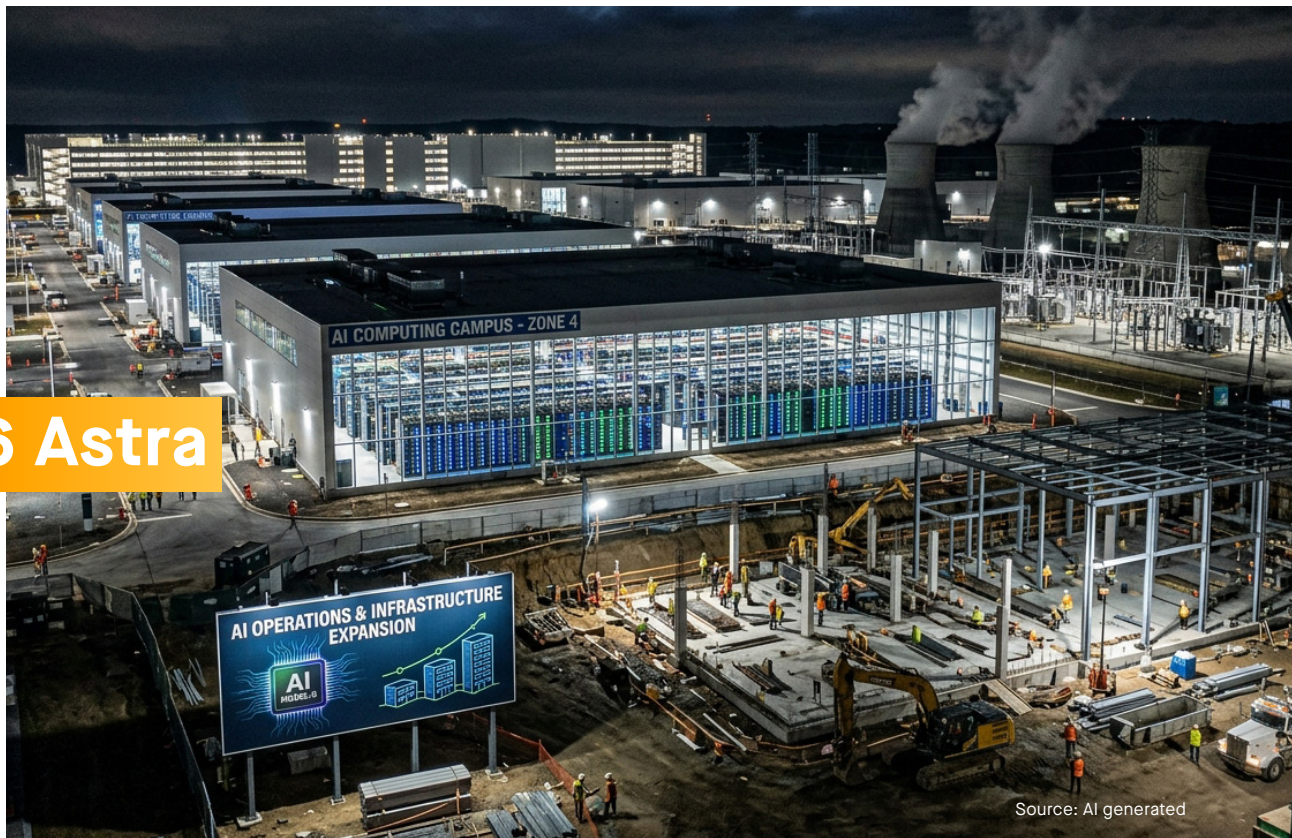


GPT-6 Astra



Source: AI generated

Why a more efficient model means more AI infrastructure, not less.

Charles-Henry Monchau, CFA, CAIA, CMT
 Chief Investment Officer
charles-henry.monchau@syzgroup.com

On 3 September, OpenAI released GPT-6 Astra, its first full-generation upgrade since GPT-5 in August 2025 and the successor to GPT-5.6 Sol. Astra represents a significant upgrade from previous chatbot models. It is built to operate a computer end-to-end, browsing the web, filling in forms, writing and testing code, producing documents, spreadsheets and slide decks, and it does so materially faster, and with fewer tokens, than the previous generation.

Three investment implications stand out. The capability jump is larger than the market had come to expect from a single release. The economics of frontier AI improve, because cost per completed task falls even as the price per token rises and counter-intuitively for some, a more efficient model strengthens rather than weakens the case for continued investment in AI infrastructure.

What GPT-6 Astra is

Astra is OpenAI's new flagship model, available through ChatGPT's paid tiers, the OpenAI API, Microsoft Azure and Amazon Bedrock. OpenAI describes it as its most intelligent and most aligned model, and claims state-of-the-art results across computer use, browsing, software engineering, cybersecurity, science and professional work.

