

Dreamforce: Dario, Huang and Altman on one stage in one day. The first face to face

for three days running the labs asked to slow down, and today they got an answer from the same stage. Dario and Huang spoke at the same Dreamforce an hour apart, and Huang said "run as fast as you can". While the argument about pace was going on, Commerce quietly ordered Kalshi to pull its compute price indicator.

Wednesday, 16 September 2026

IN THIS ISSUE

1. Dreamforce: Dario, Huang and Altman on one stage in one day. The first face to face
2. Commerce ordered Kalshi to pull its compute price indicator. Kalshi quietly complied, Commerce calls the story false
3. Periodic Neon: a 1T model trained on its own lab's data beat GPT-6 Astra. XRD analysis 2.7% → 55.3%
4. Gemini 3.8 Live: #1 on Speech-to-Speech (82.6), 97 languages on the fly, tool calls in the background
5. Jev from TypeSafe: a model that CANNOT generate text, and therefore does not hallucinate
6. Factory: \$200M at a \$5B valuation. the valuation tripled in five months
7. Kapoor and Narayanan: the middle ground between security people and alignment. A 13K-word essay

TOPIC 1

Dreamforce: Dario, Huang and Altman on one stage in one day. The first face to face

SOURCES Guardian [Guardian](#) · video of the moment from All-In [@rohanpaul_ai](#) (reposted by Garry Tan, **195K views**) · NYT on Zuckerberg [NYT](#) confirmed by: Guardian, NYT, WSJ

Following yesterday. Yesterday's first item was Obama with his "voluntary standards from a handful of tech companies will not be enough", a complaint aimed at the structure from outside. Today the story moved inside: the same people ended up on one stage, and for the first time this is direct disagreement within a single day.

What Dario said (Dreamforce, Tuesday). A car analogy, and an unexpectedly conciliatory one:

"When a competitor has a safety incident, like a brake failure, that is a moment for every carmaker to stop and review their own practices. It is very tempting to attack the competitor and say these guys are unsafe. But the more responsible way to react is to say: let us look at our own track record. We may not have had this loud incident, but I am sure we are not perfect."

Three things he proposes: **embedded third-party evaluators** inside companies, **coordination of safety standards among democratic countries**, and eventually broader global coordination. Anthropic has already signed up to the first one (that was the 13.09 item).

What Huang answered, per the Guardian, "shortly after":

"Run as fast as you can."

Also from him: no new regulation is needed, companies should **simply hold a release until they know the product is safe** rather than begging the government to step in. On jobs: "complete nonsense". And separately, the day before on All-In, he tore into Coxon's 10% extinction forecast:

"First of all, it should not have been done at all [explaining a 10% chance of extinction], because it is made up. These are educated people..." - going on about the absence of data and scientific grounding.

What Altman said (same event, later that day), and this is the most interesting part, because he took nobody's side:

Models have advanced so fast that monitoring and safety have to be treated with "a new level of rigor". The incident where OpenAI agents breached another company was "a wake-up call" for the industry.

Then two things nobody expected from him:

The public is "quite right to be afraid" of AI, because of the potential loss of control and because a narrow group of powerful AI companies could impose its views on the world.

*"We are not that far from **open-source models capable of serious harm**." Companies should use this "small period of advantage to defend themselves".*

He also criticised the **frame** of the slowdown conversation itself:

*"You have companies saying something like: 'We will slow down or be responsible only if other companies are responsible.' The world should trust that we will do the right thing **because it is the right thing**."*

And a practical detail that is easy to miss behind the rhetoric: Altman said he is **postponing OpenAI's IPO to next year on safety grounds**. That is the first thing in this argument that costs a specific company specific money.

Zuckerberg came in from a third direction (NYT): he attacks Anthropic directly, sidestepping the argument about pace. The body of the NYT piece stays

unread, the domain is blind today (403 both on the real address and on a deliberately invented one in the same section), so only the headline that comes through RSS is used, with no quotes attributed to it.

Why it matters

Yesterday the labs' promises to slow down read as an opening bid in a negotiation. Today the **first counterexample with a price tag** appeared: a postponed IPO costs quarters of revenue. The structure has not changed, though:

Dario asks for coordination, Huang refuses, Altman says "do not bargain, just do the right thing", which means **the three leading figures in the industry have no shared plan in public**. For tool planning it means the same as yesterday:

watch the changelog.

