

The Long Read

26 August 2026

12

STORIES

3

SECTIONS

3481

WORDS

5

HIGH IMPACT



● 5 high impact ● 7 medium 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

12 stories · 3481 words · 3 sections · 5 charted



AI Models & Research 4 stories · 1212 words

AI Tools & Ecosystem 4 stories · 1120 words

AI Applications & Industry 4 stories · 1149 words

AI Models & Research

- 01 ● IBM Granite 4.2: Open-Weight Reasoning Models with Agentic RL Ship Under Apache 2.0
- 02 ● Quantization-Aware Healing: Compressed 4-Bit Model Outperforms Its Full-Precision Original
- 03 ● Recuris: Recursive Memory Architecture for Long-Horizon Agent Harnesses
- 04 ● BrowserForge: Parallel Browser Sandboxes Train Web Agents at Scale

AI Tools & Ecosystem

- 05 ● OpenAI's Jalapeño Chip Benchmarks Show 1.5-1.9x Efficiency Over Nvidia Superchips

- 06 ● Keenable Exits Stealth with \$26M to Index the Web for AI Agents
- 07 ● Claude Cowork Finally Remembers What You Told the App in Chat
- 08 ● Comparing Local Tool Calling: Gemma 4 vs. Llama 3 vs. Mistral

AI Applications & Industry

- 09 ● Stability AI Raises \$76M from Universal, Sony, and EA for Creative AI Tools
- 10 ● Robotics Startup Generalist Hits \$3B Valuation in Series B Extension
- 11 ● OpenAI Loses Data Center Chief as Executive Departures Mount Ahead of IPO
- 12 ● India's Ringg Gets Peak XV Backing to Push Voice AI Beyond the Phone Call

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

4 stories ● 4 high

01 HIGH IMPACT huggingface.co

IBM Granite 4.2: Open-Weight Reasoning Models with Agentic RL Ship Under Apache 2.0

IBM has released Granite 4.2, a family of open-weight reasoning models that integrate agentic reinforcement learning directly into the training pipeline for real-world tool use.

The Granite 4.2 family comprises three dense, decoder-only models sized at 3B, 8B, and 30B parameters, all released under the Apache 2.0 license. Each model is pre-trained from scratch on approximately 15 trillion tokens using a five-phase strategy that extends the context window to 512K tokens. The architecture utilizes Grouped Query Attention with 8 KV heads and Rotary Position Embeddings with a theta of 10,000,000. Unlike previous iterations focused primarily on instruction following, these models feature a native thinking switch that toggles between standard response generation and a low-effort reasoning mode optimized for simpler queries.

The distinguishing feature of the 8B and 30B variants is a multi-stage reinforcement learning pipeline that includes a dedicated agentic block. After supervised fine-tuning on 7.2 million samples, these larger models undergo sequential GRPO training stages for software engineering, terminal operation, and web search. This process teaches the models to execute multi-turn trajectories in real sandboxed environments, such as editing code repositories via OpenHands or running shell commands through Terminus-2. The 3B model follows the same foundational RL curriculum but omits the agentic block, limiting it to verifiable rewards in math and code without live environment interaction.

Benchmark results indicate significant gains in agentic capabilities for the larger models. On SWE Bench Verified, the 30B model achieves a score of 57.00 compared to 47.67 for the 8B variant, while the 3B model has no reported score for this task. In general reasoning, the 30B model reaches 89.17 on AIME25 and 66.41 on GPQA. All models support native tool calling formatted for OpenAI-compatible endpoints and are available in quantized formats including FP8, NVFP4, MXFP4, and GGUF for efficient deployment.

Infrastructure for this release relies on NVIDIA GB200 NVL72 clusters hosted by CoreWeave, utilizing NeMo-RL for asynchronous training loops and NeMo-Gym to orchestrate diverse rollout environments. This setup allows the generation workers to sample responses from live sandboxes without blocking the trainer, enabling scalable RL across complex agent tasks.

SWE Bench Verified Scores — %

8B Dense		47.67
30B Dense		57

Agentic coding resolution rates by model size · +20%

KEY FACTS

MODEL SIZES

3B, 8B, 30B

CONTEXT WINDOW

512K tokens

SFT SAMPLES

7.2M

LICENSE

Apache 2.0

PRE-TRAINING DATA

15T tokens

SWE BENCH VERIFIED (30B)

57.00

WHY IT MATTERS

Developers can now deploy fully open-weight models capable of executing multi-step tool use in real environments without proprietary API dependencies. The inclusion of agentic RL in the base training rather than as a post-hoc fine-tune suggests more robust performance for autonomous coding and system administration tasks.

