs (such as the lack of perceptual inputs) that are not shared by all models.
- **Stochastic parrots:** This challenge says that LLMs do not represent the world because they aren't trained to do so. In response, we'll argue that (i) the challenge overgeneralizes to threaten human cognition; (ii) the challenge ignores that LLMs could represent the external world as a means to predicting the next word; (iii) the challenge ignores the possibility that representation is an *emergent capability* of LLMs; and (iv) the challenge is incompatible with the structure of internal representations identified by interpretability research.
- **Memorization:** This challenge claims that LLMs do not represent the world because their behavior can be explained through the alternate theory that they simply memorize text. In response, we'll argue that LLMs possess the ability to generalize robustly from their training data, and then make correct predictions about new questions.

We can think of each of these three challenges as targeting one of the conditions on representation we discussed earlier. The sensory grounding challenge claims that representation requires specific kinds of causal connections between the external world and internal LLM states. The memorization challenge claims that representation requires robust models and reasoning abilities, which LLMs supposedly lack. The stochastic parrots challenge claims that representation requires the right kind of functional origin, rooted in training objectives.

### 3.1 Sensory Grounding

One skeptical challenge concerns sensory grounding. This skeptic says that LLMs don't represent the world because they lack sensory grounding and merely use text: a system that is trained only on form [such as an LLM] would fail a sufficiently sensitive test [for intelligence], because it lacks the ability to connect its utterances to the world (Bender and Koller 2020). Pure (text-only) LLMs only see patterns of text. They do not have any outside sensory input. Therefore, according to this challenge, they can't break the syntactic circle and connect any of the symbols they see to the outside world. As Harnad (1990) defined it, the symbol grounding problem is one of how the semantic interpretation of symbols can be intrinsic to the system, rather than just parasitic on the meanings in our heads.

One version of the symbol grounding problem involves causal theories of mental representation. For some theories of mental representation, the right sort of causal connection between the outside world and internal states is required for such states to count as representational in the first place. According to these theories (including Stampe (1977) and Fodor (1987)), mental states represent the world when they are causally connected to the world in

the right way.<sup>16</sup> For example, my perceptual experiences of cats tend to be caused by cats, and therefore represent them. In addition, my desire to eat ice cream tends to cause me to eat ice cream, and therefore represents ice cream. Causal theories of representation can explain, for example, why a photograph of an identical twin represents one twin and not the other, despite resembling each twin perfectly.

Naive causal conditions may make trouble for LLM representation. LLMs have activations that are related to cats. But these activations are not directly caused by cats in the way that our perceptual system directly responds to cats in our environment. Instead, LLMs have learned about cats through training on text, and this text itself has been caused by cats.

There are two ways to address causal skepticism about LLM representation. The first option is to appeal to long causal chains. The second option is to appeal to pluralism about causal structure. Let's consider each in turn.

The first response appeals to long causal chains. In some sense, LLM activations related to cats are not caused directly by cats, in the way that retinal stimulation is directly caused by cats. Instead, LLM activations related to cats are caused by text about cats. But cat text is itself caused by cats. In this way, there is a causal chain from cats to LLM activations about cats. So it isn't clear that even naive causal conditions on representation block LLMs from genuine reference.

The second response appeals to pluralism about causal structure. In fact, any causal theory of representation needs to be pluralistic about the types of causal correlation that facilitates representation. Stampe (Stampe 1977) gives an example of barometers: the barometric representation is caused by a drop in air pressure, and the drop in air pressure causes the storm, and in this way the barometer represents the storm. While the barometric reading represents the storm, the storm is neither a cause of the barometric reading nor one of its effects. This is one of the ways in which informational theories of representation improve on naive causal theories: they can allow for a wide range of causal relations to produce genuine representation.

Another version of the sensory grounding challenge is illustrated by The Octopus Test, a thought experiment from Bender and Koller (2020). Imagine two humans communicating with one another through telegraphic cables across two islands, and a hyper-intelligent octopus that eavesdrops on their communication. In the beginning, the two humans mostly make small talk and describe their environments, which the octopus cannot see. However, because the octopus is hyperintelligent, it can pick up on various patterns in their communication. It has never seen trees, but it knows the patterns of discourse around the word tree and can do a good job of mimicking the other person. However, imagine that one of the people invents a catapult, and talks about it to their partner.

The skeptical challenge claims that the octopus cannot represent the catapult, or any other object that only exists on the islands rather than in the octopus's own environment. Because the catapult is novel to the discourse, the octopus won't be able to understand what it is or make predictions about how it works. And if this is the case, it suggests that LLMs cannot represent the world either.

---

16. The precise details of the causal theory differ with each particular adherent. Stampe proposed that a mental state represents the proposition that *p* iff under optimal conditions, it is causally correlated with *p*. Stampe's notion of optimal conditions was then spelled out in terms of biologically normal conditions that guarantee the well-functioning of the mechanisms that produce the mental state. Fodor defended a different, asymmetric dependence theory. For Fodor, a concept *C* represents an object *O* when the concept is lawfully causally correlated with that object, and any other correlations between *C* and other objects are asymmetrically explained by the connection between *C* and *O*. For example, my concept ZEBRA is caused by zebras, but can also be caused by cleverly painted mules. But when ZEBRA is caused by a painted mule, this itself is explained by the underlying causal connection between ZEBRA and zebras.

We ourselves don't have a strong judgment about what the octopus can represent in this thought experiment. But we think one fruitful way of resolving the question is to consider various candidate conditions on representation. First, it is clear that the octopus carries information about the catapult, since it can reliably predict who will say what about the catapult. But, second, it is unclear whether the octopus has internal representations of the catapult that causally influence its predictions. It is also unclear whether the octopus has a series of internal representations of the catapult whose internal structure or pattern of similarity mirrors patterns in the catapult. Moreover, it is unclear whether the octopus evolved to track the catapult through some kind of optimization process such as natural selection. For all of these reasons, we think it is quite unclear whether the octopus represents the catapult.

