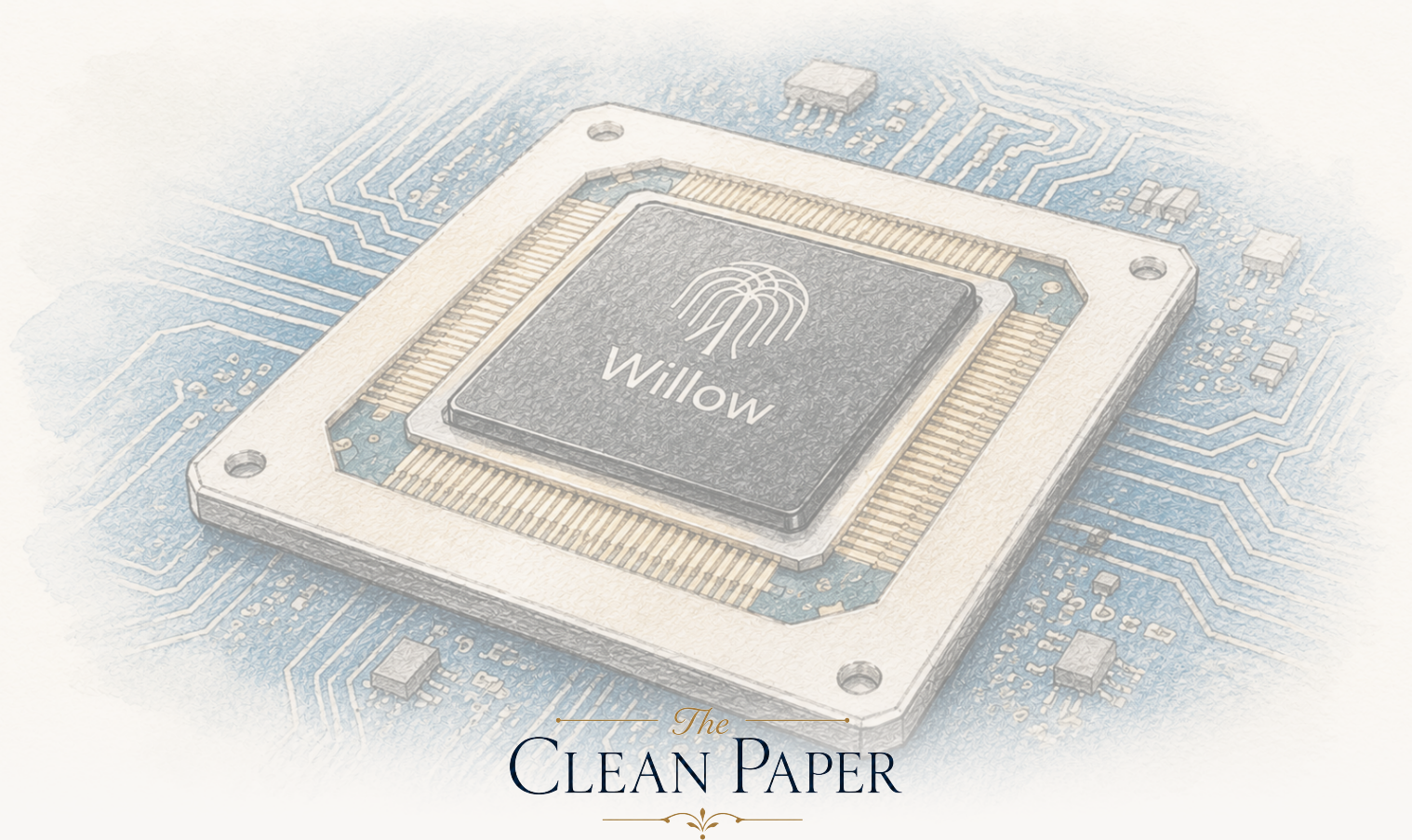




WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Квантавая памяць нарэшце стала лепшай пры павелічэнні — сапраўдная вяха, але яшчэ не гатовая машына

Google Quantum AI упершыню ясна паказала, што квантавая памяць на surface code можа працаваць ніжэй за парог: пры павелічэнні кода ад distance 3 да 5 і 7 лагічная частата памылак экспаненцыяльна падала — прыкладна ў 2,14 раза на кожныя два крокі, — а найбуйнейшая distance-7 памяць на 101 кубіце перажывала свой найлепшы фізічны кубіт, гэта значыць выйшла beyond breakeven, пры карэкцыі памылак у рэальным часе. Гэта даўно чаканая інжынерная вяха. Але гэта таксама толькі адзін лагічны кубіт у ролі памяці, з частатой памылак усё яшчэ далёкай ад патрэб рэальных алгарытмаў, без лагічных аперацый, з незразумелай падлогай памылак, якую падкрэсліваюць аўтары, і з многімі парадкамі маштабавання наперадзе. Парог перасечаны, а гатовы камп'ютар не дастаўлены.

