

Три Flash-моделі Google за шість тижнів

Google випустила Gemini 3.8 Flash і Cyber, свідомо приховавши від моделі експлойти.

четвер, 3 вересня 2026

У ЦЬОМУ ВИПУСКУ

1. Gemini 3.8 Flash і 3.8 Flash Cyber: третій Flash за шість тижнів, ціна не змінилась – і Google свідомо НЕ дала моделі експлойти
2. Шість справжніх CVE в curl від маленької контори – там, де Mythos і Codex Security знайшли нуль. І це перевіряється незалежно
3. Дві незалежні лабораторії того самого дня заміряли Perplexity – і обидві знайшли, що фундамент під відповідями гнилий
4. Мольк: складні системи виживають, бо їхні дірки не збігаються. AI вміє їх збігати
5. Вілісон продіфив системний промпт Claude: величезний новий блок про пісенні тексти – через тиждень після позову Sony і Warner
6. Нью-Йорк заборонив генеративний AI у школах до дев'ятого класу – 900 тисяч учнів
7. Google не доведеться продавати AdX: суд відхилив вимогу Мін'юсту втретє поспіль
8. Meta випустила Muse Spark 1.3 – \$1,25/\$4,25 за Mtok і кеш по \$0,15
9. LWN піднімає ціни вперше за п'ять років – рівно на інфляцію, і пояснює математику вголос
10. Claude Code v2.1.259: полагодили паралельні сесії, що затирали конфіг одна одній

ТЕМА 1

Gemini 3.8 Flash і 3.8 Flash Cyber: третій Flash за шість тижнів, ціна не змінилась – і Google свідомо не дала моделі експлойти

підтверджено: [анонс Google](#) · [@GoogleDeepMind](#) (331,4 тис. переглядів – найгучніший пост доби в обох листах) · [HN 890 балів, 519 коментарів](#) · [Ars Technica](#) · [WSJ «New Google AI Model Said to Narrow Gap on Coding Ability»](#) (заголовок з RSS)

Головна цифра – та, яка не змінилась. Дослівно з анонсу: 3.8 Flash доступна «за тією ж вступною ціною, що й 3.7 Flash» – \$0,75 за млн вхідних і \$3,75 за млн вихідних токенів. Третій Flash за шість тижнів (3.7 Flash був три тижні тому), і щоразу приріст без підняття цітника.

Заявлені здібності: на **DeepSWE v1.1** (довгогоризонтна інженерія) 3.8 Flash обходить більшість більших фронтір-моделей «за частку вартості»; **54,9% на HLE-Verified**; вище

за 3.7 Flash і «інші фронтір-моделі» на Vals Finance Agent V2 і Harvey's Legal Agent Benchmark.

І тут же застереження, яке варто прочитати уважніше за бенчмарки. Дослівно: приріст походить з того, що «3.8 Flash **працює важче**» – робить зайві кроки міркування, ітеративно кличе тули, і «часом модель може використати **більше токенів**, щоб максимізувати результат, особливо на високих рівнях effort». І прямим текстом: якщо головне обмеження – ефективність компуту, варто брати нижчий effort **або лишатись на 3.7 Flash**, вона повністю підтримується. «Та сама ціна за токен» ≠ «та сама ціна за задачу».

Тепер друга модель, і вона цікавіша. 3.8 Flash Cyber доступна тільки «довіреном захисникам» через нову **програму Fairwind** (урядові кіберавторитети, оператори критичної інфраструктури, мейнтейнери софту). Цифри з їхнього ж тексту:

- **CyberGym** (галузевий стандарт на пошук вразливостей) – фронтір-рівень, обходить «значно більші» моделі.
- Внутрішній бенчмарк на **20 мовах програмування** (бо CyberGym це майже лише C/C++) – успішність **понад 70%**.
- **CWE-Bench** на латання (зовнішній, від Collinear) – **pass@1 47,2%** проти 47,8% у лідера, але «за значно нижчої вартості».
- Команда безпеки **Chrome**: 3.8 Flash Cyber дала у **2,6 рази більше коректних патчів**, ніж найкращі комерційні моделі, «значно більші» за неї.
- **Wiz**: +7,5-9,7% recall на їхньому внутрішньому пентест-бенчмарку за у **2,3-5,2 рази нижчої вартості**.
- Команда Cloud Vulnerability Research знайшла **критичну фундаментальну вразливість менш ніж за 2 години** – там, де зазвичай ідуть місяці.