Finally, it is unclear whether the octopus can reason well about the catapult in a wide range of counterfactual scenarios, in the kinds of ways that would produce genuine world models. The claim that the octopus simply could not in principle represent the catapult in such a way ignores a crucial hypothesis. When an LLM (or octopus) receives text, if it is sufficiently capable, it can form hypotheses over what sort of processes generated such a text.

Just as scientists form hypotheses about microscopic phenomena responsible for observed patterns, advanced octopuses could make educated guesses about how the world works based on the patterns and types of text they receive. For example, if it reads lots of texts about machine learning, it might form certain hypotheses about certain types of researchers and computational infrastructure existing in the real world. From the under-water version of Wikipedia, it can make guesses about the sorts of processes and agents that would generate Wikipedia articles and what sort of underlying reality would result in such strings of text.

While direct empirical evidence for LLMs (and octopuses) forming explicit hypotheses about the world is limited, the possibility remains theoretically plausible. If advanced LLMs can form hypotheses about how the world works, then the symbols they receive could connect to the outside world, even if they do not have the capability of receiving anything other than text as input. Future research should aim to uncover more about the internal mechanisms of LLMs and their ability to form and use such hypotheses about the world.

Indeed, (Shai et al. 2024) provides empirical support for this kind of hypothesis formation in LLMs. As we discussed above, they demonstrated that when transformers are trained on sequences generated by hidden processes, they develop internal representations that capture the underlying generative structure, even when that structure is only indirectly observable through its outputs. The geometric arrangement of these representations in the model's residual stream matches theoretical predictions about optimal belief states, suggesting that LLMs can indeed form sophisticated hypotheses about hidden causal processes from observed patterns alone. This provides concrete evidence that LLMs can connect symbols to underlying reality even without direct sensory access to that reality.

If advanced LLMs can form such hypotheses about how the world works, then the symbols they receive could connect to the outside world, even if they do not have the capability of receiving anything other than text as input. The challenge for future research is to better understand how these internal mechanisms of hypothesis formation and representation operate across different domains and types of hidden structure.

Another problem for the sensory grounding challenge is its fragility. Even if skeptics can successfully argue that a pure LLM is unable to represent reality, their victory may only be a Pyrrhic one. The transformer architecture, which underlies these models, is designed for sequence prediction generally and is versatile enough to be applied to various modalities, including computer vision.

Vision-Language Models (VLMs), such as GPT-4o and Gemini, integrate both visual and textual data, processing them within a shared embedding space. Unlike systems that separately handle text and images and then attempt to combine their outputs, VLMs use a single, unified transformer model capable of processing and understanding both modalities simultaneously.

For instance, an image of a cat and the word cat would be represented in the same space, allowing the model to draw connections between the visual and textual representations. This shared embedding space helps bridge the gap between sensory input and linguistic representation. By incorporating visual data, VLMs gain a form of sensory grounding that pure text-based LLMs lack.

The architecture of VLMs is fundamentally similar to that of standard LLMs, relying on the same principles of self-attention and sequence modeling. This similarity undermines the argument that the transformer architecture itself is incapable of genuine representation. If VLMs, which are built on the same foundational architecture, can achieve sensory grounding and representational capabilities, it suggests that the perceived limitations of LLMs are not inherent to the architecture but rather a consequence of the input modalities used.

For all of these reasons, we do not consider the sensory grounding challenge a decisive threat to LLM representation.

### 3.2 Stochastic Parrots

Another skeptical challenge to LLM representation is associated with the stochastic parrots argument (Bender et al. 2021). This view contends that LLMs, despite their impressive outputs, are merely sophisticated statistical pattern matchers rather than systems capable of genuine understanding or representation. The core claim of the stochastic parrots argument is that LLMs are trained solely to predict the next word in a sequence, without any true comprehension of the content they generate. According to this view, LLMs don't represent or reason about the world; instead, they simply reproduce patterns from their training data in a statistically sophisticated but fundamentally meaningless way.

Proponents of this view often point to cases where LLMs produce fluent but nonsensical or contradictory outputs. For instance, an LLM might confidently assert a false statement or agree with contradictory premises in different conversations. These behaviors, they argue, reveal that LLMs lack genuine understanding and are merely parroting patterns from their training data.

One way to understand the stochastic parrot challenge is in terms of the teleosemantic conditions on representation we discussed earlier. Those conditions required that the internal states of LLMs have the function of representing features of the external world. This function would itself need to emerge from some kind of evolutionary, selective process.

Here, we'll lay out four challenges for the stochastic parrots view: The first challenge is overgeneralization. The argument that systems trained on pattern recognition can't develop genuine understanding potentially overgeneralizes. Human cognition, for instance, involves significant pattern recognition and statistical learning, yet we don't deny humans the capacity for genuine understanding. Similarly, other AI systems trained on pattern recognition (like computer vision models) are often considered to represent features of the world. The key question is not whether a system is trained on patterns, but whether it develops more sophisticated capabilities as a result of this training.

The second challenge is representation as a means to an end. While LLMs are indeed trained to predict the next word, representing the world could be an efficient means to this end. For example, to accurately predict words in a physics textbook, it would be helpful for



an LLM to develop some internal understanding of physics concepts. Our earlier discussions of probing studies and world models provide evidence that LLMs do develop structured internal representations that go beyond simple pattern matching (Herrmann and Levinstein 2024). Again, this can be unpacked in terms of teleosemantic requirements on representation. The idea would be that the training process of predicting the next word selects for the ability to track features of the external world.

The third challenge is emergent capabilities. Recent research suggests that LLMs can exhibit capabilities that were not explicitly part of their training objective. For instance: a) Few-shot learning: Brown et al. (2020) demonstrated that GPT-3 can perform new tasks with just a few examples, despite not being explicitly trained for this ability. This few-shot learning suggests a form of rapid adaptation that isn't easily explained by simple pattern matching. b) In-context learning: Wei et al. (2022) showed that LLMs can learn to perform new tasks from instructions and examples provided in the input prompt, without any change to their weights. This ability to learn within the context of a single forward pass challenges the notion of mere pattern reproduction. c) Chain-of-thought reasoning: Wei et al. (2022) also demonstrated that prompting LLMs to generate step-by-step reasoning significantly improved their performance on complex reasoning tasks, suggesting a capacity for structured thinking beyond simple pattern matching. d) Emergence of coding abilities: Chen et al. (2021) found that large language models trained on natural language can develop unexpected coding abilities, despite not being explicitly trained on programming tasks. e) Zero-shot task generalization: Kojima et al. (2022) showed that LLMs can solve novel tasks they weren't explicitly trained on when prompted to think step by step, demonstrating a form of reasoning capability.

