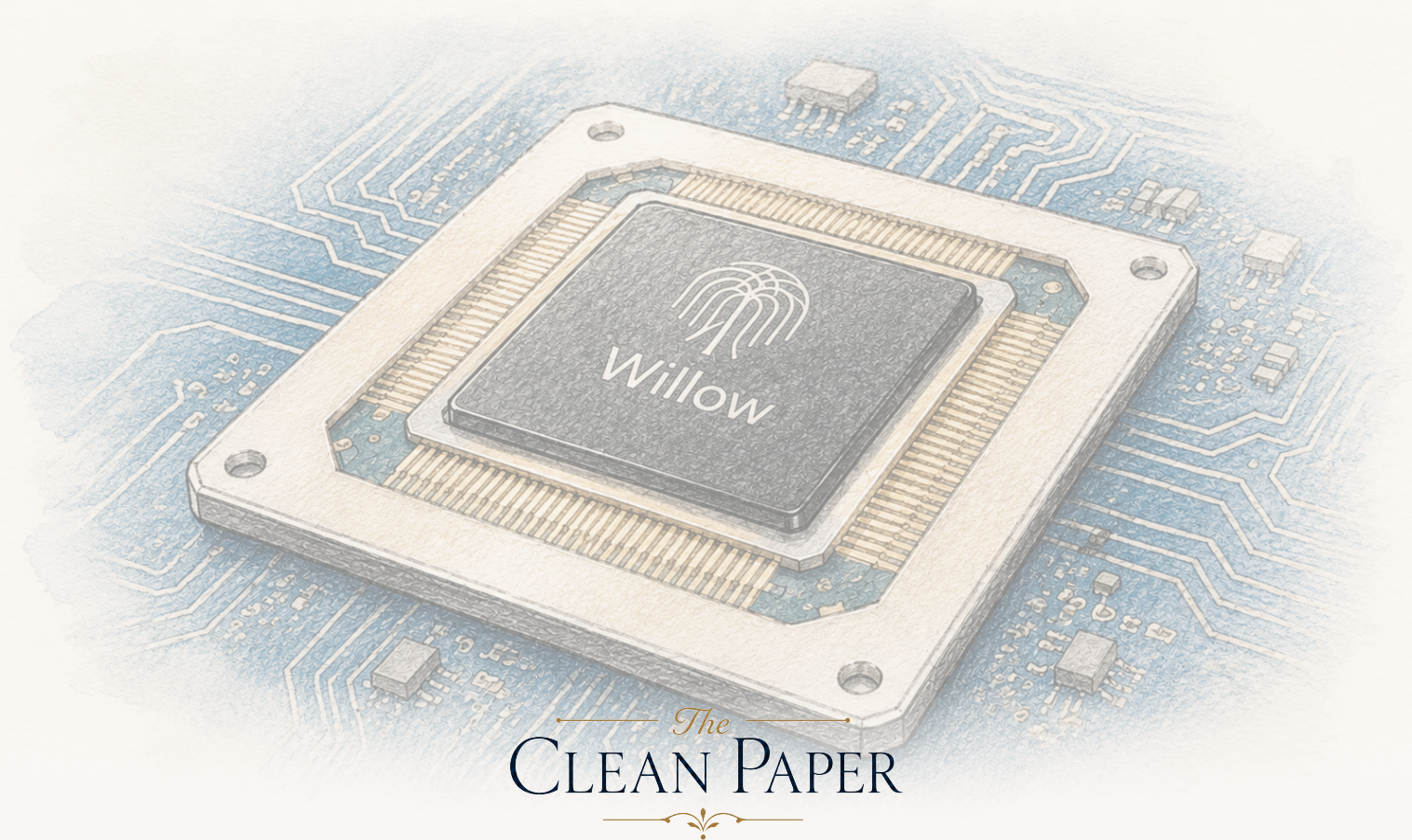




WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Квантова памет най-сетне стана по-добра, когато нарасна — истинският пробив и машината, която това не е

