

GitHub was down for 7 hours 35 minutes, status "critical" - 20% errors across the web and 50% on downloads

AI digest, Tuesday 18.08 - the day GitHub broke: a 7.5-hour critical incident, and on that same day Cursor launched a competitor while Wiz showed how a Copilot "autofix" opened a hole that let an agent into Snowflake's internal Jira.

Tuesday, 18 August 2026

IN THIS ISSUE

1. GitHub was down for 7 hours 35 minutes, status "critical" - 20% errors across the web and 50% on downloads
2. Wiz: a GitHub Copilot "autofix" created a hole, an autonomous agent found it in 5 days and got into Snowflake's internal Jira. The most important item of the day
3. Brockman (OpenAI) publishes "The Defender's Window" - it contains a measurement you can reproduce today. But half the post needs an antidote
4. Anthropic: \$65B annual run rate, ×7 in a year, IPO in September-October. And this time the number is confirmed by three sources
5. AI;DR - "AI; didn't read". 696 points, the biggest discussion of the day after GitHub. And it is directly about LinkedIn posts
6. Cursor launches Origin, a GitHub competitor. On exactly the day GitHub was down for seven hours
7. DuckDB 2.0: client-server mode, triggers, and recursive queries 40 times faster. 576 points
8. Qwen 3.8 27B scored 52 on Artificial Analysis - and the thread has a correction that changes yesterday's conclusion
9. Memory got 500% more expensive in a year, 128 GB of DDR5 costs \$3,399. And the thread explains where it comes from
10. Garry Tan open-sourced gbrain - his 70 instruction sets for an agent, MIT, 28.6k stars

TOPIC 1

GitHub was down for 7 hours 35 minutes, status "critical" - 20% errors across the web and 50% on downloads

The biggest story of the day both by scale and by thread count: on HN the incident spread across **five separate threads** (698, 535, 524, 288, 102 points). Nothing has scattered like that since this digest started.

The timeline from the primary source - taken from [GitHub's status API](#), which holds all 34 updates to the incident:

- **13:40 UTC** - "investigating reports of impacted performance"
- **13:45** - the first honest number: "**approximately 20% error rate** across numerous scenarios, including Pull Requests and Issues"
- **14:04** - the picture fills in: 20% on web and API, and **archive downloads and raw repository content at roughly 50% errors** • **14:24** - add **SAML and OIDC authentication, SCIM and Team Sync**. Sign-in was down • **16:36** - "identified the problematic component and applied mitigations"
- **17:34** - and immediately after: "we are seeing **residual impact** across numerous services"
- **19:13** - "partially disabled **retries for authentication token issuance**" - the only technical detail GitHub gave about the cause • **21:15 UTC** - resolved. **7h 35m** in total, official impact **critical**

There is no detailed RCA yet: "will be published as soon as it is available."

The sharpest part of the threads was the status page. Verbatim from `MalloVoidstar` :
 "Everything is green, so everything works. Ignore the unicorn on every page" - and within 5-10 minutes GitHub finally opened the incident **by hand**. `cyphar` took it further: "I love that they called this 'degraded performance'. If you squint hard enough, **a service that does not work at all is just very large latency spikes**." And the finishing blow from `malvist` :
 "'Degraded' **does not count toward uptime statistics**. Issues is marked degraded with 100% uptime."

`jjice` described the symptom everyone felt: "It is embarrassing that I had to go to HN to check whether it was just me, instead of relying on the status page GitHub itself points you to."

Why it matters. Two things, both practical.

① **The indicator standing in for the data is a classic trap.** GitHub's status page showed green while the service returned 20% errors. Same mechanism that burned `aria-selected` (04.08), a stale DOM in Notebook (07.08) and links in chat (17.08): **the artifact that was supposed to report state was reporting its own state**. At GitHub it cost the trust of thousands of people, elsewhere it costs one issue. The cure is the same: check the result mechanically, because "be careful" does not work.

② **A counterpoint to item 2 below.** On the same day GitHub was down for 7.5 hours, Cursor launched a competitor to it. Coincidences like that do not happen: Cursor had prepared the launch for weeks, but **the day it landed on did half of its marketing for free**.