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

Момант, калі крывая нарэшце выгнулася ў правільны бок

Трыццаць гадоў квантавая карэкцыя памылак трымалася на абяцанні, якое яшчэ ніколі не выконвалася настолькі чыста. Тэорыя кажа: калі размеркаваць адну адзінку квантавай інфармацыі — адзін *лагічны* кубіт — паміж мноствам шумных фізічных кубітаў і калі гэтыя фізічныя кубіты дастаткова якасныя, то даданне *яшчэ большай* іх колькасці павінна рабіць лагічны кубіт *лепшым*, а памылкі павінны экспаненцыяльна зніжацца па меры павелічэння кода. Умова схаваная ў «калі»: ніжэй за крытычны парог шуму большая колькасць кубітаў дапамагае; вышэй за яго яна толькі дадае шум. Усе папярэднія эксперыменты знаходзіліся не па той бок гэтай мяжы або не паказвалі трэнд дастаткова чыста. Павелічэнне кода рабіла сітуацыю горшай, а не лепшай.

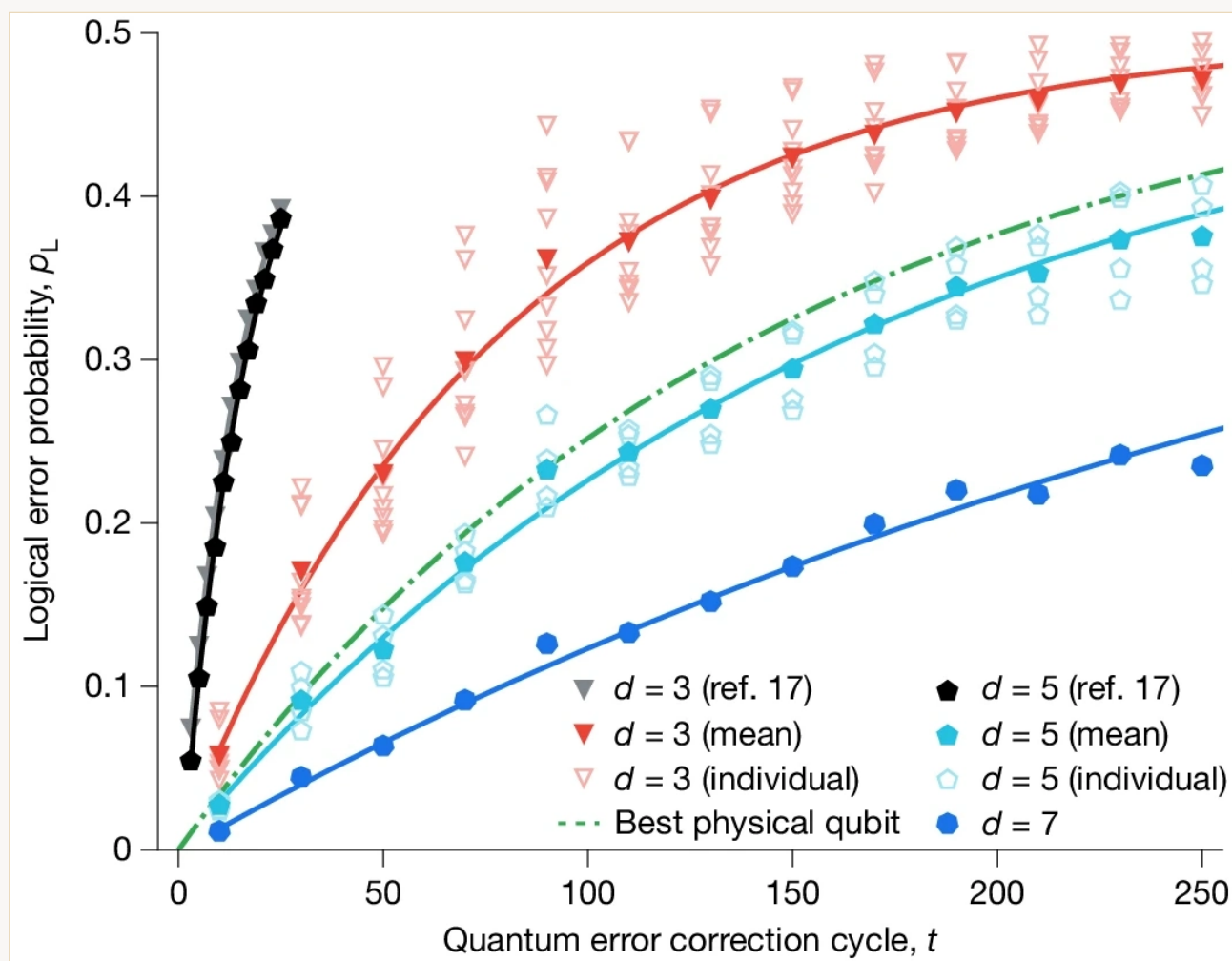
У снежні 2024 года Google Quantum AI паведаміла пра першую ясную дэманстрацыю іншага рэжыму. На Willow, новым пакаленні звышправодных працэсараў кампаніі, даследчыкі пабудавалі памяць з surface code на адлегласцях кода 3, 5 і 7 і ўбачылі, як лагічная частата памылак *падае* кожны раз, калі код становіцца большым — у $\Lambda = 2,14 \pm 0,02$ раза на кожнае павелічэнне адлегласці на два. Найбуйнейшая сістэма — distance-7 памяць на 101 кубіце — утрымлівала лагічны кубіт з памылкай $0,143\% \pm 0,003\%$ за цыкл карэкцыі і, што стала асобным загаловам унутры загалова, жыла *даўжэй за найлепшы фізічны кубіт*, з якога была пабудавана, у $2,4 \pm 0,3$ раза. Гэта называюць «beyond breakeven» — выхадам за кропку бясстратнасці, — і гэта першы выпадак, калі ўвесь апарат карэкцыі памылак на гэтым абсталяванні сапраўды акупіў уласныя накладныя выдаткі.