Google Quantum AI показва за пръв път ясно, че квантова памет със *surface code* може да работи под прага: когато кодът нараства от *distance 3* на *5* и *7*, логическата честота на грешки спада експоненциално (около $2,14$ пъти на всеки две стъпки), а най-голямата памет — *distance 7* със 101 кубита — надживява най-добрия си физически кубит, т.е. минава *beyond breakeven*, при корекция на грешки в реално време. Това е дългоочакван инженерен етап. Но това е и един логически кубит, използван като памет, с грешка все още далеч от нужната за реални алгоритми, без изпълнени логически операции, с необяснен праг на остатъчни грешки, който самите автори отбелязват, и с още много порядъци мащабиране напред. Преминат праг, не доставен компютър.

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

Моментът, в който кривата най-сетне се огъна в правилната посока

В продължение на трийсет години квантовата корекция на грешки се опира на обещание, което никога не е било ясно изпълнено. Теорията казва, че ако разпределите една единица квантова информация — един *логически* кубит — върху много шумни физически кубити и ако тези физически кубити са достатъчно добри, тогава добавянето на *още* от тях трябва да направи логическия кубит *по-добър*, като грешките спадат експоненциално с нарастването на кода. Уловката е в „ако“: под критичен праг на шума повече кубити помагат; над него повече кубити просто добавят повече шум. Всеки предишен експеримент е бил от грешната страна на тази линия или не е успявал ясно да покаже тенденцията. Увеличаването на кода е правело нещата по-лоши, не по-добри.

