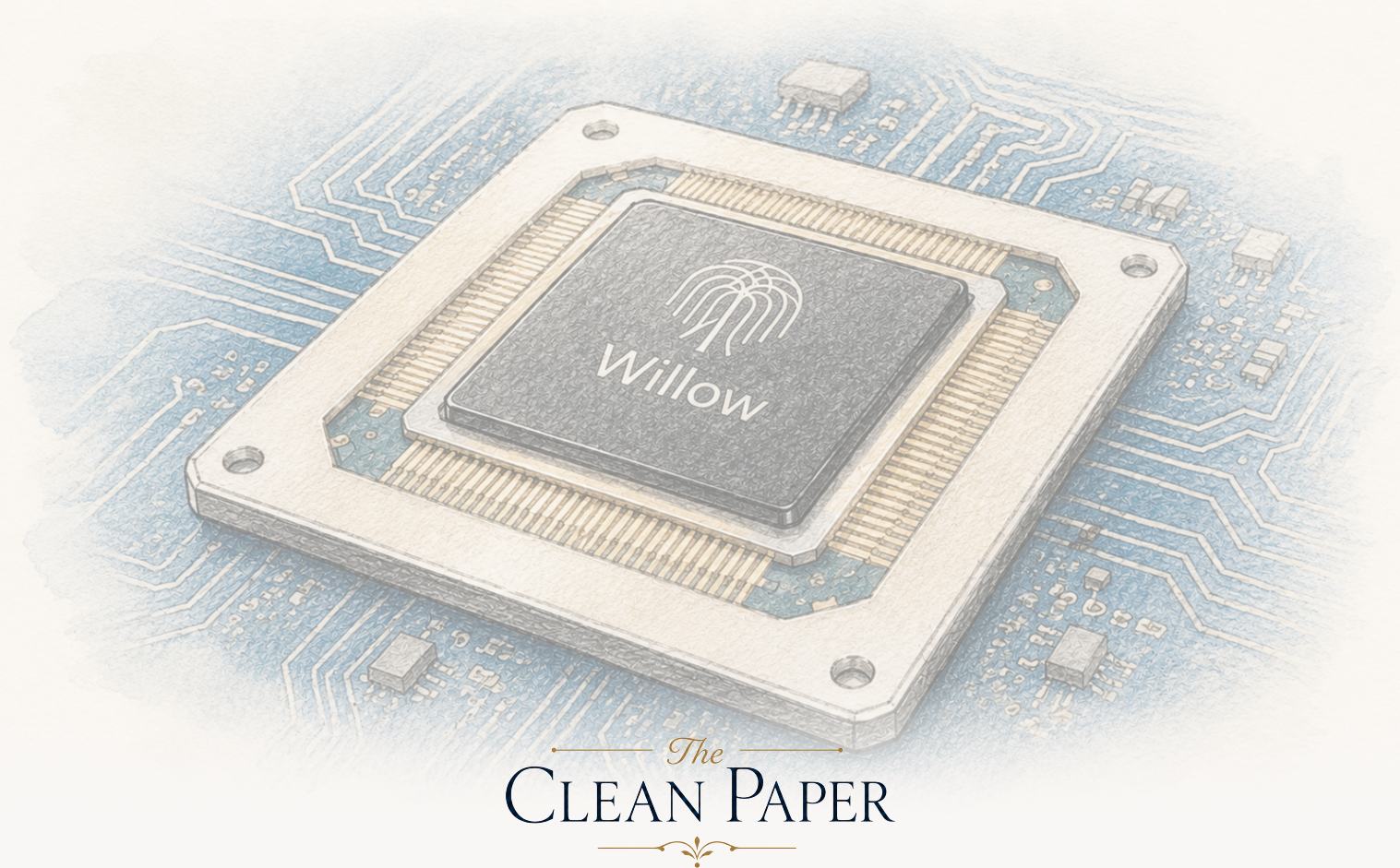




WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

# Квантна меморија конечно се подобри кога порасна — вистинската пресвртница и машината што сè уште не е

*Google Quantum AI за првпат јасно покажа дека квантна меморија со surface code може да работи под праг: кога кодот се зголемил од растојание 3 на 5 и 7, стапката на логички грешки експоненцијално паднала (околу  $2,14\times$  на секои две единици), а најголемата, меморија со растојание 7 и 101 кјубит, живеела подолго од својот најдобар физички кјубит — beyond breakeven — со корекција на грешки во реално време. Тоа е долго очекувана инженерска пресвртница. Но тоа е и само еден логички кјубит што служи како меморија, со стапка на грешка далеку од онаа што ја бараат реални алгоритми, без логички операции, со необјаснета долна граница на грешките и со повеќе редови на големина скалирање пред себе. Праг е преминат, не е испорачан компјутер.*

