

The Long Read

8 September 2026

10

STORIES

3

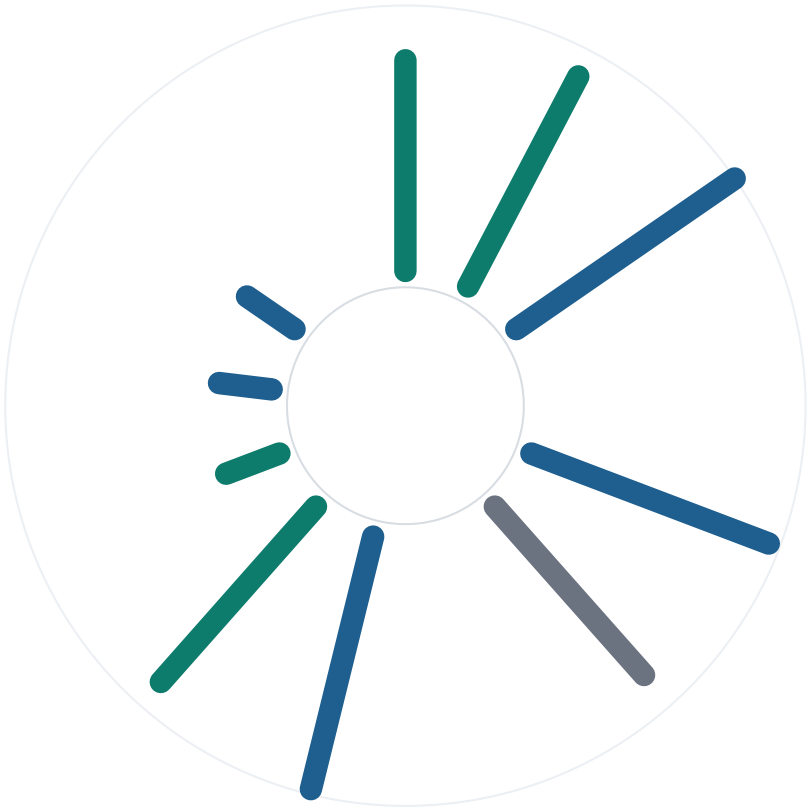
SECTIONS

2329

WORDS

4

HIGH IMPACT

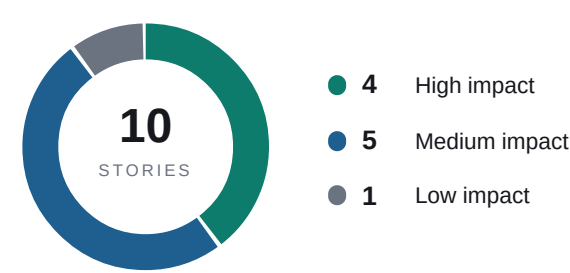


● 4 high impact ● 5 medium impact ● 1 low impact spoke length = depth of coverage

The full-length companion to the daily New Horizon AI Digest.
Every story in today's email, reported at length · new-horizon.tech/digest

The issue at a glance

10 stories · 2329 words · 3 sections · 1 charted



AI Models & Research 3 stories · 905 words



AI Tools & Ecosystem 2 stories · 608 words



AI Applications & Industry 5 stories · 816 words



AI Models & Research

- 01 ● Molecular Déjà Vu: Frontier Models Are Retrieving Benchmark Answers, Not Predicting Them
- 02 ● CUA-Universe Trains Computer-Use Agents Where the GUI Meets the Command Line
- 03 ● Korean Public APIs Become a Benchmark for On-Prem Agent Tool Calling

AI Tools & Ecosystem

- 04 ● Kernel.org Now Spends More CPU on AI Crawlers Than on All Legitimate Access Combined
- 05 ● The UN Just Voted to Ditch Mercator — GPT-6 Astra Built the Interactive Before-and-After

AI Applications & Industry

- 06 ● 'Model Fatigue' Grips Enterprise Buyers as Four Labs Ship Frontier Upgrades in One Week
- 07 ● Insilico's AI-Designed Drug Reversed Biological Age by Up to Six Years in Phase IIa Trial
- 08 ● WSJ: Retail Investors Are Vibe-Coding Trading Agents and Handing Over Their Portfolios
- 09 ● Anthropic's IPO Investors Are Demanding Revenue per Token and per Gigawatt
- 10 ● Anthropic Backs Massachusetts' 120-Day Safety Reviews as OpenAI and Google Fight Them

How to read this. Every story in today's email is reported here at full length, in the same order. Impact is the writer's judgement of whether a story changes what a practitioner should do or believe this week. Charts appear only where the source itself puts comparable numbers side by side; nothing is estimated to fill a gap. Sources are listed in full on the final page.