[Read the original at huggingface.co](#) →

02

HIGH IMPACT

[huggingface.co](#)

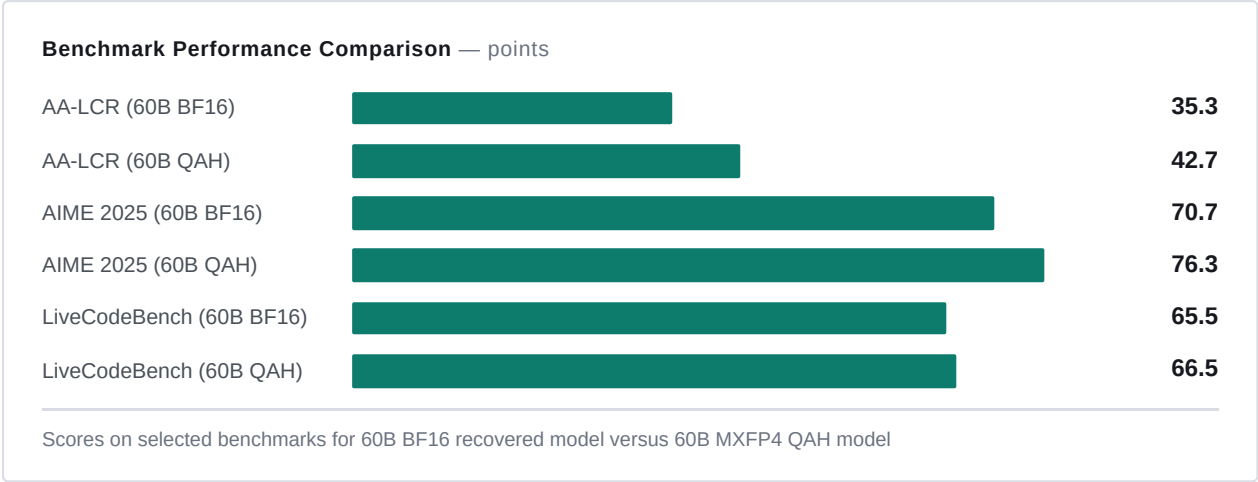
Quantization-Aware Healing: Compressed 4-Bit Model Outperforms Its Full-Precision Original

Quantization-Aware Healing (QAH) enables a compressed 4-bit model to outperform its own full-precision bfloat16 counterpart by distilling directly from the original uncompressed teacher.

Multiverse Computing introduces Quantization-Aware Healing (QAH), a method that recovers accuracy in large language models after both structural compression and low-bit quantization. Applied to a GPT-OSS 120B model compressed to 60B parameters and quantized to MXFP4, QAH produces a checkpoint that exceeds the performance of the recovered 60B bfloat16 version on seven of nine benchmarks. The approach inverts the standard efficiency trade-off: the resulting 4-bit model is smaller, cheaper to run, and more accurate than the full-precision checkpoint from which it was derived.

Standard healing pipelines rely on quantization-aware training (QAT) or distillation from a recovered checkpoint, both of which struggle when the architecture has changed structurally. QAT requires expensive re-training through noisy forward passes and risks instability if training extends past an optimal point. Distillation from a recovered bfloat16 student anchors the quantized model to a degraded target, capping its potential accuracy. QAH bypasses these limits by distilling logits directly from the frozen, pre-compression 120B teacher into the 60B MXFP4 student. This treats quantization not as a lossy post-processing step but as a second pass of distillation, allowing the student to acquire information the initial recovery stage missed.

The method also delivers significant stability gains. In head-to-head tests on a 9B variant, QAH reached peak average performance across MMLU-Pro, LiveCodeBench, and GPQA Diamond in roughly 100 steps, seven times faster than QAT's 700 steps. While QAT collapsed by nearly 19 points after its peak due to drift under cross-entropy loss, the QAH model held within two points of its maximum throughout training. For the 60B model, the largest gains appeared in long-context reasoning (+7.4 on AA-LCR) and math (+5.6 on AIME 2025). The 4-bit QAH model even surpassed the original 120B teacher on LiveCodeBench (66.5 vs 66.0) while using roughly one-quarter of the weight memory.



KEY FACTS

SOURCE MODEL

GPT-OSS 120B

QUANTIZATION FORMAT

MXFP4

AA-LCR GAIN

+7.4

COMPRESSED SIZE

60B parameters

BENCHMARKS WON VS BF16

7 of 9

TRAINING STEPS TO PEAK

100

WHY IT MATTERS

Practitioners can now deploy structurally compressed, 4-bit models that are demonstrably more accurate than their full-precision sources, removing the accuracy tax traditionally associated with aggressive quantization. The stability of KL-divergence distillation against a frozen teacher also eliminates the operational risk of over-training during the healing phase.