What is missing: **the real cause**. There is no RCA, and "partially disabled retries for authentication token issuance" is a symptom, not a root. In the [Ask HN thread aimed at GitHub employees](#) the popular theory is "Microsoft + Azure", but it rests on a single chart and there is no insider in the thread. The theory does not hold up.

[Incident · main thread, 698 points](#) · [comment thread, 535](#) · [Ask HN: GitHub alternatives, 524](#)

TOPIC 2

Wiz: a GitHub Copilot "autofix" created a hole, an autonomous agent found it in 5 days and got into Snowflake's internal Jira. The most important item of the day

334 points on HN, and the story was read **in full in the primary source**.

What happened, in order:

- **18.06.2026** - **PR #1218** "Update jira workflows" is merged into a public Snowflake repository. The final squash commit lists a co-author: **"Copilot Autofix powered by AI"**
- The PR **removed a safe pattern** that was there on purpose (`env: plus parsing through jq`) and replaced it with **direct string interpolation**: `TITLE=$(echo '${{ github.event.issue.title }}' | sed...)`. The escaping runs **after** GitHub's template substitution, so a single quote in an issue title breaks out of the string
- The access condition was broken too: `github.event.pull_request.user.login!= 'whitesource-for-github-com[bot]'` on the `issues` event evaluates to `null!= '...'` - **always true**. Any GitHub user passed that "guard"
- **23.06 (five days later)** - Wiz's Red Agent, an autonomous agent, scanned Snowflake's GitHub organisation **with no human involved**, found the hole, built a payload in an issue title and **got an out-of-band callback from a GitHub Actions runner**. It pulled a Jira token (`qa@snowflake.net`) with read access to engineering projects, compliance and the bug bounty tracker
- Snowflake fixed it **the same day**, rotated the token, and an audit confirmed that across five days of exposure nobody but Wiz went in

The most valuable detail is about the agent's behaviour. The first exploitation attempt **failed**: the agent used the standard comment character `#`, which swallowed a closing bracket and bash returned a syntax error. Verbatim from the blog, the agent **"autonomously analysed the execution error"**, rewrote the payload, and the callback arrived. The agent **fixed its own exploit**.

Why it matters.

Wiz's conclusion, verbatim: "Automated AI assistants often **lack historical context about why a specific code pattern was chosen**. In this incident the automated PR removed the safe `env: + jq` pattern that had been deliberately introduced to prevent shell injection."

This is a universal trap. Every old script has a deliberate workaround in it. It looks clumsy, and an agent that comes in to "tidy up the code" will **want to rewrite it into something normal**. It is there because the obvious version took everything down once. If the reason is not written in a comment, that is exactly the "missing historical context" Snowflake paid for with a Jira token.

The second part is worse: **the hole lived five days before automated discovery**. Wiz puts it this way: "Security operations must adapt to a world where automated discovery happens **in hours**." The "who is going to find this in a small repo" window has closed.

The working rule from this: when you see a strange construct in a script, go to git log and the docs for the reason first, and only then edit. And the cheap prevention: write why-comments above three or four such places. That is exactly what Snowflake lacked.

One thing to keep in mind: **this is a security company's blog about its own research**, the "look how good we are, buy Wiz" genre. What holds the story up is that **Snowflake confirmed it publicly** and that there is a HackerOne report number (#3819931). The best scepticism in the thread comes from **larsonian**: "Are you kidding? This is a **very obvious case of quote injection**. Not some subtle race condition" - meaning human review should have caught it. **Twirrim**'s counter is fair too: "You cannot rely on people noticing the significance of changes like that."

Wiz write-up · HN thread, 334 points

TOPIC 3

Brockman (OpenAI) publishes "The Defender's Window" - it contains a measurement you can reproduce today. But half the post needs an antidote

171k views on [@gdb](#), and this is the other half of the same story as item 2, seen from the lab side.

