

The next AI acceleration may come from more inference, not recursive self-improvement

AI News Digest

2026-09-20

The next AI acceleration may come from more inference, not recursive self-improvement

By AI News Digest • September 20, 2026

The day’s clearest capability signal is a shift toward inference-time scaling and agent swarms rather than demonstrated recursive self-improvement. The digest also tracks the physical, legal, financial, and local-compute constraints that are beginning to shape how far that acceleration can go.

The near-term acceleration thesis

Agent swarms are scaling; recursive self-improvement is still an open claim

A new Interconnects analysis describes frontier labs—especially OpenAI and Anthropic—as already using thousands of concurrent agents, while warning that rapidly scaling inference-time compute should not be confused with recursive self-improvement (RSI). Its baseline is “lossy self-improvement”: more agents and more compute can accelerate clearly stated, verifiable work, but exponential resource costs, diminishing returns from parallel agents, and hardware and political bottlenecks remain. [1]

The practical version of that thesis is visible in a current inference-scaling tutorial. The base model scored 15.2% on the reported task, versus about 48% for a reasoning variant and 40% with chain-of-thought prompting; adding top-p sampling, chain-of-thought, and self-consistency raised the result to 52% with five or ten samples, but took roughly six times as long. The reasoning model reached 55% with the same techniques, and the tutorial’s conclusion is that self-consistency is useful when accuracy matters—not as a default operating mode. [2]

That is the important distinction for the next phase of competition: capability

gains may increasingly come from allocating more inference and parallel attempts to each problem, with corresponding cost and latency tradeoffs, rather than from an unexplained jump in peak intelligence. The Interconnects analysis says the clearest internal automation gains so far are in software engineering, log monitoring, planned experiments, and other routine research support, while Anthropic’s cited system-card language reports no clear acceleration beyond the current rate of progress. [1]

Physical-world deployment

A reported Anthropic wet lab raises a harder safety boundary

A report circulating in the monitored feed, attributed to Reuters, says Anthropic has quietly established a Bay Area “wet lab” to push Claude into real-world biology experiments. The report contains no operational details or primary announcement, so it is best treated as a signal rather than confirmation of a particular capability. [3]

Emad Mostaque responded that AI-automated labs should be barred from viral- or pathogen-related work because, in his view, they would otherwise conduct gain-of-function research. The significance is the governance boundary: once models are connected to physical laboratory workflows, oversight has to cover experiment authorization, equipment access, monitoring, and shutdown—not only the model’s text outputs. [4]

Constraints on the race

AI “pacing” is now an antitrust allegation

The Associated Press reported that a new lawsuit claims Anthropic, OpenAI, SpaceXAI, and Google illegally agreed to slow AI development. The allegation is unproven, but it turns a debate usually framed around safety and competition into a legal claim about market conduct. [5]

Gary Marcus said the stated motivation was that slowing development would reduce the value consumers receive from paid AI subscriptions; he separately suggested liability concerns around future models may be part of the real incentive to slow down, while calling the lawsuit’s broader premise misguided. Those are interpretations, not established facts, but they expose the tension now surrounding “pacing”: slowing may protect consumers or reduce liability, while also changing the economics of the model race. [6, 7]

Data-center finance is becoming an AI bottleneck

A post quoting the *Financial Times* says about \$18 billion of loans tied to a New Mexico data center leased to Oracle entered stressed territory, with investors worried that local backlash could derail the company’s broader AI-infrastructure build-out. [8] Gary Marcus, while cautioning that it might not be this specific

case, warned that a development like it could trigger a cascade through the AI-infrastructure “house of cards.” [9]

One stressed financing package is not evidence of a sector-wide credit event. It does show, however, that scaling AI capacity depends on local political consent and financing structures as much as on chips, power, and model demand.

Research path

Compile a language task once, then run it locally

The *Program-as-Weights* preprint proposes a 4B compiler that turns a natural-language function into parameter-efficient adapters for a frozen 0.6B Qwen3 interpreter. Its abstract reports that the resulting local program matches direct prompting of Qwen3-32B while using roughly one-fiftieth the inference memory and running at 30 tokens per second on a MacBook M3. [10]

A follow-up paper, *Compile by Training*, reports 83.6% semantic accuracy on FuzzyBench-Hard—a subset where the fast compiler produced no exact matches—at roughly a minute of compile time rather than seconds. The authors’ broader proposition is more consequential than either number: recurring fuzzy text functions can become small, reusable, versioned, offline artifacts instead of repeated calls to a large remote model. These are preprint-reported results, but they point to a practical route toward lower-cost, more private local AI for fixed workflows. [11]

Sources

1. Why I still haven’t bought into true RSI
2. Build A Reasoning Model From Scratch 4: Inference Scaling 1 (Temperature, Top-p, Self-Consistency)
3. X post by @Polymarket
4. X post by @EMostaque
5. X post by @AP
6. X post by @GaryMarcus
7. X post by @GaryMarcus
8. X post by @carlquintanilla
9. X post by @GaryMarcus
10. Program-as-Weights: A Programming Paradigm for Fuzzy Functions
11. Compile by Training: Turning Natural-Language Specifications into Local Neural Functions