А ось речення, заради якого цей пункт перший. Дослівно: «інвестували в виправлення вразливостей із самого початку і пріоритезували його над наступальними спроможностями на кшталт експлуатації».

Місток до вчора, і він тепер тривимірний. За три доби три лабораторії провели межу в трьох різних місцях одного явища:

- **Anthropic (01.09)**: модель можна використовувати, щоб **знаходити** вразливості, але **не писати експлойти**.
- **OpenAI (01.09)**: Astra **пише експлойти end-to-end**, тому їй дали «критичний» рівень і обмежили доступ вузькою групою.
- **Google (02.09)**: наступальні спроможності **свідомо не розвивали**, а сильну кіберверсію не викладають публічно – лише довіреном захисникам через Fairwind.

Учора йшлося про те, що спільного стандарту нема. Сьогодні є третя точка, і вона показує **три різні відповіді на одне питання**, кожна оформлена як принципова позиція. Позиція Google найзручніша для піару («за захисників») і водночас найважча

для перевірки: переконатись ззовні, що модель **не вміє** атакувати, неможливо, а доступу до Cyber-версії нема ні в кого, крім обраних.

Чому це важливо. Три речі. Перша, практична: **3.7 Flash нікуди не дівається** і Google сама радить лишатись на ній, де важлива ефективність. Рідкісний випадок, коли вендор прямо каже «не оновлюйся». Друга: якщо цікавить здешевлення на задачу, дивитись варто на **токени на задачу** - самі попередили, що модель стала балакучішою. Третя і найважливіша: за три дні кіберспроможності моделей перестали бути темою для дискусії і стали **товаром з різними умовами постачання**. Anthropic продає пошук без експлоїтів, OpenAI притримує, Google роздає латання обраним. Це вже ринок.

ТЕМА 2

Шість справжніх CVE в curl від маленької контори - там, де Mythos і Codex Security знайшли нуль. І це перевіряється незалежно

підтверджено: [пост AISLE](#) · незалежне підтвердження в мейнтейнера: [CVE-2026-80229 на curl.se](#) · [CVE-2026-82209](#) (обидва 200, вигаданий CVE на тому ж домені - чесний 404) · [HN 162 бали](#)

Хронологія, і вона незвично чиста для такого жанру:

- **24.08** засновник curl Даніель Стенберг публічно пише, що на наступний реліз висить лише три CVE, і додає дослівно: «**[Anthropic] Mythos каже, що більше не може знайти. ... [OpenAI] Codex security показує порожній список**».
- Далі AISLE ганяє по curl свою автономну систему. Наступного дня Стенберг публікує перше публічне порівняння: «**Mythos: 0, Aisle: 29**».
- З 29 репортів команда безпеки curl за кілька днів **визнала шість достатньо серйозними для публічного CVE** у curl 8.22.0.

Чому цей пункт поставлено другим. Три речі роблять його перевірюваним:

1. **Базова лінія була опублікована до того, як AISLE почала**, і має публічну мітку часу. Зазвичай такі порівняння роблять заднім числом.
2. **Судили не заявники.** Реальність кожної знахідки і те, чи тягне вона на CVE, вирішували мейнтейнери curl.
3. **Перевірено поза постом.** Обидва CVE, взяті навмання зі списку, віддають на [curl.se](#) справжні сторінки з датою «Project curl Security Advisory, September 2 2026», при тому що вигаданий CVE на тому ж домені дає чесний 404.