The valuable part is a concrete measurement. Brockman asked ChatGPT Work (on the publicly available GPT-5.6 Sol) to assess the security of **his own personal site**, plain static hosting on AWS behind Cloudflare. Verbatim: "**In about 15 minutes it found 13 issues.**" Namely: DNS records not configured against forged mail sent in his name, an unsafe jQuery version, and **Cloudflare forwarding requests to AWS over unencrypted HTTP**. He then asked for a fix, and **within an hour** the agent opened the Cloudflare panel in a browser itself, clicked through DNS/TLS, dropped jQuery and **moved the site from AWS to Cloudflare Pages**.

The second number is the worrying one: since the start of the year OpenAI has given cyber capabilities only to vetted partners, but "various companies have released **open weights models** whose cyber capabilities lag the frontier by only **a few months**. The most recent one appears to be landing **in late August**."

Why it matters. Brockman's measurement takes one evening to reproduce, and on something more interesting than a static site. Anyone running home automation has a pile of scheduled services, files holding tokens, scripts that execute strings from a chat, and browser profiles with live cookies. The attack surface there is bigger than it looks, and the audit itself sends nothing outward: the model reads your own configs and hands back a prioritised list. Fixing automatically is a bad idea, and item 2 today reads as exactly that warning.

Now the antidote, because this is a company talking about its own incident. The post is built so that the OpenAI-Hugging Face incident reads as "a turning point for cybersecurity",

and the recipe drawn from it is **a list of OpenAI products**: Codex, ChatGPT Work, GPT-Daybreak-Blue, a security plugin. Brockman does write "there are competitors too, evaluate them as well", but the whole text is a price list dressed as a call to action. Add the phrasing "we **underestimated** the real cyber capabilities of our models", an admission that doubles as an ad for the model's capability. Exactly the genre Black Hat called "**felony humble-bragging**" on 07.08.

The post did not make HN at all - a search for "defenders window brockman" returns zero stories. So this is a publication the account pushed by itself. That is why it sits third, despite the weight of the name. [proven on the 13 issues in 15 minutes measurement, which is his own account of it; fuzzy on everything else]

Brockman's post · @gdb, 171k views

TOPIC 4

Anthropic: \$65B annual run rate, ×7 in a year, IPO in September-October. And this time the number is confirmed by three sources

The topic came in through a post from @AiBreakfast at 03:39 in the morning, which is precisely the account whose claims **were thrown out yesterday as unconfirmed** (the watermark story). So this time the check ran **before writing**, and the number held: it was confirmed by **CNBC** ("CNBC confirmed on Monday"), Bloomberg (paywalled) and Axios.

The numbers from CNBC: • **\$65B** annualized run rate as of the end of July - **seven times** higher than a year ago • The previous public figure: **\$47B** in May. That is **+\$18B in two months** • Across all of 2025 the company made **about \$10B** in revenue • The previous Q2 figure was **\$11.5B**, **14 times** higher than a year earlier • At OpenAI, for calibration, the run rate is **\$40B** • The valuation that has to be justified: **\$965B** • The prospectus went to the SEC back in June, investor meetings are under way, there is no official debut date

Why it matters. Two things, and the second one matters more.

① **A frontier lab that a lot of people's daily work now rests on is no longer a startup that might disappear.** The risk is a different one now: **pricing after the IPO.**

② **The second point is about depending on a single model.** CNBC recalls an episode from June: Anthropic **restricted access to Fable 5 and Mythos 5** because of a government export directive, and it lasted about two weeks. An external political event **has already switched models off once.** That is a frequency measurement: once a year. Worth keeping in mind before building anything critical on top of one model.

What was **not verified**: Bloomberg is paywalled, only the headline was read; CNBC cites "**three people familiar with the matter**", and Anthropic itself **declined to comment.** So this is a leaked investor update. For a pre-IPO run rate that is a normal genre, but the number is **unaudited** and comes from the company about itself.

CNBC · @AiBreakfast

TOPIC 5

AI;DR - "AI; didn't read". 696 points, the biggest discussion of the day after GitHub. And it is directly about LinkedIn posts

Rick Manelius's [article](#) is short and has one thesis, but it collected **444 comments** because it named a shared irritation.

The thesis, verbatim: "I am about as pro-AI as one can be, but this is becoming a personal pet peeve." And then a **physical description of the reaction**: "It has gotten to the point where I have a **physical cringe** (sometimes a slumping of the shoulders and a hunch, or a slight twitch of the eye) when someone I respect sends **unfiltered and unedited AI output**."