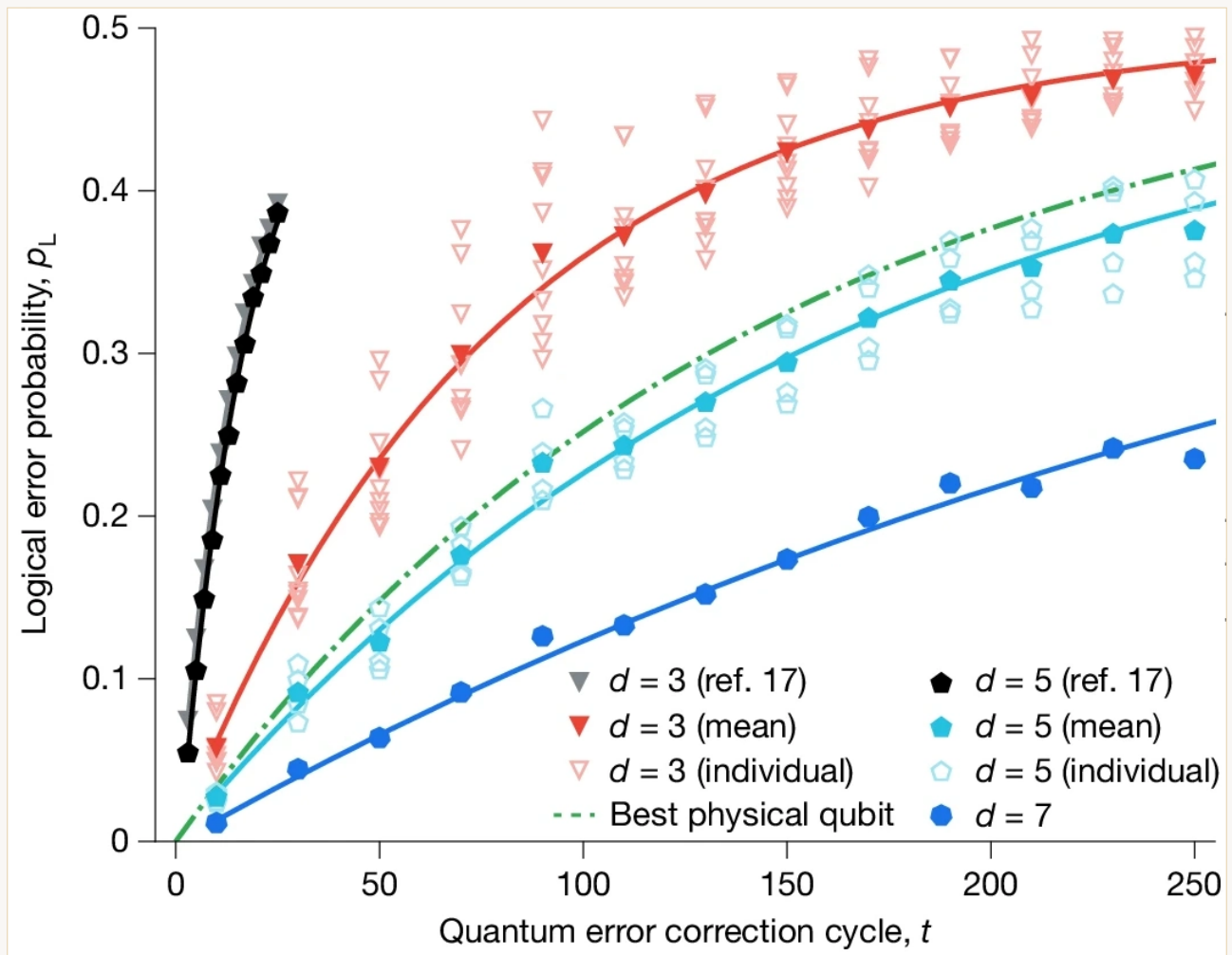
През декември 2024 г. Google Quantum AI съобщи първата ясна демонстрация на другия режим. На Willow, най-новото им поколение свръхпроводящи процесори, те изграждат памети с повърхностен код с разстояние на кода 3, 5 и 7 и наблюдават как логическата честота на грешки *спад*а всеки път, когато кодът стане по-голям — с фактор $\Lambda = 2.14 \pm 0.02$ за всяко увеличение на разстоянието на кода с две стъпки. Най-голямата им памет, разстояние 7 със 101 кубита, задържа логически кубит с грешка $0.143\% \pm 0.003\%$ на цикъл на корекция и — заглавието вътре в заглавието — оцелява *по-дълго от собствения си най-добър физически кубит*, с фактор 2.4 ± 0.3 . Това се нарича „отвъд точката на изравняване“ и е първият път, когато целият апарат на корекцията на грешки се изплаща на този хардуер.

Това е истински етап и си струва да сме точни какъв. Това е доказателство, че мащабирането вече върви в правилната посока. Не е работещ квантов компютър и статията не твърди, че е.

Какво означават „повърхностен код“, „разстояние на кода“ и „под прага“

Логически кубит е една защитена единица квантова информация, кодирана върху много физически кубити. **Повърхностният код** е конкретен начин това кодиране да се направи върху 2D решетка, където допълнителни „измерващи“ кубити постоянно проверяват за грешки, без да нарушават съхраняваната информация. **Разстоянието на кода** d не е физическо разстояние: това е най-малкият брой правилно разположени грешки, които могат да повредят логическия кубит, без кодът да ги забележи. По-голямо d означава по-голям и по-устойчив участък — използва повече физически кубити (приблизително $2d^2 - 1$) и коригира повече едновременни грешки, до $(d - 1)/2$. Така трите размера, проверени тук — разстояние 3, 5 и 7 — коригират 1, 2 и 3 едновременни грешки и използват приблизително 17, 49 и 97 физически кубита; паметта разстояние 7, която Google изгражда, използва 101 — малко над този учебникарски минимум. „**Под прага**“ е ключовият израз. Корекцията на грешки помага само ако

физическата честота на грешки е под критична стойност; там всяко увеличение на разстоянието на кода потиска логическата честота на грешки експоненциално. Факторът на потискане Λ измерва това — $\Lambda > 1$ означава, че увеличаването на кода помага, и колкото по-високо е Λ , толкова по-добре. Google съобщава $\Lambda \approx 2.14$, което означава, че всяко увеличение на разстоянието на кода с две стъпки намалява логическата честота на грешки приблизително наполовина. Фактът, че Λ е уверено над 1, е целият резултат.