Гэта сапраўдная вяха, і важна дакладна сказаць, якога тыпу. Гэта доказ таго, што маштабаванне цяпер ідзе ў правільны бок. Гэта не гатовы квантавы камп'ютар, і артыкул гэтага не сцвярджае.

Што азначаюць «surface code», «distance» і «below threshold»

Лагічны кубіт — адна абароненая адзінка квантавай інфармацыі, закадаваная ў многіх фізічных кубітах. **Surface code** — канкрэтны спосаб такога кадавання на двухмернай рашотцы, дзе дадатковыя «вымяральныя» кубіты пастаянна правяраюць наяўнасць памылак, не разбураючы захаваную інфармацыю. **Адлегласць кода** d — не фізічная адлегласць: гэта мінімальная колькасць правільна размешчаных памылак, здольных сапсаваць лагічны кубіт так, каб код гэтага не заўважыў. Большае d азначае большы і ўстойлівейшы ўчастак кода: ён расходзе больш фізічных кубітаў — прыкладна $2d^2 - 1$ — і выпраўляе больш адначасовых памылак, да $(d - 1)/2$. Таму тры правярэння тут памеры — адлегласці 3, 5 і 7 — выпраўляюць адпаведна 1, 2 і 3 адначасовыя памылкі і тэарэтычна патрабуюць прыкладна 17, 49 і 97 фізічных кубітаў; пабудаваная Google distance-7 памяць выкарыстоўвала 101 — крыху больш за гэты падручнікавы мінімум. **«Below threshold»** — ключавая фраза. Карэкцыя памылак дапамагае толькі тады, калі фізічная частата памылак ляжыць ніжэй за крытычнае значэнне; у гэтым

рэжыме кожнае павелічэнне адлегласці экспаненцыяльна душыць лагічную частату памылак. Каэфіцыент падаўлення Λ вымярае менавіта гэта: $\Lambda > 1$ азначае, што павялічваць код карысна, і чым большае Λ , тым лепш. Google паведамляе $\Lambda \approx 2,14$, гэта значыць кожнае павелічэнне адлегласці на два прыкладна ўдвая памяншала лагічную частату памылак. Сам факт, што Λ упэўнена вышэй за 1, і ёсць цэнтральны вынік.



Як лагічная памылка назапашваецца на працягу цыклаў карэкцыі для памяці з адлегласцямі 3, 5 і 7 — зверху ўніз. Лінія, на якую варта глядзець, — зялёная пункцірная: *найлепшы адзінкавы фізічны кубіт* на чыпе. Distance-7 памяць — сіняя, самая ніжняя — назапашвае памылку павольней за гэтую лінію, таму закадаваны кубіт жыве даўжэй, чым найлепшы фізічны кубіт, з якога ён пабудаваны: «beyond breakeven», у $2,4\times$ раза. Гэта вынік пра час жыцця; само падаўленне ніжэй за парог ($\Lambda = 2,14$ пры росце кода ад distance 3 да 5 і 7) утрымліваецца ў лічбах у тэксце.