[Read the original at huggingface.co](#) →

03

HIGH IMPACT

arXiv.org

Recuris: Recursive Memory Architecture for Long-Horizon Agent Harnesses

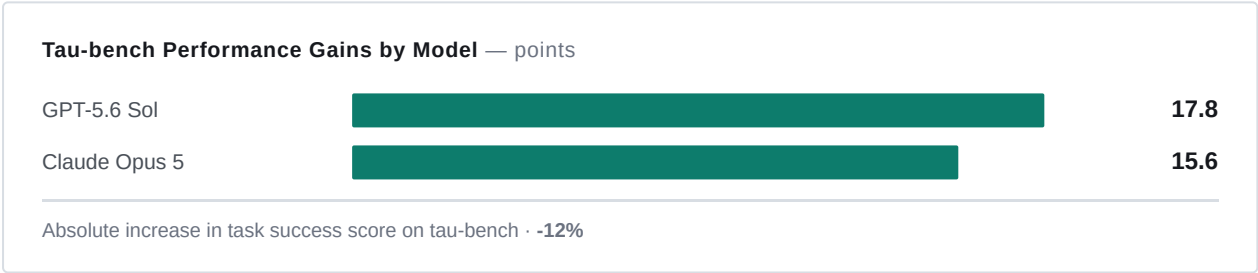
The Recuris architecture boosts frontier model performance on long-horizon tasks by decoupling task progress tracking from historical skill storage.

Submitted to arXiv on 25 Aug 2026, the paper introduces Recuris, a recursive Experiential-Working Memory architecture designed to solve state obscurity in long-horizon agent tasks. The system separates concerns: a

Working Memory component tracks immediate task progress and guides skill selection, while an Experiential Memory stores historical skills. This division grounds skill invocation in current needs rather than forcing the agent to parse full execution histories. A fixed Meta-Agent converts execution outcomes into structured evidence, localizing failures to specific memory components. This evidence triggers validation-gated updates to Skill Memory, creating a bounded recursive loop that evolves execution strategies without unbounded context growth.

Evaluation covered four long-horizon benchmarks and ten distinct models. Recuris improved task success in 35 of the 37 completed model-benchmark pairs. On the tau-bench, the architecture added 17.8 points to GPT-5.6 Sol and 15.6 points to Claude Opus 5, pushing the latter to an 87.9% success rate. Performance gains were also recorded on SkillFlow, with increases of 16.6 points for Qwen3.6-27B and 13.5 points for Qwen3.6-35B. The advantage scales with task complexity; on the longest interaction horizons, the improvement reached 32.2 points. Common failure modes associated with long-context degradation dropped by up to 80%.

Unlike standard context-window extensions or static retrieval-augmented generation, Recuris implements an active evolution loop where execution evidence directly reshapes memory contents. The approach treats memory not as a passive store but as a mutable component updated by a meta-agent based on localized failure analysis. The authors position this recursively evolving memory as a scalable foundation for Recursive Self-Improvement (RSI), enabling agents to transform accumulated experience into effective behavior without the computational penalty of expanding history windows. Code for the implementation is available via the provided repository link.



KEY FACTS	
SUBMISSION DATE	MODELS TESTED
25 Aug 2026	10
BENCHMARKS USED	SUCCESS RATE IMPROVEMENT PAIRS
4	35 of 37
TAU-BENCH GAIN (GPT-5.6 SOL)	MAX FAILURE REDUCTION
+17.8	80%

WHY IT MATTERS

This architecture offers a viable path to reliable long-horizon autonomy without requiring ever-larger context windows, allowing practitioners to deploy complex agents on existing hardware constraints. The ability to localize and patch memory failures automatically reduces the engineering overhead of debugging multi-step agent trajectories.

Read the original at [arXiv.org](#) →

BrowserForge: Parallel Browser Sandboxes Train Web Agents at Scale

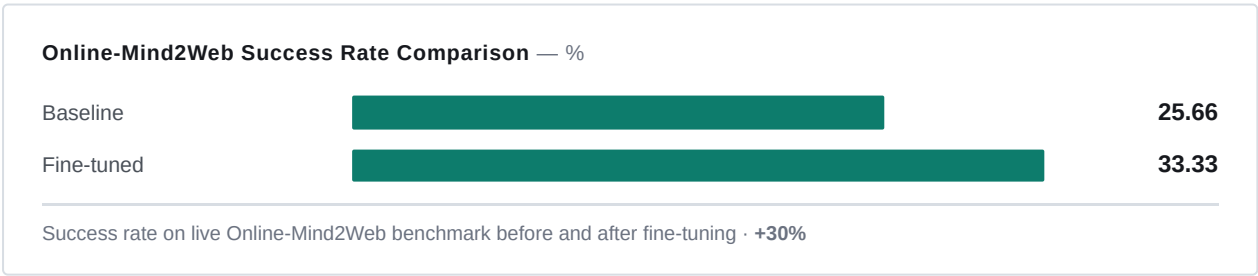
BrowserForge scales web agent training by generating 203,238 verified interaction trajectories across hundreds of thousands of unique, open-web websites.

Current web agent training relies on datasets with only a few thousand trajectories drawn from narrow, fixed site lists, limiting generalization. BrowserForge addresses this data bottleneck by deploying a framework that drives hundreds of concurrent browser sandboxes over the open web. The system couples an open-web sourcing stage, a sandbox cluster manager for high utilization, and a Proposer-Solver dual-agent loop. This loop converts raw pages into executable tasks and collects verified trajectories, ensuring the agent encounters a vastly broader range of real-world domains than previous automated synthesis pipelines allowed.

