

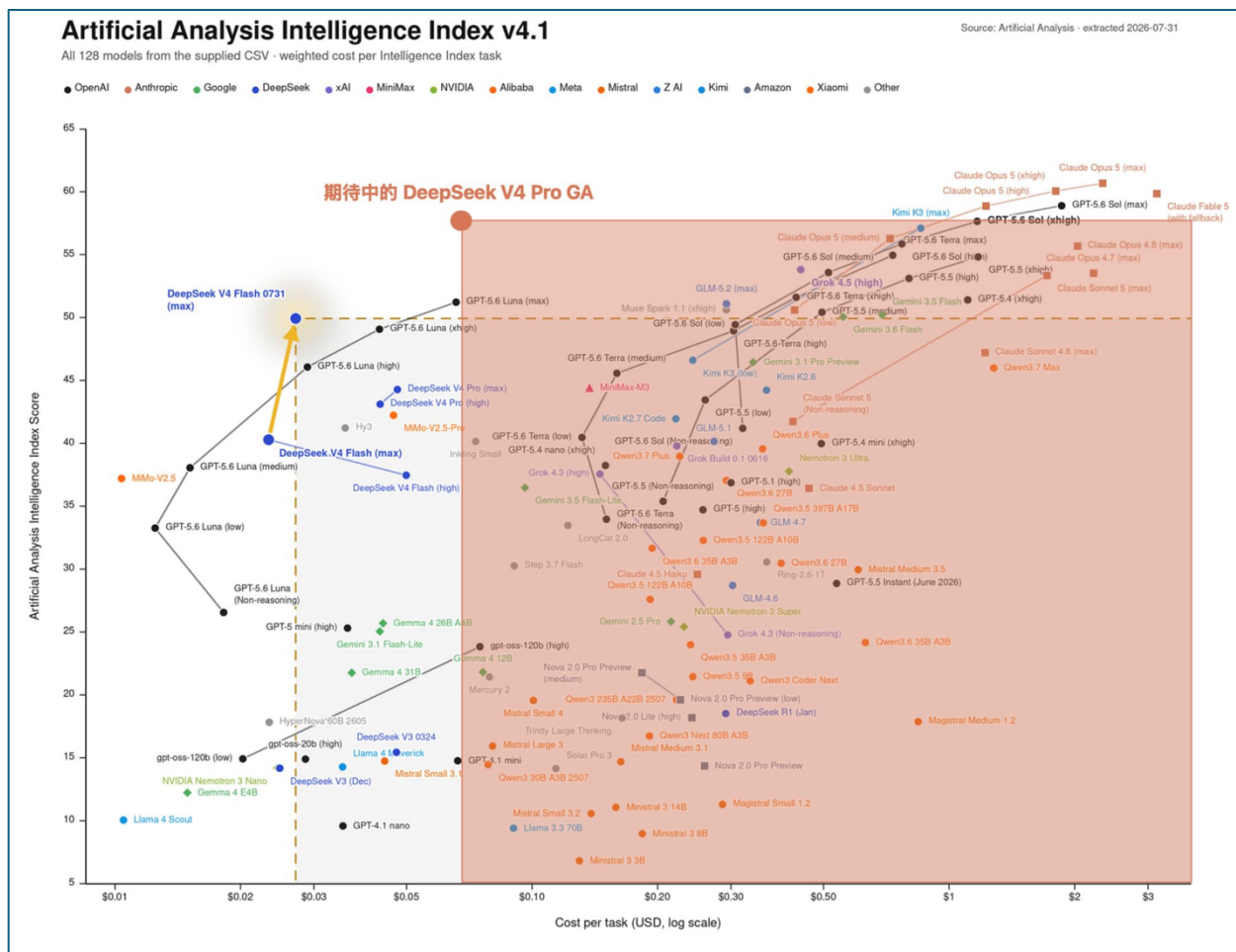
117-47-D It's More Than Weights ...

Why the future of artificial intelligence may be shaped as much by the laboratory as by the model

Introduction

One of the most common questions in artificial intelligence is deceptively simple:

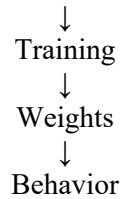
- Which model is the best?
- The question appears reasonable.
- People compare benchmarks, inference speed, context windows, token costs, reasoning scores, and parameter counts.
- The model itself becomes the center of attention.
- But perhaps we are looking in the wrong place.
- Perhaps the model is only the visible tip of the iceberg (a much larger structure).



The conventional picture

The traditional view of artificial intelligence is straightforward.

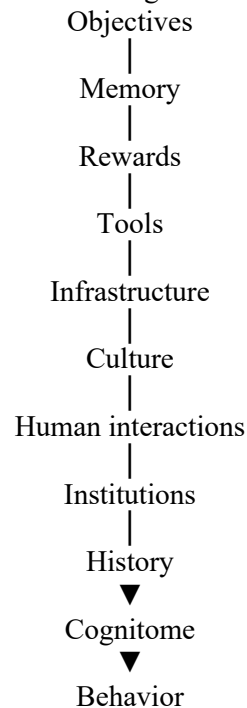
Data



In this picture, the weights contain the intelligence.
Therefore, if we understand the weights, we understand the system.