Как логическата грешка се натрупва през циклите на корекция за паметите разстояние 3, 5 и 7 (отгоре надолу). Линията, която трябва да се следи, е зелената прекъсната — *най-добрият единичен физически кубит на чипа*. Паметта разстояние 7 (синя, най-ниска) натрупва грешка по-бавно от тази линия, така че кодираният кубит надживява най-добрия физически кубит, от който е изграден — „отвъд точката на изравняване“ с фактор 2.4×. Това е резултатът за живота на паметта; самото потискане под прага ($\Lambda = 2.14$ при нарастване на кода от разстояние 3 към 5 и 7) е в числата в текста.

Какво са направили авторите

- Изграждат памети с повърхностен код върху два чипа Willow: процесор със 105 кубита, който изпълнява кодовете разстояние 3, 5 и 7 за теста на мащабирането (най-големият е памет разстояние 7 със 101 кубита, от които 49 кубита за данни), и процесор със 72 кубита, който изпълнява памет разстояние 5 с декодер в реално време плюс кодове с повторение с голямо разстояние на кода.
- Измерват как логическата грешка *на цикъл* се променя, когато увеличават разстоянието на кода от 3 на 5 и 7, и извличат фактора на потискане Λ .
- Сравняват живота на логическия кубит с най-добрия отделен физически кубит на същия чип, за да проверят за „точката на изравняване“.
- Изпълняват кода разстояние 5 с **декодер в реално време** — класически хардуер, който интерпретира проверките за грешки толкова бързо, колкото се генерират — за до един милион цикъла, за да покажат, че корекцията може да смогва на машината.
- Разширяват по-прости **кодове с повторение** до разстояние 29, за да търсят редки, дълбоки източници на грешки, които поставят долна граница на представянето.

Какво са установили

- **Кодът е под прага.** Логическата грешка на цикъл намалява с $\Lambda = 2.14 \pm 0.02$ за всяко увеличение на разстоянието на кода с две — чисто експоненциално потискане, поведението, което теорията обещава и което никой процесор не е показвал категорично досега.
- **Паметта разстояние 7 достига $0.143\% \pm 0.003\%$ грешка на цикъл** и живее 2.4 ± 0.3 пъти по-дълго от най-добрия си физически кубит — отвъд точката на изравняване.
- **Декодирането в реално време смогва.** Декодерът има средна латентност 63 микросекунди при разстояние 5 спрямо време на цикъла 1,1 микросекунди, устойчиво през един милион цикъла — корекцията на грешки работи на живо, не само в анализ след експеримента.
- **Остава рядък, дълбок източник на грешки.** В тестовите с код с повторение представянето накрая е ограничено от корелирани изблици от грешки приблизително **веднъж на час** (около един на всеки 3×10^9 цикъла), които задават дъно на грешката близо до 10^{-10} и чийто произход според авторите **все още не е разбран**.

Какво не доказва това

- Това **не е** квантов компютър, който извършва изчисления. Това е квантова *памет*: тя съхранява и защитава един логически кубит. Не изпълнява логически операции (гейтове) между логически кубити и не изпълнява алгоритъм.