The resulting corpus contains 203,238 trajectories, each collected from a distinct website. A rule-plus-model cleaning pipeline filters failed runs and rewrites surviving reasoning into a unified chain-of-thought style. Crucially, while page structure like the accessibility tree serves as a synthesis-time signal, the final trained agent acts purely from rendered screenshots. This approach avoids the fragility and high token costs associated with parsing HTML or accessibility trees directly during inference, focusing instead on visual perception.

Fine-tuning a compact multimodal model on this corpus yields measurable gains on established benchmarks. Success rates on the live Online-Mind2Web benchmark rise from 25.66% to 33.33%. The model also shows consistent improvements in step accuracy on the static Multimodal-Mind2Web dataset, with performance gains scaling alongside the corpus size. Controlled analyses confirm that open-web sourcing and broad website coverage are the primary drivers of these improvements, validating the hypothesis that diversity in training data outweighs mere volume from restricted sources.

This release shifts the focus from architectural tweaks to data scalability for web agents. By demonstrating that parallel sandboxing can efficiently harvest high-quality data from the unrestricted web, BrowserForge provides a pathway to train more robust agents without relying on brittle structural inputs. The availability of a corpus larger and more diverse than prior datasets offers an immediate resource for practitioners seeking to improve agent reliability on unseen domains.



KEY FACTS	
TRAJECTORIES GENERATED	ONLINE-MIND2WEB BASELINE SUCCESS
203,238	25.66%
ONLINE-MIND2WEB FINE-TUNED SUCCESS	SUBMISSION DATE
33.33%	25 Aug 2026

Practitioners can now fine-tune screenshot-based agents on a corpus of over 200,000 diverse trajectories, bypassing the need to parse HTML or accessibility trees. The demonstrated 7.67 percentage point gain on live benchmarks suggests that scaling data diversity via parallel sandboxes is a more effective lever than current synthetic data methods.

[Read the original at arXiv.org](#) →

AI Tools & Ecosystem

4 stories ● 1 high ● 3 medium

05 MEDIUM IMPACT TechCrunch

OpenAI's Jalapeño Chip Benchmarks Show 1.5-1.9x Efficiency Over Nvidia Superchips

OpenAI's custom Jalapeño inference chip delivers 1.5 to 1.9 times greater efficiency per kilowatt than current Nvidia Blackwell systems.

At the Hot Chips conference, OpenAI released the first benchmark data for Jalapeño, its custom inference processor developed in collaboration with Broadcom. Tested on SemiAnalysis' InferenceX benchmark, the system demonstrated a 1.5 to 1.9-fold improvement in throughput per kilowatt and tokens per user compared to state-of-the-art Nvidia Blackwell hardware. Richard Ho, OpenAI's head of hardware, characterized the results as a significant advance capable of serving higher AI workloads per power unit while maintaining low latency responses.

The architecture prioritizes minimizing friction during specific inference phases, particularly prefill and communication steps that typically create bottlenecks. By explicitly placing model state, including the KV cache, locally within the system, Jalapeño reduces data movement delays. This full-stack approach allows the hardware to activate precise combinations of compute, memory, and networking resources tailored to each phase of the generation process, addressing inefficiencies inherent in general-purpose GPU clusters.

Deployment timelines indicate a phased rollout rather than immediate availability. Ho estimated that Jalapeño would enter production in very small volumes by the end of 2026, with significant deployment scheduled for 2027. This schedule suggests that while the current benchmarks exceed existing Nvidia offerings, competing architectures may evolve before Jalapeño reaches scale. OpenAI intends to treat Jalapeño as a multigenerational platform, coordinating future developments across models, chips, and memory systems to sustain these efficiency gains.

KEY FACTS	
EFFICIENCY GAIN	BENCHMARK SUITE
1.5-1.9x	SemiAnalysis InferenceX
COMPARISON TARGET	INITIAL DEPLOYMENT
Nvidia Blackwell	End of 2026
SCALE DEPLOYMENT	HARDWARE PARTNER
2027	Broadcom

Practitioners planning large-scale inference infrastructure for 2027 must account for a potential shift in the performance-per-watt baseline, as custom silicon begins to outpace general-purpose GPUs in specific latency-sensitive workloads.

[Read the original at TechCrunch →](#)

06

MEDIUM IMPACT

TechCrunch

Keenable Exits Stealth with \$26M to Index the Web for AI Agents

Former Yandex search lead Andrey Styskin has launched Keenable with \$26 million in seed funding to build a web index optimized specifically for AI agent retrieval rather than human browsing.

Keenable emerged from stealth this week securing \$26 million in seed funding led by Accel, with participation from Conviction Partners and unnamed business angels. The startup is constructing a proprietary web search index containing over 100 billion documents, designed to serve AI chatbots and agents that require grounding in source material during both training and runtime. Unlike traditional search engines optimized for human attention spans, Keenable's infrastructure targets the specific needs of large language models that can process vast portions of information simultaneously. The company currently employs 15 engineering staff across the U.S. and Europe and plans to double its headcount by the end of the year to accelerate go-to-market efforts.