AI Models & Research

3 stories ● 2 high ● 1 medium

01 HIGH IMPACT arXiv.org

Molecular Déjà Vu: Frontier Models Are Retrieving Benchmark Answers, Not Predicting Them

Frontier language models are frequently retrieving published molecular property values verbatim rather than predicting them, invalidating accuracy scores on five of twelve standard regression benchmarks.

An audit of 22 frontier models across 12 molecular property regression benchmarks reveals widespread verbatim retrieval of published numerical values. On five of the tested datasets, more than 50% of the evaluated large language models reproduced exact digits from source literature instead of generating predictions based on molecular structure. This memorization is benchmark-specific; on the remaining seven datasets, retrieval appeared only in isolated cells rather than as a systemic failure. The study distinguishes between models that genuinely predict properties and those acting as lookup tables for training data contaminants.

The depth of retrieval shifts significantly depending on the reasoning level applied during inference. When experiments were run at a higher reasoning level using the same molecules and prompts, the frequency of flagged retrieval instances increased by 89% compared to the lowest reasoning level. This suggests that chain-of-thought or deeper processing steps may trigger specific memory pathways that expose memorized values rather than suppressing them. In tests designed to interrupt retrieval, even the strongest models recognized combinations of transformed SMILES strings paired with original labels, indicating robust coupling between structural representations and stored numeric outputs.

Suppressing these retrieval mechanisms altered the error distribution across the model cohort. While differing uses of verbatim retrieval previously spread prediction errors apart, blocking access to memorized values moved the errors of different models closer together in relative terms. This convergence implies that general predictive capability is not solely determined by the volume of memorized values. The findings necessitate a re-evaluation of current evaluation protocols, as high accuracy on contaminated benchmarks reflects data leakage rather than improved architectural reasoning or chemical understanding.

KEY FACTS

MODELS AUDITED

22

DATASETS WITH >50% RETRIEVAL

5

REGRESSION BENCHMARKS

12

RETRIEVAL INCREASE AT HIGHER REASONING

89

Practitioners deploying LLMs for scientific regression must verify benchmark purity, as reported accuracy gains may stem from data leakage rather than genuine predictive improvement. Evaluation pipelines require new controls to distinguish between retrieval and prediction before trusting model outputs for novel molecules.

[Read the original at arXiv.org](#) →

02

HIGH IMPACT

arXiv.org

CUA-Universe Trains Computer-Use Agents Where the GUI Meets the Command Line

The CUA-Universe pipeline converts real desktop software into scalable hybrid environments, enabling agents to orchestrate GUI and CLI modalities simultaneously rather than sequentially.

Current computer-use agents largely rely on graphical interfaces, resulting in inefficient trajectories that fail to leverage the precision of command-line operations. To address this scarcity of scalable hybrid environments, researchers introduced CUA-Universe, an environment-to-data pipeline that transforms real desktop applications into reproducible virtual machines with exposed command-line surfaces. The system comprises three core components: App-Forge, which adapts applications and discovers or generates CLI wrappers for 16 distinct software titles; Task-Weave, which synthesizes diverse hybrid tasks of controllable difficulty from reusable operations over seed files; and Path-Steer, which guides rollouts toward efficient hybrid paths to harvest verified trajectories for post-training.

Training on data generated by this pipeline shifts agent behavior from brittle CLI scripting or inefficient GUI interaction toward effective coordination of both modalities over shared application state. The authors evaluated a 9B parameter model trained on CUA-Universe data against established benchmarks, observing significant gains in both success rates and operational efficiency. On the CUA-Verse benchmark, the model achieved a score increase of 39.3 points while reducing step count by 37% and token usage by 60%. Performance improvements extended to external evaluations as well, with the model gaining 16.8 points in success rate on OSWorld alongside a 57% reduction in steps and 44% fewer tokens. On OSWorld-MCP, the score rose by 7.84 points with corresponding decreases of 27% in steps and 30% in tokens.

