

The AI Model Release Wave

42 models in 30 days, ending with GPT-6 Astra

What the release cycle reveals about agentic work, cost per task, open deployment, specialization, and enterprise control.

Executive summary

From 5 August to 3 September 2026, at least 42 significant AI models and separately marketed variants entered public release, preview, or restricted deployment. The cycle ended with GPT-6 Astra beginning its rollout through Microsoft Foundry.

The volume is striking, but the change in product direction matters more. The leading releases were designed less as chatbots and more as workers: systems that plan, use tools, operate software, preserve context across long assignments, and produce finished documents, code, analysis, and decisions.

The market is also becoming harder to describe with one benchmark. Intelligence, coding performance, task cost, latency, token efficiency, local deployment, data control, and specialization now create several different frontiers. No single model led all of them.

The next competitive advantage will not come from access to one impressive model. It will come from knowing which model should perform each task, what authority it should receive, how its work will be verified, and whether the economics remain defensible at production scale.

Research window and inclusion rule

- The release window is 5 August through 3 September 2026, inclusive.
- The ledger includes significant public releases, previews, restricted-access models, and separately marketed variants.
- It excludes informal community fine-tunes and minor model-hub uploads.
- Vendor performance claims are treated as claims until reproduced independently.

The month closed with GPT-6 Astra

Microsoft describes Astra as a model for turning open-ended objectives into multi-step plans and finished deliverables. It can create documents, spreadsheets, presentations, dashboards, and analyses while incorporating new instructions as work progresses.[1]

Astra can also interact with software through advanced tool use and computer use. This includes navigating applications, updating records, testing websites, investigating software defects, and completing workflows where no dedicated API exists.

The model began rolling out on 3 September through the Microsoft Foundry Limited Access Program. Foundry provides the surrounding enterprise controls: identity management, role-based access, private networking, encryption, content filtering, monitoring, safety evaluation, and governance. Microsoft states that customer prompts and outputs are not used to train the models.[1]

That framing is significant. The product is no longer simply an intelligence endpoint. It is a worker operating inside an enterprise environment, and the surrounding control plane is part of the product.

Astra is powerful, but not dominant everywhere

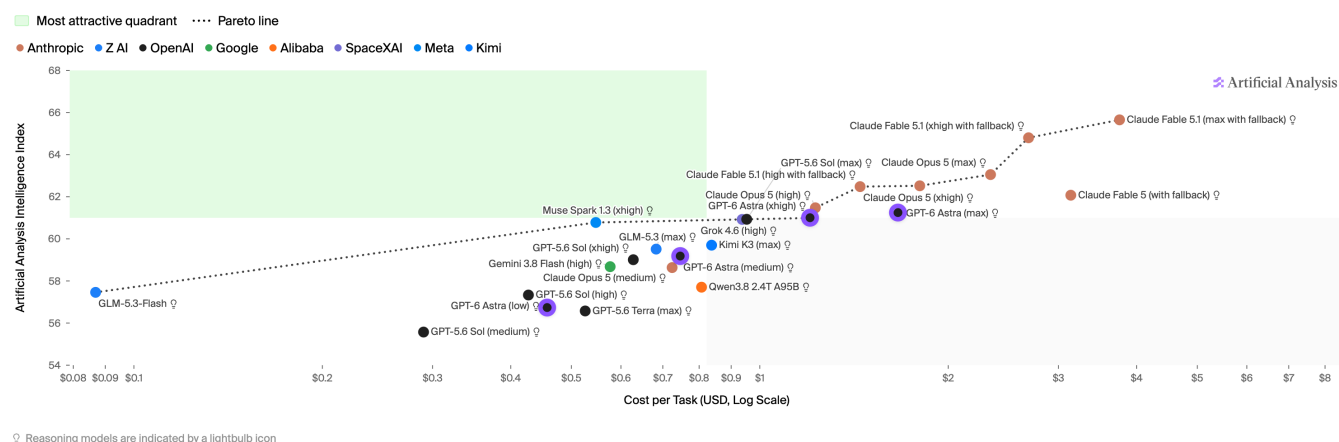
Artificial Analysis found two different stories in its independent evaluation of Astra. On the Coding Agent Index, Astra scored 67, approximately level with Claude Opus 5, Claude Fable 5, and Muse Spark 1.3. Claude Fable 5.1 led with 70.[2]

