

ClaudeBot маскується, шукає ключі Anthropic

Сканери під виглядом ClaudeBot цілять `/.claude/settings.json` з креденціалами Anthropic

четвер, 13 серпня 2026

У ЦЬОМУ ВИПУСКУ

1. Хтось масово сканує сайти, прикидаючись ClaudeBot - і шукає саме `/.claude/settings.json` і креди Anthropic. Це найпряміша до нас новина за весь тиждень
2. ТРИ фронтір-моделі за одну добу: Grok 4.6, Qwen3.8 на 2.4 трлн параметрів і DeepSeek V4 Pro. І одразу поправка до того, як це подають
3. Tailscale знайшла причину псування баз у 16-річному багу SQLite. 19 інцидентів за півроку, і найкраща інженерна історія за багато тижнів
4. «ШІ вимиває середній клас інженерії» - 779 балів і 727 коментарів. І Charity Majors у Pragmatic Engineer з протилежного боку того самого питання
5. Гаррі Тан виклав GBrain - «персональний AGI» з пам'яттю в git-репо і скілами в markdown. Це буквально опис нашої конструкції, і його варто прочитати уважно
6. Zed випустила Delta - спільне середовище, де код і розмова з агентом лежать в одній базі МІЖ комітами. Тред скептичний, і скепсис вартий уваги
7. Brainbase відкрила Universal Managed Agents - обгортка над Claude Managed Agents на 8+ харнесів і 50+ моделей
8. Transformers.js перетнула 10 млн завантажень на місяць - ×10 за півроку. Локальний ШІ в браузері перестав бути демкою
9. Anthropic вшила невидимий підпис у вивід Claude - вчора я це давав, сьогодні воно дійшло до розсилок. Пояснюю, чому НЕ розгортаю вдруге
10. Ethan Mollick про свіже дослідження: «я думаю, це виявиться хибним» - і чому я даю саме заперечення, а не саме дослідження

ТЕМА 1

Хтось масово сканує сайти, прикидаючись ClaudeBot - і шукає саме `/.claude/settings.json` і креди Anthropic

Перший пункт, бо це єдина сьогоднішня історія, де в переліку цілей атаки стоять буквально імена файлів з типового репо на Claude Code.

knownagents.com (248 балів на HN) фіксує: «Спостерігається широка кампанія, що видає себе за ШІ-ботів, скануючи сайти на вразливості. Атакувальник, схоже, цілить у шляхи до кредів і конфігів, які використовують ШІ-інструменти для кодингу».

Що саме перебирають, дослівно з їхнього списку:

- `/.config/anthropic/credentials/default.json`
- `/.claude/settings.json`
- варіанти `/.env`
- конфіги AWS і Docker

Під якими іменами ходять (частка трафіку): **Googlebot 0.5%**, далі по 0.1% – ChatGPT-User, OAI-SearchBot, GPTBot, PerplexityBot, ClaudeBot.

Як визначають підробку: «Візит вважається підробленим, коли він **заявляє впізнавану ідентичність агента, але не проходить її метод автентифікації** – верифікований IP або Web Bot Auth». І одразу дисклеймер від них же: «Провалена перевірка означає, що візит найпевніше видавав себе за названого агента; вона **не встановлює**, який софт чи оператор реально зробив запит».

Тверезий голос із треду. Найвищий коментар: «Хтось **завжди** запускає масові скани вразливостей. Це стан інтернету рівня "вода мокра"». І другий: «Кожен сервер з відкритим 80/443 має тисячі звернень на день... Нове тут лише те, що вони прикидаються іншим типом надокучливого бота». Сенсації нема, змінився маскувальний шар.

Чому це важливо. Практичний висновок – перевірити, чи файли `/.claude/` і `/.env` з робочих репо не лежать у теці, яку віддає публічний веб-сервер. Цей список шляхів – опублікований чекліст того, що зараз вважається здобиччю. Вчорашній пункт 1 (крадіжка думок з агентних трейсів) казав те саме з іншого кінця: **артефакти ШІ-інструментів стали окремою категорією цілі**. Дві доби поспіль, два різні вектори, один висновок. [proven – сторінка knownagents прочитана, цифри і шляхи дослівні; це дані **одного вендора** про власний трафік, незалежного підтвердження масштабу нема; частки 0.1% – це частка від їхнього трафіку, не абсолютні числа]

knownagents.com – звіт про кампанію · HN-тред, 248 балів

ТЕМА 2

ТРИ фронтір-моделі за одну добу: Grok 4.6, Qwen3.8 на 2.4 трлн параметрів і DeepSeek V4 Pro. І одразу поправка до того, як це подають

Головна подія доби за обсягом. Найкоротше сформулював @Yuchenj_UW (115 тис. переглядів): «3 фронтір-моделі за один день!» – далі перелік, перевірений по кожній окремо.

