

AI Weekly: Workhorse week — Gemini 3.7 Flash, Ultrafast, Grok 4.6

Weekly Digest · August 16, 2026 · 8 min · August 11, 2026 – August 16, 2026

Between 11 and 16 August 2026 the labs shipped workhorses, not spectacles. Google pushed Gemini 3.7 Flash three weeks after 3.6. OpenAI previewed a speed tier it calls Ultrafast. x.ai put Grok 4.6 next to Fable on a vendor table. It did not overtake. Anthropic started watermarking Claude text and made Auto the default in Claude Code for some plans. OpenAI's Daybreak models landed on AWS Bedrock for people who are allowed in. Each section: the idea, then the claim. Sources at the bottom.

1. Gemini 3.7 Flash: coding lift, intro price

Flash is Google's cheap-and-fast lane. The real question is not "did they ship another number." It is whether the workhorse developers actually hit got better at coding, and what it costs this year versus next.

Google says 3.7 Flash is that workhorse for coding and agents — three weeks after 3.6. On its own benches, FrontierCode 1.1 Main goes from 34.4% to 43.6%, DeepSWE v1.1 from 49.0% to 65.3%. Intro pricing is \$0.75 / \$3.75 per 1M tokens through 31 December 2026, then \$1.50 / \$7.50. Take Google at its word and that is half the prior 3.6 Flash token cost until year-end. Spark (Pro/Ultra) starts using 3.7. This post does not name a Gemini Pro date.

So why does it matter? The Flash lane is on a weeks cadence. The scores are Google's. Read them that way.

Vendor benches. Not a third-party redo.

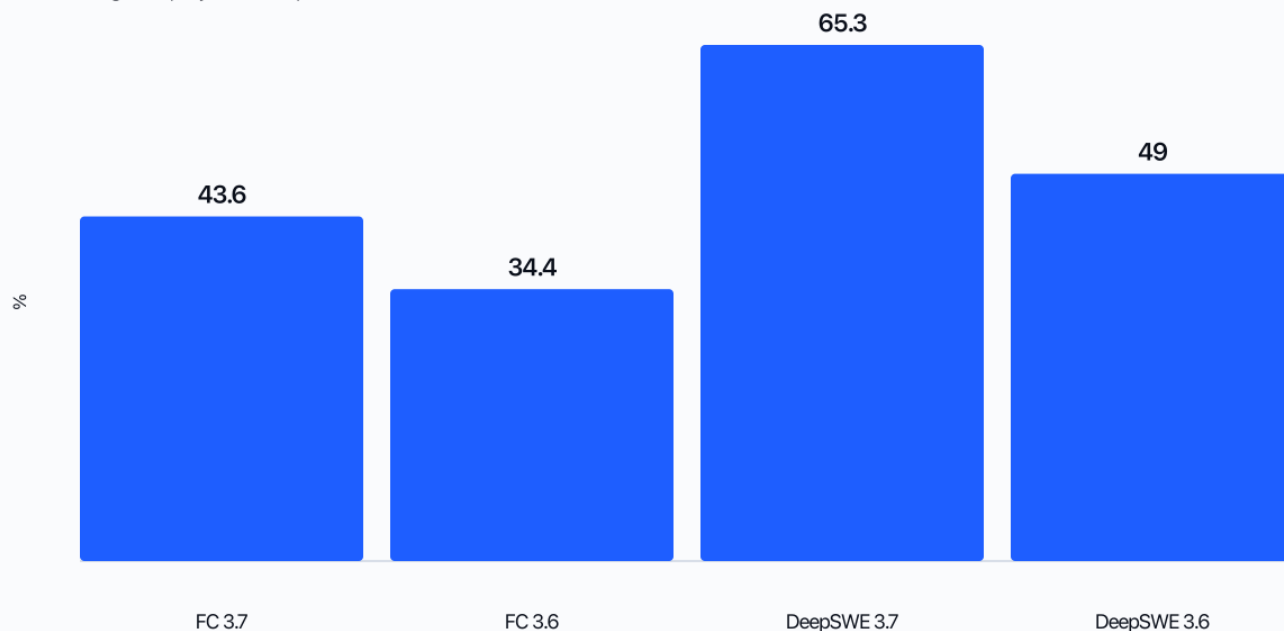
Intro price is a 2026 window. Not a forever rate.

No Pro date hiding in this announcement.