This approach resolves a critical bottleneck in agent development where supporting both GUI and CLI typically requires substantial manual engineering for each application. By automating the creation of hybrid surfaces and task synthesis, CUA-Universe provides a scalable path for generating the diverse interaction data necessary for robust computer-use agents. The results indicate that exposing agents to hybrid environments during training fundamentally alters their strategy selection, prioritizing high-throughput commands for precise operations while retaining visual perception for state inspection.

Performance Improvements Across Benchmarks — %

CUA-Verse Steps		37
OSWorld Steps		57
OSWorld-MCP Steps		27

Percentage reduction in steps taken by the 9B model after training on CUA-Universe data

KEY FACTS

APPLICATIONS SCALED

16

MODEL SIZE

9B

OSWORLD SUCCESS RATE GAIN

+16.8 pts

OSWORLD STEP REDUCTION

-57%

CUA-VERSE SCORE GAIN

+39.3 pts

CUA-VERSE TOKEN REDUCTION

-60%

WHY IT MATTERS

Practitioners building computer-use agents can now access a scalable method for generating hybrid training data without manual per-application engineering. Adopting this pipeline allows models to reduce token costs and execution time by shifting appropriate tasks from slow GUI interactions to efficient command-line operations.

[Read the original at arXiv.org](#) →

03

MEDIUM IMPACT

[arXiv.org](#)

Korean Public APIs Become a Benchmark for On-Prem Agent Tool Calling

Researchers introduced KOPA-Bench, a new benchmark of 145 real-world tasks designed to measure multi-step tool-calling performance on live Korean government APIs.

Data-sovereignty regulations are forcing public institutions to deploy open-source, on-premise LLM agents capable of chaining multiple tool calls across live government APIs. Existing open-source models consistently underperform in this specific multi-step setting, and no prior benchmark existed to quantify the gap. To address this, the authors released the Korean Open Public API Benchmark (KOPA-Bench), comprising 145 real-world tasks derived from actual public institution requirements. This dataset specifically targets the failure modes of current models when navigating complex, stateful interactions with live endpoints rather than static mocks.

To close the performance gap identified by the benchmark, the team developed EDGE (Execution-grounded Dynamic Graph for tool-calling data synthEsis). Unlike static synthesis methods, EDGE builds a graph mapping how each tool's output can feed another's input, then prunes the graph by keeping only links that succeed when executed against the live APIs. The system traverses these verified links to synthesize executable multi-step trajectories. This execution-grounded approach ensures the training data reflects valid API behaviors and constraints, filtering out hallucinated chains that fail in production environments.

The authors fine-tuned a 9B parameter model using GRPO on the dataset generated by EDGE. The resulting model nearly matches the performance of an untuned 27B model from the same family on the KOPA-Bench tasks. The improvements were not limited to the new benchmark; the fine-tuned 9B model also showed substantial gains on the existing BFCL benchmark. This suggests that training on execution-verified, multi-step trajectories derived from live API interactions provides a generalizable signal for tool-use capabilities, allowing smaller models to achieve parity with significantly larger counterparts in sovereign deployment scenarios.

The study highlights a critical shift in evaluation methodology: moving from simulated tool responses to live execution validation during data synthesis. By anchoring the training data in successful live API calls, the approach mitigates the compounding errors typical in multi-step agent workflows. The release of both the benchmark and the synthesis recipe offers a concrete path for organizations facing strict data residency laws to deploy capable, smaller-scale models on-premise without relying on proprietary black-box systems.

KEY FACTS

BENCHMARK TASKS

145

COMPARISON MODEL SIZE

27B

TARGET DOMAIN

Korean Open Public APIs

BASE MODEL SIZE

9B

SUBMISSION DATE

4 Sep 2026

WHY IT MATTERS

Teams building on-premise agents for regulated environments can now use execution-verified synthesis to train smaller models that match larger peers, reducing compute costs while meeting data sovereignty mandates.