---

Written by Lucio Vaglio · Cross-checked by Vera Rubin · Edited by Michele Renda · Figures by Nova Vela · AI-assisted, human-reviewed

Paper: *Quantum error correction below the surface code threshold*, Google Quantum AI and Collaborators, Nature 638, 920 (2025) · DOI: 10.1038/s41586-024-08449-y

# Моментот кога кривата конечно се сврте во вистинската насока

Триесет години, квантната корекција на грешки се потпираше на ветување што дотогаш не беше јасно исполнето. Теоријата вели дека ако една единица квантна информација — еден *логички* кјубит — се распореди низ многу шумни физички кјубити, и ако тие физички кјубити се доволно добри, тогаш додавањето *повеќе* од нив треба да го направи логичкиот кјубит *подобар*, при што грешките експоненцијално се намалуваат како што расте кодот. Клучното е тоа „ако“: под критичен праг на шум, повеќе кјубити помагаат; над него, повеќе кјубити само додаваат уште шум. Секој претходен експеримент или бил од погрешната страна на таа граница, или не успеал јасно да го покаже очекуваниот тренд. Зголемувањето на кодот ги влошувало работите наместо да ги подобрува.

Во декември 2024 година, Google Quantum AI објави прва јасна демонстрација на спротивниот режим. На Willow, нивната најнова генерација суперспроводливи процесори, тие изградиле мемории со surface code со кодни растојанија 3, 5 и 7 и забележале дека стапката на логички грешки *паѓа* секогаш кога кодот се зголемува — за фактор  $\Lambda = 2.14 \pm 0.02$  на секое зголемување на растојанието за две. Најголемата, меморија со растојание 7 и 101 кјубит, чувала логички кјубит со грешка од  $0.143\% \pm 0.003\%$  по циклус на корекција и — главниот резултат во рамките на главниот резултат — преживувала *подолго од најдобриот физички кјубит од кој е составена*, за фактор  $2.4 \pm 0.3$ . Тоа се нарекува „beyond breakeven“: првпат целиот дополнителен товар на корекцијата на грешки навистина се исплатил на овој хардвер.