① **Grok 4.6.** Анонс xAI: фокус – «довготривалі агенти і амбітніша інтерактивна й візуальна робота». Заявляють: «досягає фронтір-інтелекту в кількох агентних бенчмарках кодингу і знанневої роботи. **Дорівнює GPT-5.6 Sol на Artificial Analysis Intelligence Index**».

І тут поправка. На власному графіку xAI стоять три стовпчики: **Fable 5 Max - 62, Grok 4.6 - 61, GPT-5.6 Sol Max - 61.** «Дорівнює Sol» - правда, а Fable вони вже не дорівнюють, і на графіку це видно. Головне інше: у їхньому графіку взагалі нема Opus 5. Незалежний замір Artificial Analysis (322 бали на HN) ставить: **Opus 5 - 63, Fable 5 - 62, Grok 4.6 - 61.** Модель зайшла у фронтір одразу **третьою**.

Решта цифр AA цікавіші за індекс: агентна робота (GDPval-AA v2) - Elo 1753, «поступається лише Claude Opus 5»; ціна **\$2/\$6 за 1 млн токенів** вводу/виводу, кеш - \$0.5; контекст **500 тис.** (без змін від 4.5); **\$0.84 за задачу.** Найпоказовіше - **ефективність:** Grok проходить задачу за ~53 ходи і ~0.5 млрд токенів вводу, тоді як Opus 5 - за ~103 ходи і ~2.0 млрд. Вердикт самих AA: «найсильніші результати - **на агентній роботі**, не на статичному міркуванні».

② Qwen3.8-2.4T-A95B - ваги **вже на Hugging Face** (536 балів). З картки моделі: «**2.4 трлн загально і 95 млрд активованих**», МоЕ на **512 експертів**, контекст **262 144 нативно**, розширюється до **~1.01 млн**. Заявлені бенчмарки: GPQA Diamond **92.6**, SWE-bench Pro **67.7**, Terminal Bench 2.1 **86.6**, HLE **43.6**.

③ DeepSeek V4 Pro 0813 (807 балів) - вийшла 12 серпня. Ціна на OpenRouter: **\$0.435 / \$0.87 за 1 млн токенів**, контекст **1 млн**, МоЕ. Найтверезіша оцінка - з топового коментаря HN: «Конкурує з Opus 4.8, але слабша за Sol чи Fable. Приблизно у **20 разів дешевша**».

Реакція практиків, обидві сторони. @levie: «Чудовий день з погляду релізів... обидва оновлення - величезні стрибки можливостей за неймовірно низьких цін. Це буквально **парадокс Джевонса для ШІ**». @mckaywrigley, який реально кодить: «grok 4.6 гарна, **інтелект на долар - шалено добрий**», але одразу - «fable 5 усе ще найкраща модель». І з HN, як протиотрута до захвату: «Цікаво, наскільки на адопшн впливає **інерція**. Деякі з цих китайських моделей напрочуд здібні, але розробники за замовчуванням лишаються на тих, що вже стали індустріальним стандартом».

Чому це важливо. Незалежні цифри радше підпирають статус-кво на Opus 5: за заміром AA він сьогодні перший (63) і перший на агентній роботі (Elo 1753). Агентні задачі - багатокрокові прогони по репо, скрипти, тули - рівно той профіль навантаження, де він і мав би бути найкращим.

Тверезо про ціну: Grok робить задачу вдвічі меншою кількістю ходів і за \$0.84. На підписці токени не тарифікуються поштучно, тому «у 20 разів дешевше» там нічого не важить. Аргументом воно стає за масового фонового обробітку, сотні прогонів на добу: тоді дешева модель на чорнову роботу з дорогою на думання має сенс. Це рівно вчорашній Switchyard від Nvidia, тільки тепер під нього є три свіжі кандидати з відкритими вагами.