The fourth challenge is internal representations. As we've discussed earlier in this paper, interpretability research provides evidence that LLMs develop rich internal representations of concepts and relationships. This structured internal knowledge is difficult to reconcile with the stochastic parrots view. It's worth noting that the stochastic parrots argument raises important questions about the nature of understanding and representation. Even if we reject the strong claim that LLMs are merely stochastic parrots, we might consider intermediate positions. For instance, LLMs might combine aspects of statistical pattern matching with more sophisticated representational capabilities. Ultimately, the stochastic parrots argument highlights the need for careful empirical investigation of LLM capabilities and limitations. While the behavior of LLMs can sometimes be consistent with sophisticated pattern matching, the evidence we've reviewed throughout this paper suggests that LLMs do develop meaningful internal representations and capabilities that go beyond simple parroting.

### 3.3 Memorization

Another skeptical challenge concerns memorization. According to this challenge, we don't need to posit genuine reasoning in LLMs, because we can explain their behavior in another way. Instead of reasoning about questions, LLMs simply memorize great quantities of data and then use shallow heuristics to generalize to new prompts. They are trained on billions of sentences, and in the course of this training, they simply store the answers to a large number of questions. These answers are then returned in response to prompts, without any genuine reasoning taking place.

Whether memorization threatens LLM representation depends on what is required for representation. If representation is simply a matter of carrying information, then memorization is perfectly compatible with genuine LLM representation. On the other hand, if

representation requires the presence of robust models or of reasoning about the world that matches folk psychology, then memorization may rule out representation.

In fact, there is a lot of evidence that LLMs do not merely memorize the answers to questions they are asked. Instead, LLMs seem to possess the ability to generalize robustly from their training data and then make correct predictions about new questions. They do not simply use shallow heuristics.

Return to the case of modular addition. As we discussed above, Nanda et al. (2023) found that during training, the LLM learns modular addition in multiple steps: an initial period of overfitting, based on memorization, followed by a transition to a general solution to the problem. In the initial phase of training, the LLM does engage in memorization. But after using memorization, the LLM learns a more general algorithm for computing the answer. After this algorithm is learned, the LLM then removes its memorization components. As evidence for this claim, Nanda et al. found that in the initial period of training, the model achieved 100% accuracy on the training data, but low accuracy on the testing data. After 10,000 epochs of training, the model learned how to actually perform the task, and achieved high accuracy on the testing data. Overall, this suggests that LLM skeptics are missing out on much of the rich structure of LLM reasoning.

Similar abilities to generalize were found in the Othello experiment. Every game of Othello starts with one of four moves. Each initial move creates a quadrant of Othello game space. To train Othello, the experiments only included training data from 3 of the 4 quadrants. But they found that the model was equally successful at playing Othello in the omitted quadrant of game space.<sup>17 18</sup>

#### 4. Action and Folk Psychology

While we've argued that LLMs have mental representations, the question remains: do they have a robust folk psychology? To address this, we need to consider whether LLMs have beliefs about the world, desires they aim to satisfy, and intentions that guide their actions. This question is complex, involving issues of stability, coherence, and goal-directed behavior.

To approach this question, we'll examine two influential philosophical perspectives on folk psychology: interpretationism and representationalism. Each offers a different framework for understanding mental states and presents different challenges when applied to LLMs.

---

17. Our focus in this paper is on whether LLMs have a folk psychology. This question is worth distinguishing from another skeptical target: the question of whether LLM *outputs* are meaningful. Here, the key question is whether text produced by LLMs has meaning. This is a different question than whether LLMs have a folk psychology. Compare: we can imagine a human being who has beliefs and desires, but who does not know how to speak French. If they started saying French sentences out loud phonetically, there would be an interesting question to ask about whether these sentences are meaningful in their mouth. But this question is a separate one from whether the speaker has a psychology. In practice, much of the debate about LLM outputs has itself been connected to debates about LLM folk psychology. For example, one skeptical challenge to the meaningfulness of LLM outputs starts from the premise that LLMs lack communicative intentions. This is itself a skeptical premise about LLM folk psychology. In response, (Mandelkern and Linzen 2024) have suggested that LLM psychology may not be necessary for LLM outputs to be meaningful, as long as LLMs are part of our linguistic community. (See also Mollo and Millière (2023). Others have suggested that we think of LLMs as more like *libraries* rather than speakers, which could also allow outputs to be meaningful (for discussion, see (Gopnik 2022)). By contrast, (Lederman and Mahowald 2024) argue that some kinds of LLM outputs can only be meaningful if LLMs are genuine speakers. Our focus in this paper is the question of whether LLMs have a psychology, rather than the question of whether LLM outputs are meaningful.

