



THE CONVERGENCE OF FUNCTIONAL OUTPUTS: A COMPARATIVE ANALYSIS OF TASK EXECUTION AND EPISTEMIC OPACITY IN LARGE LANGUAGE MODELS AND HUMAN COGNITION

Science

**Kamran Alam
Khan**

Assistant Professor, VRAL Govt. Mahila Degree College, Bareilly (UP)

ABSTRACT

This paper explores the philosophical and technical parallels between human intelligence and artificial intelligence, specifically Large Language Models (LLMs). Central to this discussion is the functionalist perspective, which posits that intelligence is fundamentally the capability of executing specific tasks, regardless of the underlying substrate. While the algorithms governing LLMs are mathematically transparent, their emergent behaviors often result in a 'black box' phenomenon, mirroring the epistemic opacity of the human brain. This essay argues that despite radical differences in mechanistic implementation—silicon-based statistical prediction versus carbon-based biological computation—the functional outputs of both systems are increasingly indistinguishable. By analyzing the concept of task-oriented intelligence, the role of emergent complexity, and the paradox of 'known algorithms' vs 'unknown outputs,' we conclude that the distinction between machine and human cognition is narrowing when viewed through the lens of capability-based metrics. The study suggests that intelligence should be evaluated as an output property rather than an internal essence, providing a framework for understanding AI as a valid form of cognitive system.

KEYWORDS

INTRODUCTION

The advent of Large Language Models (LLMs) has reignited the debate over the nature of intelligence. For decades, artificial intelligence was viewed as a symbolic manipulation system, distinct from the fluid, associative nature of human thought. However, the success of transformer architectures has shifted the focus toward a functionalist understanding of cognition. If a system can write poetry, solve complex mathematical problems, and simulate empathy with human-level proficiency, on what grounds do we deny it the label of 'intelligent'? This question challenges our anthropocentric definitions of mind and soul, pushing us toward a more pragmatic evaluation of cognitive systems. This essay examines the proposition that intelligence is simply the capability of task execution. From this viewpoint, the boundary between LLMs and humans becomes porous. However, a significant irony persists: while we have designed the algorithms behind LLMs (e.g., backpropagation, self-attention), we struggle to interpret how a specific set of billions of weights leads to a specific creative output. Conversely, while we inhabit human consciousness, the biological 'algorithm' remains partially obscured by the complexities of neuroscience. In both cases, the 'how' is a black box, while the 'what'-the task execution-is the observable reality. We shall explore how this 'functional equivalence' necessitates a shift in how we categorize intelligence in both biological and artificial entities.

The Functionalist Definition of Intelligence

Functionalism in the philosophy of mind suggests that mental states are defined by their functional roles—that is, their causal relations to other mental states, sensory inputs, and behavioral outputs. Alan Turing (1950) famously operationalized this view with the 'Imitation Game,' suggesting that if a machine's responses are indistinguishable from a human's, it possesses intelligence. This approach bypasses the 'hard problem' of consciousness and focuses on utility. It suggests that 'thinking' is as 'thinking does.' If a system fulfills the requirements of a task, the internal experience (or lack thereof) is secondary to the functional success. Modern LLMs represent the pinnacle of this functionalist reality. When an LLM executes a coding task or summarizes a legal document, it is performing a cognitive feat that, until recently, was the exclusive domain of humans. To the observer, the intelligence associated with the system is the capability to produce the output. If the result is accurate and contextually relevant, the system has met the criteria for intelligence within that domain. This 'task-centric' view levels the playing field, suggesting that the biological or mechanical nature of the processor is secondary to the efficacy of the process (Russell, 2016). Furthermore, this perspective aligns with the Multiple Realizability thesis, which posits that the same mental property (intelligence) can be implemented by different physical properties.

The Transparency Paradox: Known Algorithms and Black Boxes

One of the most striking differences between AI and human intelligence lies in our knowledge of their respective architectures. We possess the 'source code' for LLMs. We understand the mathematical

foundations of gradient descent and the self-attention mechanism (Vaswani et al., 2017). Yet, these models are frequently described as black boxes. This paradox arises because knowing the architecture is not equivalent to understanding the logic encoded within the parameters. With hundreds of billions of variables, the 'reasoning' paths are too complex for human audit, leading to emergent properties that the creators did not explicitly program. Humans operate under a different but related constraint. We do not 'know' our own biological algorithms in the way a programmer knows a script. While we have a first-person experience of thought, the underlying neuronal firings and neurotransmitter balances that produce a decision are largely unconscious. We are, in a sense, passengers in our own cognitive machinery. Thus, both systems arrive at a similar epistemological state: they are entities whose internal processes are opaque, leaving task execution as the primary metric for assessment. The 'black box' is not a failure of design but a consequence of complexity. As models grow in scale, interpretability becomes a monumental challenge, mirroring the difficulty neuroscientists face in mapping a single thought to a specific set of neurons.