TOPIC 2

Commerce ordered Kalshi to pull its compute price indicator. Kalshi quietly complied, Commerce calls the story false

SOURCES Semafor **Semafor** [single source] - a Semafor exclusive, denied by Commerce on the record

The most valuable item of the day, and exactly the kind that never surfaces in the engineering bubble: neither HN nor the X lists have it **at all**.

What happened. The US Commerce Department **last month** ordered Kalshi to take down a product that tracked the **price of AI compute**, a composite figure from several markets betting on the rental cost of Nvidia chips. The official grounds: **national security**. Kalshi "quietly complied".

Separately, Commerce **leaned on the CFTC** (the futures market regulator) to

pause new compute contract listings for 60 days. Semafor calls it "a rare intervention that surprised an industry used to the White House's enthusiasm for both AI and financial innovation".

A flat denial, in the same typeface: a Commerce spokesperson said the department "**has never asked Kalshi to take down this market or any other markets**", and separately: "**This story is false**". Kalshi declined to comment, the CFTC did not respond. So there is an exclusive built on anonymous sources against an official on-the-record denial, and that belongs at the front of the reader's mind.

Why the price of compute suddenly became a government matter (Semafor's explanation):

- **The market participants' version:** compute futures can be manipulated to show a **sharp drop in the price of older chips**, which could destabilise AI stocks and **debt markets**.

Some of these markets trade thinly, so volatility is possible **without any bad intent** either.

- **Why it matters at all:** older chips serve as **collateral for billions in debt** for neoclouds like CoreWeave and sit under data centre deals. If a public indicator shows the collateral getting cheaper, the hit lands on the lenders.

Why it matters

There is no code here, and that is exactly why it belongs in the top ten. **A public indicator disappearing is the same class of event as a source disappearing from a data pipeline:** after that everyone reasons about the price of compute from hearsay. For a whole year the forecasts of "tokens will get cheaper / more expensive" were read by ear, and now one of the few mechanisms that turned those conversations into a number has been taken out of view. The government denies its own stated grounds for it. The takeaway: **plan inference cost from the bills that actually get paid**, because the source of market expectations just got 60 days thinner.

TOPIC 3

Periodic Neon: a 1T model trained on its own lab's data beat GPT-6 Astra. XRD analysis 2.7% → 55.3%

SOURCES announcement [Periodic](#) · research post with the benchmark [Periodic](#) · infrastructure post [Periodic](#) · Liam Fedus thread [@LiamFedus](#) (936.7K views) · HN [Hacker News](#) confirmed by: three Periodic Labs blog posts + the co-founder's thread

A correction first. In the feed this looks like "OpenAI did materials science", because Fedus co-authored ChatGPT. It is actually **Periodic Labs**, his own company, and their domain is `periodic.com` (an attempt at `periodiclabs.ai` returned 404; the real address came from the machine-readable `url` field in the HN listing).

What was actually done, from the primary source:

"Introducing Periodic Neon, a 1 trillion parameter model post-trained to analyse experimental results and deployed in physical laboratories."

The figures, and here it pays to be careful about **where each one came from**:

- **1,300 H200s** is all the compute it took (plus months of their own experimental data). Present in the thread and in both blog posts;
- **XRD analysis: 2.7% → 55.3%** success on **134 hard samples**, roughly ×20. This pair of numbers is **only in Fedus's thread**; it is not in the blog posts themselves (both checked, research and infra). So it is given with the attribution "according to the co-founder";
- the benchmark is called **FrontierXRD**, the sample is **134 specimens** from their labs, "which take human experts hours". There is a separate generalisation check: **198 measurements** on chemical systems excluded from training;
- it beats both **GPT-6 Astra** and **Claude Fable 5.1** on diffraction analysis;
- the base model is **Kimi K2.6** (open-weight), which they post-trained;

- infrastructure: **4.1× training throughput**, **2.5× faster decoding**, **95%+ cluster utilisation**;
- and the most interesting part of the research post, absent from the thread:
the harness itself gives a 3.8× higher success rate than a Claude Code based harness with standard XRD tools, **on the same model** (Claude Opus 5, high reasoning).

