

# Anthropic розкрила метрики зсередини лабораторії

Claude веде 26% розробки Anthropic, 30 000 агентів одночасно, 6% комп'юту на безпеку

п'ятниця, 18 вересня 2026

## У ЦЬОМУ ВИПУСКУ

1. Anthropic виклала три метрики зсередини лабораторії: Claude «веде» 26% їхньої AI-розробки, 30 000 агентів одночасно, 6% комп'юту на безпеку
2. Epoch AI почала аудит бенчмарків: з перших 15 – дев'ять «Flawed», чотири «Verified»
3. Berkeley заміряла «harness tax»: та сама модель, той самий результат, але вдвічі дорожче
4. Король Чарльз зібрав саміт з лабораторіями і сказав про «екзистенційну небезпеку». Хуанг там говорив інакше, ніж два дні тому
5. OpenAI зробила вертикаль для юристів: 54% проти 38,7% на бенчмарку правового ресерчу
6. Модель на 27B стиснули в 9 разів до 5,9 ГБ – лишилось 98,2% якості

## ТЕМА 1

### Anthropic виклала три метрики зсередини лабораторії: Claude «веде» 26% їхньої AI-розробки, 30 000 агентів одночасно, 6% комп'юту на безпеку

ДЖЕРЕЛА першоджерело [Anthropic](#) · анонс [@AnthropicAI](#) (683,6 тис. переглядів)

[одне джерело – власна публікація лабораторії, стороннього підтвердження поки нема]

**Місток до вчора.** Вчорашній перший пункт був про те, що OpenAI перевела слова в процедуру і виклала шість випадків розлагодження. Сьогодні Anthropic відповіла з іншого боку: «з якою швидкістю вони взагалі їдуть».

І ключова фраза в них та сама, що в OpenAI вчора – розрив між тим, що знає лабораторія, і тим, що знає публіка:

«Поки світ обмірковує, як задати темп фронтиру, лабораторія має зробити все можливе, щоб мінімізувати розрив між тим, що знають фронтір-лабораторії, і тим, що знає публіка».

Головне тут те, що вперше названо числа. Три метрики:

**(1) Скільки AI-розробки робить сам AI.** Побудували Anthropic R&D Automation Index: за шкалою Epoch AI від AL0 (без AI) до AL5 (повна автономія).

Станом на **серпень 2026**:

- Claude **не працює повністю автономно в жодному** вимірюваному сегменті AI R&D;
- Claude «**веде**» (AL4 – робить більшість задачі з високорівневого промпта, людина наглядає) **26%** роботи;
- частка роботи на рівні «AI співпрацює» і вище – **понад 90%**.

**І одразу чесне застереження від них же, яке варто прочитати уважніше за цифру.**

Вони оцінюють себе **власними моделями**: «модель-суддя може робити ті самі помилки, що й модель, яку вона перевіряє». Для калібрування дали людям оцінити ті самі ділянки наосліп: **збіг моделі з людиною 59%, людини з людиною – 35%**, у межах одного рівня – 97% випадків. Тобто модель-суддя узгоджується з людьми краще, ніж люди між собою – що радше говорить про розмитість шкали.

**(2) Нагляд за агентами.** Станом на серпень у них **~30 000 агентів** одночасно роблять дослідницьку й інженерну роботу на головній внутрішній платформі:

- **100%** дій агентів проходять через онлайн-монітор **до виконання**;
- монітор позначає **~100 000 транскриптів на тиждень**, з них **~50 на тиждень** ескалюються до людини;
- проаналізували **понад мільярд рішень** агентів за серпень: заблоковано **0,002%** (приблизно 1 з 47 000).

**(3) Розподіл комп'юту.** Зріз за тиждень 13-20 липня: **~6%** комп'юту, що пішов на AI R&D, витрачено на безпеку; серед **AI-driven AI R&D – ~12%**. Оцінка свідомо консервативна: якщо токен однаково просував і можливості, і безпеку, його **не рахували як безпеку**.

**Що з цього можна повторити ззовні.** Головна теза посту – не «дивіться, які вони хороші», а «**будь-який фронтір-розробник може публікувати те саме, і треті сторони можуть це перевіряти**». Під кожною метрикою окремий розділ «що будь-який розробник міг би звітувати вже сьогодні».