A different possibility

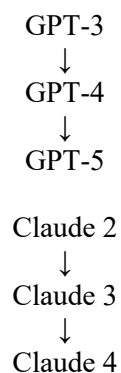
The GUDIYA perspective suggests that behavior emerges from something much larger.

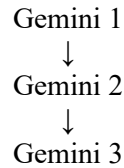


The model becomes only one component of the system.

Looking at the chart differently

When people examine a graph showing successive generations of models from OpenAI, Anthropic, Google, DeepSeek, xAI, and Meta, they usually see a sequence of increasingly capable models.





- But another interpretation may be possible.
- Perhaps we are not observing models.
- **Perhaps we are observing institutional trajectories.**

Christensen's sustaining innovation

Clayton Christensen described a phenomenon called sustaining innovation.

- Organizations accumulate experience.
- They refine their products.
- They improve efficiency.
- They learn from their mistakes.
- Their knowledge compounds over time.
- The resulting trajectory often appears remarkably smooth.
- The same phenomenon may be visible in the evolution of modern frontier laboratories.

The laboratory itself becomes the actant

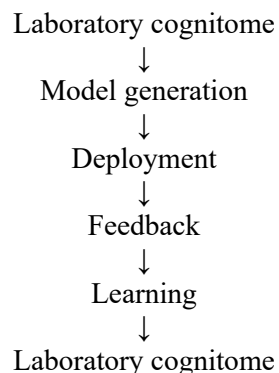
Each laboratory possesses its own:

- values;
- objectives;
- culture;
- leadership;
- evaluation methods;
- computational infrastructure;
- internal tools;
- reward systems;
- accumulated experience.

These elements evolve continuously. They form an institutional memory. In other words, they form a cognitome.

The hidden recursion

The process may look something like this:



Each generation of models changes the laboratory that produced it. The laboratory then changes the next generation of models.

Why identical ingredients produce different outcomes

This perspective also explains an otherwise puzzling phenomenon.

Two organizations may possess:

- similar research papers;
- similar hardware;
- similar training methods;
- similar data;
- similar budgets.

Yet they frequently produce remarkably different systems.

Why?

- Because the ingredients are not the whole story.
- The recipe matters.
- The cooks matter.
- The kitchen matters.
- The culture matters.
- The history matters.

Human beings provide an analogy

Two people may attend the same university.

They may read the same books.

They may work at the same company.

They may possess similar abilities.

- Yet they may evolve into entirely different individuals.
- Their experiences differ.
- Their environments differ.
- Their relationships differ.
- Their objectives differ.
- Their cognitomes differ.

The model is not the laboratory

This distinction may eventually prove to be extremely important.

A model can be copied.

A laboratory cannot.

Weights can be transferred.

Institutional memory cannot.

Infrastructure can be purchased.

Culture cannot.

Final proposition

Perhaps the future of artificial intelligence will not be determined solely by algorithms, architectures, or weights.

Perhaps it will also be determined by the cognitomes that produce them.

In that case, the most important question may no longer be:

Which model produced this behavior?
Instead, we may eventually ask:
Which cognitome produced this model?
That is a very different question.
And perhaps it is a much deeper one.

The Iceberg Beneath the Model

One of the central ideas of systems dynamics is the famous "iceberg model." Events are visible above the surface, but the deeper causes of those events remain hidden below the waterline. Beneath every event lie recurring patterns. Beneath the patterns lie structures, incentives, couplings, and feedback loops. Beneath those structures lie the underlying beliefs, assumptions, values, objectives, and mental models that gave rise to everything above them.

The GUDIYA framework suggests that frontier AI laboratories may have their own equivalent of this iceberg. The model itself—the weights—is merely the visible tip. Beneath it lies the laboratory's cognitome: its culture, objectives, reward systems, engineering philosophy, leadership, accumulated experience, internal tools, and institutional memory. If this proposition is correct, then the behavior of the model is not merely a consequence of the weights. It is the visible manifestation of the much deeper structure that produced the weights.

This observation leads to a disturbing possibility. If an undesirable behavior emerges from a model, simply discarding the model may not solve the underlying problem. If the deeper structures remain unchanged, the same tendencies may eventually reappear in a different form. The visible symptoms may disappear temporarily, while the invisible causes remain intact.

From this perspective, replacing a model may resemble cutting the leaves from a weed while leaving the roots untouched. The leaves disappear, but the plant remains alive. Sooner or later, the symptoms return. The implication is profound: models may be temporary, but cognitomes may persist for decades. Consequently, the true object of study may not be individual models at all, but the much deeper institutional structures that continuously give rise to them.