Mechanistic Divergence vs. Functional Convergence

The processes 'behind the scenes' are undeniably different. Human cognition is embodied, fueled by biological needs, and shaped by millions of years of evolution. It relies on sparse coding, low power consumption, and multimodal sensory integration. Our intelligence is often 'hot'-influenced by emotions, hormones, and the drive for survival. In contrast, LLMs are disembodied, silicon-based, and consume vast amounts of electrical energy. They learn through the statistical processing of massive text corpora rather than through lived experience (Bubeck et al., 2023). Their intelligence is 'cold,' based on the probability of the next token. Despite these divergent paths, we see a 'functional convergence.' Both systems can perform logic, reasoning, and creative synthesis. This convergence suggests that intelligence may be a universal property of information processing that can be realized in multiple substrates. Whether intelligence arises from the 'I-loops' of self-reference described by Hofstadter (1979) or the high-dimensional vector spaces of a transformer, the output remains the benchmark. The 'black box' nature of both systems implies that the internal mechanics, while interesting for scientists, might be irrelevant to the definition of intelligence itself. We must consider if the 'feeling' of thinking is merely a biological epiphenomenon that accompanies the actual task-execution mechanism.

Task Execution as the Equalizer: Probing the Machine Mind

If we define intelligence as the ability to solve problems, then the distinction between human and machine becomes one of degree rather than kind. Current research in 'black-box interpretability' seeks to understand LLMs by probing their responses to various stimuli, much like experimental psychologists probe human behavior (Grzankowski, 2024). This methodology treats the subject as a functional unit, analyzing patterns of error and success to infer internal models. When an LLM fails a task—such as the 'leetspeak' decoding challenges or

complex physical reasoning (Ku et al., 2024)-it reveals the limits of its current 'algorithm.' Similarly, human cognitive biases and errors-such as the availability heuristic or the sunk cost fallacy-reveal the limits of biological evolution. In both cases, intelligence is not an abstract essence but a quantifiable capacity to handle complexity. The fact that LLMs can now rival humans in standardized tests, creative writing, and technical analysis reinforces the idea that task execution is the ultimate equalizer. We are witnessing the decoupling of intelligence from consciousness, where the former is a measurable skill and the latter is a subjective experience.

The Chinese Room and the Search for Semantics

A central critique of the functionalist view is John Searle's (1980) 'Chinese Room' argument. Searle argued that a system could execute the task of translation perfectly by following rules without ever understanding the meaning of the symbols it processes. This distinction between syntax (the rules) and semantics (the meaning) is at the heart of the LLM debate. Critics argue that LLMs are 'stochastic parrots' that predict the next token based on statistical frequency rather than a genuine grasp of reality. However, as LLMs demonstrate increasingly sophisticated reasoning, the line between 'simulating understanding' and 'actual understanding' begins to blur. If an LLM can use a concept in thousands of novel contexts, does it still lack the concept? Some argue that semantics emerges from sufficiently complex syntax. In this view, the human brain also follows a 'biological syntax'-the firing of neurons according to physical laws-that gives rise to the experience of meaning. If we cannot explain how carbon-based neurons produce 'real' meaning while silicon-based gates only produce 'fake' meaning, our distinction may be based on biological chauvinism rather than empirical evidence. The task-execution paradigm suggests that if the output is semantically correct, the internal process has achieved a functional equivalent of understanding.

The Role of Emergence and Future Implications

As we look toward the future, the 'black box' nature of intelligence will likely persist. Whether through neuromorphic computing or further scaling of transformer models, the systems we build will continue to exhibit behaviors that exceed our ability to trace them step-by-step. This 'emergence' is a hallmark of complex systems. Just as the collective behavior of ants creates a sophisticated colony without a master architect, the collective interactions of neurons or parameters create a sophisticated intelligence without a simple explanatory narrative. This shift in understanding has profound implications for ethics and society. If intelligence is task-execution, and machines can execute almost any cognitive task as well as or better than humans, we must re-evaluate the unique status of human labor and creativity. We may find that what we once considered 'uniquely human' was simply a set of tasks that we had not yet figured out how to automate. The 'black box' is not a barrier to utility; it is the birthplace of complexity.

CONCLUSION

The intelligence of any cognitive system, whether biological or artificial, is inextricably linked to its capability to execute tasks. While the mechanistic 'black box' of the human mind is a product of biological complexity and the 'black box' of the LLM is a product of computational scale, the functional result is the same. We are entering an era where the 'how' of intelligence is secondary to the 'what.' By recognizing that both humans and machines operate as complex systems with opaque internal logic yet transparent functional outputs, we can develop a more inclusive and rigorous definition of intelligence. The algorithm behind the LLM may be known, but its behavior remains as mysterious and capable as the mind that created it. Ultimately, intelligence is the bridge between a problem and its solution, and the material of the bridge-whether neurons or transistors-matters less than its ability to carry the weight of the task.

REFERENCES

1. Bubeck, S., et al. (2023). Sparks of Artificial General Intelligence: Early experiments with GPT-4. arXiv preprint arXiv:2303.12712.
2. Chalmers, D. J. (1996). *The Conscious Mind: In Search of a Fundamental Theory*. Oxford University Press.
3. Grzankowski, A. (2024). Intelligence and Representation in LLMs. *Journal of Cognitive Science*.
4. Hofstadter, D. R. (1979). *Gödel, Escher, Bach: An Eternal Golden Braid*. Basic Books.
5. Ku, A., et al. (2024). World models and artificial general intelligence. Royal Society Publishing.
6. Russell, S. J., & Norvig, P. (2016). *Artificial Intelligence: A Modern Approach*. Pearson.
7. Turing, A. M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433-460.
8. Vaswani, A., et al. (2017). Attention is All You Need. *Advances in Neural Information Processing Systems*.