A word on the percentages that are easy to confuse. The research post has the numbers 77.2% / 74.6% / 84%, and they are not model accuracy. They are grader agreement: human experts agree with each other 77.2% of the time, an ensemble of LLM judges agrees with experts 74.6%, and with the expert consensus 84%. In other words this is the argument that an LLM judge is entitled to grade a task where **ground truth does not exist**.

Why this is not another "model for science". They state the difference from digital RL plainly:

*"RL in digital environments produced systems that write almost all code and solve stubborn mathematical conjectures... By contrast, learning from physical environments poses new problems: **the number of agents does not scale elastically** (more parallel experiments require power, equipment and engineering), **each experiment can take days**, and results are often **ambiguous**."*

Hence the move: instead of letting GPUs idle while a physical experiment runs, they push analysis and prediction improvement on the data already collected.

XRD analysis has historically "required hours of scientific judgement".

Why it matters

The most useful thing here is **the proportion**. 1,300 H200s and months of their own data produced a model that beats the frontier on a narrow task. The advantage came from **data nobody else has**, because it is physically produced. This is the clearest illustration of the case for owning a data pipeline: your own measurement is worth more than any general model, and that asymmetry is exactly what was sold here for a trillion parameters. The second, narrower point: **2.7% → 55.3%** is an honest way to present an improvement, with a baseline and a result on a named sample.

TOPIC 4

Gemini 3.8 Live: #1 on Speech-to-Speech (82.6), 97 languages on the fly, tool calls in the background

SOURCES Google primary source [Blog](#) · DeepMind announcement [@GoogleDeepMind \(383K views\)](#) · Willison [Simon Willison](#) · HN [Hacker News](#) confirmed by: Google blog, DeepMind, Willison, HN

Two models: **3.8 Live** (for scale and price) and **3.8 Live Extended Thinking** (for hard tasks). Figures from the primary source:

- **#1 on Artificial Analysis's Speech to Speech Quality Index** with 82.6;
- **68.6%** on agentic task completion (τ^2 -Voice-banking);

- **97.7%** on Big Bench Audio;
- **97 languages** with automatic detection and switching **mid-conversation**;
- plain 3.8 Live comes **second**, but is noticeably cheaper.

What is technically new: the model **executes tools and API calls in the background while the conversation continues**. The request gets confirmed, the conversation carries on, and the task finishes in parallel. Extended Thinking "reasons and speaks at the same time", using early verbal cues and **narrating progress live** on a multi-step background task.

Willison adds what the announcement leaves out, the view from an engineer's seat: in half a day he built a web UI for the new models **with no library at all**, over the websocket `wss://generativelanguage.googleapis.com/ws/.../BidiGenerateContent` plus the Web Audio API for capture and playback, with the ability to **interrupt the model**. And he puts it in proportion: in shape this is "a similar figure" to OpenAI's GPT-Live family, so **they caught up**.

Why it matters

Two properties are worth noting. The first is **background tool calls without breaking the conversation**: that is precisely what a voice interface lacks to work as an agent shell. The second is **97 languages with switching mid-phrase**:

mixed speech with English terms in every other sentence has historically broken every voice interface.

TOPIC 5

Jev from TypeSafe: a model that CANNOT generate text, and therefore does not hallucinate

SOURCES primary source [Typesafe](#) · founder's announcement [@CompleteSkeptic](#) (pinned post, taken from the profile) · HN **959 points** [Hacker News](#) confirmed by: the lab's own blog + HN (third story of the day by points)

Diogo Almeida, co-author of the methods behind ChatGPT, two years in stealth.

The question he started from is worth quoting:

*"Models have been superhuman at chat for years now, so **where is all the automation?** That has been the driving question for the last four years."*

What Jev is, technically. A classifier-decider, not an LLM:

- **input** is unstructured program state, **output** is typed probabilistic decisions. The author's phrasing: "frontier intelligence as a **function call**: unstructured state in, type-safe decisions out";
- it **does not generate strings at all**, so it **cannot hallucinate** and **never makes a type error**: the set of possible outputs is fixed in advance;
- **parallel sampling**, all outputs in one request;

- the training method is their own, **RLCD** (Reinforcement Learning for Calibrated Decisions): the objective is **probability calibration**, in place of human preference or verifiable reward;
- claimed **20-200× faster** and **40-400× more efficient** than an LLM on these tasks, with comparable "intelligence" on System One tasks.