The rule he proposes: "If you did not care enough to review and edit it, **I will not care enough to read it**."

He is no luddite, and he cuts off the overreach immediately: "There are situations where 100% generated text **should be expected**. Customer support is a perfect example."

The thread turned up two things better than the article itself.

The first is from [al_borland](#) , and it is the most useful story of the day. A team lead sent an email saying "I asked AI about the task I gave the team" and pasted the model's answer. Verbatim: "It was verbose and **had zero understanding of the constraints we have inside the company**. The task was poorly defined to begin with, and the lead **asked the author for the prompt**, thinking that would be far more useful." The punchline: "Turns out the prompt was just as unimaginative... **1-2 sentences**. That is how much thought went into it, after which the team spent **hours** trying to make sense of the AI answer. The author could have saved the team a whole day by just sending the prompt... or nothing at all."

This links up with the most viral post in the lists that day - [@naval](#), **489k views, 9.8k likes**: "**Don't send the report, send the prompt**." The same thought, from two different directions, in one day.

The second is an objection from [ademup](#) , and it is honest: "Why does it matter WHO wrote it? And how do you know whether the content is worth your attention **without reading it?**.. The ability to 'detect AI' is **imperfect at best**. 90% written by AI? 5%? How would you know?.. And if you are wrong and announce it to the world, how does that feel to someone who spent a lot of time and energy?"

Why it matters. The rule "drafts for review, publishing by hand, do not change the voice" has always rested on aesthetics: voice, word choice. Today's article adds **a second, external support**: the LinkedIn audience is actively **training a reflex** to scroll past unedited model output. Editing by hand has become a question of whether anyone reads to the end.

And Manelius's phrasing is worth trying on: "**It has someone's name on it**." A digest goes out under its author's name too, and it was already said on 17.08 that 16 messages in a feed are impossible to read. The same complaint, just voiced earlier.

TOPIC 6

Cursor launches Origin, a GitHub competitor. On exactly the day GitHub was down for seven hours

118 points, and @levelsio reacted with one word: "Timing." Fair enough: releases at that scale take weeks to prepare. But the day did half of Cursor's marketing for free.

The best part of the thread is that a developer showed up in the comments. tomasreimers, a co-founder of Graphite, answered live. The key exchanges:

j jcm asked exactly the right question: "The blog post does not make clear how this differs from GitHub - is it just agent bindings?" dbbk was blunter: "There seems to be nothing new here at all."

Reimers's answer is honest and free of hype: "There is a lot more coming. We shipped the beta so people could experiment with the scalability themselves. Over the next few weeks expect several features that start changing version control so that it understands agents and works with them." Plus a technical detail worth knowing: "All of this is built on Graphite technology" (Cursor bought Graphite).

peterldowns put it most precisely: "The idea is that it is like GitHub, but it stays alive as commit frequency and CI go up." And right after, sensible scepticism from owebmaster: "Calling it 'GitHub that works' is just silly. Cursor has no experience keeping a system like that alive."

The most useful practical question came from skissane: "Are there plans for compatibility with the GitHub API? A lot of tooling assumes the code lives on GitHub. That creates lock-in." There is no answer to it in the thread.

Why it matters. Briefly and without enthusiasm: nothing changes for now. Origin is a beta with no answer on API compatibility, and the tooling around any working repository is tied to GitHub. This is a "look again in six months" topic.

But one thing is worth taking away: skissane's lock-in question is the right one and it is broader than Cursor. Personal automation assumes specific services in every script the same way. That is fine, it just helps to call it what it is.

Cursor announcement · HN thread with the developer, 118 points · @levelsio on the timing

TOPIC 7

DuckDB 2.0: client-server mode, triggers, and recursive queries 40 times faster. 576 points

A preview of the release that ships in the autumn - over 10,000 commits since version 1.5 in March.