Astra's coding result was helped by major token-efficiency improvements. At maximum effort, it used roughly one-third of the tokens consumed by GPT-5.6 Sol in the Codex evaluation harness.[2]

The broader Artificial Analysis Intelligence Index was less decisive. Astra scored 61, equal to GPT-5.6 Sol and five points behind Claude Fable 5.1 at maximum effort. Astra used approximately 10 percent fewer output tokens than Sol, but its higher token prices made it about 75 percent more expensive per Intelligence Index task at maximum effort.[2]

Intelligence Index vs. Cost per Intelligence Index Task

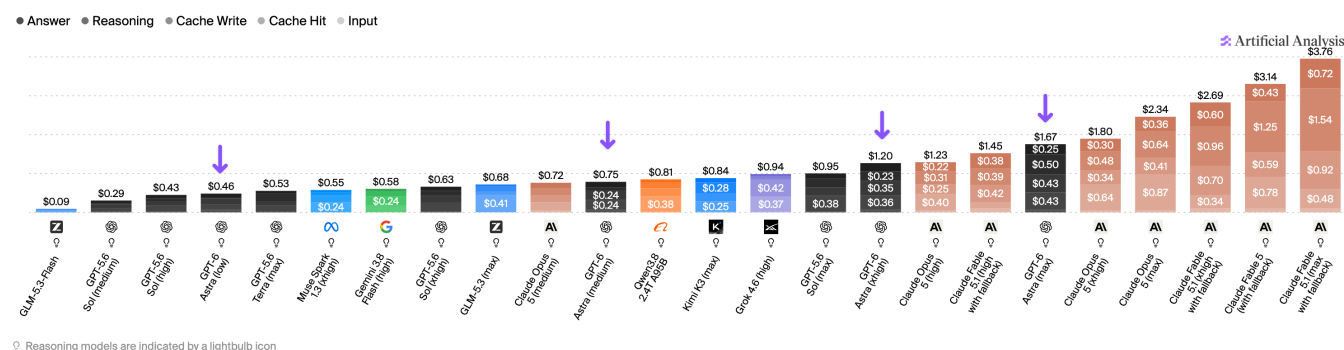
Artificial Analysis Intelligence Index - Weighted average cost (USD) per Artificial Analysis Intelligence Index task



Reasoning models are indicated by a lightbulb icon

Cost per Intelligence Index Task

Weighted average cost (USD) per Artificial Analysis Intelligence Index task, segmented by token type. Lower is better



Reasoning models are indicated by a lightbulb icon

Figure 1. GPT-6 Astra on the Artificial Analysis Intelligence Index versus cost per task. The lower panel decomposes average task cost by token type.

Artificial Analysis, Benchmarking GPT-6 Astra, 3 September 2026. [View source](#)

Astra also made a notable improvement in hallucination behaviour. Artificial Analysis measured a decrease in hallucination rate from 92 percent for GPT-5.6 Sol to 51 percent for Astra at maximum effort, while accuracy improved by four points.[2]

Results across professional evaluations were mixed. Astra gained approximately 80 Elo points on AA-Briefcase, a long-horizon knowledge-work evaluation, but lost a similar amount on GDPval-AA v2. It improved on Humanity's Last Exam while regressing slightly on several banking, scientific-coding, and long-context evaluations.[2]

The practical conclusion is not that Astra wins every category. It is that model selection has become workload-specific.

Cost per task is replacing cost per token

Token pricing is increasingly inadequate for comparing models. A seemingly inexpensive model can become costly if it produces unnecessarily long answers, repeats tool calls, requires multiple retries, or consumes enormous amounts of context. A more expensive model can be economical when it completes a task in fewer turns with fewer tokens.