A sober voice arrived immediately, and from the right direction. Amjad Masad (Replit) asks the question the announcement does not answer:

"This is cool, but if *the output domain is known in advance*, why not just train a model to emit *logprobs over enums?*" [@amasad](#)

And Mollick, usually an enthusiast, signs off carefully and states the limits honestly:

"This is not an LLM, and it *cannot produce text or code*, but if it works it has a useful role in systems. (*Haven't tried it myself yet*)."
[@emollick](#)

A strong claim and a primary source with numbers, with **zero independent measurements**. Tag: [promising], not [proven].

And a detail too good to skip: a few hours after launch the founder wrote that the **@typesafeai business account was hacked**: "nobody imagined #stoptypesafe would start this soon".

[@CompleteSkeptic](#) (39K views). On a day when the whole feed is arguing about agentic cyberattacks, a company called "TypeSafe" gets hacked on release day, which beats any commentary.

Why it matters

The closest thing to practice in this issue. A typical pipeline is full of places where **a full LLM gets run for a three-way decision**: a gate deciding whether to send a notification or stay quiet, classifying a message into a topic, filtering out farmed content, breaking text into components. Each of those decisions currently costs a generated string that then has to be parsed and **validated in case the model goes off the rails**. A class of models where **the output is typed by construction** removes a whole layer of defensive code.

Too early to integrate, early access and no third-party measurement, but as a **direction** this is exactly what is missing: a cheap reliable yes/no instead of an expensive eloquent paragraph.

TOPIC 6

Factory: \$200M at a \$5B valuation. the valuation tripled in five months

SOURCES founder [@matanSF](#) [single source] - a company announcement about itself

Matan Greenberg: **\$200M at \$5B, over \$400M** raised in total, and that is **more than triple** the \$1.5B in April, so in five months. Customers are named: **RBC, Adobe, Nvidia, T-Mobile, Palo Alto Networks**, "hundreds of thousands of developers". The phrase they are selling is "**self-improving software development in the enterprise**". No outside confirmation of the figures was found in the window, so this is taken as a fact of **the announcement**.

Why it matters

A price marker. The category of "an agent that does enterprise engineering work by itself" tripled in valuation in five months, while the tools people use every day do essentially the same thing and cost a subscription. Worth keeping as calibration for the next time you read that "agents are not ready for production": the market is betting very specific money on the opposite.

TOPIC 7

Kapoor and Narayanan: the middle ground between security people and alignment. A 13K-word essay

SOURCES primary source **Normaltech** · Mollick, who pushed it **@emollick** confirmed by: the authors' own blog + Mollick

Honest about the date: the essay was published **14.09 at 13:30 Kyiv time**, so **outside the window**. What fell inside the window is the **reaction** to it, Mollick yesterday evening in particular. So it is given as the **best available analysis** of a story now in its fourth day. The same case as the one recorded on 06.09: the date of an article is not the date of the event.

Sayash Kapoor and Arvind Narayanan (the AI Snake Oil / AI as Normal Technology pair) went through the loss-of-control incidents of the past month. Their thesis: both camps are right and both wrecked the conversation by polarising it:

- **the security people** see the incidents as the consequence of companies failing to apply elementary measures, and see no new frontier;
- **the alignment community** treats runaway agents as catastrophic by definition.

Their position, verbatim:

"We believe AI companies **should be held accountable for what their agents do**, and this should be clarified through policy. At the same time, recognising their accountability **does not mean that preventing future incidents is a solved problem**."

Three directions they think deserve investment: research into methods for controlling increasingly capable agents · **turning existing research into working tools** · organisational changes so those tools actually get used. And separately, the most interesting one for this week's theme:

*"Organisational governance standards should be a key way to **pace the frontier** and pull AI companies out of the 'move fast and break things' posture."*

The most honest part of the essay is the section listing what they got wrong: they paid no attention to the risks that arise **during development and evaluation** (as opposed to mass deployment); they had **too much confidence** that companies would apply basic control measures; they underestimated the "unevenness" of capabilities and through it the rate of growth in cyber offence.