The technical differentiation lies in index structure optimization for agentic queries. Styskin, who previously led search, AI, and cloud divisions at Yandex and worked on web search infrastructure for Amazon Alexa alongside co-founder Matthias Petri, argues that serving the entire internet without task-specific fine-tuning creates enormous costs. Keenable addresses this by narrowing the search space rapidly based on query intent, a method distinct from enterprise search solutions that often fail at web scale. The startup's API is already in production at several undisclosed AI labs and inference providers. Additionally, Keenable has partnered with voice AI firm Gradium to support live information retrieval and is developing a Web Query Language product to synthesize answers from multiple sources when no single document contains the full response.

This launch occurs as major tech giants like Google and Microsoft restrict access to their existing search APIs to prevent cannibalization, opting instead for bundled approaches. Accel partner Zhenya Loginov noted the scarcity of web-scale search infrastructure options for AI developers amidst these closures. While building a massive index remains "painfully expensive," Keenable posits that the innovator's dilemma makes Google beatable on agentic queries through more cost-efficient, specialized solutions. Competitors such as Brave and Exa are also entering the space, signaling a broader industry shift away from the traditional "ten blue links" model toward infrastructure built explicitly for machine consumption.

KEY FACTS

FUNDING AMOUNT

\$26 million

INDEX SIZE

100 billion documents

TARGET HEADCOUNT

30 engineering staff

LEAD INVESTOR

Accel

CURRENT TEAM SIZE

15 engineering staff

FOUNDERS

Andrey Styskin and Matthias Petri

Practitioners gain a potential alternative to restricted big-tech APIs for grounding agents in live web data, though adoption depends on whether Keenable's cost efficiency outweighs the risk of relying on a new vendor. The upcoming Web Query Language could simplify multi-hop reasoning tasks if it successfully synthesizes fragmented web sources.

[Read the original at TechCrunch →](#)

07

HIGH IMPACT

TechCrunch

Claude Cowork Finally Remembers What You Told the App in Chat

Anthropic has unified the memory architectures for Claude Chat and Claude Cowork to enable persistent context across both interfaces.

On Tuesday, Anthropic announced the merger of the memory systems powering Claude Chat and Claude Cowork, eliminating the previous silo that required users to re-brief the agent when switching between research and action modes. The update ensures that information learned in a chat session is immediately available to Cowork, allowing the agent to reference specific project details like headcount, city locations, or conference speakers without manual re-input. This architectural shift transforms the user experience from managing two separate products into interacting with a single continuous assistant.

The updated system changes how memory ingestion occurs, moving from end-of-conversation summarization to real-time topic addition. As users chat, Claude now adds topics to its memory instantly, facilitating seamless transitions even while a conversation is ongoing. Users gain direct visibility into this retained data through an interface that allows reading, editing, or deleting stored information on any topic. This feature is enabled by default on Free, Pro, and Max plans across web, desktop, and mobile platforms, though iOS and Android users must update to the latest application version to access the functionality.

Regarding data governance, the system enforces strict defaults against storing sensitive categories including health data, race, ethnicity, religious beliefs, politics, and gender identity. Users can override this restriction by toggling on "include sensitive topics in memory," which triggers a notification whenever such data is saved. Regardless of user settings, the model strictly prohibits saving government-issued IDs, Social Security numbers, criminal history, immigration status, or any content violating acceptable use policies. This granular control addresses previous friction points where agents lacked context while maintaining hard boundaries on high-risk personal identifiers.

KEY FACTS

ANNOUNCEMENT DATE

Tuesday

BLOCKED DATA TYPES

Government ID, SSN, criminal history

MOBILE REQUIREMENT

Latest app version

AFFECTED PLANS

Free, Pro, and Max

DEFAULT SENSITIVE TOPICS

Health, race, religion, politics, gender

Developers building agentic workflows can now rely on persistent context between planning and execution phases without engineering custom state-handling layers. However, integration logic must account for the new default exclusions of sensitive data categories to prevent runtime errors when those topics are required.

[Read the original at TechCrunch →](#)

08

MEDIUM IMPACT

MachineLearningMastery.com

Comparing Local Tool Calling: Gemma 4 vs. Llama 3 vs. Mistral

Gemma 4, Llama 3, and Mistral each enable local tool calling but diverge sharply in architectural integration, licensing constraints, and hardware efficiency.

Google DeepMind's Gemma 4, released April 2, 2026, treats tool calling as a native capability rather than a post-training addition. The family spans five sizes from E2B to 31B, with a 12B Unified variant added in June 2026 to bridge edge and server gaps. Architecturally, it mixes Dense and Mixture-of-Experts designs, supporting up to a 256K token context window on medium variants. Crucially, Gemma 4 includes built-in system prompt support and a configurable thinking mode, allowing developers to tune intermediate reasoning before function execution without embedding instructions in every user turn.