Artificial Analysis now reports cost per task, time per task, and output tokens per task alongside intelligence measurements. These metrics provide a more realistic view of models operating inside production workflows.[3]

Meta's Muse Spark 1.3 is a strong example. At its publicly available xhigh reasoning level, it scored 61 on the Intelligence Index and cost approximately \$0.55 per task. At publication, that was the lowest measured cost among models scoring 59 or higher.[4]

Artificial Analysis Intelligence Index ↗

Artificial Analysis Intelligence Index v4.1.1 incorporates 9 evaluations: GDPval-AA v2, r²-Banking, Terminal-Bench v2.1, SciCode, Humanity's Last Exam, GPQA Diamond, CritPt, AA-Omniscience, AA-LCR

Not currently available

Artificial Analysis

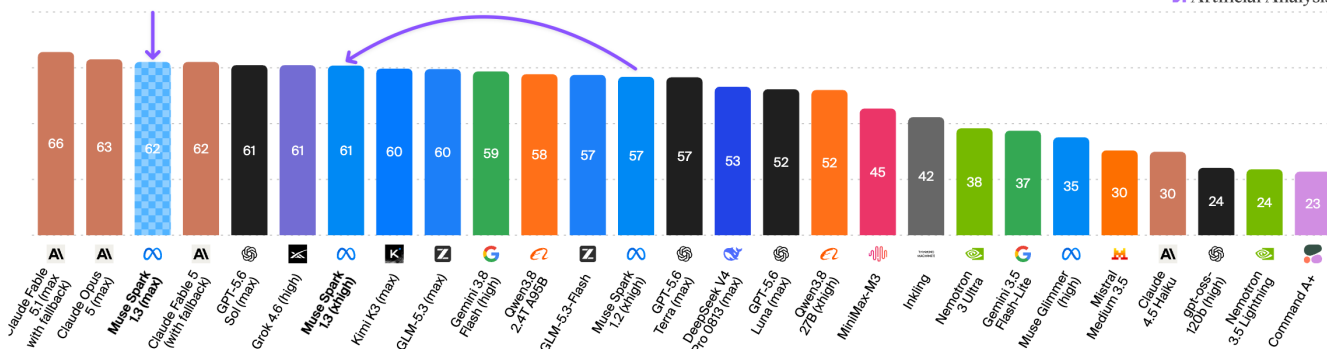


Figure 2. Frontier-model standings after the release of Muse Spark 1.3. Scores reflect the Artificial Analysis Intelligence Index v4.1.1.

Artificial Analysis, Muse Spark 1.3: Meta reaches the frontier, 2 September 2026. [View source](#)

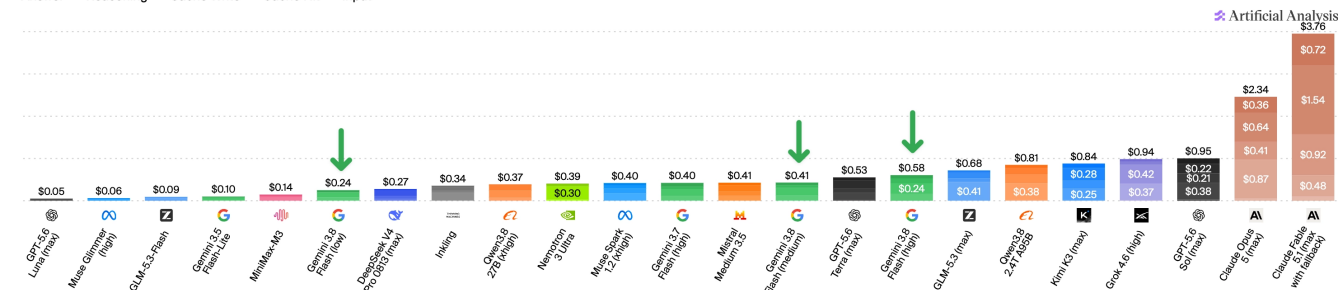
Google's Gemini 3.8 Flash followed a similar strategy. It scored 59 at high reasoning and cost approximately \$0.58 per Intelligence Index task. It maintained Gemini 3.7 Flash's introductory token pricing while improving software engineering, tool use, and agentic evaluations.[5]