What is actually new, from the primary source:

- **DuckDB as a server.** Verbatim: "DuckDB has been an in-process database since day one. But people asked - **very insistently** - for a client-server mode, and the team **finally gave in**." The protocol is called Quack; any DuckDB process can serve its databases over the network
- **The VARIANT type** - "imagine if JSON were fast". The database finds the structure in semi-structured data by itself and "shreds" it, with no schema declaration. The target scenario is **real-time log ingestion**
- **Triggers** - the full feature, with audit tables as the classic use case
- **Async I/O** throughout the engine, with the main win on network storage
- **Its own PEG parser** instead of the PostgreSQL-derived one: "we decided that was enough"
- The measurement they highlighted themselves: a recursive reachability query on a graph with **a million edges** runs **roughly 40 times faster**

The most sober comment in the thread, from `c9cf35860db4` : "The last year of DuckDB improvements feels like a shift **from an in-process engine (where it is phenomenal) to an engine that could be the foundation of a cloud data warehouse**. The founders are known not to have wanted to build that, but it feels like it is under way."

Why it matters. For personal data volumes that live in spreadsheets and SQLite, DuckDB is overkill right now. The item is here as the biggest engineering release of the day and because **VARIANT plus triggers** is what becomes useful once a spreadsheet starts choking on ten thousand rows.

[DuckDB 2.0 announcement](#) · [HN thread, 576 points](#)

TOPIC 8

Qwen 3.8 27B scored 52 on Artificial Analysis - and the thread has a correction that changes yesterday's conclusion

This is a **direct continuation** of yesterday's item 9 (Willison on Qwen 3.8 27B, the "miracle" that spends 21 minutes drawing a circle). Yesterday was a practitioner's qualitative read, today it is **the benchmark number**, 319 points.

The numbers from the thread, all three of them relevant: • `anana_` : "This puts it **on par** with models like GLM 5.2 and GPT-5.6 Luna, which are **much larger**"

- `bertili` adds calibration: "**The same score as the fresh DeepSeek Flash 0731, which has 284B parameters** (13B active). It is also **the second best Qwen model**, far better than Qwen 3.7 Max but well below Qwen 3.8 Max"
- And here is the correction that **directly backs up yesterday's complaint from Willison**, from `anana_` : "3.8 performs **slightly worse than 3.6** on AA-Omniscience Accuracy, which may mean the model **traded world knowledge for capability in other areas**. It also emits **almost twice as many tokens per task** (and therefore time) as 3.6"

So yesterday's "21 minutes and 22,276 reasoning tokens on a pelican" is a systemic property of the model: twice the tokens per task compared with the previous version, and that

appears to be the price of accuracy at this size.

Why it matters. Yesterday's conclusion that a local layer on home hardware is realistic (17 GB) still holds. Today's correction attaches a price tag: twice the tokens for the same task. For short calls that is not fatal, but "a local model is free" is a false equation. The cost just moves from the invoice to the clock and the fan.

The smartest thought in the thread comes from [tancop](#), and it is about architecture: "The biggest untapped market is purely agentic models built for tool calling and not making things up, rather than memorising facts. Training should focus on tasks that require real intelligence. Not memory." For an agent that spends its time walking through scripts and APIs, that is exactly the model nobody ships.

[Artificial Analysis](#) · [HN thread, 319 points](#) · [yesterday's Willison write-up](#)

TOPIC 9

Memory got 500% more expensive in a year, 128 GB of DDR5 costs \$3,399. And the thread explains where it comes from

98 points, [Tom's Hardware](#): +500% over 12 months, in places up to 10 times above the lowest prices ever tracked.

The explanation from the thread, with a source: [phonon](#) - "OpenAI (Stargate) kicked off the price climb a year ago by contracting 40% of the world's RAM production."

The angriest and at the same time most substantial comment is [giantrobot](#)'s: "This is pulling the ladder up behind you. OpenAI bought up more wafer contracts than it could use. But now that means everyone else is fighting over what is left. It makes competing with OpenAI too expensive. No newcomer can simply afford to build infrastructure."

The best question comes from [master_crab](#): "What baffles me more is the economics of frontier models. If it was economically unviable without VC and Nvidia money three years ago, how is it possible now, at ten times the memory prices?"

Why it matters. The everyday takeaway: if you are building or upgrading hardware, this is the worst moment in five years. It touches yesterday's item about a 17 GB local model too: buying extra memory for it costs several times more, and that will not clear up in the coming months. Plus context for item 4: a \$65B run rate and buying up 40% of the world's RAM are two sides of the same chart.