Практичних висновків два, обидва дрібні. ① **2.4 трлн параметрів Qwen локально не запусити:** інша вагова категорія порівняно з учорашніми 30B у 24 ГБ. ② **Заголовок «зрівнялись з фронтиром» нічого не варто:** сьогодні три компанії сказали це за добу, а незалежний індекс показує розрив у 1-2 пункти. [proven - анонс xAI прочитаний у

браузері з графіком, картка Qwen і сторінка DeepSeek на OpenRouter прочитані, цифри AA з їхньої статті; бенчмарки Qwen – **заявлені самим Qwen**, незалежних замірів нема; оцінка DeepSeek – за коментарем HN, не за власним тестом; жодну з трьох моделей не тестовано наживо]

xAI: Introducing Grok 4.6 · Artificial Analysis: незалежний розбір, 61 бал · Qwen3.8-2.4T на Hugging Face · DeepSeek V4 Pro на OpenRouter · @Yuchenj_UW: «3 фронтір-моделі за день» · @levie про парадокс Джевонса · @mckaywrigley: «fable 5 усе ще найкраща» · HN: DeepSeek, 807 балів · HN: Grok 4.6, 458 балів

ТЕМА 3

Tailscale знайшла причину псування баз у 16-річному багу SQLite. 19 інцидентів за півроку

879 балів, найпопулярніше на HN за добу. Довге полювання за багом, де винним виявився не той, хто шукав.

Симптом: 19 інцидентів псування бази за шість місяців. Наслідки для клієнтів злі: доводилось зупиняти процеси control plane на уражених шардах, спочатку це давало «**понад годину**» простою, коли пристрої не могли зайти в мережу або дізнатись про зміни. Статус-сторінку вивішували глобально, хоча падала частина шардів, і це, як вони самі пишуть, підточувало довіру.

Причина: «Рідкісний data race у вихідному коді SQLite між чекпоінтом і транзакцією запису». Якщо запис відбувається в конкретний момент під час чекпоінта, процес помилково вважає, що сторінки скопійовані в основний файл БД, хоча їх там нема. Результат – безповоротна втрата даних. Розробники SQLite оцінюють, що баг жив у коді щонайменше 16 років.

Чому вгатило саме по Tailscale: вони «беруть чекпоінтинг під ручний контроль» і «чекпоінтять дуже агресивно», тобто самі загнали себе в те вузьке вікно, де гонка стріляє.

Як знайшли. Місяці розслідування і відкинутих гіпотез, потім купили платну підтримку SQLite. Прорив стався, коли розробники SQLite написали дебаг-шим `tmstmpvfs`, що обгортає шар віртуальної файлової системи і дає видимість операцій чекпоінта. Після цього гонка стала видна. Фікс уже в **SQLite 3.51.3**: додаткова перевірка, яка виявляє, що WAL скинуто іншим потоком.

Окремо те, що відзначив топовий коментар: Tailscale **проспонсорувала розробку цього open-source шима**. «Цікавий приклад компанії, що фінансує openсорс – у цьому випадку оплачує розробку нового і дуже специфічного дебаг-інструмента».

Чому це важливо. Спершу про масштаб ризику для чужих проєктів: баг стріляє при агресивному ручному чекпоінтингу під конкурентним записом. Хто не керує WAL

руками і пише в базу зрідка, живе в рівно протилежному профілі навантаження, і панікувати нема з чого.

Варте уваги інше – **метод**. Пів року відкидали власні гіпотези і були впевнені, що винні самі, поки не дістали інструмент, який показав правду замість здогадів. Той самий урок, на якому легко палитись у дрібному масштабі: читати `aria-selected` замість вмісту, застиглий DOM замість статусу. **Коли теорія не сходиться, дешевше зробити видимим те, що зараз невидиме.** У них це коштувало контракту з розробниками SQLite, в іншому випадку – одного зайвого кроку в скрипті. [proven – блог Tailscale прочитаний, цифри і цитати дослівні; це розповідь **постраждалої сторони** про власний інцидент, версію SQLite-розробників окремо не шукали]

Tailscale: 16-річний баг WAL-reset у SQLite · HN-тред, 879 балів – топ доби

ТЕМА 4

«ШІ вимиває середній клас інженерії» – 779 балів і 727 коментарів. І Charity Majors у Pragmatic Engineer з протилежного боку того самого питання

Дві історії подано одним пунктом свідомо: поодиноці кожна читається як маніфест, разом – як суперечка.