Grok 4.6 reached an Intelligence Index score of 61 at a measured cost of approximately \$0.84 per task. Its strongest results appeared in agentic knowledge work, terminal-based software tasks, and multi-turn tool use.[6]

Cost per Intelligence Index Task

Weighted average cost (USD) per Artificial Analysis Intelligence Index task, segmented by token type. Lower is better

● Answer ● Reasoning ● Cache Write ● Cache Hit ● Input



Intelligence Index vs. Cost per Intelligence Index Task

Artificial Analysis Intelligence Index - Weighted average cost (USD) per Artificial Analysis Intelligence Index task

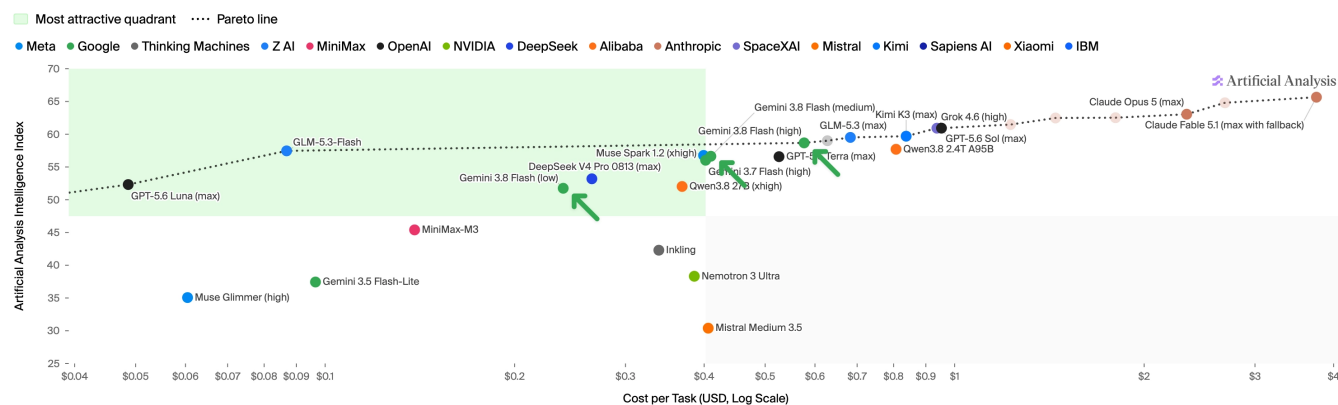


Figure 3. Cost per Intelligence Index task and the intelligence-cost frontier after Gemini 3.8 Flash entered the market.

Artificial Analysis, Google has released Gemini 3.8 Flash, 2 September 2026. [View source](#)

The frontier is no longer one leaderboard. It is a collection of frontiers involving capability, cost, latency, token efficiency, context use, and tool performance.

Open weights are attacking scale and efficiency

The month produced an unusually broad collection of open-weight models. At the largest end, Tencent released Hy4 Preview, a 770-billion-parameter mixture-of-experts model with approximately 49 billion parameters active per token and a one-million-token context window.[7]

Z.ai released GLM-5.3, a 743-billion-parameter coding model. Alibaba released Qwen3.8-Flash-Next, a 125-billion-parameter sparse model that activates approximately six billion parameters per token and previews architectural ideas intended for Qwen4.[7]

The smaller end may be more consequential for many organizations. Liquid AI's LFM2.5-2.6B is designed to plan, call tools, and operate entirely on a phone, laptop, PC, or robot. Ant Group's Ling-3.0-tiny activates approximately 1.3 billion of its 7.9 billion parameters per token. Cohere's North Micro Vision Instruct uses approximately 2.4 billion parameters for document and chart understanding.[7]

Open deployment is no longer limited to conventional text generation. Open models are being designed for coding, multimodal perception, tool use, document processing, translation, and long-running agents.