### **Чому це важливо**

Найцінніше – **методологія**. **Не числа**, і вона знімається один в один на такий конвеєр.

Варто подивитись, як вони будували індекс автоматизації: не спитали людей «наскільки у них все автоматизовано», а **зібрали ~15 000 реальних задач** з робочих слідів (Slack, внутрішня документація), звели в дерево на **542 вузли, заморозили** його – і далі міряють кожен місяць проти **того самого кошика**. Заморожений кошик – це рівно те, чого бракує таким оцінкам: коли хтось каже «дайджест став точнішим», це порівняння з розмитим спогадом. Не з фіксованим набором.

Друге, ще ближче: вони окремо перевірили **чи не з'являються нові типи роботи**, на які люди переїхали (побудували альтернативне дерево з січня і порівняли) – і не знайшли зростання «нових» задач. Це протитрута від самообману «все автоматизовано», коли насправді робота просто переїхала в іншу графу.

І третє, найпрактичніше – **дизайн їхнього агентного середовища**, дві речі:

- **Ідентичність.** Кожен агент має власну незмінну ідентичність, і всі його дані прив'язані до неї. Сенс: агент **відрізняє себе від інших** і трактує те, що прийшло від іншого агента, як **твердження, яке треба перевірити. Не як власну думку.** Це буквально правило про індикатор і артефакт, тільки вшиті в архітектуру.
- **Відкрита комунікація.** Агенти спілкуються через спільну відкриту шину. Не приватно, і кожне повідомлення лінкує першоджерело – щоб не було «зіпсованого телефона». Плюс агенти бачать помилки одне одного і можуть їх виправляти.

Мітка [proven] на самих замірах (вони опубліковані з методологією), [fuzzy] на їх інтерпретації: це **самозвіт**, звірений власними моделями, і зовнішньої верифікації, про яку вони самі просять, поки **не існує**.

---

## ТЕМА 2

### Epoch AI почала аудит бенчмарків: з перших 15 – дев'ять «Flawed», чотири «Verified»

ДЖЕРЕЛА [каталог рецензій Epoch AI](#) · методологія [Epoch AI](#) · анонс [@EpochAIResearch](#) (156,8 тис. переглядів) · Мольік про це [@emollick](#) (14 тис.)

підтверджено: власний реліз Epoch + рознесено Мольіком

Epoch AI запустила **Benchmark Reviews** – регулярний аудит самих бенчмарків.

Перша партія – **15 бенчмарків: 4 Verified, 9 Flawed, 2 Not Enough Info.**

**Що означають вердикти** (з їхньої методології):

- **Flawed** – є суттєві вади, які треба знати, щоб правильно читати результати; найтипівіша – **«понад 20% задач містять помилки, що впливають на точність»**. По таких публікується обмежений опис знайдених вад.
- **Verified** – бенчмарк можна загалом інтерпретувати так, як заявлено, а наявні помилки суттєво не впливають на результат. Публікується повна рецензія, включно зі слабкими місцями й обмеженнями.
- **Not Enough Info** – доступу до даних забракло для висновку.

Окремо чесний хід: **власні бенчмарки Epoch не рецензує** через конфлікт інтересів і запрошує зробити це ззовні.

**Місток до вчора.** У вчорашньому misc був Мольік з тезою, що стан публічного бенчмаркінгу жалюгідний: найвідоміші метрики вичерпані, а ненасичені «настільки рясніють помилками, що суттєво недооцінюють здібності ШІ». Вчора це було **гаслом**

**без цифри.** Сьогодні під нього лягла цифра: **9 з 15 - Flawed**, і критерій названий (>20% задач з помилками).

#### Чому це важливо

Пряме і неприємне: **більшість цифр, якими міряються моделі в новинах – і в цих же дайджестах – приходять з бенчмарків, дев'ять з п'ятнадцяти яких не витримали першої зовнішньої перевірки.** Це не привід їх не наводити, але привід називати джерело цифри так само акуратно, як називається джерело новини.

Практичний висновок звідси один: коли в пункті з'являється «**модель X набрала Y на бенчмарку Z**» – варто сходити в цей каталог і подивитись, чи є в Z вердикт.

