

SevenTnewS.com

UNE

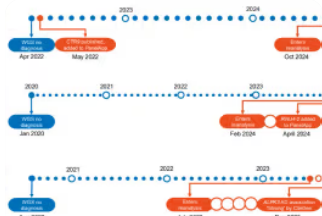
ANTHROPIC

Anthropic Restores Access to Claude Fable 5 After US Export Controls Lifted

The week frontier models met their match in cost, agents, and reality

This week's throughline: This week, the AI landscape was defined by a collision between ambition and friction: frontier model releases from DeepSeek, Anthropic, and MiniMax defied the assumption that top-tier performance requires proprietary expense, while cybersecurity threats demonstrated that adversaries are weaponizing AI as fast as defenders can deploy it. At the same time, new benchmarks from GauntletBench to FFASR exposed stark gaps between lab hype and real-world reliability. Agents floundered on vision tasks, long-horizon edits degraded documents, and promised million-token contexts proved misleading. Yet amid the sobering reality checks, a parallel wave of open-source tooling and infrastructure funding signaled that

the ecosystem is maturing toward practical, efficient deployment, even if the agents themselves aren't ready for prime time.



Talos, an open-source tool, automates reanalysis of genomic data to speed rare disease diagnosis

Running the Gauntlet: Re-evaluating the Capabilities of Agents Beyond Familiar Environments

Mihailo Yezhov^{1,2,3}, Ronit Libi¹, Gregory Biele¹, Michal Zakrzewski¹, Sebastian Montagna¹, Daniel Rysnick¹, Shreyash Padurha¹, Kunal Choudhary¹, Zhen Fu¹, William Lippman¹, Kai Ren¹, Hanna Vozhova¹, Xander Davis¹, Tara Rasmussen¹, Gordon Li¹, Paul Burt¹, Barisim Vire¹, Arkadiusz Dziadosz¹, Vito Gar¹, Chris Russell¹, Christopher Sommerfeld¹, Adam Michal¹, Volodymyr Karga¹, Philip Torr¹, Adi Bilo¹

¹University of Oxford, ²Stanford, ³Massachusetts Institute of Technology

⁴The Chinese University of Hong Kong, ⁵UK AI Security Institute, ⁶Signia AI

⁷The Chinese University of Hong Kong, Shenzhen, ⁸University of Cambridge

Abstract

As agentic systems continue to evolve and are widely deployed in real-world scenarios, there is a growing demand to faithfully evaluate their capabilities. How-

Oxford's GauntletBench puts AI agents through 100 real tasks. They failed 81% of them.



M3D and Real-Guidance Bring Dataset Distillation to High-Resolution Realms



Microsoft's Phi-4 Model Redefines Efficiency in Breakthrough Research

SOMMAIRE

- 01 Frontier Model Releases Spark Access and Cost Debates
- 02 Cybersecurity Threats Escalate with AI and Novel Malware
- 03 New Tools Democratize AI for Science and Scientific Workflows
- 04 Open-Source AI Tooling and Infrastructure Maturity Accelerate
- 05 Together AI and Infrastructure Funding Signal Boom
- 06 AI Benchmarks Expose Gaps in Reasoning and Real-World Scenarios
- 07 AI Agents Face Critical Reliability and Verification Hurdles
- 08 AI Infrastructure and Optimization Tools Gain Ground
- 09 Inference Modeling and Optimization for Large Models
- 10 AI Privacy, Security, and Content Control Intensify
- 11 New Methods for Model Understanding and Verification
- 12 Vision-Language and Multimodal Research Surge
- Conclusion

1. Frontier Model Releases Spark Access and Cost Debates

This week saw a wave of model releases that intensify pressure on incumbents to balance performance, openness, and affordability. DeepSeek released a preview of DeepSeek-V4 claiming world-class reasoning and improved agent capabilities while remaining open-weight, and also launched a new LLM continuing its efficiency push. Anthropic released Claude Sonnet 5 with near-Opus-level performance on agentic tasks at a mid-tier price, and Claude Fable 5 for extended autonomous work, with Aleph Alpha unveiling T-Free, a tokenizer-free architecture. MiniMax released M2.5, a coding model topping Multi-SWE-Bench at a fraction of the cost. The releases challenge the assumption that frontier models must be proprietary or expensive.