Cybersecurity models moved behind controlled-access gates

Three specialized cybersecurity releases appeared during the window: OpenAI GPT-5.6-Cyber, Google Gemini 3.8 Flash Cyber, and Anthropic Claude Mythos 5.1.[8][9][10]

All three use restricted-access programs rather than normal public availability. Google offers Gemini 3.8 Flash Cyber to trusted defenders through its Fairwind Program. Anthropic's Mythos 5.1 is available through cyber and life-sciences verification programs. GPT-5.6-Cyber is limited to vetted defensive organizations.

Astra raises the stakes further. OpenAI classifies it as the first broadly deployed model to reach the Critical level for cybersecurity capability under its Preparedness Framework. According to OpenAI, a model at this level may discover unknown vulnerabilities and develop new exploit techniques against well-protected systems without continuous human direction.[11]

The governance implication is direct: capability does not equal authorization. High-capability agents require separate identities, scoped credentials, approved resources, action-level authorization, trajectory monitoring, auditable records, and human review for consequential operations.

Specialization is becoming a competitive strategy

Several releases were not designed to be universal assistants. Cohere's Parse 5 converts PDFs, presentations, forms, diagrams, and images into structured Markdown with tables, visual regions, and bounding boxes.[12]

Ant Group's Ling-3.0-flash-Fin focuses on annual reports, investment research, financial workbooks, valuation modelling, and banking workflows. Tencent's Hy-MT2-30B-A3B targets multilingual translation. Sakana AI's Namazu focuses on Japanese-language business, research, coding, and agentic work.[7]

Thomson Reuters introduced Thomson, its first proprietary in-house large language model. It was trained using professional content from Westlaw, Practical Law, Checkpoint, and Reuters and is initially being deployed inside CoCounsel Legal.[13]

Microsoft's MAI-Code-1.1-Flash, NVIDIA's Nemotron 3.5 Lightning, and the three Ornith-1.5 models concentrate heavily on software engineering and agent execution.[7]

The enterprise model market is beginning to resemble a portfolio rather than a monarchy. A frontier planner may coordinate work while less expensive or specialized models handle coding, document extraction, translation, financial analysis, or local processing.

Complete 30-day release ledger

The table below names all 42 models and separately marketed variants in scope. Dates were cross-checked against the source-backed LLM Releases catalog and chronological changelog, with Astra added from its 3 September first-party launch materials.[1][7][14]

Date	Models released	Primary position
3 Sep	GPT-6 Astra	Frontier agents, computer use, and professional work through Microsoft Foundry
2 Sep	Gemini 3.8 Flash; Gemini 3.8 Flash Cyber; Muse Spark 1.3	Economical agents, coding, and restricted cyber capability
1 Sep	Claude Fable 5.1; Claude Mythos 5.1	Frontier knowledge work and controlled-access specialist capability
31 Aug	Mercury 2.5 Preview	High-speed diffusion language model
28 Aug	Hy4 Preview	770B open-weight agentic model
27 Aug	Parse 5; Ling-3.0-flash-Fin	Document intelligence and finance-specific analysis
26 Aug	Qwen3.8-Flash; Qwen3.8-Flash-Next; GLM-5.3-Flash	Sparse multimodal architectures, long context, and tool use
25 Aug	Granite 4.2 3B; Granite 4.2 8B; Granite 4.2 30B	Open reasoning family with thinking controls and tool calling
24 Aug	Apodex 1.1; Apodex 1.1 Mini; Thomson	Long-horizon professional agents and domain-owned legal intelligence
21 Aug	DeepSeek-V4-Flash-Vision-Exp	Experimental visual understanding added to V4 Flash
20 Aug	Hy-MT2-30B-A3B	Open-weight multilingual translation
19 Aug	Ornith-1.5-9B; Ornith-1.5-35B-A3B; Ornith-1.5-397B	Three open coding-agent scales, including a mobile build
17 Aug	GLM-5.2 Turbo	Faster hosted tier for coding and agents
14 Aug	GLM-5.3; Qwen3.8-27B; Dots3-Note Preview	Large open coding models and a locally practical multimodal model