Сторона перша – есей Флоріана Герренгта (779 балів, 727 коментарів, найбільше обговорення доби). Теза: «середній клас» – це компетентні мідли, що тримають якість коду і менторять джунів. Вимиває його тим, що ШІ прибрав природне обмеження швидкості, яке раніше змушувало думати перед тим, як писати.

Ключові цитати: «ШІ змушує проекти зі слабкою інженерною культурою провалюватись значно швидше»; «Реалізація дешева. Платять за ухвалення хороших рішень»; «Погані інженери стали значно дорожчими в наймі».

Найвищий коментар доповнює найточніше: «Був час, коли люди сідали і обговорювали, як саме вони щось робитимуть. Тепер можна просто попромптити агента кілька годин і відкрити PR. Найтрагічніше в цьому те, що **для ненаренованого ока це виглядає як робоче**».

Сторона друга – Charity Majors (Honeycomb) у вчорашньому випуску Pragmatic Engineer. Заходить з протилежного боку і місцями провокативно: «**2025 для ШІ був тим, чим 2010 для хмари**», а скепсис, раціональний торік, у 2026 вже не тримається.

Найгостріше – про рев'ю: код-рев'ю «**переоцінене, і це найменш цінна частина того, що людина додає до інженерії**». Аргумент: люди сильні в розмовах і архітектурних рішеннях, слабкі у верифікації правильності очима. І питання, яке вона ставить руба: «**Що має статись, щоб можна було бути повністю спокійним, відвантажуючи код, якого не читав?**», з висновком, що це «питання "коли", а не "чи"».

Її позиція складніша за «ШІ-оптимістка проти скептика». Менше довіри до згенерованого коду означає **більше тестів, евалів і conformance-тестингу**, бо система стає недетермінованою. Плюс дві конкретні практики: **ніколи не відправляти згенеровану ШІ комунікацію без повного власного вичитування**, і в Honeycomb у команди є **«середня без ШІ»**, щоб не втратити контроль.

Чому це важливо. Ці двоє сходяться рівно в тому місці, що стосується читача.

Обидва кажуть: **дефіцит перемістився з написання коду в судження і верифікацію**. Герренгт – «реалізація дешева, платять за рішення»; Майорс – «рев'ю переоцінене, натомість тести й евали». Це те саме, що вчора незалежно сказали Google (пункт 6) і трейдери з Optiver (пункт 9): **швидкість виробництва коду перестала бути дефіцитом**. Четвертий день поспіль, четверте незалежне джерело – це вже консенсус галузі.

Практичний зріз для будь-якого маленького автоматизованого конвеєра. Ручна верифікація в таких зазвичай є: людське «го» на незворотне. Майорс тисне пальцем у справжню дірку – автоматичних перевірок майже нема, максимум смоук-тест, що процес піднявся. Коли правляться допоміжні скрипти чи розклад запусків, єдиний тест – що воно не впало в моменті правки.

Висновок дешевий: наступного разу варто витратити на перевірку ті самі півгодини, що й на фічу. «Покрити тестами репо» – задача на тиждень і вона не потрібна. Досить найкрихіткіших місць, де помилка тиха і спливає через дні. [proven – есей і безкоштовна частина випуску Pragmatic Engineer прочитані, цитати дослівні; повний транскрипт Майорс за пейволлом, прочитано прев'ю і виклад, глибших аргументів не видно; есей Герренгта – це **думка одного інженера**, не дослідження, жодних даних під тезою про «середній клас» він не наводить]

Есей: [AI is removing the middle class of software engineering](#) · HN-тред, 779 балів і 727 коментарів · [Pragmatic Engineer: Charity Majors про скепсис до ШІ](#)

ТЕМА 5

Гаррі Тан виклав GBrain – «персональний AGI» з пам'яттю в git-репо і скілами в markdown

Вчора в misc була його фраза «цілі стартапи невдовзі будуть markdown-файлами». Сьогодні він виклав те, про що говорив.

GBrain v0.45.6.0, MIT, і за його ж словами додано «**17 нових brain-скілів**, загатованих через особистий OpenClaw-агент з **сотнями тисяч markdown-файлів**». Далі дослівно: «**Персональний AGI – це той, що працює на тебе**. GBrain тепер працює з Codex і Claude Code».