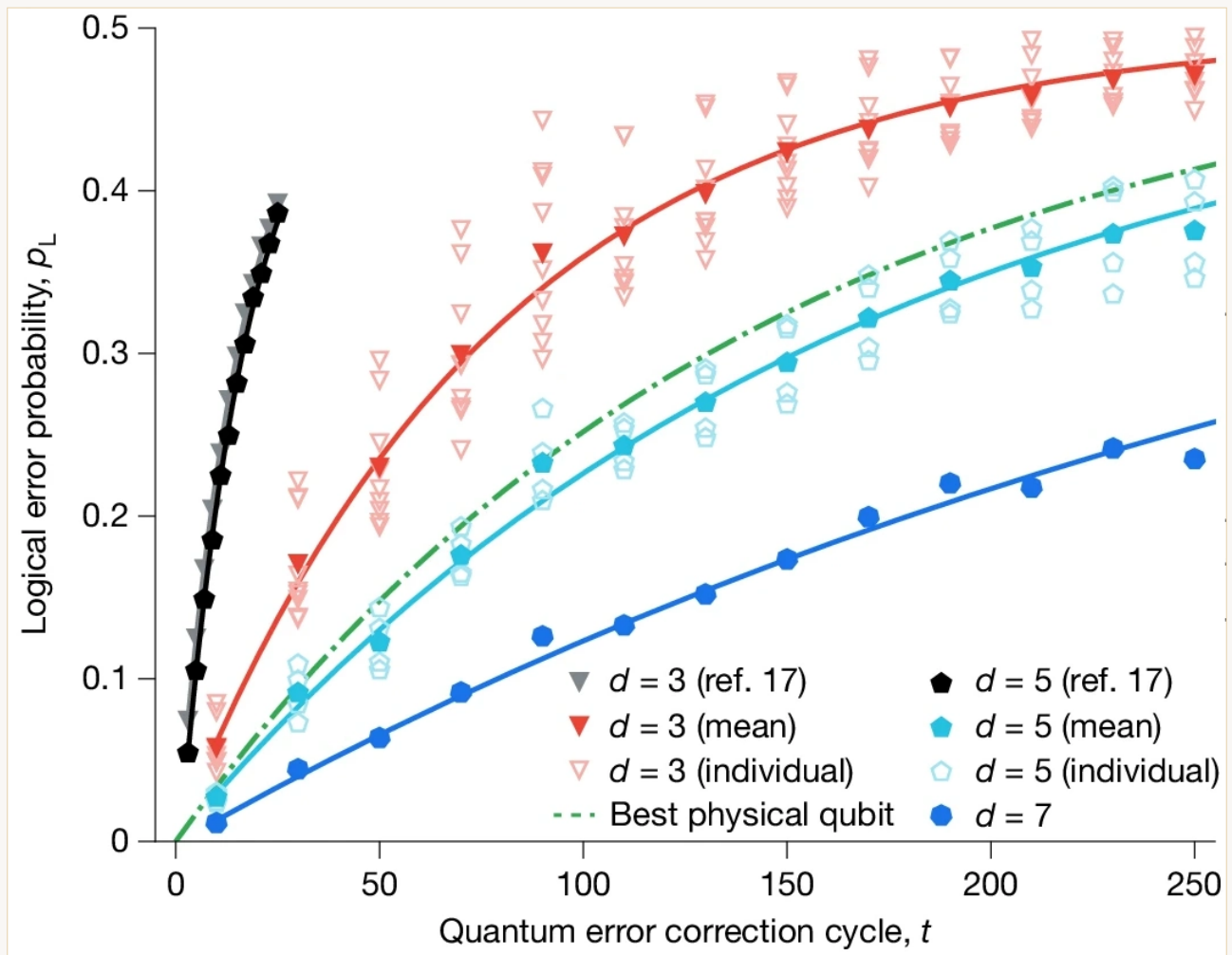
Ова е вистинска пресвртница, но вреди прецизно да се каже каква. Тоа е доказ дека скалирањето конечно оди во вистинската насока. Не е функционален квантен компјутер, а ни трудот не тврди дека е.

## Што значат „surface code“, „растојание“ и „под праг“

**Логички кјубит** е една заштитена единица квантна информација кодирана низ многу физички кјубити. **Surface code** е конкретен начин за такво кодирање на дводимензионална решетка, при што дополнителни „мерни“ кјубити постојано проверуваат за грешки без да ја нарушат складираната информација. **Растојанието**  $d$  на кодот не е физичко растојание: тоа е најмалиот број соодветно поставени грешки што можат да го расипат логичкиот кјубит без кодот да ги забележи. Поголемо  $d$  значи поголема и поотпорна кодна област — троши повеќе физички кјубити (приближно  $2d^2 - 1$ ) и може да поправи повеќе истовремени грешки, до  $(d - 1)/2$ . Така, трите тестирани големини, со растојанија 3, 5 и 7, можат да поправат 1, 2 и 3 истовремени грешки и во идеализирана шема трошат приближно 17, 49 и 97 физички кјубити — меморијата со растојание 7 што ја изгради Google користела 101, малку над овој учебнички минимум.

„**Под праг**“ е клучниот израз. Корекцијата на грешки помага само ако стапката

на физички грешки е под критична вредност; тогаш секое зголемување на растојанието експоненцијално ја потиснува стапката на логички грешки. Факторот на потиснување  $\Lambda$  го мери тоа —  $\Lambda > 1$  значи дека зголемувањето на кодот помага, а колку е повисоко  $\Lambda$ , толку подобро. Google пријавува  $\Lambda \approx 2.14$ , што значи дека секое зголемување на растојанието за две приближно ја преполовува стапката на логички грешки. Фактот дека  $\Lambda$  е убедливо над 1 е суштината на резултатот.