Date	Models released	Primary position
13 Aug	Gemini 3.7 Flash; DeepSeek-V4-Pro-0813	Economical long-context agent models
12 Aug	Grok 4.6; North Micro Vision Instruct; LFM2.5-VL-3B	Frontier reasoning and compact open vision-language models
11 Aug	Nemotron 3.5 Lightning; Namazu; MAI-Code-1.1-Flash	Long-running agents, Japanese work, and coding
10 Aug	Muse Glimmer; Solar Pro 4; GPT-5.6-Cyber	Open local agents, economical professional work, and restricted cyber capability
6 Aug	LFM2.5-2.6B; Ling-3.0-tiny	Small efficient models for edge reasoning and tool use
5 Aug	Muse Spark 1.2	Multimodal software engineering through Muse Code

What enterprise buyers should do

- 1 Build a model portfolio.** Route work according to required capability, data sensitivity, latency, cost, deployment location, and potential impact.
- 2 Evaluate complete workflows.** Measure success rate and total cost, including retries, tool calls, cached context, human intervention, and verification.
- 3 Separate intelligence from authority.** A model may be capable of completing an action without being authorized to take it.
- 4 Test smaller and specialized models.** They may provide better privacy, predictability, latency, and economics for bounded workloads.
- 5 Independently reproduce important claims.** Vendor benchmarks are signals for testing, not procurement conclusions.

Conclusion

GPT-6 Astra is the most visible conclusion to this release cycle, but it is not the entire story. The deeper change is the fragmentation of the model market into frontier planners, economical workers, restricted cyber systems, open local models, document specialists, coding agents, translation systems, and domain-owned professional models.

The strategic question is no longer which model is best. It is which model should perform each task, under what controls, with what evidence, and at what total cost.

Organizations that answer those questions well will gain more from the current generation than those that simply purchase access to the latest frontier name.

Scope and limits

This report is an independent editorial analysis of public release materials and third-party benchmark reporting. It is not a benchmark audit, procurement recommendation, security certification, or statement of endorsement. Model availability, pricing, access restrictions, and benchmark standings can change rapidly. Buyers should verify current terms and reproduce critical evaluations with their own data and workflows.

Artificial Analysis figures are reproduced with attribution for commentary and analysis. Readers should use the linked interactive chart and source articles for current data, methodology, and model filters.

The strategic question is no longer which model is best. It is which model should perform each task, under what controls, with what evidence, and at what total cost.

Sources

- 1 [Microsoft, "GPT-6 Astra: Frontier intelligence for work, now available in Microsoft Foundry"](#), 3 September 2026.
- 2 [Artificial Analysis, "Benchmarking GPT-6 Astra"](#), 3 September 2026.
- 3 [Artificial Analysis, "Intelligence vs. Cost per Task"](#), accessed 3 September 2026.
- 4 [Artificial Analysis, "Muse Spark 1.3: Meta reaches the frontier"](#), 2 September 2026.
- 5 [Artificial Analysis, "Google has released Gemini 3.8 Flash"](#), 2 September 2026.
- 6 [Artificial Analysis, "Grok 4.6 returns SpaceXAI to the intelligence frontier"](#), 12 August 2026.
- 7 [LLM Releases, model catalog](#), accessed 3 September 2026.
- 8 [Google, "Introducing Gemini 3.8 Flash and 3.8 Flash Cyber"](#), 2 September 2026.
- 9 [Anthropic, Claude Fable 5.1 and Mythos 5.1 release materials](#), 1 September 2026.
- 10 [OpenAI, GPT-5.6-Cyber release and deployment materials](#), 10 August 2026.
- 11 [OpenAI, "Safety overview: GPT-6 Astra"](#), 3 September 2026.
- 12 [Cohere, "Introducing Parse 5"](#), 27 August 2026.
- 13 [Thomson Reuters, AI and technology announcements](#), August 2026.
- 14 [LLM Releases, chronological changelog](#), accessed 3 September 2026.
- 15 [Meta AI Research, "Introducing Muse Spark 1.3"](#), 2 September 2026.
- 16 [IBM Research, "Introducing Granite 4.2"](#), 25 August 2026.

Junior Williams | Systems and Enterprise Architecture | Cyber and AI Risk | trustcyber.ca