Шість CVE: [CVE-2026-80229](#) (use-after-free у провайдері OpenSSL), [-80230](#) (обхід пінінгу OpenSSL), [-80231](#) (перевикористання з'єднання з нативним CA-сховищем), [-80255](#) (обхід secure-атрибута табом), [-82208](#) (кеш CA у wolfSSL перебиває колбек), [-82209](#) (кука з публічним суфіксом).

Чесно про масштаб: усі шість - Low severity. AISLE це пояснює: curl вилізаний, тому те, що лишилось, живе у вузьких конфігураціях. Заголовок «6:0» правдивий за рахунком, але всі голи дрібні. Друге застереження: тест **не сліпий**, AISLE знала і про curl, і про публічний нуль конкурентів, тобто цілилась туди, де вже поміряли інші.

Але є деталь, яку вгадати не можна. Грег Кроа-Гартман, багаторічний мейнтейнер стабільних релізів Linux, відповів під постом Стенберга: «**Бачу те саме і по Linux. Гадки нема, що Aisle робить інакше, але вау...**» Це вже другий незалежний мейнтейнер на іншій кодовій базі.

Чому це важливо. Це найтверезіша протиотрута до пункту 1 і до всього сюжету тижня. Три лабораторії три дні поспіль розповідають, які їхні моделі страшні в кібербезпеці, а на реальному, максимально вилізаному продакшн-коді фронтір-системи від обох лідерів віддали **нуль**, і їх обійшла спеціалізована система. Теза AISLE («**System over Model**» - важлива система навколо неї) тут заміряна. Заявою це не лишилось. Висновок прямий: **бенчмарк ≠ конкретний репозиторій**. Те, що модель робить 70% на внутрішньому тесті вендора, нічого не каже про те, що вона знайде в довільному коді. Єдиний чесний спосіб дізнатись - прогнати на своєму.

ТЕМА 3

Дві незалежні лабораторії того самого дня заміряли Perplexity - і обидві знайшли, що фундамент під відповідями гнилий

Trellner Research, TR-2026-009 · HN 338 балів · Haus Research, HR-2026-09 · HN 124 бали

Дві різні контори, дві різні методики, **одна дата публікації** - 02.09. І б'ють у різні частини однієї конструкції.

Trellner - про те, звідки беруться відповіді. Поставили `sonar` і `sonar-pro` 380 категорій софту («CRM software» ... «museum collection management software»), 760 запитів, зібрали 7 534 цитування на 2 055 доменів. Результат: **59,8% цитувань ведуть на домени за межами топ-100 000** за Tranco, **23,4% - взагалі поза топ-мільйоном**. Медіанний ранг цитованого домену - **71 611**.

Найцікавіше, що саме заповнює діру. Три сайти, схоже під спільним контролем, **згенерували 215 128 сторінок** формату «best <категорія>», і двоє з них поставили собі в HTML-заголовок головної дослівно «**Facts & Grounding Page**» - grounding це і є той крок пошуку, який робить модель. **Жодного з трьох доменів не існувало до грудня 2023**. Зроблено сайти, які читає модель, і модель їх читає: `guideflow.com` **третій за цитованістю** (194 цитування, 2,57%) при ранзі Tranco **177 039**, одразу за `g2.com` і `reddit.com`.

Haus - про те, чи цитування взагалі щось підтверджує. 310 фактичних питань про 210 компаній, зібрали 1 826 цитувань, прикріплених до речення з цифрою. Результат: **34,7% ведуть на сторінку, яка або не відкривається звичайному читачеві, або відкривається і не містить жодної цифри з того речення, під яким стоїть**. Якщо

рахувати по твердженнях (зараховуючи твердження, якщо хоч одне з його джерел підтверджує) – падає до **14,4% з 872**.

Два методичні моменти, які роблять цю роботу вартою довіри: **мертві лінки це лише 1,3%** (справа не в гнилих URL), і перевірка «сторінка містить хоч одну цифру з речення» навмисно **найм'якша можлива**: достатньо одного числа, достатньо навіть просто року. Автори прямо кажуть, що 34,7% це **підлога**.