#### 4.1 Interpretationism

The most radical view is interpretationism. The idea behind interpretationism is that folk psychology is for explaining behavior. All that is required to have a folk psychology is for the system to behave in sufficiently complex ways best explained by appeal to folk psychological states. When this happens, the system has beliefs and desires: the system's desires are the goals promoted by its actions, and its beliefs are the views about the world that are required to be true in order for its actions to promote its goals. Whether it has internal states of a certain kind is not directly relevant to the question of folk psychology. Prominent interpretationists include Dennett (1981) and Davidson (1984).<sup>19</sup>

The previous conditions we've explored could allow for a separation between mental representation in general and belief/desire psychology in particular. LLMs could satisfy informational, causal, structural, and teleosemantic requirements on content, even if their behavior is too disorganized and happenstance to count as possessing beliefs and desires. Importantly, for interpretationists, folk psychological states come as a package deal. As (Stalnaker 1984) puts it:

Belief and desire ...are correlative dispositional states of a potentially rational agent. To desire that P is to be disposed to act in ways that would tend to bring it about that P in a world in which one's beliefs, whatever they are, were true. To believe that P is to be disposed to act in ways that would tend to satisfy one's desires, whatever they are, in a world in which P (together with one's other beliefs) were true. (p. 15)

---

18. The connection between belief and learning may offer another route to LLM skepticism. The philosopher Grace Helton has argued that in order to genuinely have beliefs, those beliefs must be able to change in response to evidence (Helton 2018). Helton's argument has two premises. First, you have a belief only if you are obligated to change that belief in response to strong counter-evidence. Second, you are obligated to do something only if you are able to do so. From these premises, it follows that LLMs have beliefs only if they are able to change their beliefs in response to strong counter-evidence. But one might worry that after training, LLMs can no longer learn or change their beliefs. If Helton's premises are correct, it follows that LLMs do not have beliefs.

The initial skeptical concern here is that after training, LLMs no longer learn or change their beliefs. The idea is that all of the learning that an LLM engages in occurs during training, when the weights of the LLM's neural network gradually change in response to feedback. But once a user interacts with the LLM, for example using ChatGPT, the weights are frozen, and so the LLM no longer learns.

This skeptical response fails, because of in-context learning. In-context learning refers to the ability of a language model to use the context provided within a given prompt or conversation to generate relevant and coherent responses (Brown et al. 2020). In general, in-context learning can refer to any way in which the model aptly adapts to the context of the prompt: e.g., adapting the right style, responding to corrections, or maintaining continuity over a conversation.

For our purposes, we can focus on few-shot learning, which is a form of in-context learning. In few-shot learning, you can provide examples within your prompt, and the model will use these to understand the type of response you're looking for. For example, if you want GPT-3 to predict nationalities, you might just ask it about Marie Curie's nationality. But with a few-shot prompt, you instead give it Einstein and Gandhi's nationality, and then ask it Marie Curie's nationality. Often, LLMs will perform better after seeing a few examples, rather than answering a question zero shot. This suggests that the LLM learns from the initial examples (Xie et al. 2021).

Xie et al. (2021) found that in-context learning in LLM satisfies many features of Bayesian inference. In particular, they hypothesized that in few-shot learning, there's a hidden concept that explains each of the examples in the prompt and that should be used to generate the response. They studied how a perfect Bayesian would guess the right response given this hypothesis and compared it to how GPT-2 performed. GPT-2 learned quickly, like a Bayesian, and got better and better with more examples.

Although LLMs of today do not retain the information learned in-context in other inference cycles, they do update their beliefs within a given inference cycle. Humans likewise often forget some of their beliefs—albeit not always—so we do not see a special reason to deny that LLMs can learn and change their minds.

Philosophers often distinguish dispositional from occurrent beliefs. This distinction is particularly interesting in the setting of LLMs. In LLMs, dispositional representations would be stored in the underlying weights of the neural net, while occurrent representations would be tokened by the activations in response to prompts. When LLMs engage in on-line learning, their occurrent representations change in response to evidence. In this way, Helton's argument could allow that LLMs possess occurrent beliefs, even if they lack dispositional beliefs. Rather, they would possess unchanging dispositional representations that do not respond to evidence, and so are not subject to rational obligations, even though they could guide action and inference.

19. Closely related to interpretationism is dispositionalism. On a dispositionalist picture, an agent has a belief that P if it is disposed to behave as if P. See, e.g., (Marcus 1990).

So, even if LLMs have something like mental representations of the world, they might not have desires or the right kinds of behavioral dispositions required for folk psychological states.

Do large language models satisfy interpretationist conditions on mental representation? There are at least two reasons for skepticism. The first challenge concerns LLM affordances: what actions an LLM can perform. Pure LLMs do not have access to robotic bodies. Instead, they simply produce text (or probability distributions over tokens). But at first glance these text interactions aren't naturally suited for performing complex actions.

There are at least two good responses to the problem of affordances. First, the problem of affordances is fragile. Some LLMs today do have access to robotic limbs. For example, Google's Palm-E system integrates an LLM with a robotic limb, which can perform actions after being prompted by a user with text (Driess et al. 2023), and see Gemini Robotics for newer models in this vein). When asked to pick up a bag of potato chips from the counter, the robotic limb can skillfully sort between different objects and pick up the chips.

Second, text production provides a rich enough space of alternative outcomes to count as a full-fledged action space. In today's digital world much human action takes place in text. When we imagine a human being whose life is confined entirely to text-based actions, we see no barrier to such a being possessing beliefs and desires. To see this point in greater detail, consider the use of LLMs in game environments. Mei et al. (2024) studied the behavior of LLMs in social cooperation games. These kinds of environments allow LLMs to formulate complex plans. They found that GPT-3 and GPT-4 would adjust their behavior throughout the game:

In games with multiple roles (such as the Ultimatum Game and the Trust Game), the AIs' decisions can be influenced by previous exposure to another role. For instance, if ChatGPT-3 has previously acted as the responder in the Ultimatum Game, it tends to propose a higher offer when it later plays as the proposer, while ChatGPT-4's proposal remains unchanged. Conversely, when ChatGPT-4 has previously been the proposer, it tends to request a smaller split as the responder. (p. 17)

This kind of behavior suggests that LLMs have the ability to use game theory to navigate environments involving other agents.

More recently, we have seen the rise of LLM-based *agents*—LLMs that plan and execute multi-step tasks with access to tools like code editors, web browsers, and file systems. This has provided striking new evidence relevant to the question of action dispositions. Coding agents, for example, can now accomplish complex software engineering tasks that require sustained goal-directed behavior over many steps: decomposing a high-level objective into subgoals, writing and debugging code across multiple files, running tests, and iterating when plans fail. For an interpretationist, this behavior seems naturally described in terms of beliefs (e.g., about the current state of the codebase), desires (to satisfy the user's specification), and intentions (the current plan of action). Indeed, coding agents will often make an explicit bullet-point plan and then follow it. The action space is far richer than the game-theoretic environments discussed above, and the plans are considerably longer-horizon.