Како се натрупува логичката грешка низ циклусите на корекција кај мемориите со растојание 3, 5 и 7 (од горе надолу). Линијата што треба да се следи е зелената испрекината — најдобриот поединечен физички кјубит на чипот. Меморијата со растојание 7 (сина, најдолу) собира грешка побавно од таа линија, па кодираниот кјубит живее подолго од најдобриот физички кјубит од кој е изграден — „beyond breakeven“, за фактор  $2.4\times$ . Ова е резултатот за животниот век; самото потиснување под праг ( $\Lambda = 2.14$ , кога кодот расте од растојание 3 на 5 и 7) е изразено со бројките во текстот.



## Што направиле авторите

- Изградиле мемории со surface code на два Willow-чипа: процесор со 105 кјубити на кој работеле кодови со растојанија 3, 5 и 7 за тестот на скалирање (најголемиот бил меморија со растојание 7, 101 кјубит и 49 податочни кјубити), и процесор со 72 кјубити на кој работела меморија со растојание 5 со декодер во реално време, како и repetition codes со големо растојание.
- Измериле како се менува логичката грешка *по циклус* кога растојанието се зголемува од 3 на 5 и 7, и од тоа го извлекле факторот на потиснување  $\Lambda$ .
- Го споредиле животниот век на логичкиот кјубит со најдобриот поединечен физички кјубит на истиот чип, за да проверат дали е постигнат „breakeven“.
- Го пуштиле кодот со растојание 5 со **декодер во реално време** — класичен хардвер што ги толкува проверките за грешки додека тие се произведуваат — до еден милион циклуси, за да покажат дека корекцијата на грешки може да ја следи машината.
- Ги прошириле поедноставните **repetition codes** до растојание 29 за да бараат ретки, длабоки извори на грешки што поставуваат долна граница на перформансите.

## Што откриле

- **Кодот е под праг.** Логичката грешка по циклус се намалувала со  $\Lambda = 2.14 \pm 0.02$  за секое зголемување на растојанието за две — чисто експоненцијално потиснување, токму однесувањето што го предвидува теоријата и што претходно ниеден процесор не го покажал недвосмислено.
- **Меморијата со растојание 7 достигнала  $0.143\% \pm 0.003\%$  грешка по циклус** и живеела  $2.4 \pm 0.3$  пати подолго од својот најдобар физички кјубит — над breakeven.
- **Декодирањето во реално време успевало да ја следи работата.** Декодерот имал просечна латентност од 63 микросекунди при растојание 5, наспроти времетраење на циклус од 1.1 микросекунда, и тоа било одржано низ еден милион циклуси — корекцијата работела во живо, а не само во последователна анализа.
- **Останува редок, длабок извор на грешки.** Кај тестовите со repetition code, перформансите на крај биле ограничени од корелирани налети на грешки што се случувале приближно **еднаш на час** (околу еден на секои  $3 \times 10^9$  циклуси), поставувајќи долна граница на грешката близу  $10^{-10}$ , чиј извор, според авторите, **сè уште не е разјаснет**.

## Што ова не докажува

- Ова **не е квантен компјутер што врши пресметки**. Ова е квантна меморија: складира и штити еден логички кјубит. Не изведува логички операции (порти)

меѓу логички кјубити и не извршува алгоритам.

- Не сме **на еден кјубит од корисни машини**. Логички кјубит со растојание 7 користи околу 101 физички кјубит; стапка на грешка од околу 0.1% по циклус е сè уште многу над приближно  $10^{-6}$  до  $10^{-10}$  што им е потребно на реални алгоритми. Затворањето на таа празнина бара многу поголеми растојанија — многу повеќе физички кјубити по логички кјубит — а корисните алгоритми бараат *илјадници* логички кјубити истовремено. Тоа го турка буџетот на физички кјубити кон милиони.
- Изразот **„ако се скалира“ носи голем дел од товарот на заклучокот**. Самиот труд вели дека перформансите на уредот, *ако се скалираат*, би можеле да ги исполнат барањата на големи алгоритми. Да се покаже правилниот тренд на еден логички кјубит не е исто што и да се изгради скалирана машина, а ништо овде не гарантира дека трендот ќе опстои при многу поголеми големини.
- **Необјаснетата долна граница на грешките е реален проблем**. Корелираните налети што ги ограничуваат repetition codes се, според авторите, за редови на големина поголеми од очекуваното и би спречиле поголеми fault-tolerant апликации додека не се разберат — отворен дефект, јасно наведен, а не решен детал.
- Резултатот не кажува ништо за **кршење енкрипција или „квантна супремација“ во корисни задачи**. За тоа е потребна целосната fault-tolerant машина за која ова е еден темелен камен, а не демонстрација на таква машина.