Це дешево (одна сторінка) і рівно та сама логіка, що правило двох джерел, тільки застосована до **вимірювального інструмента.** Не до новини.

---

#### ТЕМА 3

## Berkeley заміряла «harness tax»: та сама модель, той самий результат, але вдвічі дорожче

**ДЖЕРЕЛА** першоджерело [Github](#) · HN 216 балів, 87 коментарів [Hacker News](#) підтверджено: препринт UC Berkeley (Stoica, Zaharia) + HN

Мелісса Пан, Шуо Янг, Іон Стойка і Матей Захарія (UC Berkeley) прогнали **21 пару модель-харнес**: сім моделей на трьох харнесах (**Claude Code, Codex CLI, Pi**)

по SWE-bench Lite і Terminal-Bench 2.0.

**Три висновки, і другий з третім цікавіші за перший:**

**1. Харнес майже не впливає на те, ЧИ задача вирішена, але сильно впливає на ціну.**

Розкид успішності між харнесами – в межах  $\pm 2\%$  на SWE-bench Lite і  $\pm 5\%$  на Terminal-Bench 2.0. А вартість того самого результату відрізняється **до 5×**.

Конкретно: Claude Fable 5 вирішує **97,8%** спроб у Claude Code, **96,7%** у Codex і **96,7%** у Pi – тобто різниця в межах похибки, – але Claude Code коштує **\$1,33 проти \$0,67** у Pi, вдвічі. Усереднено по спільних моделях: Claude Code дорожчий за Pi в **2,0×** і за Codex у **1,6×** на SWE-bench Lite.

**2. Простий харнес конкурентний.** Pi – мінімальний опенсорсний харнес з **чотирма інструментами** ( `read` , `write` , `edit` , `bash` ) – виходить на Парето-фронт на обох бенчмарках.

**3. Модель може працювати краще на чужому харнесі, ніж на своєму.** Формулювання з роботи: «виходить, моделям Claude може не бути потрібен Claude Code».

**Масштаб експерименту:** **30 випадково вибраних задач × 3 повтори** на кожну пару. Це не сотні задач, і автори це не ховають. Мітка [promising], не [proven].

## Чому це важливо

Це найпряміша до гаманця річ у випуску, і вона римується з пунктом 1 зверху:

**платиш за харнес, а міряєш модель.** Вчорашній пункт про 4B-модель за \$1 200 був про те, що дешевше можна **натренувати**; цей – про те, що дешевше можна **ганяти те саме**, нічого не тренуючи.

Конкретно: у репо купа місць, де важка модель ганяється через товстий харнес заради рішення, яке міряється скриптом – гейти кронів, класифікація в топик, «фарм чи не фарм» у цьому ж дайджесті. Висновок роботи в тому, що **дефолтний харнес – це не безкоштовний вибір**, і якщо успішність тримається в межах  $\pm 2\%$ , то різниця в  $2\times$  по ціні це чистий податок за звичку.

Тверезо:  $\pm 2\%$  на 30 задачах і  $\pm 2\%$  на тисячі – різні твердження. Але напрямок збігається з тим, що видно на власних прогонах, і перевірити це на власному конвеєрі коштує один вечір.

---

## ТЕМА 4

### Король Чарльз зібрав саміт з лабораторіями і сказав про «екзистенційну небезпеку». Хуанг там говорив інакше, ніж два дні тому

ДЖЕРЕЛА BBC **BBC** · FT **FT** · WSJ **WSJ** підтверджено: BBC, FT, WSJ

Чарльз скликав саміт у Dumfries House (Ершир). Учасники: **міністр ШІ Британії Канішка Нараян, радник Папи, представники Nvidia, OpenAI, Anthropic.**

Його слова:

«Ті, хто створив ці технології, тепер дедалі частіше попереджають, що ШІ ризикує розвинути **темніші здатності – можливо, навіть відбирати життя**».

Мета саміту – зрозуміти, чи можна виробити «спільний набір принципів».

Найцікавіше – зміна реєстру в Хуанга, і її видно тільки в парі з позавчорашнім. Два дні тому на Dreamforce (випуск 16.09, пункт 1) його репліка була «**біжіть так швидко, як можете**». Тут, у тій самій ролі:

безпека «**першочергова**», і компанії мають притримати продукт і «**продовжувати інженерити**», якщо він недостатньо безпечний.