Што зрабілі аўтары

- Пабудавалі surface-code памяць на двух чыпах Willow: 105-кубітным працэсары, які запускаў коды distance-3, -5 і -7 для тэсту маштабавання — найбуйнейшая канфігурацыя была 101-кубітнай distance-7 памяццю з 49 кубітамі даных, — і 72-кубітным працэсары, які запускаў distance-5 памяць з дэкодарам рэальнага часу і repetition codes вялікай адлегласці.
- Вымералі, як лагічная памылка за цыкл змяняецца пры павелічэнні адлегласці кода ад 3 да 5 і 7, і атрымалі каэфіцыент падаўлення Λ .
- Параўналі час жыцця лагічнага кубіта з найлепшым асобным фізічным кубітам на тым жа чыпе, каб праверыць «breakeven».
- Запусцілі distance-5 код з **дэкодарам рэальнага часу** — класічным абсталяваннем, якое інтэрпрэтуе праверкі памылак па меры іх з’яўлення, — да мільёна цыклаў, каб паказаць, што карэкцыя здольная працаваць у тэмпе самой машыны.
- Пашырылі больш простыя **repetition codes** да distance 29, каб шукаць рэдкія глыбінныя крыніцы памылак, якія задаюць ніжнюю мяжу прадукцыйнасці.

Што яны выявілі

- **Код працуе ніжэй за парог.** Лагічная памылка за цыкл памяншалася ў $\Lambda = 2,14 \pm 0,02$ раза на кожнае павелічэнне адлегласці на два — чыстае экспаненцыяльнае падаўленне, менавіта тыя паводзіны, якія абяцала тэорыя і якія ніводзін працэсар раней не дэманстраваў настолькі адназначна.
- **Distance-7 памяць дасягнула $0,143\% \pm 0,003\%$ памылкі за цыкл і жыла ў $2,4 \pm 0,3$ раза даўжэй,** чым яе найлепшы фізічны кубіт — beyond breakeven.
- **Дэкадаванне ў рэальным часе паспявала.** Для distance 5 сярэдняя затрымка дэкодара складала 63 мікрасекунды пры працягласці цыклу 1,1 мікрасекунды і стабільна працавала на працягу мільёна цыклаў — карэкцыя выконвалася ўжывую, а не толькі ў аналізе пасля эксперыменту.
- **Застаецца рэдкая глыбінная крыніца памылак.** У тэстах repetition code прадукцыйнасць урэшце абмяжоўвалі карэляваныя ўсплёскі памылак прыкладна **раз на гадзіну** — каля аднаго выпадку на 3×10^9 цыклаў, — што стварае падлогу памылак каля 10^{-10} ; паходжанне гэтых усплёскаў, паводле аўтараў, **пакуль незразумелае**.

Чаго гэта не даказвае

- Гэта **не** квантавы камп’ютар, які выконвае вылічэнні. Гэта квантавая *памяць*: яна захоўвае і абараняе адзін лагічны кубіт. Яна не выконвае лагічныя аперацыі — gates — паміж лагічнымі кубітамі і не запускае алгарытмы.
- Гэта **не** азначае, што да карысных машын застаўся адзін кубіт. Лагічны кубіт distance-

7 расходуе каля 101 фізічнага кубіта; частата памылак каля 0,1% за цыкл усё яшчэ нашмат вышэйшая за прыкладна 10^{-6} – 10^{-10} , неабходныя рэальным алгарытмам. Закрываючы гэты разрыву патрабуе значна большых адлегласцяў — істотна больш фізічных кубітаў на адзін лагічны, — а карысным алгарытмам патрэбныя *тысячы* лагічных кубітаў адначасова. Фізічны бюджэт для такога маштабу сыходзіць у мільёны.