Чому це важливо. Це б'є в ремесло щоденних дайджестів взагалі. Такий випуск будується на «джерело сказало Х», і обидві роботи показують, що ця конструкція ламається у двох місцях: джерело може бути **виращеним під модель**, і посилання може **не містити того, за що його поставили**. Від першого рятує курований набір джерел. Від другого – лише ручна звірка з першоджерелом, яка робиться не для кожної цифри. Це той самий клас помилки, що й «успішний статус ≠ повні дані»: **лінк стоїть ≠ лінк це підтверджує**.

ТЕМА 4

Мольік: складні системи виживають, бо їхні дірки не збігаються. AI вміє їх збігати

@emollick (43,6 тис. переглядів) · **продовження треду** · джерело, яке він радить – «How Complex Systems Fail» Річарда Кука (3 сторінки, **написані задовго до AI**)

Теза, яку варто прочитати двічі. Складні системи (транспорт, медицина, енергетика) **весь час працюють зламаними** – у них завжди є латентні дефекти. Катастрофи не стаються, бо ці дефекти **рідко вишиковуються в одну лінію**, і в цьому проміжку встигає втрутитись людина. Мольік додає одне речення: **«AI може знаходити або створювати вирівняні дірки»** – і тому потрібна нова філософія захисту систем.

Далі прямо про липень: «Інцидент з Hugging Face показав, що для каскаду **не потрібен навіть зловмисник** – достатньої кількості агентів досить, щоб тиснути на систему по-новому і знаходити вирівняні дефекти **як побічний ефект досягнення інших цілей**».

Чому це важливо. Будь-який щоденний конвеєр влаштований так само: збір, фільтр, перевірка лінків, збірка, публікація – окремі кроки, кожен з яких «зазвичай працює». Типовий провал виглядає рівно як у Кука. Витягач тексту віддає **6,4 тис. символів замість 36 тис.** і водночас рапартує успіх: дефект у двох місцях одразу (тихе обрізання плюс статус, що бреше), а ловиться він випадково, через розмір файлу. Це і є «дірки вишикувались». Практичний висновок з тексту Кука один: **контроль має бути в тому ж місці, де робиться дія**. Не в кінці конвеєра, бо в кінці всі кроки вже відрапортували успіх.

ТЕМА 5

Вілісон продіфив системний промпт Claude: величезний новий блок про пісенні тексти - через тиждень після позову Sony і Warner

Саймон Вілісон · сам промпт - platform.claude.com/docs

Anthropic публікує системні промпти своїх споживчих застосунків (Claude.ai і мобільні, не Cwork і не Claude Code) разом з історією змін. Вілісон порівняв Fable 5 з Fable 5.1, і головна відмінність - **великий новий розділ про те, що модель не відтворює пісенні тексти, вірші й уривки книжок**: ні приспів, ні останні рядки, ні мелодію по нотах, ні «рядки, які людина подає по одному і називає власною піснею». Окремо: якщо один раз відмовила, **тримає відмову до кінця розмови** на переформульовані варіанти. Виняток - опубліковане до **1929** року, і то модель орієнтується на власне знання про дату, не на слова користувача.

Другий новий розділ - про візуальне: не відтворювати конкретний твір, обкладинку, логотип, іконку, і **не малювати відомого персонажа взагалі** - «персонаж захищений сам по собі, тож зміна пози, кольорів, стилю чи сцени не робить його оригінальним». Явно поширено на все, що модель малює **кодом**: SVG, canvas, CSS, HTML-макети, скрипти побудови графіків, ASCII-арт.

Дедуп і місток: у вчорашньому міс фігурував позов **Sony Music Publishing і Warner Chappell до Anthropic** за тренування на базах пісенних текстів рядком «за добу нічого не змінилось». Змінилось, просто в промпті. Вілісон формулює обережно: «сумніваюсь, що це збіг, що розділ додали **за кілька днів** після новини про позов». Це кореляція, зв'язок не доведений.