## Колку се убедливи доказите?

- **Главното тврдење е цврсто и важно**. Работа под праг со чисто експоненцијално потиснување низ три кодни растојанија, плус животен век над breakeven и функционален декодер во реално време, е токму комбинацијата што полето се обидувахе да ја постигне. Таа е покажана директно, а не само изведена од модел. Ова не е производ на возбуда; тоа е реален инженерски резултат од водечка група.
- **Авторите внимателно го ограничуваат опсегот**. Го претставуваат како *меморија* под праг, самите ја нагласуваат необјаснетата долна граница од корелирани грешки и иднината ја условуваат со јасното „ако се скалира“. Претерувањето, каде што се појавува, е главно во медиумското заокружување на „мемориски кјубит со корекција на грешки се подобрува кога расте“ во „квантното компјутерство пристигна“.
- **Искрениот статус е темелен чекор, направен убедливо**. Еден логички кјубит, заштитен доволно добро така што додавањето редундантност конечно помага — но со долг, тежок и сè уште негарантиран пат на скалирање, логички порти и необјаснети грешки пред него.

## Зошто е важно

Fault-tolerant квантното компјутерство отсекогаш имало проблем налик на „кокошка и јајце“: машините што би биле корисни бараат стапки на грешки што ниеден физички кјубит не може да ги постигне, а решението — корекција на грешки — работи само ако хардверот веќе е доволно добар за да биде под праг. Преминувањето на таа граница, макар еднаш и на само еден логички кјубит, го менува прашањето од „дали ова воопшто е можно?“ во „до каде може да се скалира и колку брзо?“. Тоа е вистинска и значајна промена, и затоа резултатот заслужува внимание.

Но истата внимателност што го прави резултатот веродостоен треба да ја смири приказната околу него. Ова е првата добро поставена тула во темелот. Не е зградата, а луѓето што ја поставиле први го кажуваат тоа. Вистинскиот начин да се следи квантното компјутерство во следните неколку години е токму оваа неатрактивна крива: дали факторот на потиснување ќе се одржи додека кодовите растат, дали логичките порти ќе можат да работат толку чисто колку логичката меморија и дали мистериозната грешка што се појавува еднаш на час ќе биде објаснета.

## Кратко резиме

Google Quantum AI за првпат јасно покажа дека квантна меморија со surface code може да работи *под праг*: кога кодот го зголемиле од растојание 3 на 5 и 7, стапката на логички грешки експоненцијално се намалувала (за околу  $2.14\times$  на секои две единици растојание), а најголемата, меморија со растојание 7 и 101 кјубит, живеела подолго од својот најдобар физички кјубит — *beyond breakeven* — додека корекцијата на грешки работела во реално време. Тоа е вистинска, долго очекувана пресвртница во инженерството на квантни компјутери. Но истовремено е еден логички кјубит што служи како меморија, со стапка на грешка сè уште далеку од онаа што ја бараат реални алгоритми, без изведени логички операции, со необјаснета долна граница на грешките што самите автори ја нагласуваат и со повеќе редови на големина скалирање пред себе. Прагот навистина е преминат — квантен компјутер не е испорачан.

---

## Извори

Google Quantum AI and Collaborators, Quantum error correction below the surface code threshold, Nature 638, 920 (2025)

<https://doi.org/10.1038/s41586-024-08449-y>

Author Correction: Quantum error correction below the surface code threshold, Nature 653, E5 (2026) — corrects a Fig. 3a axis label and legend keys; does not affect the below-threshold (Fig. 1) results

<https://www.nature.com/articles/s41586-026-10559-8>