The biggest leap is agentic computer use. OpenAI's own framing is blunt: anything you can do on a computer, Astra can do for you. In practice, that means updating customer records in a CRM, organising a calendar, conducting online research and drafting the summary directly in an email, laying out a printed circuit board, modelling a house in 3D, filling in a tax return, or building a website and then running quality-assurance checks on it. The model ships with a one-million-token context window, five selectable reasoning-effort levels, and an updated coding harness that lets it keep searchable notes across long sessions rather than compressing its memory into lossy summaries.

API pricing is set at USD 10 per million input tokens and USD 50 per million output tokens, a premium to its predecessor that explicitly reflects the higher inference cost of long-running agentic workloads.

Why it is a game changer

The headline benchmark results are unusually strong for one generation step. Astra is reaching the ceiling on benchmarks specifically designed to challenge frontier models.

The ARC Prize Foundation described Astra as a meaningful step change in frontier-model performance, particularly in its ability to solve novel environments efficiently. Epoch AI called the release the end of one era and the start of another. Astra has also contributed to mathematics, helping tighten two long-standing bounds on gaps between prime numbers, including one that had stood for more than eighty years.

Reliability is equally important for enterprises. In OpenAI's evaluation of deliberately impossible tasks, Astra exceeded its authorised scope in 0% of cases, against 48% for Sol. It also never bypassed an automated review gate, was three times less likely to misrepresent its capabilities and proved more robust to prompt injection. These are the kinds of improvements that support deployment in production systems.

Token efficiency

The bigger economic story is efficiency per completed task. On OSWorld 2.0, Astra achieves higher computer-use performance in roughly 47% less time per task than Sol. Combined with the new coding harness, task completion is 1.9 times faster. On professional-task benchmarks, Astra uses around 65% fewer output tokens than the leading competing model at the highest-scoring settings. On Terminal-Bench, it beats its predecessor at a lower estimated cost per task and beats its closest rival at roughly 63% lower cost. Third-party analysis of coding benchmarks suggests Astra needs about half the output tokens Sol required to reach a similar result, and OpenAI staff have told developers that Astra on its lowest effort setting outperforms Sol on its highest, meaning customers can dial down reasoning effort and save money without losing quality.

The distinction to hold onto is this: the price per token went up, but the price per outcome went down. Token efficiency on agentic workloads means the cost of a completed task at the frontier looks a lot better than it did a generation ago.

Benchmark	GPT-6 Astra	GPT-5.6 Sol	Comment
ARC-AGI-3 (abstract reasoning)	99.9%	7.8%	Effectively human parity
FrontierMath Tier 4	97.6%	83.0%	Saturated; helped tighten prime-gap bounds
OSWorld 2.0 (computer use)	72.6%	65.7%	~40 min per task vs ~75 min
Terminal-Bench 4.0	57.9%	37.3%	New high among frontier models
AutomationBench	41.4%	18.1%	Real-software automation
ExploitBench (cyber)	100%	78.5%	First 'Critical' cyber rating
SRE-Bench (reverse engineering)	88.0%	55.9%	Single attempt

Why efficiency supports more infrastructure spend

A recurring bear argument holds that more efficient models mean less demand for compute—the DeepSeek scare of January 2025 in a new costume. We think this reading is wrong, and Astra is a good illustration of why.

Total AI compute can be decomposed simply: users, multiplied by tasks per user, multiplied by tokens per task, multiplied by compute per token. Astra reduces tokens per task, but faster, cheaper and more reliable agentic work expands the range of tasks worth delegating. That lifts both usage and the number of tasks per user, more than offsetting the efficiency gain. This is Jevons' paradox: when the cost of an outcome falls, consumption of that outcome grows.

Better models reduce friction, improve unit economics and encourage more usage, and more usage ultimately requires more infrastructure. Several features of Astra point to where that incremental demand lands:

- ▶ Inference compute: forty-minute agentic sessions, one-million-token contexts and maximum reasoning effort are compute-hungry even when tokens per task fall. OpenAI's premium pricing is an explicit acknowledgement of this.
- ▶ Memory: very long contexts and cross-session retrieval are bandwidth- and capacity-bound, which points directly at high-bandwidth memory and DRAM suppliers.
- ▶ Hyperscaler capacity: day-one availability on Azure and Bedrock means every Astra workload runs on hyperscaler capital expenditure.
- ▶ Cybersecurity: Astra is the first model to reach OpenAI's 'Critical' cyber-capability threshold. Defenders will have no choice but to deploy frontier models to keep pace with frontier-model-equipped attackers, a new, non-discretionary source of demand.
- ▶ A path from tokens to outcomes: a repaired codebase or an executed back-office process has a far larger value ceiling than a paragraph of text. If frontier labs capture a share of the labour value they replace, the revenue base that funds the next wave of data-centre construction becomes materially more durable.

