

---

Subject: Re: Artificial Intelligence

Posted by [Wayne Parham](#) on Sat, 24 May 2025 21:59:55 GMT

[View Forum Message](#) <> [Reply to Message](#)

---

Thanks for your kind words. Honestly, I'm just kind of recording my thoughts because I think they might be useful, if even only marginally. There are some really smart folks out there studying this stuff, but right now, I think we're kind of in a "retrograde" position in our learning curve. We've had such success with Large Language Models (LLM), that I think we've sort of forgotten some of the earlier intellectual exercises on our path.

Once we pushed past single-layer Perceptrons to a few layers, and then onto multiple-payer convolutional networks, we were able to do some pretty good image processing and image recognition. This moved onto real-time feeds and language through LLMs, and once there, we seemed to try to fit everything onto that kind of technology. It does do some remarkable things. But we seem to have become so excited with LLMs that we try to fit everything into that approach.

It's like once we got to that point, we became so enamored with the technology that we stopped looking elsewhere. We began to see LLMs as the universal tool, the crescent wrench of the digital domain.

LLMs depend on huge databases of tokenized words, and the relationships between words - mostly their closeness in many sentences - is what is used for searches. It works, but it's inefficient and somewhat more inaccurate.

In the old days, documents were searched on perfect matches. This was too limiting so fuzzy methods were also employed. In addition to queries for perfectly equal matches, we added a "like" search query for near matches. That helped, but it was still somewhat limiting. And now that we can use huge bodies of LLM tokenized data for searches, we essentially have a fuzzy search of a huge corpus of data.

This works very well for what it is. It is useful. We can do some incredible stuff with this technology. But the machines don't know a thing. They don't know up from down, inside from outside, equal from inequal. They know words, even many formulas. But they do not know concepts.

Back in the 1970s and 1980s, I was influenced heavily by the idea that machines needed to be able to do analogical reasoning in order to proceed onto other more advanced "thought." The idea I took from that is our machines need to be able to understand concepts before anything else. I still think that's true.

There is a very cool research group called DeepMind and they've done some remarkable things. Truly remarkable. They do it by combining LLMs with other approaches that can self-learn. It's an excellent approach, because fairly sophisticated concepts can be taught, and this approach takes it to another level.

But I cannot help but wonder if we're doing it fundamentally wrong, backwards sort of. Instead of helping the LLM by using genetic algorithm techniques to self-learn specific higher-level concepts,

why not instead have the concept-learning be done at the absolute lowest level?

When I think about how animals learn - including early man - it must start with primitive concepts. There is obviously no language at this level. Thought is pure concept recognition.

There is the fight-or-flight mode. The concept is "danger." It's a thought like an exclamation point.

Then there's the thought - the concept - of what caused the danger. Predator. Tiger. Lion. It's an image of the creature. Might be a memory of an image of the creature tearing into another creature at the end of a fight, devouring it.

The concept of wound. Of death. The concept of hunger. Maybe a concept of sharpened pole to be used to impale. A fish. Impaling a fish. Eating it. Hunger quenched.

Or maybe of a hanging fruit. Eating it. Some taste good. Some satisfy hunger. Some are poisonous. So the concept of sickness. Danger of poison. Death from that, an adjacent concept to death from wound.

All these kinds of concepts are essentially causality maps. They're self-constructed simplifying explanations of things experienced. They start as mental images or video-like sequences. An experience is perceived, and a quick decision/conclusion is formed that explains the experience. The images and sounds may be remembered, or that memory may fade over time. But the concept that was learned is stored, and through repeated similar experiences, the concept is reinforced in memory.

The simplifying explanation of a concept also makes it very efficient to store in memory and to process. It doesn't require all memories of past events to be recalled. They can be allowed to fade. The reinforcement by new events that are similar strengthens the concept in memory without needing to store and process all details of past events.

Language itself is a concept. And in describing things using language, other concepts can be explored. Concepts about concepts. Concepts in the past, concepts in the future. Time is a concept too.

All these concepts are very low-level information objects. They don't arise out of a word soup, as some LLM enthusiasts argue. That's why I think we're in a retrograde position now. Our understanding is an eddy current.

I think we need to go back to low-level analogical reasoning. We need to develop neural network topologies that can explore concepts. We need to nudge them into learning those low-level concepts. And then assign words to the concepts. Instead of having tokenized words, we need tokenized concepts. We need concepts vectors.

In a concepts vectoring approach, the embeddings are similar, in that each concept is ranked by its similarity to other concepts. Concepts that are similar or that are often used together would have closer vectors. So the approach is the same as with other modes.

I personally think this is a very important next step. Instead of training on words, sounds, images or other symbols or representations of concepts, we need to train in the concepts themselves. We then associate those concepts with descriptive words, images and sounds. Only then can we expect true understanding.

I also think that once the core modality is concepts, that's when the machine will gain a sense of self and a sense of its place in its environment. It will be able to develop new concepts and concept associations, and be capable of learning through its experience.