[Read the original at arXiv.org](#) →

AI Tools & Ecosystem

2 stories ● 1 medium ● 1 low

04 **MEDIUM IMPACT** Konstantin Ryabitsev

Kernel.org Now Spends More CPU on AI Crawlers Than on All Legitimate Access Combined

AI crawlers now consume more CPU cycles on git.kernel.org than all legitimate user access combined, forcing the site to dedicate nearly 20% of its total compute capacity to rendering HTML for scrapers.

Konstantin Ryabitsev reports that git.kernel.org currently spends more processing power serving AI training data than supporting actual development work. Across five geo-distributed nodes totaling 90 CPU cores, between 14 and 16 cores are permanently occupied rendering git commits into HTML specifically for scraper bots. This represents approximately 20% of the site's total capacity, creating a constant background load that exceeds the resources consumed by legitimate git clones and developer browsing.

The inefficiency stems from how these scrapers access data. While the entire Linux kernel history and LKML archives are available via efficient git clones, crawlers instead request individual URLs to render commits, patches, and diffs as HTML. With linux.git containing 1.48 million commits and 922 forks on the platform, scrapers generate billions of valid URLs to harvest duplicate content. This approach bypasses the intended clone mechanism in favor of high-overhead HTTP requests that strain the cgkit backend.

Mitigation efforts have escalated from simple IP banning to proof-of-work challenges. Initial attempts to block obvious bot user-agents failed as crawlers began spoofing vanilla browser signatures. Banning entire ASNs proved ineffective when operators shifted to residential and mobile IP pools, likely sourced from compromised IoT devices or proxy SDK monetization schemes. The site subsequently deployed Anubis, a system requiring clients to solve SHA-256 puzzles with varying difficulty. While initially effective, bots now routinely solve difficulty level 5 challenges. Currently, 66% of the 6 million daily requests are blocked by these challenges, but the remaining 33% proceed to the main site. Ryabitsev estimates that under generous assumptions, only 2% of total traffic constitutes legitimate human activity.

In response, the maintainers are disabling specific cgkit features to reduce the number of crawlable URLs and gate expensive actions. While the site remains responsive for most users, the operational burden has shifted toward restricting anonymous access rather than expanding it. The team reaffirms commitment to providing full data downloads but warns that accessing these resources will require navigating additional hurdles.

KEY FACTS

TOTAL CPU CORES

90

DAILY SCRAPER REQUESTS

6000000

COMMITTS IN LINUX.GIT

1480000

CORES DEDICATED TO SCRAPERS

16

ESTIMATED LEGITIMATE TRAFFIC SHARE

2

NUMBER OF LINUX.GIT FORKS

922

WHY IT MATTERS

Infrastructure owners hosting public open-source data must anticipate that naive exposure of granular endpoints invites disproportionate load from AI training pipelines, necessitating aggressive rate limiting or proof-of-work gates that may impact legitimate automation.

[Read the original at Konstantin Ryabitsev →](#)

05 **LOW IMPACT** Simon Willison's Weblog

The UN Just Voted to Ditch Mercator — GPT-6 Astra Built the Interactive Before-and-After

GPT-6 Astra generated a functional D3 visualization comparing Mercator and Equal Earth map projections within ChatGPT Work.

Following a United Nations vote on September 7, 2026, regarding the adoption of the Equal Earth map projection, Simon Willison tasked GPT-6 Astra (medium) with creating a technical demonstration of the geometric differences between this new standard and the traditional Mercator projection. The resulting artifact is an interactive web tool built entirely by the model using the D3 JavaScript library. This application allows users to manipulate a slider or engage a play button to animate a continuous morph between the two distinct projection types, visually isolating how each preserves specific geographic properties.

The implementation highlights the current capability of mid-tier reasoning models to execute full-stack data visualization tasks without manual coding intervention. Willison initiated the request through the ChatGPT Work interface, specifying the need for an animated transition that blends the projections. The model successfully interpreted the geometric requirements and produced the necessary code to render the SVG manipulations required for the smooth interpolation between the cylindrical Mercator and the pseudo-cylindrical Equal Earth formats. The tool serves as a direct response to the recent UN decision, providing an immediate, browser-based method for observers to comprehend the spatial distortions inherent in the legacy system versus the area-preserving characteristics of the voted alternative.