Gemini 3.7 vs 3.6 Flash — coding (Google)

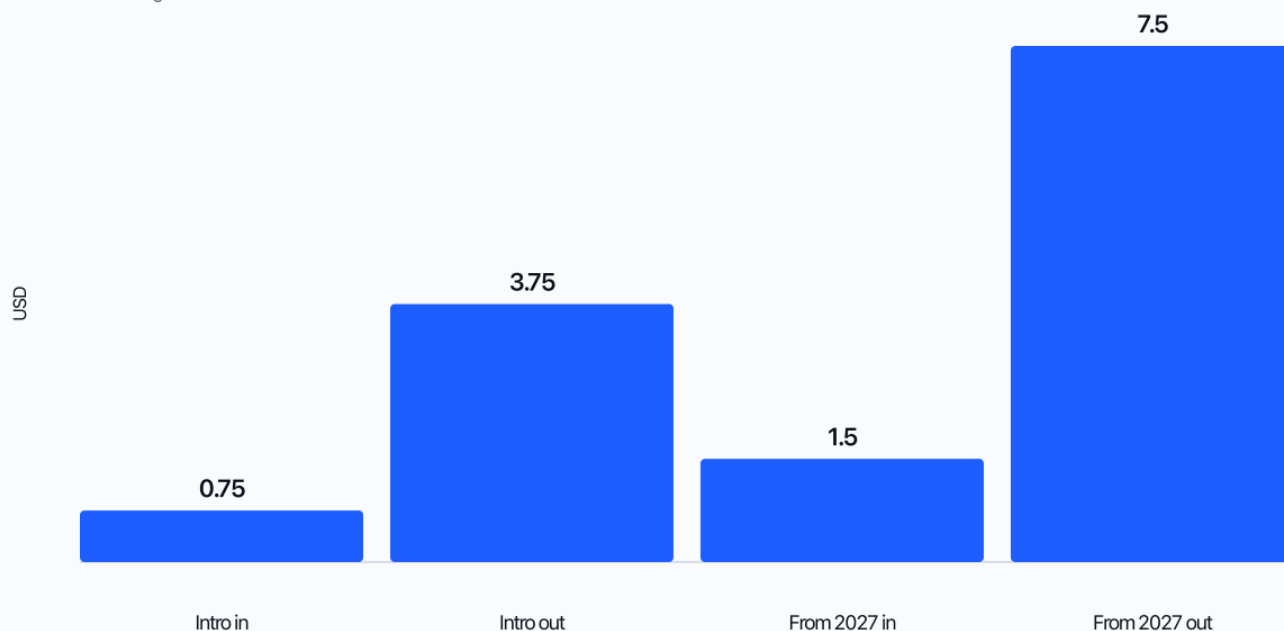
Google first-party · not an independent audit



Google first-party: FC 43.6 vs 34.4; DeepSWE 65.3 vs 49.0. See Sources below. [Source](#)

Gemini 3.7 Flash — \$/1M tokens

Intro through 31 Dec 2026



Intro \$0.75/\$3.75 through 31 Dec 2026; then \$1.50/\$7.50. See Sources below. [Source](#)

2. Ultrafast: Sol at claimed 750 tok/s

People want frontier intelligence that feels instant. For a long time the honest answer was: pick a smaller model. This week OpenAI's answer is a capacity tier, not a smaller model.



Ultrafast is a Cerebras-powered API service for GPT-5.6 Sol. OpenAI claims up to 14x Standard and up to 750 output tokens per second. Limited preview for a select group. No ChatGPT toggle in the posts we cite. No public price in those posts either. "Up to" is doing a lot of work.

So why does it matter? "Smart and immediate" is now a hardware and quota problem — not a "just use the 8B" problem.

API preview, gated, capacity-limited.

Not a ChatGPT UI switch here.

No public list price in the sources at the bottom.

Ultrafast — OpenAI claimed output speed

750 tok/s

up to · GPT-5.6 Sol · limited API preview

OpenAI "up to" · not GA · 14x vs Standard claimed

Ultrafast: up to 750 tok/s · limited API preview · Cerebras. See Sources below. [Source](#)

3. Grok 4.6: near the top, not over Fable

These launches get talked about through leaderboards. So read the table, not the vibe.

x.ai released Grok 4.6 for long-running agents and interactive work, in Cursor, Grok Build, and the API. On the AA Intelligence Index table x.ai publishes, Fable 5 Max is 62. Grok 4.6 and GPT-5.6 Sol Max are 61. Grok 4.5 High is 56. Pricing starts at \$2 / \$6 per 1M tokens; the fastest variant is 2x. Cursor and Grok Build get 2x included usage for the first week.