However, several features of LLM agents complicate straightforward belief-desire attribution. First, agent goals are *externally specified*: the agent's objectives are set by a user's prompt rather than self-generated, making the interpretationist's attributed desires parasitic on explicit instructions. Second, current agents are *ephemeral*: each session begins from a blank slate with no persistent memory, so there is no stable, enduring belief set that accumulates across interactions in the way we would expect of a genuine believer. Third, as Mil-ière (2025) has argued, the behavioral dispositions produced by current alignment methods

(RLHF and related techniques) are *shallow*—they create prompt-sensitive behavioral tendencies rather than genuine normative deliberation. Agent behavior is shaped simultaneously by pretraining, fine-tuning, system prompts, and user instructions, with no clear unified goal hierarchy. If an agent’s apparent goals are better described as a patchwork of competing dispositional sources than as a coherent desire set, the interpretationist case is weakened. These considerations do not settle the question—the impressive multi-step competence of current agents is hard to dismiss—but they suggest that the question of LLM action dispositions remains genuinely open.

A second reason for skepticism about action is instability. Shanahan and others have noted that LLM outputs are very sensitive to prompting. If you slightly change the way you ask a question, the LLM can start to behave very differently. This makes it hard to see the LLM as taking a wide range of means to promote a unified goal (Shanahan et al. 2023).

We think instability is potentially the most serious threat to LLM folk psychology. We see two potential responses worthy of further development. First, one might concede that LLMs are relatively unstable but say that each prompting session with an LLM produces a different agent with its own beliefs and desires. For example, if you ask an LLM to write poetry for you on one day and start a new conversation asking it to play chess with you the second day, you in effect have two separate agents. The chess-playing instance has no interest in poetry, but it does have the goal of beating you at chess. When you give it goals like writing poetry or playing chess or coding, it does a good job. Those tasks are very difficult. So, the thought goes, the best explanation is that the beliefs and desires of the LLM change from inference cycle to inference cycle, but the behavior of particular instances is still fruitfully explained with beliefs and desires and intentions.

Second, one might reject the claim of instability. Although the behavior looks unstable, the different behavior of LLMs in different prompting sessions might derive from a stable underlying goal, such as pleasing the user. The different outputs of LLMs in different sessions might result from a disconnect between what the model believes and what the model says. In other words, instability in model outputs is evidence of lying, not evidence that the system lacks beliefs.<sup>20</sup>

Above all, we suggest that more research needs to be done about how stable LLM outputs are to a range of different prompting environments. For example, one fruitful project would be to explore how LLM strategies in game environments change as a result of different prompting conditions. This would allow some measure of whether it forms coherent plans that are robust to a range of perturbations.

## 4.2 Representationalism

Representationalism, advocated by philosophers like Jerry Fodor and Fred Dretske, holds that having mental states requires having internal representations with appropriate functional roles. On this view, to have a belief that *p*, a system must have an internal state that represents *p* and plays the right causal role in the system’s cognitive economy.

There is cause for optimism on the representationalist picture. As we’ve argued at length, some LLM mental states can have truth-conditions. Furthermore, LLMs even appear to have some sorts of world models and structured representations of conceptual domains.<sup>21</sup>

---

20. For an argument that RLHF helps ground textual meaning through the goal of pleasing the user (Mollo and Millière 2023).

However, even on representationalist pictures like Fodor's, beliefs have to play the right role in the larger system and generally cannot be divorced entirely from desires. Even if LLMs have rich internal representations, it's not clear that these play the right kind of role in generating behavior to count as beliefs or desires. The issue of instability resurfaces here—if internal representations don't stably guide behavior across different prompts, can they really count as beliefs?

Ultimately, we think matters are easier for the representationalist than the interpretationist. While we have some behavioral evidence in the case of LLMs, behavioral evidence is much more limited for LLMs than it is for humans. However, we have perfect internal access to LLMs, and we have much to discover about the role various representations play in LLMs' cognition. Therefore, as we come to understand more about how LLMs think, it will become more obvious for representationalists whether they have folk psychological mental states. Some key open questions for attributing folk psychological states on either picture, then, include:

- **Action and planning:** How can we best understand LLM action given their limited affordances?
- **Stability:** How can we reconcile the apparent instability of LLM outputs with the need for stable beliefs and desires?
- **Goal-directedness:** Do LLMs have anything analogous to enduring goals or values?

Addressing these questions will require a combination of philosophical analysis and empirical investigation. While there's evidence that LLMs have sophisticated internal representations and can exhibit complex, apparently goal-directed behavior, significant questions remain about whether they possess full-fledged folk psychological states. Resolving these questions will be crucial for understanding the capabilities, limitations, and potential moral status of these increasingly ubiquitous AI systems.

## References

- Abdou, M., Kulmizev, A., Hershcovich, D., Frank, S., Pavlick, E., & Søgaard, A. (2021). *Can language models encode perceptual structure without grounding? A case study in color*. <https://arxiv.org/abs/2109.06129>
- Adams, F., & Aizawa, K. (2021). Causal theories of mental content. In E. N. Zalta (Ed.), *The Stanford encyclopedia of philosophy* (Fall 2021). Metaphysics Research Lab, Stanford University. <https://plato.stanford.edu/archives/fall2021/entries/content-causal/>
- Alain, G., & Bengio, Y. (2018). *Understanding intermediate layers using linear classifier probes*. <https://arxiv.org/abs/1610.01644>
- Ameisen, E., Lindsey, J., Pearce, A., Gurnee, W., Turner, N. L., Chen, B., Citro, C., Abrahams, D., Carter, S., Hosmer, B., Marcus, J., Sklar, M., Templeton, A., Bricken, T., McDougall, C., Cunningham, H., Henighan, T., Jermyn, A., Jones, A., ... Batson, J. (2025). Circuit tracing: Revealing computational graphs in language models.