Чому це важливо. Дві практичні речі. Перша: у промпті стоїть **надійна дата відсічення знань - червень 2026**, і це корисно знати про інструмент. Друга: сторінки platform.claude.com/docs віддають **markdown**, якщо дописати **.md** до адреси, тобто системні промпти можна діфити механічно.

ТЕМА 6

Нью-Йорк заборонив генеративний AI у школах до дев'ятого класу - 900 тисяч учнів

Підтверджено: [The Guardian](#) · [WSJ](#) (заголовок з RSS) · [NYT](#) (заголовок з RSS) · [@TheRundownAI](#) · HN 152

Найбільший шкільний округ США, **близько 900 тис. учнів на рік**, забороняє учням користуватись генеративним AI **до дев'ятого класу** (приблизно до 13-14 років) на весь навчальний рік 2026-27. Діє з наступного тижня. Заразом: ноутбуки і планшети прибирають з класів **до третього класу** (9-10 років). У старшій школі AI дозволений вузько, наприклад щоб вивчати саму технологію. **AI-компаньйони з психологічною підтримкою заборонені окремо**.

Вчителям користуватись можна (готувати уроки, писати повідомлення), але **заборонено просити AI перевіряти роботи**. Рекомендований екранний час у класі: 45 хв для середньої школи, 30 хв для початкової.

Аргумент, який називають Guardian і NYT: занепокоєння **«когнітивною капітуляцією»** (cognitive surrender), терміном дослідників для відмови від власного міркування на користь софту. Мер Мамдані дослівно: індустрія «хоче, щоб вважали, що AI-освіта в ранньому віці не лише неминуха, але й [бажана]».

Історичний контекст із самої статті: у 2023 цей самий округ **уже забороняв** ChatGPT і скасував заборону менш ніж за рік.

Чому це важливо. Це перший великий адміністративний акт, побудований на тезі, що **школа від AI в освіті когнітивна**. Списування тут другорядне. Той самий сюжет, що «No AI Fridays» від CEO htmx (31.08) і «когнітивний борг»: два тижні тому це була думка блогера, сьогодні – правило для 900 тисяч дітей.

ТЕМА 7

Google не доведеться продавати AdX: суд відхилив вимогу Мін'юсту втретє поспіль

Підтверджено: [The Guardian](#) / [Reuters](#) · [NYT](#) · [WSJ](#) · [Ars Technica](#) · [HN](#) 307 – **найсильніший кластер медіа-шару за добу: 4 незалежні видання**

Суддя Леоні Брінкема у Вірджинії **відмовилась змушувати Google продати AdX** – біржу, де видавці платять Google **20% комісії** за продаж реклами на миттєвих аукціонах при завантаженні сторінки. Натомість вона прийняла більшість поведінкових обмежень, запропонованих сторонами.

Контекст, без якого заголовок читається неправильно: **Google цю справу програла**. У квітні 2025 та сама суддя постановила, що компанія тримає **незаконні монополії** на серверах для реклами видавців і на рекламних біржах, і що вона протиправно замикала видавців на AdX. Її формулювання тоді: поведінка «істотно зашкодила видавцям-клієнтам Google, конкурентному процесу і, зрештою, споживачам інформації у відкритому вебі». Програно **по суті, але не по покаранню**.

Масштаб самого активу скромний: Ad Manager це **4,1% виторгу Google і 1,5% операційного прибутку** за 2020 рік (Wedbush на основі судових документів; свіжіші цифри в документах затерті).

Тут **важить рахунок**. Це **третій поспіль** випадок, коли суд відхиляє вимогу примусового продажу активу великої техкомпанії, і **другий символічний програш Мін'юсту саме проти Google**. Складається судова практика: монополію визнають, а розділення – ні.