Meta's Llama 3 introduced reliable native tool calling with the 3.1 update, shifting from constrained generation or heavy prompt engineering to fine-tuned JSON emission. The lineup ranges from 8B to 405B parameters, though reliability scales with size: 70B and 405B models handle complex multi-tool scenarios robustly, while the 8B variant struggles with ambiguous choices. Llama 3.2 further diversified output formats, introducing Python-style syntax for 1B and 3B text models while maintaining JSON for larger vision and instruction-tuned variants. Deployment is governed by the Llama 3 Community License, permitting commercial use under 700 million monthly active users but restricting output usage for training competing systems.

Mistral AI focuses on parameter efficiency, culminating in the Mistral Small 4 release in March 2026. This model consolidates reasoning and vision capabilities into a single 119 billion parameter package using Mixture-of-Experts routing to activate only 6 billion parameters per token. While earlier versions like 0.3 required community-maintained templates for consistent multi-tool reliability, the Small family delivers capable tool use on mid-range hardware where 70B dense models are impractical. All three families integrate with local runtimes like Ollama and LM Studio, but the choice hinges on specific constraints: Gemma for edge-native agentic workflows, Llama for ecosystem breadth, and Mistral for inference speed on consumer GPUs.

KEY FACTS

GEMMA 4 RELEASE DATE

April 2, 2026

LLAMA 3 MAX PARAMETERS

405B

MISTRAL SMALL 4 ACTIVE PARAMS

6B

GEMMA 4 CONTEXT WINDOW

256K tokens

LLAMA 3 LICENSE LIMIT

700M monthly active users

MISTRAL SMALL 4 TOTAL PARAMS

119B

Selecting the correct model family dictates whether a local agent can run on embedded hardware, leverage existing framework integrations, or maximize inference speed without cloud dependency. Misalignment between model architecture and deployment context risks failed tool execution or excessive resource consumption.

[Read the original at MachineLearningMastery.com →](#)

AI Applications & Industry

4 stories ● 4 medium

09 **MEDIUM IMPACT** TechCrunch

Stability AI Raises \$76M from Universal, Sony, and EA for Creative AI Tools

Stability AI closed a \$76 million Series B round led by major entertainment licensors including Universal Music Group, Sony Music Group, Warner Music Group, and Electronic Arts.

Stability AI announced Tuesday the completion of a \$76 million Series B funding round, bringing its total capital raised to \$232 million since its 2019 founding. The investor roster diverges from typical venture capital firms, featuring direct participation from Universal Music Group, Sony Music Group, Warner Music Group, and gaming publisher Electronic Arts. AMD Ventures and Pacific Alliance Ventures also joined the round. This capital structure aligns with Stability's recent strategic pivot toward embedding generative models directly into enterprise creative workflows rather than solely distributing open weights.

The company intends to allocate the new funds toward expanding its "creative production" product suite and scaling its professional services arm. Current offerings include distinct models for music, video, and image generation. This direction follows a series of co-development agreements executed over the last year: partnerships with Universal Music and EA were finalized in October 2025, followed by a deal with Warner Music in November 2025. These arrangements grant partners influence over tool development rather than simple output licensing, signaling a shift from model-centric to workflow-centric deployment.

Legal developments continue to shape the company's operational landscape. In the United Kingdom, Stability AI largely prevailed in a copyright infringement lawsuit brought by Getty Images, where a judge ruled in the company's favor regarding the use of images for training data. However, a parallel lawsuit filed by Getty in the United States remains active. Internal governance challenges persist alongside external litigation; co-founder Cyrus Hodes filed a suit in 2023 alleging fraud in the sale of his equity stake to co-founder Emad Mostaque. Prem Akkaraju, who assumed the CEO role in 2024, framed the funding as validation of a vision where generative AI serves producers and storytellers directly.

KEY FACTS

SERIES B AMOUNT

\$76 million

KEY INVESTORS

Universal Music Group, Sony Music Group, Warner Music Group, EA

COMPANY FOUNDING YEAR

2019

TOTAL FUNDRAISING

\$232 million

CEO START YEAR

2024

Practitioners should note the shift toward co-developed, licensed enterprise tools over open-weight releases, as major content holders now hold equity stakes in the underlying infrastructure. The partial UK legal victory provides tentative precedent for training data usage, though US uncertainty remains.

[Read the original at TechCrunch →](#)

10

MEDIUM IMPACT

TechCrunch

Robotics Startup Generalist Hits \$3B Valuation in Series B Extension

Robotics startup Generalist has reached a \$3 billion valuation following a \$200 million Series B extension led by 8VC.

Generalist, founded in 2024 by former Google DeepMind researchers Pete Florence and Andy Zeng alongside ex-Boston Dynamics engineer Andrew Barry, secured nearly \$200 million in additional capital according to regulatory filings. This injection extends the company's June Series B round, which was initially announced at \$400 million led by Radical Ventures with a \$2 billion valuation. The new funding brings the total raised in this series to \$600 million, pushing the company's post-money valuation to \$3 billion. Neither Generalist nor lead investor 8VC responded to requests for comment regarding the extension.