---

21. We might require that LLMs internally distinguish between true claims and false claims in a way that goes beyond representation for the LLM to count as having beliefs. In particular, we might require that they somehow systematically tag sentences as true or false and use this tag in their master algorithm to determine what text to output? We don't yet have a full answer to this question. Some (Burns et al. 2023), (Azaria and Mitchell 2023) have argued that LLMs do make an internal distinction between truth and falsity. However, others, e.g., (Levinstein and Herrmann 2024) claim these studies are flawed. (Herrmann and Levinstein 2024) argue for a representationalist account of belief for LLMs but maintain that current empirical evidence of whether LLMs actually have beliefs is inconclusive.

- Transformer Circuits Thread*. <https://transformer-circuits.pub/2025/attribution-graphs/methods.html>
- Andrews, K. (2013). Reading people or reading minds: A precis of do apes read minds? Toward a new folk psychology. *CAPE Studies in Applied Philosophy and Ethics Series*, 1, 140–151.
- Azaria, A., & Mitchell, T. (2023). *The internal state of an LLM knows when it's lying*. <https://arxiv.org/abs/2304.13734>
- Beckmann, P., & Queloz, M. (2026). Mechanistic indicators of understanding in large language models. *Philosophical Studies*. <https://doi.org/10.1007/s11098-026-02513-1>
- Belrose, N., Schneider-Joseph, D., Ravfogel, S., Cotterell, R., Raff, E., & Biderman, S. (2023). *LEACE: Perfect linear concept erasure in closed form*. <https://arxiv.org/abs/2306.03819>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? . *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On meaning, form, and understanding in the age of data. In D. Jurafsky, J. Chai, N. Schluter, & J. Tetreault (Eds.), *Proceedings of the 58th annual meeting of the association for computational linguistics* (pp. 5185–5198). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.463>
- Berglund, L., Tong, M., Kaufmann, M., Balesni, M., Stickland, A. C., Korbak, T., & Evans, O. (2024). The reversal curse: LLMs trained on “A is B” fail to learn “B is A.” *The Twelfth International Conference on Learning Representations (ICLR)*. <https://openreview.net/forum?id=GPKTIktAok>
- Boghossian, P. A. (1990). The status of content. *Philosophical Review*, 99(2), 157–184. <http://doi.org/10.2307/2185488>
- Bricken, T., Templeton, A., Batson, J., Chen, B., Jermyn, A., Conerly, T., Turner, N., Anil, C., Denison, C., Askell, A., Lasenby, R., Wu, Y., Kravec, S., Schiefer, N., Maxwell, T., Joseph, N., Hatfield-Dodds, Z., Tamkin, A., Nguyen, K., ... Olah, C. (2023). *Towards monosemanticity: Decomposing language models with dictionary learning*. Transformer Circuits Thread. <https://transformer-circuits.pub/2023/monosemantic-features>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.
- Burns, C., Ye, H., Klein, D., & Steinhardt, J. (2023). Discovering latent knowledge in language models without supervision. *The Eleventh International Conference on Learning Representations (ICLR)*. <https://openreview.net/forum?id=ETKGubyohcs>
- Cappelen, H., & Dever, J. (2025). *Going whole hog: A philosophical defense of AI cognition*. <https://arxiv.org/abs/2504.13988>
- Chaigneau, S. E., Barsalou, L. W., & Sloman, S. A. (2004). Assessing the causal structure of function. *Journal of Experimental Psychology: General*, 133(4), 601–625.
- Chan, E. (2025). Evaluating logical reasoning ability of large language models. *Preprints*. <https://doi.org/10.20944/preprints202504.1933.v1>

- Chen, M., Tworek, J., Jun, H., Yuan, Q., Pinto, H. P. de O., Kaplan, J., Edwards, H., Burda, Y., Joseph, N., Brockman, G., Ray, A., Puri, R., Krueger, G., Petrov, M., Khlaaf, H., Sastry, G., Mishkin, P., Chan, B., Gray, S., ... Zaremba, W. (2021). *Evaluating large language models trained on code*. <https://arxiv.org/abs/2107.03374>
- Christiano, P. F., Leike, J., Brown, T. B., Martic, M., Legg, S., & Amodei, D. (2017). Deep reinforcement learning from human preferences. *Advances in Neural Information Processing Systems*, 30. <https://proceedings.neurips.cc/paper/2017/hash/d5e2coadad503c91f91df240dcd4e49-Abstract.html>
- Da, J., & Kasai, J. (2019). Cracking the contextual commonsense code: Understanding commonsense reasoning aptitude of deep contextual representations. *arXiv Preprint arXiv:1910.01157*. <https://arxiv.org/abs/1910.01157>
- Davidson, D. (1984). *Inquiries into truth and interpretation*. Oxford University Press.
- Dennett, D. C. (1981). *The intentional stance*. MIT Press.
- Dretske, F. I. (1981). *Knowledge and the flow of information*. MIT Press.
- Driess, D., Xia, F., Sajjadi, M. S. M., Lynch, C., Chowdhery, A., Ichter, B., Wahid, A., Tompson, J., Vuong, Q., Yu, T., Huang, W., Chebotar, Y., Sermanet, P., Duckworth, D., Levine, S., Vanhoucke, V., Hausman, K., Toussaint, M., Greff, K., ... Florence, P. (2023). *PaLM-E: An embodied multimodal language model*. <https://arxiv.org/abs/2303.03378>
- Ettinger, A. (2020). What BERT is not: Lessons from a new suite of psycholinguistic diagnostics for language models. *Transactions of the Association for Computational Linguistics*, 8, 34–48. [https://doi.org/10.1162/tacl\\_a\\_00298](https://doi.org/10.1162/tacl_a_00298)
- Fodor, J. A. (1975). *The language of thought*. Harvard University Press.
- Fodor, J. A. (1987). *Psychosemantics: The problem of meaning in the philosophy of mind*. MIT Press.
- Forbes, M., Holtzman, A., & Choi, Y. (2019). Do neural language representations learn physical commonsense? *arXiv Preprint arXiv:1908.02899*. <https://arxiv.org/abs/1908.02899>
- Giulianelli, M., Harding, J., Mohnert, F., Hupkes, D., & Zuidema, W. (2021). *Under the hood: Using diagnostic classifiers to investigate and improve how language models track agreement information*. <https://arxiv.org/abs/1808.08079>
- Gopnik, A. (2022). Children, creativity, and the real key to intelligence. *APS Observer*, 35.
- Graham, G., Horgan, T. E., & Tienson, J. L. (2007). Consciousness and intentionality. In M. Velmans & S. Schneider (Eds.), *The blackwell companion to consciousness* (pp. 468–484). Blackwell.
- Gurnee, W., & Tegmark, M. (2024). Language models represent space and time. *The Twelfth International Conference on Learning Representations (ICLR)*. <https://openreview.net/forum?id=jE8xbmvFin>
- Harding, J. (forthcoming). Operationalising representation in natural language processing. *British Journal for the Philosophy of Science*. <https://doi.org/10.1086/728685>
- Harnad, S. (1990). The symbol grounding problem. *Physica D: Nonlinear Phenomena*, 42(1–3), 335–346.
- Heathwood, C. (2016). Desire-fulfillment theory. In G. Fletcher (Ed.), *The routledge handbook of the philosophy of well-being* (pp. 135–147). Routledge.
- Helton, G. (2018). If you can't change what you believe, you don't believe it. *Noûs*, 54(3), 501–526. <https://doi.org/10.1111/nous.12265>
- Herrmann, D. A., & Levinstein, B. A. (2024). *Standards for belief representations in LLMs*. <https://arxiv.org/abs/2405.21030>