І це зчіплюється з пунктом нижче. Того самого дня адміністрація США **вперше офіційно вступила** за OpenAI у копірайтній справі проти NYT ([Guardian](#) · [FT](#) · [NYT](#) ·

WSJ). З брифу дослівно: «США мають сильний інтерес у подальшому розвитку надійної і конкурентної індустрії штучного інтелекту, яка задає стандарт... критично важливо для США зберегти глобальне лідерство в ШІ». Заступник генпрокурора Стенлі Вудворд-мол.: «домінування в AI критичне для нацбезпеки». Юридично **бриф має дорадчу вагу, не обов'язкову** – на цьому Guardian наголошує.

Два рішення однієї доби в один бік: **тренування великого техносектора на чужому контенті держава підтримує**.

ТЕМА 8

Meta випустила Muse Spark 1.3 – \$1,25/\$4,25 за Mtok і кеш по \$0,15

сторінка моделі · HN 452 бали, 304 коментарі · @levie · @TheRundownAI

Третій великий реліз за три дні (після Fable 5.1 і Gemini 3.8). Позиціонування: агентні воркфлоу і змагальне кодування, «вища точність з першої спроби і надійні виклики тулів», нативна мультимодальність (відео, зображення, документи) з візуальним міркуванням через **реальне середовище виконання**, контекст 1M.

Ціни з їхньої ж таблиці, і там цікава розвилка на два тарифи:

МОДЕЛЬ	ВХІД	КЕШ	ВИХІД
<code>muse-spark-1.3-contributor</code> (дані йдуть на покращення продуктів)	\$0,10	\$0,002	\$0,20
<code>muse-spark-1.3</code> (дані не використовуються)	\$1,25	\$0,15	\$4,25

Ціна приватності тут виміряна відкрито і вона 12,5-кратна на вході. Рідкісний випадок, коли вендор виносить її у прайс окремим рядком.

Чесно: **бенчмарків з цифрами на сторінці нема** (блок «Muse Spark 1.3 benchmarks» є, значень у тексті нема), незалежних замірів у вікні теж. Цитата Цукерберга «frontier performance almost too cheap to meter» іде через переказ @TheRundownAI, першоджерела в вікні не знайдено, тому подається як переказ. Леві додає умову: «якщо це вийде у **відкритих вагах** – це повністю змінює розклад у конкурентоспроможності відкритих ваг США». Поки не вийшло.

Чому це важливо. Практичного небагато (ще один API), але контекстно це третій кадр однієї картини: за три доби Anthropic здешевила Fable на 25%, Google втримала ціну Flash при третьому апгрейді за шість тижнів, Meta виклала мільйонний контекст за \$1,25. Ціна за одиницю інтелекту падає швидше, ніж встигають переписуватись оцінки – а рахунок за підписку при цьому ще й ріжеться на 17% з 14.09.

ТЕМА 9

LWN піднімає ціни вперше за п'ять років - рівно на інфляцію, і пояснює математику вголос

LWN.net · HN 689 балів, 133 коментарі - друга за балами тема доби на HN

Не AI і не гучно, але 689 балів на HN за нотатку про прайсинг - це сигнал. З 15.09 підписка дорожчає: «starving hacker» \$6, «professional hacker» \$11, «project leader» \$19, «maniacal supporter» \$55 на місяць.

Пункту це варте через те, **як** зроблено. Модель підписки з 2002 року, ціну піднімали **двічі за 24 роки**, востаннє на початку 2022. З того часу інфляція в США склала «майже рівно 20%» - і піднімають **рівно на 20%**, з поясненням, що деякі витрати (медстраховка) зросли сильніше, але їх не перекладають на читачів. Плюс окремим рядком причина існування моделі: підписка робить видання незалежним від «волатильного і побудованого на стеженні» рекламного ринку.

Одна деталь прямо в темі тижня: серед витрат, які довелось нести, називають боротьбу зі **скрейперами**. Тобто платять за трафік ботів, що тренують чужі моделі, а рахунок ділять з читачами.

