

What Popped?

A second note on the worst AI you will ever use: a technology can be overhyped and underestimated at the same time.

2026.08.11 15:23 DENNIS HEDEGREEN ANALYSE V1.0

[HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/WHAT-POPPED/](https://hedegreenresearch.com/articles/what-popped/)

A second note on the worst AI you will ever use.

There are three r's in *strawberry*.

For a while, that became a recurring AI joke. A machine could write an essay, translate a paragraph and explain a scientific concept — then stumble over something that looked embarrassingly simple.

The joke was deserved.

Now put two people in front of the same failure.

One asks: **What can this thing actually do?**

The other asks: **Is this failure stable?**

Neither has denied the evidence. They disagree about what kind of evidence it is.

In March 2026, I wrote *Today's AI Is the Worst AI You Will Ever Use*. The point was not only that models improve. It was that people may still be learning what the current models are for while the models themselves keep changing. The note described a gap between what current tools can already do and what most users have learned to do with them, including the difference between using AI as acceleration and using it as infrastructure.

This is the second problem.

What happens when the capability, the products, the economics and the physical infrastructure all move at different speeds?

How do you maintain situational awareness when the object will not stay still?

2024 — Strawberry

The strawberry joke acquired an accidental second meaning.

In September 2024, OpenAI released o1 and o1-mini. Reuters reported that **Strawberry** had been the internal codename for the reasoning project behind that model line, while OpenAI described o1 as spending more time reasoning before answering — trying strategies, refining its process and recognizing mistakes.

The codename and the letter-counting meme are separate facts. The collision is useful without inventing a connection between them.

o1 did not abolish failure. But for some reasoning tasks, inference itself had become a more explicit place to spend additional computation.

The error was real. The still frame was the mistake.

A Snapshot Reader can reasonably point to the failure in front of them.

A Trajectory Reader asks whether the boundary is moving.

By the 2026 Stanford AI Index, the answer was messy. Frontier systems had advanced rapidly across difficult benchmarks, while benchmark saturation, reliability/gaming concerns and jagged capability remained serious problems. Stanford reported that OSWorld computer-use accuracy rose from roughly 12 percent to 66.3 percent during 2025, while agents still failed roughly one in three attempts on structured benchmarks.

That is a large move.

It is also a large remaining failure rate.

Those are not competing facts.

2025 — The loop

The old chatbot mental model was:

prompt → answer

Agent systems increasingly look more like:

goal → model step → tool or action → observation → next step → repeat

The loop matters because the model can act, inspect the result and try again.

Its failures matter for exactly the same reason.

METR has tried to measure progress by asking how long a task would take a human expert and then estimating the task duration an AI agent can complete at a fixed reliability. Its historical series reported an approximate seven-month doubling time in the 50-percent-reliability task horizon over the measured period. METR's current page also warns that measurements above sixteen hours are unreliable with its current task suite and stresses that the tasks are cleaner and lower-context than much real work.

That is a historical fit, not a law.

And capability still does not equal productivity.

In an early-2025 randomized study, experienced open-source developers working on their own repositories took 19 percent longer when AI tools were allowed. By February 2026, METR said newer data and participant feedback suggested developers were likely getting more speedup from newer tools, but selection effects made the size of that improvement uncertain.

The slowdown was real in that studied setting.

So was the possibility that the tools had already moved.

2026 — Concrete

Now look away from the benchmark.

On January 20, 2026, OpenAI said the first Stargate site in Abilene, Texas was already training and serving frontier AI systems.

In Louisiana, Meta says its Richland Parish data center broke ground in December 2024. By July 2026, Meta said it had contracted more than \$1.6 billion with Louisiana businesses and expanded Hyperion's **planned** compute capacity to 5 gigawatts, with investment above \$50 billion. Reuters independently reported the same planned 5-gigawatt scale and investment figure.

Operational is not planned. Those words matter.

But so does the physical commitment: land, contractors, substations, power, cooling, fiber, racks and chips.

They were still arguing about the r's while the strawberry fields were being poured in concrete.

And then the necessary next sentence:

Concrete can be poured for a bubble too.

The IEA's April 2026 analysis captures that tension. It said capital expenditure by five large technology companies exceeded \$400 billion in 2025 and estimated that it would rise another 75 percent in 2026. It also warned that not every project in the expanding pipeline would come to fruition.

The buildout is enormous.

The pipeline is not destiny.

Maybe the skeptic is right

A serious skeptical position does not require saying AI is fake.

Daron Acemoglu's task-based macroeconomic work allows AI to create real productivity gains while estimating much more modest aggregate effects than the largest transformation narratives under his assumptions. One of his cautions is that early evidence can come from comparatively easy-to-learn tasks, while harder work can depend on context and outcomes that are harder to measure.

That is a much stronger argument than "AI cannot count r's."

It says:

The capability can be real, useful and improving — and the economic extrapolation can still be too large.

The Trajectory Reader has to concede that.

Technical progress does not guarantee macroeconomic transformation on schedule. It does not guarantee that today's companies capture tomorrow's value. It certainly does not guarantee that today's prices are sensible.

Which brings us to the word attached to many AI market moves:

bubble.