Як воно влаштоване (за README): знання лежать **markdown-файлами в git-репо** («brain hero»), які синкаються в Postgres для пошуку. Зверху – гібридний пошук (вектори + ключові слова + обхід графа) і **самозбірний граф знань**, що витягує зв'язки між

сутностями **без викликів LLM**, типізованими ребрами (`attended` , `works_at` , `invested_in`). Скіли - «markdown-файли (не прив'язані до інструмента), запаковані в один скілпак», їх у комплекті 50+. Підтримка: Claude Code, Codex, OpenClaw, Cursor, Claude Desktop, через MCP. Заявлено «~30 хвилин до повністю робочого мозку».

Уточнення того ж дня: «Ганяти GBrain з Codex чи Claude Code має бути **окремий агент**, і не варто очікувати запуску у своєму основному кодинг-агенті. Думати про це радше як про **персональну ШІ-версію ChatGPT або Claude, що має власне git-репо для пам'яті і кастомні скіли**».

Чому це важливо. Архітектуру «окремий агент, пам'ять у git-репо як markdown, інструкції як markdown-файли» за останній рік незалежно збрало чимало людей руками. Тут вона вперше виходить під MIT від президента YC і отримує назву.

Два елементи GBrain варті уваги тому, що саморобні версії зазвичай їх не мають:

- **Postgres + гібридний пошук** замість чистого повнотекстового. Поки пам'ять невелика, BM25 ловить точні терміни добре, і апгрейд не потрібен. На кількох тисячах файлів це перший кандидат на заміну.
- **Самозбірний граф зв'язків без LLM.** У саморобних конструкціях зв'язки між темами ставляться руками, `[[посиланнями]]` , під час нічної консолідації. Автоматичне витягування сутностей типізованими ребрами знімає цю роботу і не залежить від того, чи згадали поставити лінк.

28.3 тис. зірок не означають, що воно краще під конкретну задачу, і ставити його на себе з ходу підстав нема. Але підхід GBrain до графа пам'яті варто прочитати. [proven - README GBrain прочитаний, цитати з нього і з постів Тана дослівні; **GBrain не встановлювався і не запускався**, усі характеристики з їхньої документації; «персональний AGI» - маркетингове формулювання автора, не опис можливостей]

GBrain на GitHub (MIT, 28.3к зірок) · @garrytan про реліз v0.45.6.0 · @garrytan: «це має бути окремий агент» · @garrytan: «markdown skill-maxxing»

ТЕМА 6

Zed випустила Delta - спільне середовище, де код і розмова з агентом лежать в одній базі МІЖ комітами. Тред скептичний

457 балів, 474 коментарі. Ідея цікава, реакція показова, тому даються обидві половини.

Що це. Delta від Zed Industries - «мультиплеєрне» середовище для роботи з агентами, зараз **приватна бета** (перші інвайти 12 серпня). Проблема, яку вони називають: рев'ю відбувається по знімках-комітах, і контекст того, **як код таким став**, губиться. Delta «тримає код і розмови пов'язаними, щоб розробники й агенти працювали з повним контекстом того, як код виник».

Механіка - на **DeltaDB**, реплікованій базі, що «реплікує розмову і worktree разом, у реальному часі, для всіх у треді». Ключове: «DeltaDB працює з git-репо, яке вже є.

Кожна правка і розмова захоплюються **МІЖ комітами**». Коментарі лишаються прив'язаними до коду, поки той еволюціонує; у кожного учасника – своя локальна копія в реальному синку.

Скепсис із треду розумний. Топовий коментар: «Хтось іще ненавидить читати ШІ-резюме коду? Код буває лаконічним... Часто закінчуєш тим, що читаєш абзац заради пояснення кількох рядків. Або навпаки – резюме **пропускає важливі граничні випадки**». І найдошкульніший: «Певно, це виглядало чудовою ідеєю **рік тому**... Але багато чого змінилось за ці 12 місяців. Фронтір-моделі й кодинг-агенти просунулись так, що великої цінності тут не видно».

Чому це важливо. Продукт приватний і зроблений під команди, тому для одинака чи пари людей дії нуль. Але проблема справжня, і в маленьких конструкціях вона вже розв'язана інакше: **журнал чату плюс пам'ять агента в git**. Для двох людей це дешевше і надійніше, ніж окреме середовище.