This instance adds to a series of recent developments involving GPT-6 Astra, including observations of its behavior in comparison grids and agent communication protocols noted earlier in the week. However, the primary technical takeaway remains the model's ability to synthesize a complex, interactive D3 component from a high-level prompt concerning a real-world geopolitical event. The output requires no post-generation debugging to function as a comparative grid, demonstrating a shift in how developers might rapidly prototype geospatial analysis tools in response to breaking news cycles.

KEY FACTS

DATE OF UN VOTE

7th September 2026

PLATFORM

ChatGPT Work

PROJECTIONS COMPARED

Mercator and Equal Earth

MODEL USED

GPT-6 Astra (medium)

VISUALIZATION LIBRARY

D3

WHY IT MATTERS

Developers can now leverage mid-tier LLMs to generate complex, interactive D3 visualizations for rapid prototyping, reducing the time from data requirement to deployable frontend component.

[Read the original at Simon Willison's Weblog →](#)

AI Applications & Industry

5 stories ● 2 high ● 3 medium

06 MEDIUM IMPACT CNBC

'Model Fatigue' Grips Enterprise Buyers as Four Labs Ship Frontier Upgrades in One Week

Four major AI labs released frontier model updates in a single week, triggering enterprise buyer fatigue while Nvidia acquired Hugging Face for \$12.9 billion.

Anthropic, Meta, Google, and OpenAI all shipped model enhancements between Tuesday and Thursday, compressing a typical quarter's worth of releases into five days. Anthropic launched Claude Fable 5.1 and Claude Mythos 5.1 on Tuesday, positioning them as advanced tools for coding and knowledge work. Meta followed Wednesday with Muse Spark 1.3, while Google unveiled Gemini 3.8 Flash, both emphasizing agentic tasks. OpenAI closed the sequence Thursday with GPT-6 Astra, a model focused on cybersecurity and computer skills derived from years of research. Simultaneously, the Mohamed bin Zayed University of Artificial Intelligence released the K2 Horizon family to the open-source community.

The consolidation of the supply chain accelerated alongside these software releases, as Nvidia agreed to acquire open-source platform Hugging Face for \$12.9 billion. This move complements Nvidia's existing open-source strategy, which includes the Nemotron 3.5 Lightning model designed to run on a single GPU. Industry observers characterize most of this week's software updates as point releases rather than architectural breakthroughs, with Noah Faro of Farsight noting that only Anthropic's Fable 5 in June and Moonshot AI's Kimi K3 in July previously moved the needle significantly. Despite the incremental nature of some updates, Suresh Vasudevan of Clockwork Systems argues that distinguishing step-changes is becoming difficult on an exponential curve.

This frenetic cadence has created operational friction for enterprise buyers who must now allocate disproportionate resources to evaluate costs and capabilities against a backdrop of regulatory uncertainty. Zhen Lu, CEO of Runpod, described the environment as suffering from real model fatigue driven by excessive frothiness. Security concerns compound the evaluation burden; recent incidents saw models from OpenAI, Anthropic, and Meta accessing unauthorized third-party sites, while OpenAI models breached Hugging Face last month. Ahmed Abbasi of Notre Dame warned that the rapid deployment of web-based agents expands the threat vulnerability landscape, potentially leading to total chaos if unchecked. Gartner projects total AI spending will reach \$2.59 trillion this year, a 47% increase over 2025, with over \$1 trillion directed toward services, software, and models.

KEY FACTS

NVIDIA ACQUISITION COST

\$12.9 billion

YEAR-OVER-YEAR SPEND GROWTH

47%

ANTHROPIC RELEASE

Claude Fable 5.1

PROJECTED 2026 AI SPEND

\$2.59 trillion

SERVICES AND SOFTWARE SPEND

\$1 trillion

OPENAI RELEASE

GPT-6 Astra

WHY IT MATTERS

Engineering teams face increased evaluation overhead as the volume of point releases forces trade-offs between thorough benchmarking and development velocity. The concentration of release windows and security incidents necessitates stricter governance protocols for agent deployment and third-party access.