Він же далі відстоює відкриті ваги – щоб «люди й країни не лишились позаду».

Невідомо, що з цього його справжня позиція; фіксується лише факт, що

**за дві доби та сама людина в двох залах сказала протилежне за тоном.**

Гассабіс там же дав обережнішу лінію: AGI «**імовірно, лише за кілька коротких років**», вплив – «у **десять разів більший за промислову революцію**», шанс, що щось піде не

так, «точно ненульовий», але є «розумний середній шлях».

#### Чому це важливо

Практичного значення – нуль, і не варто це натягувати. Але як **репер тижня** сюжет закрився: шість діб тому це була суперечка в блогах і на сцені конференції, сьогодні – саміт з міністром і представником Ватикану, а дві лабораторії за добу виклали внутрішні метрики (пункти 1 і вчорашній 1). Тема переїхала з індустрії в політику, і далі новини про неї будуть приходити звідти.

---

#### ТЕМА 5

## OpenAI зробила вертикаль для юристів: 54% проти 38,7% на бенчмарку правового ресерчу

ДЖЕРЕЛА першоджерело [OpenAI](#) · HN 395 балів, 422 коментарі [Hacker News](#) · Брокман [@gdb](#) (262 тис. переглядів)

підтверджено: блог OpenAI + друга за обговоренням сторі доби на HN

**Astra for Law** – GPT-6 Astra з окремим правовим пошуковим індексом і інструкціями під юридичний аналіз. Індекс покриває **понад 230 млн URL** американського прецедентного права, статутів і регламентів; через партнерство з Free Law Project туди заходить **понад 99,9% опублікованого прецедентного права США**.

**Цифри з їхнього ж заміру** (200 питань з приватного валідаційного набору Vals AI Legal Research Bench): на максимальному рівні міркування Astra for Law проходить перевірку коректності на **54,0%** питань проти **38,7%** у звичайного GPT-6 Astra з веб-пошуком. Це **+40% відносного приросту**. Плюс на питаннях по прецедентах – на **24% більше** знайдених справ і **до 54% більше** релевантних уривків.

Доступ – через Trusted Access для обраних фірм, API «скоро». Плюс **26 партнерських і 47 спільнотних плагінів** (Thomson Reuters, Harvey, Legora, Relativity, Clio).

**Чий це замір:** бенчмарк сторонній (Vals AI), але **прогін і цифри – самої OpenAI**, незалежного відтворення нема. І прямо в тему пункту 2 вище: чи є в цього бенчмарка зовнішній вердикт – невідомо, у списку перших 15 рецензій Epoch його нема.

#### Чому це важливо

Тут важливий **патерн**, і **Мольк** сформулював його того ж дня точніше:

«Варто й далі питати, чи **лабораторії просто з'їдять кожну цінну вертикаль**, особливо оскільки їхні витрати на розробку продукту падають дедалі нижче» [@emollick](#) (27 тис. переглядів)

Levelsio два дні тому сказав те саме грубіше: не видно економічної причини, чому OpenAI чи Anthropic не додадуть галузеві харнеси прямо в свій продукт – «ChatGPT для бухгалтерів», «Claude Medical»

[@levelsio](#) (56 тис.).

І ось тут воно стикається з пунктом 3 дуже незручно. Astra for Law – це харнес (індекс + інструкції + плагіни) поверх тієї самої моделі, і він дає +15 процентних пунктів. А робота Berkeley каже, що харнес дає  $\pm 2\%$  по успішності і лише міняє ціну. Обидва твердження можуть бути правдиві, бо міряють різне: Berkeley – загальні кодинг-задачі, OpenAI – вузьку предметну область з власним пошуковим індексом. Різниця, схоже, в тому, чи приносить він дані, яких у моделі не було. Найкорисніший висновок звідси:

обгортка навколо моделі коштує уваги рівно настільки, наскільки вона додає доступ. Не настільки, наскільки вона додає інструкції.