- Hutto, D., & Ravenscroft, I. (2021). Folk psychology as a theory. In E. N. Zalta (Ed.), *The Stanford encyclopedia of philosophy* (Fall 2021). Metaphysics Research Lab, Stanford University. <https://plato.stanford.edu/archives/fall2021/entries/folkpsych-theory/>
- Karvonen, A., Chua, J., Dumas, C., Fraser-Taliente, K., Kantamneni, S., Minder, J., Ong, E., Sharma, A. S., Wen, D., Evans, O., & Marks, S. (2026). *Activation oracles: Training and evaluating LLMs as general-purpose activation explainers*. <https://arxiv.org/abs/2512.15674>
- Kassner, N., & Schütze, H. (2020). *Negated and misprimed probes for pretrained language models: Birds can talk, but cannot fly*. <https://arxiv.org/abs/1911.03343>
- Kojima, T., Gu, S. S., Reid, M., Matsuo, Y., & Iwasawa, Y. (2022). Large language models are zero-shot reasoners. *Advances in Neural Information Processing Systems*, 35, 22199–22213.
- Lake, B. M., & Baroni, M. (2023). Human-like systematic generalization through a meta-learning neural network. *Nature*, 623(7985), 115–121.
- Lederman, H., & Mahowald, K. (2024). Are language models more like libraries or like librarians? Bibliotechnism, the novel reference problem, and the attitudes of LLMs. *Transactions of the Association for Computational Linguistics*, 12, 1087–1103. [https://doi.org/10.1162/tacl\\_a\\_00690](https://doi.org/10.1162/tacl_a_00690)
- Levinstein, B. A., & Herrmann, D. A. (2024). Still no lie detector for language models: Probing empirical and conceptual roadblocks. *Philosophical Studies*, 182(7), 1539–1565. <https://doi.org/10.1007/s11098-023-02094-3>
- Li, K., Hopkins, A. K., Bau, D., Viégas, F., Pfister, H., & Wattenberg, M. (2023). Emergent world representations: Exploring a sequence model trained on a synthetic task. *The Eleventh International Conference on Learning Representations (ICLR)*. [https://openreview.net/forum?id=DeGo7\\_TcZvT](https://openreview.net/forum?id=DeGo7_TcZvT)
- Lin, J., Schonbrun, J., Karvonen, A., & Rager, C. (2024). *OthelloGPT learned a bag of heuristics*. LessWrong / AI Alignment Forum. <https://www.lesswrong.com/posts/gcpNuEZnxAPayaKBY/othellogpt-learned-a-bag-of-heuristics-1>
- Lindsey, J., Gurnee, W., Ameisen, E., Chen, B., Pearce, A., Turner, N. L., Citro, C., Abrahams, D., Carter, S., Hosmer, B., Marcus, J., Sklar, M., Templeton, A., Bricken, T., McDougall, C., Cunningham, H., Henighan, T., Jermyn, A., Jones, A., ... Batson, J. (2025). On the biology of a large language model. *Transformer Circuits Thread*. <https://transformer-circuits.pub/2025/attribution-graphs/biology.html>
- Mandelkern, M., & Linzen, T. (2024). Do language models' words refer? *Computational Linguistics*, 50(3), 1191–1200. [https://doi.org/10.1162/coli\\_a\\_00522](https://doi.org/10.1162/coli_a_00522)
- Marcus, R. B. (1990). Some revisionary proposals about belief and believing. *Philosophy and Phenomenological Research*, 50, 133–153. <https://doi.org/10.2307/2108036>
- Mei, Q., Xie, Y., Yuan, W., & Jackson, M. O. (2024). A Turing test of whether AI chatbots are behaviorally similar to humans. *Proceedings of the National Academy of Sciences*, 121(9), e2313925121.
- Meng, K., Bau, D., Andonian, A., & Belinkov, Y. (2022). Locating and editing factual associations in GPT. *Advances in Neural Information Processing Systems*, 35. [https://proceedings.neurips.cc/paper\\_files/paper/2022/hash/6f1d43d5a82a37e89b0665b33bf3a182-Abstract-Conference.html](https://proceedings.neurips.cc/paper_files/paper/2022/hash/6f1d43d5a82a37e89b0665b33bf3a182-Abstract-Conference.html)
- Millière, R. (2025). Normative conflicts and shallow AI alignment. *Philosophical Studies*, 182(7), 2035–2078. <https://doi.org/10.1007/s11098-025-02347-3>
- Millière, R., & Buckner, C. (2024a). *A philosophical introduction to language models – Part I: Continuity with classic debates*. <https://arxiv.org/abs/2401.03910>
- Millière, R., & Buckner, C. (2024b). *A philosophical introduction to language models – Part II: The way forward*. <https://arxiv.org/abs/2405.03207>