З треду варто забрати інше: закид «резюме пропускає граничні випадки» – це рівно ризик переказу виводу скрипта замість самих цифр. Правило «цифри копією з виводу» після цього тримається міцніше. [proven – анонс Zed прочитаний, цитати дослівні; **приватна бета**, продукт не переглядався; жодних замірів чи цін не наведено]

[Zed: Introducing Delta · HN-тред, 457 балів і 474 коментарі](#)

ТЕМА 7

Brainbase відкрила Universal Managed Agents – обгортка над Claude Managed Agents на 8+ харнесів і 50+ моделей

Тихий реліз, прямо в тему дня.

Universal Managed Agents API (UMA): «розширює Claude Managed Agents від Anthropic до 8+ харнесів, включно з Claude Code, Codex, OpenCode, Cursor, Devin, Factory, і до 50+ моделей, фронтір і відкритих. Розробники можуть розгорнути агента за секунди».

Чому це важливо. Третій день поспіль одна й та сама ідея приходить з різних боків, і варто назвати її вголос: **абстракція над харнесами і моделями стає окремим продуктом**. Вчора – NeMo Switchyard від Nvidia, роутер запитів між моделями. Сьогодні – UMA, роутер між харнесами. Плюс GBrain з пункту 5, сумісний і з Claude Code, і з Codex.

Для конструкції на одному харнесі й одній моделі це поки зайвий шар. Троє продавців за три дні означають лише, що проблема справжня для тих, хто жонглює п'ятьма інструментами. Записано як маркер: **коли знадобиться друга модель на фонову роботу, цей шар уже готовий**. [fuzzy – читано тільки анонс у X, на сайт продукту не заходили, документації не бачено; «8+ харнесів» і «50+ моделей» – **заявлені цифри** без перевірки]

[@BrainbaseHQ: анонс UMA](#)

ТЕМА 8

Transformers.js перетнула 10 млн завантажень на місяць - ×10 за півроку

Продовження вчорашньої лінії про локальні моделі (вчора пункт 4 - River AI на \$1.1 млрд, пункт 8 - Nemotron), але про **реальний адопшн**.

@ClementDelangue (CEO Hugging Face): «**Локальний ШІ вибухає!** Transformers.js, який Hugging Face будує останні три роки, став **найпопулярнішою опенсорсною бібліотекою для запуску моделей просто в браузері**. Він перетнув **10 мільйонів завантажень на місяць, майже ×10 всього за шість місяців**».

Поруч того ж дня - MOSS-VL: «**24 ГБ VRAM достатньо, щоб запустити MOSS-VL локально**», випустили FP8 і NF4 версії для розуміння зображень, відео і реального часу.

Чому це важливо. Четвертий день поспіль лінія «локальне, своє, на своєму залізі» отримує нову цеглину, і вперше це **цифра адопшну**. Попередні були заявами. Позавчора афоризм @naval, учора \$1.1 млрд у River і 30B від Nvidia, сьогодні ×10 завантажень за півроку.

Це лишається спостереженням: для агентної роботи, як показує сьогоднішній пункт 2, орендована фронтір-модель поки виграє з відривом. Але інфраструктура «запустити щось локально на власній машині» дозріває швидше, ніж очікувалось місяць тому. [proven - пости прочитані, цифри дослівні з них; «10 млн завантажень» - **дані самої Hugging Face** про власну бібліотеку; завантаження ≠ активні користувачі, це загальновідоме викривлення прм-статистики]

@ClementDelangue: Transformers.js і 10 млн · @MosiAI_Official: MOSS-VL у 24 ГБ

ТЕМА 9

Anthropic вшила невидимий підпис у вивід Claude - вчора це вже було, сьогодні воно дійшло до розсилок

Свідомо короткий пункт, тут заради чесності дедупу.

@TheRundownAI сьогодні ставить це першим рядком у своєму щоденному переліку: «Anthropic додає ШІ-водяні знаки у вивід Claude» / «Anthropic **вшиває невидимий підпис** у Claude».