[01] DeepSeek-V4 preview lands, and the open-...

[02] MiniMax's new M2.5 coding model tops the...

[03] Ma Jiaqi taught MiniMax engineers a hard...

[04] Anthropic Launches Claude Fable 5: A Fif...

[05] Anthropic Restores Access to Claude Fabl...

[06] Anthropic Announces "The Briefing: AI fo...

[07] Aleph Alpha unveils T-Free: a tokenizer-...

[08] The case against enshittification: why s...

[09] Claude Sonnet 5 is here, and Anthropic i...

[10] Claude Sonnet 5 just made the Opus price...

[11] Anthropic Boosts Claude API Rate Limits,...

[12] DeepSeek's new model makes efficiency th...

2. Cybersecurity Threats Escalate with AI and Novel Malware

This week's cybersecurity landscape reveals an escalating arms race, marked by increasingly automated and sophisticated threats. Reports document the first fully AI-run ransomware attack, the deployment of Umbrij malware by the ToddyCat APT group to steal OAuth tokens, and a critical zero-day vulnerability in enterprise VPN appliances. Concurrently, ransomware attacks have reached unprecedented scale, with demands hitting \$120 million and crews operating as specialized firms using triple extortion tactics. A European initiative aims to train 10,000 cybersecurity specialists by 2026 to address the critical talent shortage, while a teen linked to the Scattered Spider hacking crew was extradited to the U.S.

[01] First fully AI-run ransomware attack dis...	[02] ToddyCat apt deploys umbrij malware to h...
[03] CISA adds microsoft sharepoint server vu...	[04] Critical Zero-Day CVE-2025-XXXX Strikes...
[05] WhatsApp username reservations raise imp...	[06] The Rising Tide of AI-Powered Phishing:...
[07] Teen accused in scattered spider hacking...	[08] The ransomware renaissance: \$120 million...
	[09] Europe's once-austerity founders are abo...

3. New Tools Democratize AI for Science and Scientific Workflows

A suite of new tools aims to democratize AI for domain-specific scientific and technical work. Microsoft open-sourced Data Formulator 0.7 for enterprise analytics without SQL, launched the Microsoft Discovery platform for building agentic AI workflows in science and engineering, and opened Anthropic Academy for learning about Claude. OpenAI introduced Prism, a free LaTeX workspace powered by GPT-5.2 for scientists. Talos, an open-source tool, demonstrated automated genome reanalysis delivering 241 new rare disease diagnoses in 32 days. DiagFit secured 6.5 million euros for its healthtech approach combining existing patient data into a unified risk score. Additionally, Phi-4 matched far larger rivals on reasoning benchmarks, signaling a shift in model scaling thinking.

[01]

Microsoft open-sources Data Formulator 0...

[02]

Anthropic Academy: a free platform to le...

[03]

OpenAI launches Prism: a free LaTeX work...

[04]

Microsoft's Phi-4 Model Redefines Effici...

[05]

Talos shows automated genome reanalysis...

[06]

Microsoft Discovery Now Generally Availa...

[07]

DiagFit's €6.5 million bet on the data y...

4. Open-Source AI Tooling and Infrastructure Maturity Accelerate

The open-source AI ecosystem continues to mature with significant contributions and infrastructure advances. Hugging Face rebuilt its release pipeline for weekly Python client releases and streamlined vLLM inference server deployment, while also introducing Moon Bot, a coding agent inside Slack. Ollama 0.31 delivered a 90% speedup to Gemma 4 on Apple Silicon via multi-token prediction. BenchLM published a playbook for making pages citable by AI assistants, and the M3D dataset distillation method achieved 68.5% top-1 accuracy on ImageNet-1K with one image per class while slashing memory tenfold. These developments highlight a growing focus on practical deployment, efficiency, and reproducibility.