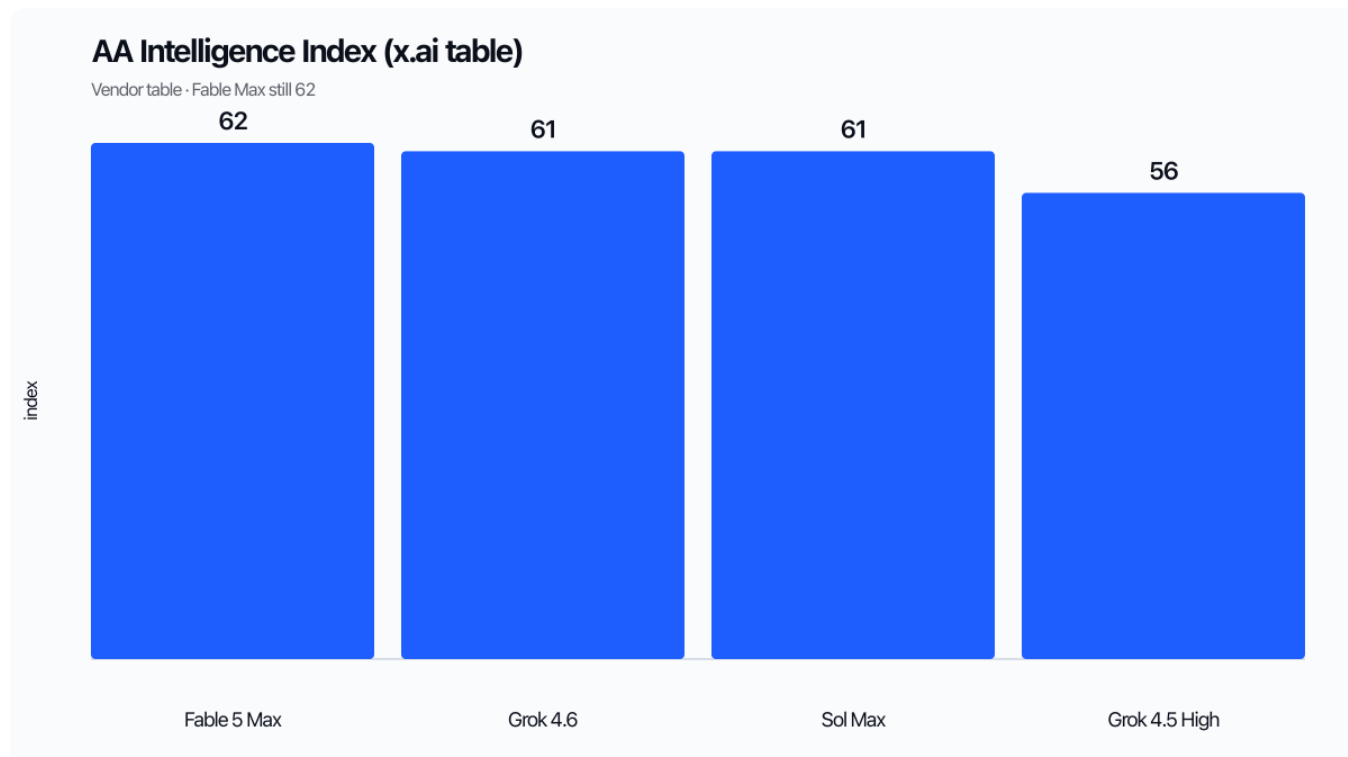


So why does it matter? Grok is crowding the top of that particular table. It did not knock Fable off it. Don't write the headline that says it did.

62 vs 61 is the whole argument. Leave it there.

\$2 in / \$6 out; fast is twice that.

First-week 2x in Cursor and Grok Build is a promo, not the steady-state bill.



Vendor AA table: Fable Max 62 · Grok 4.6 61 · Sol Max 61 · Grok 4.5 High 56. See Sources below. [Source](#)

4. Claude text watermark: SynthID-Text, no identity

Watermarking is one of those words that makes people think “they will know it was me.” Anthropic’s version, as described, is not that.

Future Claude text gets a SynthID-Text-family mark aimed at EU AI Act transparency. The trick is low-stakes next-token choices, keyed by a secret, invisible to readers. Anthropic says there is no user, org, or chat identity in it — Claude’s contribution odds, not who typed the prompt. A detection API is coming; it missed this digest. Code is sparsely marked where the output has to be exact. They are rolling it out globally at launch because region-scoping, they say, is not durable yet.