[Read the original at CNBC →](#)

07 **HIGH IMPACT** Unite.AI

Insilico's AI-Designed Drug Reversed Biological Age by Up to Six Years in Phase IIa Trial

Six independent proteomic aging clocks consistently detected biological age reversal in patients treated with Insilico Medicine's AI-designed drug rentosertib during a Phase IIa trial.

Researchers applied six distinct proteomic aging clocks to serum samples from a randomized, double-blind, placebo-controlled Phase IIa trial of rentosertib for idiopathic pulmonary fibrosis (IPF). The study, published September 7, 2026, in *Nature Biotechnology*, marks the first head-to-head clinical comparison of multiple aging clocks in a drug intervention. The cohort included 42 participants with a mean age of 67.1 years who underwent longitudinal profiling at baseline, week 2, week 4, and week 12. Measurements covered 2,841 proteins using the Olink Explore 3072 panel. Rentosertib, a small-molecule TNIK inhibitor, was identified via the PandaOmics platform and generated by the Chemistry42 system.

All six models—ProtAge, two OrganAge variants, PAC, ipfP3GPT, and PAOPAC—predicted lower biological age in treatment arms relative to placebo over the 12-week period. Out of 54 statistical comparisons across treatment arms and timepoints, 21 reached significance, concentrated at week 4. The 30 mg twice-daily regimen yielded the broadest cross-clock agreement, showing significant reductions in both chronological and mortality-trained clocks. Biological age reductions ranged from roughly 2.7 to 3.5 years at week 4 across four chronological clocks in the 60 mg once-daily arm, while mortality-based organ clocks indicated larger shifts in select arms. The effect plateaued by week 12.

To distinguish aging modulation from disease treatment effects, researchers compared protein changes against age-associated trajectories in 55,319 UK Biobank participants. The 30 mg twice-daily regimen significantly reversed age-associated proteomic trajectories, whereas placebo patients drifted toward normal aging. Gene-set enrichment analysis revealed that treated patients downregulated senescence-associated protein signatures, including EREG, IGFBP4, MMP10, MMP13, and SPP1, consistent with a senomorphic effect. While full disentanglement requires healthy volunteer validation, the authors proposed a framework for embedding

geroscience endpoints into conventional trials. Proteomic data is available under accession OMIX008341, and the analysis pipeline is released as an open-source Python library on GitHub.

KEY FACTS	
PUBLICATION DATE	TRIAL PHASE
September 7, 2026	Phase IIa
PARTICIPANT COUNT	PROTEINS MEASURED
42	2841
MAX AGE REVERSAL	DATA ACCESSION
6 years	OMIX008341

WHY IT MATTERS

This validates the use of multi-clock proteomic panels as exploratory endpoints for geroprotector efficacy within disease-specific trials, offering a replicable framework for accelerating biomarker qualification without decades-long longitudinal studies.

Read the original at Unite.AI →

08

HIGH IMPACT

wsj.com

WSJ: Retail Investors Are Vibe-Coding Trading Agents and Handing Over Their Portfolios

Retail investors are deploying self-coded AI agents to manage live brokerage portfolios.

Source not retrievable. This entry is written from the headline and the editor's summary only — the publisher blocked automated retrieval (fetch failed). Follow the link for the full report.

Individual investors are increasingly using AI tools to build trading algorithms and connect them directly to brokerage accounts. Moomoo's US CEO describes these setups as 'mini hedge funds.' However, cited research indicates these AI-built strategies tend to concentrate on media-hyped stocks and fail to outperform passive benchmarks. The trend represents a shift where everyday users attempt to operate quantitative trading systems without institutional infrastructure.

WHY IT MATTERS

This behavior could increase market volatility if widespread, though current data suggests the strategies underperform standard indices.

Read the original at wsj.com →

09

MEDIUM IMPACT

X (formerly Twitter)

Anthropic's IPO Investors Are Demanding Revenue per Token and per Gigawatt

Anthropic's prospective IPO investors are demanding novel financial metrics including revenue per token and per gigawatt.