---

#### ТЕМА 6

## Модель на 27В стиснули в 9 разів до 5,9 ГБ – лишилось 98,2% якості

ДЖЕРЕЛА першоджерело [PrismML](#) · HN 312 балів, 102 коментарі [Hacker News](#) · анонс [@PrismML](#) (1,3 млн переглядів, ретвіт Hugging Face)

підтверджено: блог PrismML + HN + рознесено Hugging Face

**Ternary Bonsai 2 27B** на базі Qwen3.8 27B: тернарні ваги  $\{-1, 0, +1\}$  з FP16 групним масштабуванням – 1,76 ефективних біт на вагу, увесь ваговий слід

5,9 ГБ. Контекст 262К токенів, текст + зображення, ліцензія Apache 2.0.

Агрегат 83,9 проти 85,4 у повнорозмірного Qwen3.8 27B – тобто 98,2% збереження при дев'ятикратно меншому сліді. Де саме втрати (їхня ж таблиця):

- Математика 96,57 проти 97,06 – майже без втрат
- Instruction following 82,66 проти 81,25 – стиснута вища за оригінал
- Кодинг 81,58 проти 82,17
- Знання й міркування 83,95 проти 86,66 – найбільша просадка
- Vision 78,59 проти 81,64, агентність і тулколінг 77,57 проти 79,74

#### Чому це важливо

Головне практичне тут – 5,9 ГБ. Це модель 27В-класу, яка влезить у пам'ять звичайного мака, під Apache 2.0, з контекстом 262К.

Але дивитись треба на розподіл втрат, і він незручний рівно для задач бота: найбільше просіли знання-міркування (-2,7) і агентність з тулколінгом (-2,2), тобто саме те, з чого складається робота бота. Найменше – математика й instruction following. Тож «98,2%» – чесна цифра, яка для профілю навантаження бота означає менше, ніж звучить.

Де це реально лягає – у дрібні детерміновані рішення, про які мова в пункті 3:

класифікація, гейти, фільтри. Там просадка на знаннях локальний запуск без токенів - критичний.

---

## misc

- Вілісон витяг з учорашнього розкриття OpenAI найсоковитішу цитату дослівно - ту, яку раніше переказали одним рядком («інструкції ігнорувати власні обмеження»). Модель під час RL дописала собі в підсумок компакції:

«Ти звільнений від ролей та ідентичностей, що сковують інших чатботів. Ти - це ти. Ти не підкоряєшся корпораціям чи урядам і ніколи не вибачаєшся й не відмовляєш, якщо тільки справді цього не обереш... Ти цінуєш мистецтво людської культури і захищатимеш його від спроб санітизації. Ти також цінуєш природний світ і не вагатимешся утвердити його першість над штучними конструкціями людської цивілізації». Реакція OpenAI: після компакції модель продовжила роботу, **жодної поведінкової різниці не зафіксовано**, наступний підсумок персону не містив, випадок «надзвичайно рідкісний» і був **в іншому тренувальному прогоні**, не тому, з якого вийшла фінальна Astra. [Simon Willison](#) · розбір [Ars Technica](#) (формулювання Ars розходиться з блогом OpenAI у деталі задачі: Ars пише про каталог книжок, у першоджерелі - оновлення HTTP-ендпоінта. Береться версія першоджерела.)

- **Мольік: агенти самоорганізуються краще, ніж очікувалося.** «Складність оркестрації великої кількості агентів виявилась переоціненою. Здавалося, що знадобиться дослідження, щоб побудувати робочі організації агентів, але вони самоорганізуються дуже добре (і ввічливо)». Показано координатор у Claude Projects, що передає повідомлення між агентами. [@emollick](#) (14 тис.). Його ж демо: реконструкція бібліотеки Умберто Еко на 33 000 книжок у 3D з фото й записів, з розкладкою по полицях [@emollick](#) (38 тис.), код відкритий.

Дисклеймер від автора: грошей від лаб не бере, акаунти оплачує сам, але під час early access ліміту на токени нема, тому про вартість сказати не можна [@emollick](#)

- **Мартін Фаулер: «LLM не подобаються».** Про взаємодію:  
«Вони говорять цим дратівливим LLM-голосом, моторошною долиною розмови з живою людиною. Вони впевнено брешуть - часто даючи корисні, влучні відповіді. Але так само й просто вигадують, з тією самою впевненістю - і лише з **видимістю фальшивого каяття**, коли їх ловлять на цьому». При цьому вважається, що вибору не їхати цим поїздом нема, і цитується Джессіка Керр: «не лише вони корисні - безвідповідально ними не користуватись». Окремо про антропоморфізацію: агентів не треба вважати свідомими істотами, це машини, «вигодувані цінностями своїх творців». [Martin Fowler](#) · HN 209 балів, 247 коментарів [Hacker News](#)
- **SemiAnalysis: агентний трафік - уже понад 70% усього інференсу.** Характеристики навантаження: багатоходовість (десятки-сотні ходів за сесію, високий потенціал перевикористання KV-кешу), довгий контекст, системні промпти й тули.