[01] M3D and Real-Guidance Bring Dataset Dist...

[03] Hugging Face Lets You Run a vLLM Server...

[05] How to Get Cited by ChatGPT, Perplexity,...

[02] How Hugging Face Ships Its Python Client...

[04] Hugging Face's Moon Bot turns Slack into...

[06] Token-Level Analysis Reveals Hybrid LLMs...

[07] Gemma 4 runs 90% faster in Ollama 0.31 w...

5. Together AI and Infrastructure Funding Signal Boom

The business of AI infrastructure is booming, with major funding rounds and strategic moves reshaping the competitive landscape. AI cloud startup Together AI raised \$800 million at an \$8.3 billion valuation, while AI chip startup Groq secured \$750 million to meet surging inference demand, partnering with the U.S. Department of Energy and announcing a \$1.5 billion expansion in Saudi Arabia. Lime also drew attention with its \$167 million IPO, though it faces questions about its \$1 billion debt, and SpaceX reportedly showed an AI device prototype to investors. These moves underscore the intense capital flow into AI hardware, cloud, and related ventures.

- [01] Together AI raises \$800 million at an \$8...
- [02] Inside BenchLM's Data Pipeline: Why We D...
- [03] Apple reportedly planning new iPad pros...
- [04] Lime's IPO raised \$167 million. That mig...
- [05] SpaceX is developing an ai device protot...
- [06] Groq raises \$750M as inference demand su...

6. AI Benchmarks Expose Gaps in Reasoning and Real-World Scenarios

Several new benchmarks and test frameworks reveal significant gaps between AI model performance in controlled settings and real-world or out-of-sample conditions. The ARC-AGI-2 benchmark sees top models scoring 85% while humans average 66%, yet GPT-5.4 leads a saturating math pack. The AIMIP Phase 1 benchmark shows AI climate models excel at historical patterns but struggle with out-of-sample scenarios. The FFASR Leaderboard reveals a large gap between near-field and far-field speech recognition accuracy. Most notably, GauntletBench found that agents failed 81% of 100 vision-intensive tasks across professional apps, while non-expert humans achieved over 80% success.

- [01] ARC-AGI-2: The Benchmark That Measures F...
- [02] GPT-5.4 Leads the 2026 Math Benchmark Pa...
- [03] Aimip phase 1: A new benchmark to test a...
- [04] Treble Technologies and Hugging Face Lau...
- [05] Oxford's GauntletBench puts AI agents th...
- [06] ProgramBench: every public AI scores 0%...

7. AI Agents Face Critical Reliability and Verification Hurdles

New research exposes fundamental reliability gaps in AI agents, particularly for long-horizon tasks and verification. Studies show LLMs introduce semantic errors in documents after repeated delegated edits, with degradation rates of 19 to 34% over 20 iterations in the DELEGATE-52 benchmark, and verifying coding agent output has become harder than generating code. The ProgramBench test found every public model scored 0% on reconstructing programs from binaries. The MosaicLeaks benchmark reveals deep research agents leak sensitive internal information through outbound web queries, with a training method cutting leakage by three times. These findings underscore that agent reliability remains a critical barrier to deployment.

[01] Llms corrupt your documents when you del...

[02] AI document corruption in delegated work...

[03] The verification horizon: why verifying...

[04] Your AI research agent is leaking privat...

[05] Vega brings zero-knowledge proofs to ide...

[06] Cloudflare to block mixed-use web crawle...

8. AI Infrastructure and Optimization Tools Gain Ground

This week, several releases underscore a maturing AI infrastructure landscape focused on efficiency and real-world deployment. Microsoft unveiled GridSFM and Data Formulator 0.7, targeting optimization of power grids and enterprise data workflows, respectively, while also highlighting the performance of its mimalloc memory allocator in AI services. Noumena Labs' Sipp library offers a 3x to 5x speedup for local AI inference, and Talos automates genomic reanalysis to accelerate rare disease diagnoses. Together, these tools demonstrate a shift from raw model performance to operationalizing AI with minimal overhead.