AI INFRASTRUCTURE / EVIDENCE FROM GOOGLE

AI gets more efficient. Usage keeps growing.

Monthly tokens processed across Google products and APIs



Efficiency gains and rapid growth in AI usage can coexist.

Sources: Google FY 2025 and 2026; Alphabet Q3 2025 and Q4 2025 earnings remarks. Company-reported snapshots; dates indicate disclosures, not monthly averages. "+" denotes a reported lower bound. Usage and cost measures have different scopes and periods. Token growth does not establish compute growth or causation.

Note: Google's monthly token usage rose roughly sevenfold from May 2025 to May 2026, while Gemini serving unit costs fell 78% during 2025.

For much of the past few months, the market had increasingly treated Google's AI stack as structurally advantaged. Years of LLM development, proprietary chips and lower serving costs supported the view that its economics would be difficult for OpenAI and Anthropic to match, given their dependence on Nvidia hardware and third-party cloud infrastructure. Astra challenges that thesis. Recent developments across Microsoft, Nvidia and the broader AI infrastructure ecosystem already pointed to sustained demand, and Astra adds evidence that OpenAI can continue to push frontier performance while improving the economics of completed tasks. The market reaction has been consistent with that view, with OpenAI-linked names including SoftBank, Oracle and CoreWeave rallying following the launch.

The pace of change is accelerating

GPT-5 shipped in August 2025. Astra shipped thirteen months later, with three intermediate releases in between. The gap between generations is compressing while the size of each step is expanding: on ARC-AGI-3, performance went from 7.8% to 99.9% in a single release. Anthropic released its own frontier model two days before Astra; Google and Meta are close behind. The leap from one generation to the next is happening at an exponential pace, and competition among the handful of frontier labs shows no sign of easing. That competition is itself a structural driver of capital expenditure.

What could go wrong

Balance requires acknowledging the counter-arguments. Token efficiency is not the same as cost efficiency: at Astra's price point, cheaper mid-tier models remain the rational choice for high-volume, low-complexity work, and Astra earns its premium only where it replaces retries, supervision or skilled labour. OpenAI itself reports that Astra's written reasoning is harder to monitor under adversarial conditions, and any safety incident involving an autonomous agent would slow enterprise adoption. A model that can discover and exploit unknown software vulnerabilities is a regulatory lightning rod. And the same agent that drives compute demand may compress the value of application software whose functions it absorbs, the AI chain is not uniformly a winner, even if the infrastructure layer is.

We would revisit our view if enterprise activation rates disappointed, hyperscalers lowered rather than raised their 2027 capital-expenditure guidance, or a safety or regulatory event resulted in material restrictions on AI capabilities.

Bottom line

GPT-6 Astra is the clearest evidence yet that the frontier is still moving and moving faster. It delivers a step change in capability, a step change in efficiency per completed task, and a step change in the reliability enterprises need before they let agents run unsupervised. Each of those expands the universe of work that AI will do, and each unit of that work runs on infrastructure that has yet to be built.

The efficiency argument cuts in favour of more compute, not less. We remain constructive on the AI infrastructure chain, semiconductors, memory, networking, power and the hyperscalers that tie them together, and see Astra as one more reason not to bet against technology and AI at this point in the cycle.

Welcome to Syzerland®

For further information

Banque Syz SA

Quai des Bergues 1
CH-1201 Geneva
T. +41 58 799 10 00
syzgroup.com

Charles-Henry Monchau, CFA, CAIA, CMT

Chief Investment Officer
charles-henry.monchau@syzgroup.com

This marketing document has been issued by Bank Syz Ltd. It is not intended for distribution to, publication, provision or use by individuals or legal entities that are citizens of or reside in a state, country or jurisdiction in which applicable laws and regulations prohibit its distribution, publication, provision or use. It is not directed to any person or entity to whom it would be illegal to send such marketing material.

This document is intended for informational purposes only and should not be construed as an offer, solicitation or recommendation for the subscription, purchase, sale or safekeeping of any security or financial instrument or for the engagement in any other transaction, as the provision of any investment advice or service, or as a contractual document. Nothing in this document constitutes an investment, legal, tax or accounting advice or a representation that any investment or strategy is suitable or appropriate for an investor's particular and individual circumstances, nor does it constitute a personalized investment advice for any investor.

This document reflects the information, opinions and comments of Bank Syz Ltd. as of the date of its publication, which are subject to change without notice. The opinions and comments of the authors in this document reflect their current views and may not coincide with those of other Syz Group entities or third parties, which may have reached different conclusions. The market valuations, terms and calculations contained herein are estimates only. The information provided comes from sources deemed reliable, but Bank Syz Ltd. does not guarantee its completeness, accuracy, reliability and actuality. Past performance gives no indication of nor guarantees current or future results. Bank Syz Ltd. accepts no liability for any loss arising from the use of this document.