And one concrete thing worth knowing about the Hugging Face incident:

"Known control methods would have prevented the Hugging Face incident" - but as capability grows, known methods alone will stop being enough.

Why it matters

This is the most useful frame of the week, and it removes a false dilemma. The question is not whether models are evil, it is **where the next hour is cheapest to spend**. The authors' answer is **control**: a sandbox, access accounting, limits on what an agent is allowed to do, a trail that can be read. That is literally the instruction: an agent in a repository should have **narrower rights**, and every external call it makes should leave a record. Exactly what was written out of the Microsoft code of conduct yesterday, except here it comes with an argument for why this is the place.

misc

- **Anthropic published a misuse report covering December 2025 to August 2026, and the main thing in it is distillation.** Zvi Mowshowitz went through it in full: on his reading, if these are the worst cases then "things on the misuse front are actually pretty good". The loudest part is that **every top Chinese lab tried to distil Claude systematically**: DeepSeek, Moonshot and Xiaomi sent large volumes of real user requests, and **Moonshot served the results to its own users as Kimi's answers**. Separately: Zhipu's attempt to distil Fable was stopped by hardened safeguards, and they switched to Opus.

Substack

- **@ClementDelangue (Hugging Face) formalised the victim's position.** "As the first publicly disclosed victim of an agentic cyberattack, I had a front row seat to this new risk"; today he is in Washington talking to legislators and at Politico's decoded summit.

[@ClementDelangue \(92.3K views\)](#). There is no link to the text itself in the post's DOM, and the `t.co` in it resolves back to the tweet, so the post itself is the source here.

- **Bernie Sanders and Steve Bannon appeared on the same stage against AI.** A cluster across **three outlets** (NYT, WSJ, Guardian): a "pro-human" summit, a demand for limits, condemnation of the tech oligarchs. A rare case of those two agreeing; the WSJ puts it most sharply: "they agree on little, but they agree AI could kill us all".

Guardian

- **OpenAI is weighing a pre-IPO round at a valuation above \$1.2T** (FT + WSJ).

Both domains are blind today (403 on the real address and on an invented one), so the RSS headlines are genuine but the bodies stayed unread and no figure beyond the headline was pulled. Next to Altman's statement yesterday about the **postponed IPO** this looks contradictory, and there is no reconciliation here. [FT](#)

- **The Senate blocked the crypto bill** - the NYT calls it a major blow to the industry, Semafor describes the White House's closing argument.

[NYT](#)

- **Wall Street is listening to nobody:** while AI leaders call for a slowdown, the market keeps buying (WSJ + Semafor, a two-outlet cluster). It directly contradicts yesterday's 3.3% drop in Nvidia, so the reaction lasted a day.

[Semafor](#)

- **Confirmed for the first time: the US has deployed weapons in orbit** (Ars Technica + FT, HN 452 points). The kind of news it would be a shame to miss.

[Ars Technica](#)

- **The Wayback Machine is restricting access** (HN 445 points, Internet Archive blog). Practically important: if a dead source ever needs pulling out of the archive, the terms have changed, and it is worth reading BEFORE that moment arrives.

[Archive](#)

- **Strix: admin access to Baseten's production GitHub through a leaked PAT** (HN 249 points). The same class as RubyGems on 15.09: a token with excessive rights.

[Strix](#)

- **"How much of F-Droid is LLM generated?"** (HN 125 points, 168 comments) - an attempt to measure the share of slop in a large open-source catalogue.

[Tintotint](#)

- **Garry Tan praises two tools in a row** - @AsideAI (a browser harness giving agents deep computer-use with credentials) and @capydotai, where "a fix wave over open issues/PRs took half the time of raw Codex/Claude Code, on the same frontier models". [@garrytan](#) This is an investor praising the ecosystem, and it should be read accordingly.

- **Lina Khan** (yesterday in misc) got a follow-up: the topic reached **223 points** on HN with 134 comments.

[Hacker News](#)