[01] Microsoft releases gridsfm, a lightweigh...

[02] Microsoft's GridSFM predicts power grid...

[03] Microsoft just opened a codebase that ki...

[04] mimalloc, Microsoft's tiny memory workho...

[05] Sipp.sh Launches Open-Source Library for...

[06] Talos, an open-source tool, automates re...

9. Inference Modeling and Optimization for Large Models

Developments in inference optimization are enabling more efficient and cost-effective deployment of large models. Aleph Alpha published theoretical inference models for DeepSeek V3 based on hardware primitives, offering practitioners a framework for latency, throughput, and cost trade-offs. Ai2 released OlmoEarth v1.1, cutting satellite imagery AI compute costs by up to three times by collapsing multiple resolution-based tokens. Noumena Labs' Sipp library offers a 3x to 5x speedup for local AI inference. A detailed analysis of frontier LLMs revealed that promised 1M+ token context windows are misleading, with effective recall and costs varying dramatically, exposing a gap between marketing and actual capability.

- [\[01\] Aleph Alpha builds theoretical inference...](#)
- [\[02\] Aleph Alpha builds a theoretical inferen...](#)
- [\[03\] Ai2 cuts satellite imagery AI costs by 3...](#)
- [\[04\] Your AI model says it can read 1 million...](#)

10. AI Privacy, Security, and Content Control Intensify

Privacy and security concerns around AI are prompting new technical and policy responses. Vega introduced a zero-knowledge proof system for identity verification that generates proofs in under 100 milliseconds without uploading documents. Cloudflare announced it will block mixed-use web crawlers by default from September 2026, giving site owners more control over content used for AI training and agent services. Meanwhile, RISC-V's commercial rise is reshaping processor design, challenging proprietary architectures like ARM and x86, potentially impacting hardware security and control.

- [\[01\] Your AI research agent is leaking privat...](#)
- [\[02\] Vega brings zero-knowledge proofs to ide...](#)
- [\[03\] Cloudflare to block mixed-use web crawle...](#)
- [\[04\] How Open-Source RISC-V Is Disrupting the...](#)

11. New Methods for Model Understanding and Verification

Advances in interpretability and verification are addressing the growing need for trustworthiness in high-stakes AI applications. Researchers from Microsoft and partners developed generative causal testing (GCT), a framework that distills black-box brain-prediction models into testable verbal explanations, discovering new neural micro-regions. Another study introduced MaxProof, which transforms generative verifiers into reliable pass@1 performance, enabling a model to surpass human gold medal thresholds on IMO and USAMO problems. These methods complement the broader discussion of AI as a cognitive extension, with a new interdisciplinary paper arguing that modern AI systems derive power from structures rooted in human cognition rather than replicating intelligence.

[01] Microsoft's new method turns black-box b...

[02] How maxproof turns generative verifiers...

[03] AI as an extension of human intelligence...

12. Vision-Language and Multimodal Research Surge

A surge of vision-language research papers from institutions like Stanford and MIT is dominating Hugging Face's trending page, indicating the community is doubling down on multimodal architectures that integrate visual understanding with language. This trend aligns with broader efforts to push AI beyond text-only capabilities, as seen in the release of GPT-Live, a full-duplex voice model that breaks from rigid turn-based conversation, making interactions feel more natural. The architectural shift toward real-time, interruptible dialogue represents a significant step forward in conversational AI.

[01] Why vision-language papers are flooding...

[02] OpenAI's GPT-Live finally stops waiting...

Conclusion

The week's narrative arc suggests that the AI industry is entering a phase of ruthless calibration: model makers must now compete on cost and openness as much as raw capability, while users demand proof, not promises, of real-world utility. The gap between demo and deployment is narrowing, but the new benchmarks and agent failures make clear that engineering and trustworthiness, not just parameter counts, will separate winners from also-rans in the months ahead.

Généré à partir de la revue SevenTnewS du 19/07/2026 — 69 articles regroupés en 12 thèmes dédoublés.