- Millikan, R. G. (1984). *Language, thought, and other biological categories: New foundations for realism*. MIT Press.
- Mollo, D. C., & Millière, R. (2023). The vector grounding problem. *Philosophy and the Mind Sciences*, 7. <https://philosophymindscience.org/index.php/phimisci/article/view/12307>
- Musker, S., & Pavlick, E. (2023). Testing causal models of word meaning in GPT-3 and GPT-4. *arXiv Preprint arXiv:2305.14630*. <https://arxiv.org/abs/2305.14630>
- Nanda, N., Chan, L., Lieberum, T., Smith, J., & Steinhardt, J. (2023). Progress measures for grokking via mechanistic interpretability. *arXiv Preprint arXiv:2301.05217*. <https://arxiv.org/abs/2301.05217>
- Neander, K. (2013). Toward an informational teleosemantics. In D. Ryder, J. Kingsbury, & K. Williford (Eds.), *Millikan and her critics* (pp. 21–40). Wiley.
- Neander, K. (2017). *A mark of the mental: A defence of informational teleosemantics*. MIT Press.
- Opie, J., & O'Brien, G. (2004). Notes toward a structuralist theory of mental representation. In H. Clapin (Ed.), *Representation in mind: New approaches to mental representation* (pp. 1–20). Elsevier.
- Patel, R., & Pavlick, E. (2022). Mapping language models to grounded conceptual spaces. *International Conference on Learning Representations (ICLR)*. <https://openreview.net/forum?id=gJcEM8sxHK>
- Perez-Cruz, F., & Shin, H. S. (2024). *Testing the cognitive limits of large language models*. Bank for International Settlements.
- Petroni, F., Rocktäschel, T., Lewis, P., Bakhtin, A., Wu, Y., Miller, A. H., & Riedel, S. (2019). Language models as knowledge bases? *arXiv Preprint arXiv:1909.01066*. <https://arxiv.org/abs/1909.01066>
- Piantadosi, S. T., & Hill, F. (2022). *Meaning without reference in large language models*. <https://arxiv.org/abs/2208.02957>
- Ravfogel, S., Elazar, Y., Gonen, H., Twiton, M., & Goldberg, Y. (2020). *Null it out: Guarding protected attributes by iterative nullspace projection*. <https://arxiv.org/abs/2004.07667>
- Rosenthal, D. M. (2005). *Consciousness and mind*. Oxford University Press.
- Russell, S. (2019). *Human compatible: AI and the problem of control*. Penguin.
- Schulte, P. (2023). *Mental content*. Cambridge University Press.
- Shai, A. S., Marzen, S. E., Teixeira, L., Oldenziel, A. G., & Riechers, P. M. (2024). *Transformers represent belief state geometry in their residual stream*. <https://arxiv.org/abs/2405.15943>
- Shanahan, M., McDonell, K., & Reynolds, L. (2023). Role play with large language models. *Nature*, 623(7987), 493–498.
- Shea, N. (2024). *Concepts at the interface*. Oxford University Press.
- Stalnaker, R. C. (1984). *Inquiry*. Cambridge University Press.
- Stampe, D. W. (1977). Towards a causal theory of linguistic representation. *Midwest Studies in Philosophy*, 2(1), 42–63. <https://doi.org/10.1111/j.1475-4975.1977.tb00027.x>
- Templeton, A., Conerly, T., Marcus, J., Lindsey, J., Bricken, T., Chen, B., Pearce, A., Citro, C., Ameisen, E., Jones, A., Cunningham, H., Turner, N. L., McDougall, C., MacDiarmid, M., Tamkin, A., Durmus, E., Hume, T., Mosconi, F., Freeman, C. D., ... Henighan, T. (2024). *Scaling monosemanticity: Extracting interpretable features from Claude 3 Sonnet*. Transformer Circuits Thread. <https://transformer-circuits.pub/2024/scaling-monosemanticity/>
- Usher, M. (2001). A statistical referential theory of content: Using information theory to account for misrepresentation. *Mind and Language*, 16(3), 331–334. <https://doi.org/10.1111/1468-0017.00172>

- Vafa, K., Chen, J. Y., Rambachan, A., Kleinberg, J., & Mullainathan, S. (2024). Evaluating the world model implicit in a generative model. *Advances in Neural Information Processing Systems*, 37. <https://openreview.net/forum?id=aVK4JFpegY>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008. <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd0531c4a845aa-Abstract.html>
- Vylomova, E., Rimell, L., Cohn, T., & Baldwin, T. (2016). Take and took, gaggle and goose, book and read: Evaluating the utility of vector differences for lexical relation learning. In K. Erk & N. A. Smith (Eds.), *Proceedings of the 54th annual meeting of the association for computational linguistics (volume 1: Long papers)* (pp. 1671–1682). Association for Computational Linguistics. <https://doi.org/10.18653/v1/P16-1158>
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Xia, F., Chi, E. H., Le, Q. V., & Zhou, D. (2022). Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35, 24824–24837.
- Xie, S. M., Raghunathan, A., Liang, P., & Ma, T. (2021). An explanation of in-context learning as implicit Bayesian inference. *arXiv Preprint arXiv:2111.02080*. <https://arxiv.org/abs/2111.02080>
- Yildirim, I., & Paul, L. A. (2024). From task structures to world models: What do LLMs know? *Trends in Cognitive Sciences*. <https://doi.org/10.1016/j.tics.2024.02.008>
- Zhong, Z., Liu, Z., Tegmark, M., & Andreas, J. (2023). *The clock and the pizza: Two stories in mechanistic explanation of neural networks*. <https://arxiv.org/abs/2306.17844>