- Това **не** означава, че сме на един кубит от полезни машини. Логически кубит разстояние 7 използва около 101 физически кубита; честота на грешка 0,1% на цикъл все още е далеч над приблизително 10^{-6} до 10^{-10} , необходими на реални алгоритми. Затварянето на тази разлика означава преминаване към много по-големи разстояния на кода — много повече физически кубити на логически кубит — а полезните алгоритми се нуждаят от *хиляди* логически кубити едновременно. Бюджетът във физически кубити стига до милиони.
- **„Ако се мащабира“ е съществено условие.** Собственото заключение на статията е, че представянето на устройството, *ако бъде мащабирано*, може да изпълни изискванията на големи алгоритми. Да покажете правилната тенденция на един логически кубит не е същото като да сте построили мащабираната машина и нищо тук не гарантира, че тенденцията ще се запази при много по-големи размери.
- **Необясненото дъно на грешките е активен проблем.** Корелираните изблици, които ограничават тестовите с кодове с повторение, са, по думите на авторите, с порядъци по-големи от очакваното и биха попречили на по-големи приложения, устойчиви на грешки, докато не бъдат разбрани — открит недостатък, заявен ясно, не решен детайл.
- Това не казва нищо за **разбиване на криптиране или „квантово превъзходство“ при полезни задачи.** Те изискват цялата машина, устойчива на грешки, за която това е основополагащ камък, а не демонстрация на такава машина.

Колко силни са доказателствата

- **Основното твърдение е солидно и важно.** Работа под прага с чисто експоненциално потискане през три разстояние на кода, плюс живот отвъд точката на изравняване и работещ декодер в реално време — това е точно комбинацията, която областта се опитва да достигне, и е демонстрирана пряко, а не изведена косвено. Това не е артефакт на медийния шум; това е реален инженерен резултат от водещ екип.
- **Авторите са внимателни с обхвата.** Те го представят като *памет* под прага, сами отбелязват необясненото дъно от корелирани грешки и поставят бъдещето зад ясно „ако се мащабира“. Преувеличението, където го има, е в околното отразяване, което закръгля „квантов кубит-памет с корекция на грешки се подобрява при увеличаване“ до „квантовите изчисления вече са тук“.
- **Честният статус е основополагаща стъпка, направена чисто.** Един логически кубит, защитен достатъчно добре, че добавянето на излишък най-сетне да помага — с дълъг, труден и все още негарантиран път на мащабиране, логически гейтове и необяснени грешки отпред.

Защо е важно

Квантовите изчисления, устойчиви на грешки винаги са имали усещане за проблема с кокошката и яйцето: полезните машини се нуждаят от честоти на грешки, които никога физически кубит не може да достигне, а решението — корекция на грешки — работи само ако хардуерът вече е достатъчно добър, за да бъде под прага. Преминването на тази линия, дори веднъж и дори на един логически кубит, променя въпроса от „възможно ли е това изобщо?“ на „докъде може да се мащабира и колко бързо?“. Това е реална и смислена промяна и затова резултатът заслужава внимание.

Но същата прецизност, която прави резултата надежден, трябва да укроти разказа около него. Това е първата тухла от фундамент, поставена добре. Не е сградата, и хората, които са я положили, са първите, които го казват. Правилният начин да се следят квантовите изчисления през следващите години е точно тази негламурна крива: дали факторът на потискане се запазва при растежа на кодовете, дали логическите гейтове могат да се изпълняват толкова чисто, колкото логическата памет, и дали мистериозната грешка веднъж на час някога ще бъде обяснена.

Кратко обобщение

Google Quantum AI показва за пръв път ясно, че квантова памет с повърхностен код може да работи *под прага*: когато кодът нараства от разстояние 3 към 5 и 7, логическата честота на грешки спада експоненциално (с около $2,14\times$ на всеки две стъпки), а най-голямата памет, разстояние 7 със 101 кубита, надживява най-добрия си физически кубит — отвъд точката на изравняване — докато корекцията на грешки работи в реално време. Това е истински, дългоочакван етап в инженерството на квантови компютри. Но това е и един логически кубит, действащ като памет, с честота на грешки все още далеч от изискваната от реални алгоритми, без изпълнени логически операции, с необяснено дъно на грешките, което самите автори отбелязват, и с път на мащабиране от много порядъци напред. Реален праг е преминат — квантов компютър не е доставен.

Източници

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>