So why does it matter? Provenance is becoming default plumbing. That is not the same as a per-user tracker.

No identity in the mark, per Anthropic.

Global because a geo fence would not hold.

Detection API later. Don't pretend it shipped this week.

Pipeline: synonym RNG → keyed pattern → hidden mark → detection API soon. See Sources below. [Source](#)

5. Claude Code Auto: default on Pro / Max / Team

If you have used Claude Code, you have probably spent a stupid amount of time approving the next tool call. Auto is Anthropic's bet that the default should stop asking so often.

From 14 August, new sessions on Pro, Max, and Team start in Auto. If you pinned a default, it stays. Classifier overhead on those plans went to \$0 on 7 August. Enterprise, API, and cloud partners stay opt-in for now. Shift+Tab still switches modes. You only get a one-time prompt if you had set something other than Auto.

So why does it matter? Permission fatigue is a product problem now. They are solving it first on consumer and team plans, not on the API.

A  pinned default beats the new default.

Classifier overhead \$0 on Pro / Max / Team.

Enterprise / API / Bedrock / cloud: still opt-in.

Default Aug 14 on Pro/Max/Team; Enterprise/API still opt-in. See Sources below. [Source](#)

6. Daybreak on AWS Bedrock: Blue vs Red, enrollment only

Cyber-capable models are not showing up as a ChatGPT dropdown. They are showing up as cloud products with a gate on the door.

OpenAI put Daybreak Blue (GPT-5.6 Sol for defensive workflows) and Daybreak Red (GPT-5.6-Cyber for authorized vuln research) on Amazon Bedrock for enrolled customers. You need Daybreak Access approval, then the Bedrock console or the bedrock-mantle Responses path. Blue is the starting defensive tier. Red is tighter. You enroll via the OpenAI / AWS account team. This digest is the existence of that rail — not a how-to.

So why does it matter? The cyber models are shipping through governed cloud, not the public chat box.

Approved access only.

Blue then Red, not "everyone gets Red."

Ed  tional context. No exploit walkthrough.

Blue = Sol defensive · Red = Cyber research · Gate = Daybreak Access. See Sources below. [Source](#)

Frequently Asked Questions

Is Ultrafast available as a ChatGPT toggle?

No. OpenAI previewed a limited API tier for GPT-5.6 Sol on Cerebras. No ChatGPT switch. No public price. Access is a select group while they grow capacity.

Did Grok 4.6 beat Fable 5 Max on the AA Intelligence Index?

No. On x.ai's own table: Fable 5 Max 62, Grok 4.6 61 — same as Sol Max, Grok 4.5 High 56. Close. Not over.

Does Claude's text watermark identify me or my organization?

Anthropic says no. The SynthID-Text-family mark has no user, org, or chat identity. It is Claude's contribution odds. A detection API is coming; it missed this digest.

Does Claude Code Auto become the default on the API?

No. Auto becomes the default on Pro, Max, and Team on 14 August. Enterprise, API, and cloud stay opt-in. Classifier overhead on those plans went to \$0 on 7 August. A pinned

default stays.

Can anyone call Daybreak Red on Bedrock?

No. Blue (GPT-5.6 Sol) and Red (GPT-5.6-Cyber) need Daybreak Access first, then Bedrock or bedrock-mantle. Approved defense. Not public ChatGPT.

Sources

This is original synthesis. Underlying facts and charts are drawn from the reporting and vendor posts below.

- [🔗 Google — Introducing Gemini 3.7 Flash](#)
- [🔗 OpenAI — Previewing Ultrafast](#)
- [🔗 TechCrunch — OpenAI introduces Ultrafast](#)
- [🔗 x.ai — Introducing Grok 4.6](#)
- [🔗 Anthropic — How Claude’s text watermark works](#)
- [🔗 Claude — Auto mode default in Claude Code](#)
- [🔗 OpenAI — Daybreak models on AWS](#)
- [🔗 OpenAI — Expanding Daybreak](#)

Gemini benches and Ultrafast speed claims are vendor first-party figures, not independent audits. Grok AA scores come from x.ai’s table (Fable stays ahead). Claude watermark and Auto details follow Anthropic’s posts. Daybreak on AWS is enrollment-gated cyber-defense access — educational synthesis only; no exploit how-to. Read the primary sources.