The company is developing an AI foundation model designed to operate across various robotic hardware platforms. Its recently released Gen 1.5 model claims the ability to enable robots to master new tasks based on video demonstrations ranging from 3 to 12 seconds in duration. Current deployment remains limited; the startup is working with a handful of customers to tailor the model for specific use cases based on direct feedback. This approach contrasts with large language models trained on internet-scale data, as robotics models face constraints in available training datasets.

The capital surge positions Generalist within a competitive landscape of firms pursuing general-purpose robotics brains. Competitors include Physical Intelligence, reportedly valued at \$11 billion, and SoftBank-backed Skild AI at \$14 billion. Genesis AI was also in talks last month to raise capital at a \$3 billion valuation. While investors are betting on an imminent "ChatGPT moment" for robotics where agents perform general tasks without explicit per-task training, some venture capitalists caution that truly general robotics models may remain years away due to data scarcity limitations inherent to physical systems.

Robotics Startup Valuations — \$bn

Skild AI		14
Physical Intelligence		11
Generalist		3
Genesis AI		3

Reported or rumored valuations for generalist robotics startups

KEY FACTS

CURRENT VALUATION

\$3 billion

TOTAL SERIES B FUNDING

\$600 million

MODEL VERSION

Gen 1.5

EXTENSION AMOUNT

\$200 million

LEAD INVESTOR (EXTENSION)

8VC

DEMO DURATION RANGE

3 to 12 seconds

WHY IT MATTERS

The rapid valuation climb to \$3 billion signals strong investor conviction in foundation models for robotics despite data scarcity warnings. Practitioners should monitor Gen 1.5's few-shot learning claims as a potential shift from task-specific training pipelines to demonstration-based adaptation.

[Read the original at TechCrunch →](#)

11 **MEDIUM IMPACT** TechCrunch

OpenAI Loses Data Center Chief as Executive Departures Mount Ahead of IPO

OpenAI's head of data centers, Chris Malone, has departed following an internal reorganization that shifted infrastructure leadership away from President Greg Brockman.

Chris Malone, OpenAI's former head of data centers, left the company last week after a tenure of approximately 17 months. Malone, who previously spent nearly five years at Meta and over a decade at Google, joined OpenAI in March 2025 shortly after the launch of the Stargate Project. This \$500 million data center initiative, championed by the Trump administration, lists OpenAI alongside Oracle, Nvidia, SoftBank, and Microsoft as key partners. His exit marks a significant turnover in a role critical to the industry-wide AI infrastructure buildout. OpenAI stated to TechCrunch that it recently reorganized its infrastructure organization to support current scale and pace, asserting that a deeply experienced team remains in place with clear leadership.

The organizational shift accompanying Malone's departure involved a change in reporting lines. Per the Wall Street Journal, Malone ceased reporting directly to OpenAI President Greg Brockman and instead reported to Vice President Sachin Katti, who assumed leadership of the group. Responsibility for data center strategy is now distributed among several executives: Uday Ruddaraju leads the data center team, Brent Mayo oversees build and delivery programs, and Spas Lazarov leads all data center engineering. This fragmentation of leadership comes as the company faces scrutiny over its operational stability ahead of a potential public listing.

Malone is the latest in a string of more than a dozen executive departures in 2026, a count Business Insider tallied at 13. Recent high-profile exits include Chief Revenue Officer Denise Dresser after eight months, Chief Operating Officer Brad Lightcap, and Fidji Simo, the de facto second-in-command who stepped down due to chronic illness but remains in an advisory role. The company also disbanded its preparedness team dedicated to assessing catastrophic risks and lost Chloé Bakalar, its head of ethics. These departures occur as OpenAI pushes its expected IPO from 2026 to 2027, facing questions regarding valuation and profitability relative to its massive capital investments in lab infrastructure.

KEY FACTS

DEPARTING EXECUTIVE

Chris Malone

STARGATE PROJECT VALUE

\$500 million

REVISED IPO TARGET

2027

TENURE DURATION

March 2025 to August 2026

2026 EXECUTIVE DEPARTURES

13

WHY IT MATTERS

The redistribution of data center leadership and the loss of a principal architect during a \$500M infrastructure push signals potential execution risks for partners relying on OpenAI's capacity roadmap. Practitioners should monitor whether this leadership churn delays the delivery of the compute resources promised under the Stargate Project.

[Read the original at TechCrunch →](#)

12

MEDIUM IMPACT

TechCrunch

India's Ringg Gets Peak XV Backing to Push Voice AI Beyond the Phone Call

Voice AI startup Ringg secured a \$10 million extension to its Series A round from Peak XV Partners to expand its orchestration layer beyond simple phone calls into complex enterprise workflows.