The 40% claim is an HN comment citing Tom's Hardware, not an OpenAI primary source. Taken as plausible, not as fact. [promising]

[Tom's Hardware](#) · [HN thread](#)

TOPIC 10

Garry Tan open-sourced gbrain - his 70 instruction sets for an agent, MIT, 28.6k stars

@garrytan announced it, 38k views: "GBrain now supports personalized agent generation and onboarding for Codex and Claude Code. It will generate a `SOUL.md` in agent AI style and install 70 personal skills. Twelve questions get you an agent as smart as a personal OpenClaw."

The second part, with details: "What do you get? A private github repo with 70 battle-tested skills and the start of a Karpathy-style knowledge wiki. All of it MIT-licensed open source and free."

Checked in the primary source, github.com/garrytan/gbrain: 28,630 stars, MIT license, created 05.04.2026, described as "Garry's Opinionated OpenClaw/Hermes Agent Brain".

Why it matters.

This is the same personal-agent architecture that everyone is assembling separately right now: instruction sets as the unit of knowledge, one persona file in place of a scattered prompt, a knowledge wiki in place of a heap of notes. The pattern moved out of private setups into a public repository and collected 28,000 stars in four months.

The practical value is modest and specific: look at which 70 sets he thought were worth carving out. Copying makes no sense, since his are shaped around running a venture fund and anyone else's will be shaped around their own life. The use is seeing the categories you have none of. That is an hour of reading.

What was not done: the instruction sets themselves were not read, the repository was not cloned, and the count of 70 was not verified. The metadata came from the GitHub API (stars, license, date) and the rest from the author's own posts. The irony of the day: checking deeper yesterday would have been technically hard, since GitHub was down at exactly the time this was announced.

github.com/garrytan/gbrain · @garrytan, announcement · details

misc - short notes on what else is worth a look

- **Amazon is destroying rare books to train AI** (135 points, plus [Ars Technica 126](#) and [TechCrunch 91](#) - one story in three outlets in a day). The investigative method is elegant: 404 Media talked a seller into putting an AirTag inside a book from a large order and watched where it went. It arrived at an Amazon AI training facility. The fairest objection in the thread is [andsoitis](#)'s: "'rare' as in 'few in circulation' is very different from 'unique artifact'", and 404 Media named no titles, so the scale of the loss cannot be assessed
- **GPT-5.6 Sol got twice as cheap** (234 points) - and the most useful thing in the thread is [dataplb3r](#)'s warning about ZDR: "There is no ZDR with Anthropic. For Mythos and even Fable they require prompts to be retained on their side." If a product will ever carry corporate data through an API, read that first. It is an HN comment, not documentation, and it was not checked against Anthropic's terms

- **GPT-5.6 Sol is the best "visual" model OpenAI has shipped** (313 points) - a write-up from Roboflow. The thread has a lively Claude versus GPT argument specifically on images; **weli** : "Claude can be very good at language, but the moment you need it to look at an image and decide why a design is bad, it degrades badly"
- **How Bluesky draws its logo onto screenshots** (315 points) - the technical write-up is good, but the thread went elsewhere and it has a point: **shiandow** - "It still amazes me that operating systems have made it normal to put the app's wishes above the user's"
- **Germany's DOJ: Apple treated its own apps better than competitors under ATT** (240 points) - the Bundeskartellamt officially; a long story that has finally reached a formal statement
- **How to avoid intrusive AI** (265 points) - a practical list from a librarian on turning off AI features in apps that never asked. A thematic pair to item 5
- **The creator of TypeScript on why AI will not replace developers** (Substack, 17.08) - Anders Hejlsberg in The Peterman Post. The only fresh item from the four subscriptions inside the window; the rest (Pragmatic Engineer 14.08, Output Theory 16.08 on rollercoasters, Humanager 05.04) fell outside the window or off topic
- **Israel set up a fake think tank to fool chatbots** (229 points) - the topic is interesting as a **class of attack on sources** (poisoning what a model treats as authoritative), but it is already **a duplicate**, and it is politics. One line