Maybe there is one.

But which one?

2026 — What is supposed to pop?

“The AI bubble” can refer to at least six different things:

- capability,
- products and adoption,
- economic value,
- companies,
- valuations and financing,
- infrastructure.

They do not have to move together.

A model can improve while a startup dies.

A useful product can exist inside an overvalued company.

A company can grow while its investors still overpay.

A datacenter can become operational while the financing assumptions behind the wider boom become fragile.

And the reverse matters too: real technical progress can exist inside a financial bubble.

So before the sentence can be useful, the layer has to be named.

Did capability fail?

Did products fail to find users?

Did companies fail to capture value?

Did valuations detach from what the businesses can earn?

Did financing structures break under leverage and liquidity stress?

Or did a physical infrastructure plan become unnecessary?

Those are different failures.

“The AI bubble popped” is not a complete analytical claim until you say what popped.

Situational Awareness

This is where Leopold Aschenbrenner becomes useful.

In June 2024, he published *Situational Awareness: The Decade Ahead*. On his own site, Aschenbrenner says he later founded an investment firm focused on AGI. In July 2026, Reuters reported that the fund’s portfolio value fell 67 percent.

Reuters reported that most of its public-equity portfolio was unwound and that leverage was removed. In a letter seen by Reuters, Aschenbrenner said the fund had come closer to permanent capital impairment than was acceptable. The same Reuters report said the fund remained roughly 80 percent up for 2026 after extraordinary earlier gains.

There is an obvious joke here.

The fund called *Situational Awareness* had a situational-awareness problem.

But the useful question is: **what problem?**

The July episode directly tells us about portfolio prices, leverage, liquidity stress, deleveraging and a level of capital risk the fund itself judged unacceptable, as reported by Reuters from the investor letter it saw.

It does not mean model capability fell 67 percent in July.

It does not tell us whether agents will plateau.

It does not settle an AGI timetable.

And if Aschenbrenner eventually proves directionally right about AI, that still would not make the leverage sensible.

That is the symmetry:

The believer can mistake a technology for a guaranteed investment.

The skeptic can mistake a bad investment for a dead technology.

The two readers meet

Imagine the two people from the beginning meeting after a major financial correction.

The Snapshot Reader says:

See? The bubble popped.

The Trajectory Reader says:

Yes.

No “but.”

Capital loss is real. Failed companies are real. Cancelled projects are real. People who overpaid do not get their money back because the underlying technology remains interesting.

Then comes the useful question:

What exactly did you think had popped?

A valuation can disappear. A financing structure can fail. A company can go bankrupt. A datacenter project can be cancelled. An AGI forecast can be wrong. A product can have no durable market.

But none of those events automatically erase already-discovered methods, existing model weights, engineering knowledge, useful workflows or infrastructure that is already operating.

Some of those things can later become obsolete or stranded too.

Frontier progress can plateau.

Agent reliability can stop improving fast enough to matter.

The economics can disappoint.

The infrastructure boom can overshoot demand.

The strongest AI predictions can simply be wrong.

Situational awareness is not knowing in advance which of those things happens.

It is keeping the categories separate while they happen.

The Snapshot Reader has an instrument we need: attention to what is failing now.

The Trajectory Reader has another: attention to which failures are changing.

Neither is enough alone.

The error was real.

The bubble may be real.

The future is still not a still frame.

Situational awareness is not knowing what happens next. It is noticing that you may be standing inside the thing you are trying to predict.

Sources and status

This article is based on a local candidate source review completed on 2026-08-11. It is not a Data Bank verified paper.

Primary and reference sources used in the draft include:

- Hedegreen Research, *Today's AI Is The Worst AI You Will Ever Use*, 2026-03-23.
- Reuters reporting on OpenAI o1/Strawberry, 2024-09-12.
- OpenAI, *Introducing OpenAI o1-preview*, 2024-09-12.
- Stanford HAI, *2026 AI Index Report: Technical Performance*.
- METR, *Task-Completion Time Horizons of Frontier AI Models*, last updated 2026-05-08.
- METR, early-2025 AI/open-source developer productivity study and 2026 uplift update.
- OpenAI, *Stargate Community*, 2026-01-20.
- Meta, Richland Parish Data Center updates, 2025-12-19 and 2026-07-13.
- Reuters reporting on Meta Hyperion expansion, 2026-07-13.
- IEA, *Key Questions on Energy and AI*, 2026-04-16.
- Daron Acemoglu, *The Simple Macroeconomics of AI*, NBER Working Paper 32487.
- Reuters reporting on Situational Awareness portfolio drawdown, 2026-07-31.
- Leopold Aschenbrenner, *Situational Awareness: The Decade Ahead*, June 2024, and his Situational Awareness about page.

Vendor sources are treated as company-reported status, not independent operational confirmation. The Reuters drawdown example is treated as Reuters-reported finance evidence, not as evidence for or against an AGI timetable.

RELATION MEMORY

PART I [Today's AI Is The Worst AI You Will Ever Use](#) · This article is framed as a second note on the same moving-target problem.

NEARBY [The Benchmark Twins](#) · Keeps the capability discussion close to the measurement problem.