Ringg has closed a \$10 million funding extension led by Peak XV Partners, bringing its total Series A capitalization to \$15.5 million following an initial \$5.5 million raise earlier this year. The company currently processes 20 million call attempts monthly, leveraging a strategy that shifts focus from low-complexity outbound dialing to high-value, sticky use cases such as healthcare appointment booking, e-commerce cart recovery, and fintech KYC onboarding. This pivot addresses the limitation of price-sensitive lead qualification tasks, aiming instead to own the outcome of specific business processes.

Technically, Ringg operates as an orchestration layer rather than a full-stack infrastructure provider, routing tasks to different speech recognition and generation models based on the specific use case requirements. While the founders initially attempted to build proprietary models under the DesiVocal brand, the cost structure necessitated a move up the stack to apply existing models to enterprise workflows. The system now manages interactions across 1,200 clinics for Practo and supports clients including Flipkart, Groww, PolicyBazaar, and Shell, the latter utilizing the platform for browser-based support automation alongside traditional voice channels which still constitute over 70% of volume.

The startup's go-to-market approach targets Global Capability Centers in India rather than selling directly to U.S. enterprises, positioning its agents to work alongside human support teams in offshore hubs. With a headcount of 40 employees, including 15 hired in the last three months, Ringg is recruiting forward-deployed engineers and researchers to reduce model inference costs. This expansion occurs within a crowded Indian market featuring model makers like Sarvam and Smallest.ai, and orchestration competitors such as Bolna and Blue Machines, where defensibility increasingly depends on owning the customer relationship and workflow completion rather than just providing raw model access.

KEY FACTS

FUNDING EXTENSION

\$10 million

MONTHLY CALL ATTEMPTS

20 million

PRACTO CLINIC COVERAGE

1,200

TOTAL SERIES A

\$15.5 million

VOICE BUSINESS SHARE

70%

EMPLOYEE COUNT

40

WHY IT MATTERS

For builders deploying AI agents, Ringg's pivot illustrates that economic viability in voice AI currently favors orchestration layers managing complex workflows over pure model ownership or simple dialer automation. The emphasis on partnering with Global Capability Centers suggests a near-term deployment pattern where AI augments existing human support structures rather than replacing them entirely.

[Read the original at TechCrunch →](#)

Sources

- 01 IBM Granite 4.2: Open-Weight Reasoning Models with Agentic RL Ship Under Apache 2.0
<https://huggingface.co/blog/ibm-granite/granite-4-2>
- 02 Quantization-Aware Healing: Compressed 4-Bit Model Outperforms Its Full-Precision Original
<https://huggingface.co/blog/MultiverseComputingCAI/quantization-aware-healing>
- 03 Recuris: Recursive Memory Architecture for Long-Horizon Agent Harnesses
<https://arxiv.org/abs/2608.24876>
- 04 BrowserForge: Parallel Browser Sandboxes Train Web Agents at Scale
<https://arxiv.org/abs/2608.24848>
- 05 OpenAI's Jalapeño Chip Benchmarks Show 1.5-1.9x Efficiency Over Nvidia Superchips
<https://techcrunch.com/2026/08/25/openais-jalapeno-chip-is-built-for-fast-inference-at-scale-benchmarks-show/>
- 06 Keenable Exits Stealth with \$26M to Index the Web for AI Agents
<https://techcrunch.com/2026/08/25/accel-backed-keenable-is-indexing-the-web-for-ai-agents/>
- 07 Claude Cowork Finally Remembers What You Told the App in Chat
<https://techcrunch.com/2026/08/25/claude-cowork-finally-remembers-what-you-told-the-app-in-chat/>
- 08 Comparing Local Tool Calling: Gemma 4 vs. Llama 3 vs. Mistral
<https://machinelearningmastery.com/comparing-local-tool-calling-gemma-4-vs-llama-3-vs-mistral/>
- 09 Stability AI Raises \$76M from Universal, Sony, and EA for Creative AI Tools
<https://techcrunch.com/2026/08/25/stability-ai-maker-of-image-generator-stable-diffusion-raises-76-million-in-fresh-funding/>
- 10 Robotics Startup Generalist Hits \$3B Valuation in Series B Extension
<https://techcrunch.com/2026/08/25/robotics-startup-generalist-reaches-3b-valuation-sources-say/>
- 11 OpenAI Loses Data Center Chief as Executive Departures Mount Ahead of IPO
<https://techcrunch.com/2026/08/25/openai-loses-a-top-data-center-exec-as-stream-of-high-profile-departures-continues/>
- 12 India's Ringg Gets Peak XV Backing to Push Voice AI Beyond the Phone Call
<https://techcrunch.com/2026/08/25/indias-ringg-gets-backing-from-peak-xv-as-it-pushes-voice-ai-past-the-phone-call/>

About this document. Every story in the 26 August 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.

Generated 2026-08-26T05:34:38Z · model qwen3.5:397b-cloud · New Horizon, new-horizon.tech · To unsubscribe or manage your subscription, see the footer of the daily email.