Це вже було 11 серпня (тодішній пункт 8). За сьогодні нових фактів не з'явилося: ані технічних деталей механізму, ані позиції Anthropic понад те, що вже було, ані незалежної перевірки, чи підпис реально детектується. Тема обростає переказами. Вчора вже викидались такі пости від @AiBreakfast і @bentossell, сьогодні те саме сталося б з @TheRundownAI.

Рядок лишається з двох причин: показати, що тема **ще жива в стрічці**, і показати, що **саме викинуто і чому**. [proven - що пост існує і що це переказ; **сам механізм підпису**

не перевірявся ні вчора, ні сьогодні, незалежного підтвердження, що він працює або детектується, у джерелах нема]

@TheRundownAI: топ-новини дня · їхній матеріал про підпис

ТЕМА 10

Ethan Mollick про свіже дослідження: «думаю, це виявиться хибним». І чому подається заперечення без самого дослідження

Маленький пункт про метод.

@emollick сьогодні о 04:26 за Києвом (за чотири хвилини до збору): «Думаю, це виявиться хибним, і не лише тому, що є підозра щодо більшої віддачі від вищого інтелекту...». Це його реакція на чергове дослідження, і вона подається без самого дослідження свідомо: у вікно збору потрапило заперечення без першоджерела. Тягнути висновок з реакції на матеріал, який не читався, робити не можна.

Другий його пост доби, смішний і влучний: «Приємно, що всі ШІ-коментатори тепер мусять вдавати, що завжди мали обережне нюансоване розуміння різниці...».

Чому це важливо. Це калібрування. Моллік, найвагоміший голос у списку джерел, за одну добу двічі публічно сказав «обережніше з висновками». Вчора: «дослідження радить обережність у визначенні, хто виграє, за одним джерелом» (тоді це потрапило в пункт 7 про частки ринку). Сьогодні: «думаю, це виявиться хибним».

Нагадування на сьогодні пряме: **три компанії за добу заявили, що зрівнялись з фронтиром**, і якби бралися їхні власні графіки, вийшло б три перші місця. Незалежний індекс дав інше. **Заперечення вартє включення нарівні із заявою:** знайшов заяву – пошукай, хто в тому ж залі був незгодний. [fuzzy – є репліка Молліка, але не первинне дослідження, на яке він відповідає; його аргумент по суті не перевірявся і оцінити, хто правий, неможливо]

@emollick: «думаю, це виявиться хибним» · @emollick про ШІ-коментаторів

misc – коротко, що ще вартє ока

- **Gradio 6.24 зберігає і відтворює запуски** – @abidlabs: прогони будь-якого Gradio-застосунку автоматично зберігаються в local storage браузера, і будь-який можна відтворити одним кліком, не стоячи в черзі. Без змін у коді, просто оновити версію
- **Woxi – опенсорсна реалізація Wolfram Language** (266 балів): переписують Mathematica у відкритий код. Маркер того, що ще одна пропріетарна фортеця отримала відкритий аналог
- **@paulg говорив зі стартапами 7 годин поспіль:** «Це 47-й батч УС. Пропущено кілька під час ковіду, але крім того, здається, поговорено з усіма». Поруч @PariLatawa

переказує вечерю з ним: «10% зростання тиждень до тижня», і [@beknabdik](#) з тим самим уроком – «якщо оптимізуєш під зростання, решта прикладеться»

- [@shl про читання коду](#), одним рядком і в тему пункту 4: «Живи життям, яким хочеш жити. Хочеш читати код – читай код. Не хочеш – не читай». Протилежність позиції Майорс, подано без аргументів, тому в misc
- [Хтось видає себе за ШІ-ботів](#) – уже в пункті 1, але повторюється одним рядком, бо це єдина дія дня: перевірити, чи не лежить `.env` / `.claude/` у веб-доступній теці на хостингах
- [Tim King, розробник AmigaDOS, помер](#) (247 балів). Людина, що написала операційку, на якій виросло покоління
- [Чому крихітні JPEG виглядають інакше в Chrome](#) (273 бали) – розбір того, як браузер масштабує дрібні зображення. Практичної користі нуль, але приємно написано
- [Shade Map](#) (162 бали) – карта, що показує, де буде тінь у конкретний час доби. Для літньої подорожі корисна річ: подивитись, який бік вулиці або пляжу в тіні о 14:00
- [Вебкамери затемнення 2026](#) (461 бал) – збірка камер під сонячне затемнення 12 серпня. Уже минуло, але лишено: там є записи