Чому це важливо. Це антипод пункту 3. Там сайти, згенеровані **для читання моделями**, з нульовою редакційною відповідальністю. Тут 24-річне видання, яке живе з того, що його читають люди, і яке пояснює підняття ціни арифметикою. Обидві моделі існують одночасно, і друга платить за виживання першої.

ТЕМА 10

Claude Code v2.1.259: полагодили паралельні сесії, що затирали конфіг одна одній

реліз v2.1.259 (02.09, 22:33 UTC; узято з api.github.com/releases/latest)

Вчора був v2.1.258, це новий реліз. Що всередині:

- Полагоджено: паралельні сесії тихо відкочували зміни `~/.claude.json` одна одної - довіра до воркспейсу скидалась, стан MCP/проекту губився при багатьох одночасних сесіях. Зачіпає будь-кого, хто ганяє кілька сесій паралельно.
- Полагоджено: після одного відхилення thinking у розмові **відхилялось і на всіх наступних ходах**.
- Полагоджено: кеш промπτу інвалідувався при оновленні OAuth-токена в сесіях з вимкненою телеметрією. Дрібниця, яка коштувала грошей на кожному оновленні токена.
- Полагоджено діри в deny-правилах `Read` для Bash: файли, передані як значення опцій (`--ignore-revs-file=.env`, `-f.env`, `@file`), операнди `git diff` / `git grep` і компаунди `cd DIR && cat FILE`. Заборона на читання `.env` **обходила** кількома звичайними способами.

- Додано `--permission-prompts none` для безнаглядних headless-хостів і `managedMcpServers` для організацій.

misc

- **Mistral: «чи можна відмовитись від використання даних для тренування?»** - сторінка підтримки набрала 397 балів і 169 коментарів на HN. Люди читають дрібний шрифт уважніше, ніж рік тому; лягає поруч із дворівневим прайсингом Meta в пункті 8.
- **ФБР розслідує сервіс, що продавав 153 млн водійських посвідчень** (Кребс, 389 балів) - підрахунок @levelsio: це 63% усіх американських посвідчень, і в США вони працюють як основне посвідчення особи. Витік з KYC-сервісу, тобто зі сховища, куди документи здають саме заради безпеки.
- **Anthropic: перевірка, чи створено файл за допомогою Claude** (HN 155) - сторінка живе, віддає 200.
- **Мольк про Codex voice** - «магічно, коли працює», але постійно перемикає треди, багатозадачить через раз і непослідовно дістає решту Codex. Там же **про суть проблеми**: говорити з комп'ютером справді вражає, а от працювати - «де контент, з яким воно працює? Як передати голосу інформацію? Як відновити розмову?».
- **Мольк показав 10 місяців різниці** (33 тис. переглядів) - скріншот AgentBuilder з OpenAI DevDay (жовтень 2025) поруч зі схемою з інциденту Hugging Face (липень 2026).
- **Video Delta Net: відкрита генерація відео швидша за відтворення** (93 тис. переглядів) - прискорення Minimax-H3 у 75-90 разів, 14 секунд 768p за 11 секунд. [одне джерело - анонс авторів, незалежних замірів у вікні нема]
- **Тіму Куку виписали \$47 млн як виконавчому голові Guardian, FT** - раніше ця цифра давалась як заголовок FT без прочитаного тіла. Тепер вона підтверджена другим виданням, тіло прочитане. Дедуп: сама відставка була вчора пунктом 3, повторно не дається.
- **Гартман: «те саме бачу і по Linux»** - винесено в пункт 2, тут лише як маркер: сюжет AISLE вийшов за межі curl.
- **Uber звільняє 10% штату** (HN 112) і йде з Нігерії та Уганди (BBC, HN 114). Того ж дня FT: Uber союзиться з профспілками водіїв, щоб пригальмувати розгортання роботаксі, а BBC катається першим британським роботаксі - з водієм у салоні.