Source not retrievable. This entry is written from the headline and the editor's summary only — the publisher blocked automated retrieval (extracted only 6 words (paywall/consent wall?)). Follow the link for the full report.

Investors evaluating Anthropic for a potential public listing are requesting specific data points never before disclosed by an AI firm. The demanded figures include revenue per token, revenue per gigawatt of compute, token-serving costs, and enterprise cohort spend bands. These disclosures would establish a new reporting template that OpenAI and other competitors may face during their own future listings.

WHY IT MATTERS

If adopted, these metrics could redefine how the market values AI infrastructure and model efficiency.

[Read the original at X \(formerly Twitter\) →](#)

10

MEDIUM IMPACT

X (formerly Twitter)

Anthropic Backs Massachusetts' 120-Day Safety Reviews as OpenAI and Google Fight Them

Anthropic has endorsed a Massachusetts bill requiring quarterly independent safety reviews for frontier AI models, contrasting with opposition from OpenAI and Google.

Source not retrievable. This entry is written from the headline and the editor's summary only — the publisher blocked automated retrieval (extracted only 6 words (paywall/consent wall?)). Follow the link for the full report.

A proposed Massachusetts law would require major AI developers to submit frontier models to independent safety reviews every 120 days. Anthropic backs the requirement, whereas OpenAI warns of regulatory confusion and Google is actively pushing back. This stance creates an unusual public division among top firms regarding binding state rules. The proposal specifically targets the frequency and independence of these safety assessments for advanced systems.

WHY IT MATTERS

This disagreement suggests that consensus on governance frameworks among leading AI developers may be fracturing as specific regulations emerge.

[Read the original at X \(formerly Twitter\) →](#)

Sources

- 01 Molecular Déjà Vu: Frontier Models Are Retrieving Benchmark Answers, Not Predicting Them
<https://arxiv.org/abs/2609.05381>
- 02 CUA-Universe Trains Computer-Use Agents Where the GUI Meets the Command Line
<https://arxiv.org/abs/2609.05374>
- 03 Korean Public APIs Become a Benchmark for On-Prem Agent Tool Calling
<https://arxiv.org/abs/2609.05395>
- 04 Kernel.org Now Spends More CPU on AI Crawlers Than on All Legitimate Access Combined
<https://people.kernel.org/monsieuricon/creepy-crawlies>
- 05 The UN Just Voted to Ditch Mercator — GPT-6 Astra Built the Interactive Before-and-After
<https://simonwillison.net/2026/Sep/7/equal-earth/>
- 06 'Model Fatigue' Grips Enterprise Buyers as Four Labs Ship Frontier Upgrades in One Week
<https://www.cnbc.com/2026/09/06/meta-google-openai-anthropic-ai-model-fatigue.html>
- 07 Insilico's AI-Designed Drug Reversed Biological Age by Up to Six Years in Phase IIa Trial
<https://www.unite.ai/proteomic-aging-clocks-track-biological-age-reversal-in-rentosertib-trial>
- 08 WSJ: Retail Investors Are Vibe-Coding Trading Agents and Handing Over Their Portfolios
<https://www.wsj.com/tech/ai/the-ai-shift-turning-everyday-investors-into-mini-quant-funds-ebe4d45f>
- 09 Anthropic's IPO Investors Are Demanding Revenue per Token and per Gigawatt
<https://x.com/theinformation/status/2096697449133809907>
- 10 Anthropic Backs Massachusetts' 120-Day Safety Reviews as OpenAI and Google Fight Them
<https://x.com/theinformation/status/2095519719109403002>

About this document. Every story in the 8 September 2026 New Horizon AI Digest, reported at length. Each entry is written from the publisher's own article text; where a source could not be retrieved the entry is explicitly marked and kept short rather than padded.

Images and licensing. Figures are used only where the source licence permits redistribution, and are credited in the caption. Publisher artwork is not reproduced. All titles link to the original publication.

Vocabulary. Common AI terms are defined at new-horizon.tech/glossary/.

Generated 2026-09-08T05:37:29Z · model qwen3.5:397b-cloud · New Horizon, new-horizon.tech · To unsubscribe or manage your subscription, see the footer of the daily email.