- Фраза «**калі маштабаваць**» сапраўды нясе вагу. Уласная выснова артыкула: прадукцыйнасць прылады, *калі яе маштабаваць*, можа адпавядаць патрабаванням вялікіх алгарытмаў. Паказаць правільны трэнд на адным лагічным кубіце — не тое самае, што пабудаваць маштабаваную машыну, і нішто тут не гарантуе, што трэнд захавецца на значна большых памерах.
- **Незразумелая падлога памылак — жывая праблема.** Карэляваныя ўсплёскі, якія абмяжоўваюць repetition code, паводле аўтараў, на парадкі большыя за чаканыя і пераацэньваюць б больш буйным fault-tolerant прыкладанням, пакуль іх не зразумеюць. Гэта адкрыты недахоп, наўпрост прызнаны ў працы, а не дробязь, якую ўжо вырашылі.
- Артыкул нічога не кажа пра **ўзлом шыфравання або «квантавую перавагу» ў карысных задачах.** Для гэтага патрэбная паўнаўстойлівая адмоўная машына, фундаментам для якой з’яўляецца гэты вынік, а не сам гэты дэманстратар.

Наколькі моцныя доказы

- **Цэнтральнае сцвярджэнне моцнае і важнае.** Праца ніжэй за парог з чыстым экспаненцыяльным падаўленнем на трох адлегласцях кода, плюс beyond-breakeven час жыцця і працоўны дэкодар рэальнага часу — менавіта тая камбінацыя, якой галіна спрабавала дасягнуць, і яна паказана непасрэдна, а не выведзена ўскосна. Гэта не артэфакт хайпу, а рэальны інжынерны вынік вядучай групы.
- **Аўтары асцярожныя з маштабам высновы.** Яны называюць вынік below-threshold memory, самі падкрэсліваюць незразумелую падлогу карэляваных памылак і прывязваюць будучыню да прыкметнага «if scaled». Перабольшанне з’яўляецца галоўным чынам у навакольным медыйным пераказе, які акругляе «абаронены ад памылак кубіт памяці палепшыўся пры павелічэнні кода» да «квантавыя камп’ютары ўжо тут».
- **Сумленны статус — фундаментальны крок, зроблены чыста.** Адзін лагічны кубіт, абаронены настолькі добра, што дадатковая залішнясць нарэшце дапамагае, — але наперадзе доўгі, цяжкі і не гарантваны шлях маштабавання, лагічных аперацый і тлумачэння неведомых памылак.

Чаму гэта важна

Адмоваўстойлівыя квантавыя вылічэнні заўсёды мелі прысмак праблемы курыцы і яйка: карысным машынам патрэбны частоты памылак, якіх не здольны дасягнуць ніводзін фізічны кубіт, а выпраўленне — карэкцыя памылак — працуе толькі тады, калі само фізічнае абсталяванне ўжо дастаткова добрае, каб апынуцца ніжэй за парог. Перасекчы гэтую мяжу хаця б адзін раз, хаця б для аднаго лагічнага кубіта, змяняе пытанне з «ці магчыма гэта ўвогуле?» на «наколькі далёка гэта можна маштабаваць і як хутка?». Гэта рэальнае і змястоўнае змяненне, таму вынік заслугоўвае ўвагі.

Але тая ж асцярожнасць, якая робіць вынік надзейным, павінна стрымліваць гісторыю вакол яго. Гэта першая добра пакладзеная цагліна падмурка. Не будынак — і людзі, якія яе паклалі, першымі гэта прызнаюць. Правільны спосаб сачыць за квантавымі вылічэннямі ў найбліжэйшыя гады — глядзець менавіта на гэтую неэфектную крывую: ці захоўваецца каэфіцыент падаўлення пры росце кодаў, ці ўдаецца выконваць лагічныя gates гэтак жа чыста, як лагічную памяць, і ці будзе калі-небудзь растлумачана тая загадкавая памылка раз на гадзінку.

Коратка

Google Quantum AI упершыню ясна паказала, што квантавая памяць на surface code можа працаваць *ніжэй за парог*: калі код павялічвалі ад distance 3 да 5 і 7, лагічная частата памылак экспаненцыяльна падала — прыкладна ў 2,14 раза на кожныя два крокі, — а найбуйнейшая distance-7 памяць на 101 кубіце перажывала свой найлепшы фізічны кубіт, гэта значыць выйшла beyond breakeven, пры гэтым карэкцыя памылак працавала ў рэальным часе. Гэта сапраўдная, даўно чаканая інжынерная вяха для квантавых камп'ютараў. Але гэта таксама адзін лагічны кубіт у ролі памяці, з частатой памылак усё яшчэ далёкай ад патрабаванняў рэальных алгарытмаў, без выкананых лагічных аперацый, з незразумелай падлогай памылак, якую падкрэсліваюць самі аўтары, і з многімі парадкамі маштабавання наперадзе. Сапраўдны парог перасечаны — квантавы камп'ютар яшчэ не дастаўлены.

Крыніцы

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>