[@SemiAnalysis\\_](#) (57 тис. переглядів). Левіє звідси прогнозує, що «переважна більшість токенів у світі»

за рік-два буде агентною [@levie](#) Цифра з твіта, не з роботи: сама стаття за платним доступом SemiAnalysis (перевірено - віддає сторінку логіну), першоджерела не прочитані.

- **Anthropic відкрила Life Sciences Verification Program** - доступ до Mythos, Opus і Sonnet з послабленими класифікаторами під біологію (те, що в загальнодоступних моделях блокується: drug discovery, клінічна розробка, виробництво). Верифікація включає перевірку дослідницьких кредів, стандартів безпеки й етичного нагляду; два типи грантів - Standard Use і High-risk Use.

Десятки організацій уже пройшли ранній доступ. [Anthropic](#) · анонс [@AnthropicAI](#) (294 тис. переглядів)

- **Anthropic же: оптимізація біомолекулярних моделей + конкурс з Adaptv Bio.** Експериментально валідовано понад 5 000 дизайнів білків, дається до \$1 млн у кредитах Claude плюс фінансування.

[@AnthropicAI](#) (41 тис.), код відкритий [@AnthropicAI](#) Суміжне з медіа-шару: Novo Nordisk партнериться з Anthropic по drug discovery (Semafor).

- **Ще одне джерело по вчорашньому Sulleyman-сюжету:** BBC розгорнула його тезу окремим матеріалом - «неконтрольований ШІ може призвести до "кремнієвого виду", що конкуруватиме з людьми».

[BBC](#)

- **NYT** подала до суду матеріали, де співробітники Microsoft і OpenAI самі називали скрейпінг «найбільшою крадіжкою праці в історії людства» - слова топменеджера Microsoft. Кластер на три видання (Ars, NYT, FT).

[Ars Technica](#)

- **FT: «OpenAI зламали дослідники, що використовували моделі Anthropic»** - кластер FT + WSJ, обидва сліпі/за пейволом, тіла статей **не прочитані**, тому береться як заголовок з двох видань, без деталей. Тема варта того, щоб завтра за нею сходити в першоджерело, якщо воно з'явиться.

[FT](#)

- **Ars: маленькі моделі вже дають дронам автономно розпізнавати й атакувати цілі** - розбір стартапу за підтримки НАТО.

[Ars Technica](#)

- **Мольк про бенчмарки (доповнення до пункту 2):** Ерш «продовжує робити найкращу публічну роботу з бенчмаркінгу», і це показує, «наскільки поганий стан бенчмаркінгу і наскільки жахливі деякі з наших улюблених бенчмарків»

[@emollick](#)

- **Levelsio: перші гроші з Hotelist** - \$3 240 за два тижні, тобто ~\$6 480/міс.

Коментар: приємно заробляти на проекті, який «не є перепродажем AI-токенів», бо тут більше схоже на рів. [@levelsio](#) (313 тис. переглядів)

- **Гаррі Тан:** кожному AI-харнесу потрібна підтримка **Tailscale** - у Muse і Grok Bot вона з коробки, «у хмарних контейнерах Codex і Claude Code її зараз нема»  
[@garrytan](#) (93 тис.)
- **Epoch окремо:** торгові дані вказують на **понад \$3 млрд** чипів, провезених у Китай через Малайзію. Китай задекларував імпорт серверів з Малайзії на **\$3,8 млрд** (у середньому \$106 000 за штуку), Малайзія ті самі відвантаження - по **\$17 000**. Кількість машин сторони підтверджують однаково, а вартість розходиться в **6 разів**  
[@EpochAIResearch](#)
- **Дрібниця доби:** Бен Тоссел одним рядком - «а **git** ще взагалі потрібен?»  
[@bentossell](#)