



The CLEAN PAPER

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

Një memorie kuantike më në fund u përmirësua ndërsa u rrit — piketa reale dhe makina që nuk është

Google Quantum AI tregoi për herë të parë qartë se një memorie kuantike me kod sipërfaqësor mund të punojë nën prag: ndërsa kodi u rrit nga distanca 3 në 5 në 7, norma e gabimit logjik ra eksponencialisht (rreth $2.14\times$ për çdo dy hapa), dhe memoria më e madhe, me distancë 7 dhe 101 kubitë, jetoi më gjatë se kubiti i saj fizik më i mirë — përtej breakeven-it — me korrigjimin e gabimeve në kohë reale. Është një piketë inxhinierike e kërkuar prej kohësh. Por është edhe një kubit i vetëm logjik që vepron si memorie, me normë gabimi ende larg asaj që kërkojnë algoritmet reale, pa operacione logjike, me një dysheme gabimi të pashpjeguar që autorët e theksojnë dhe me disa rende madhësie shkallëzim përpara. Një prag i kaluar, jo një kompjuter i dorëzuar.

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

Momenti kur kurba më në fund u përkul në drejtimin e duhur

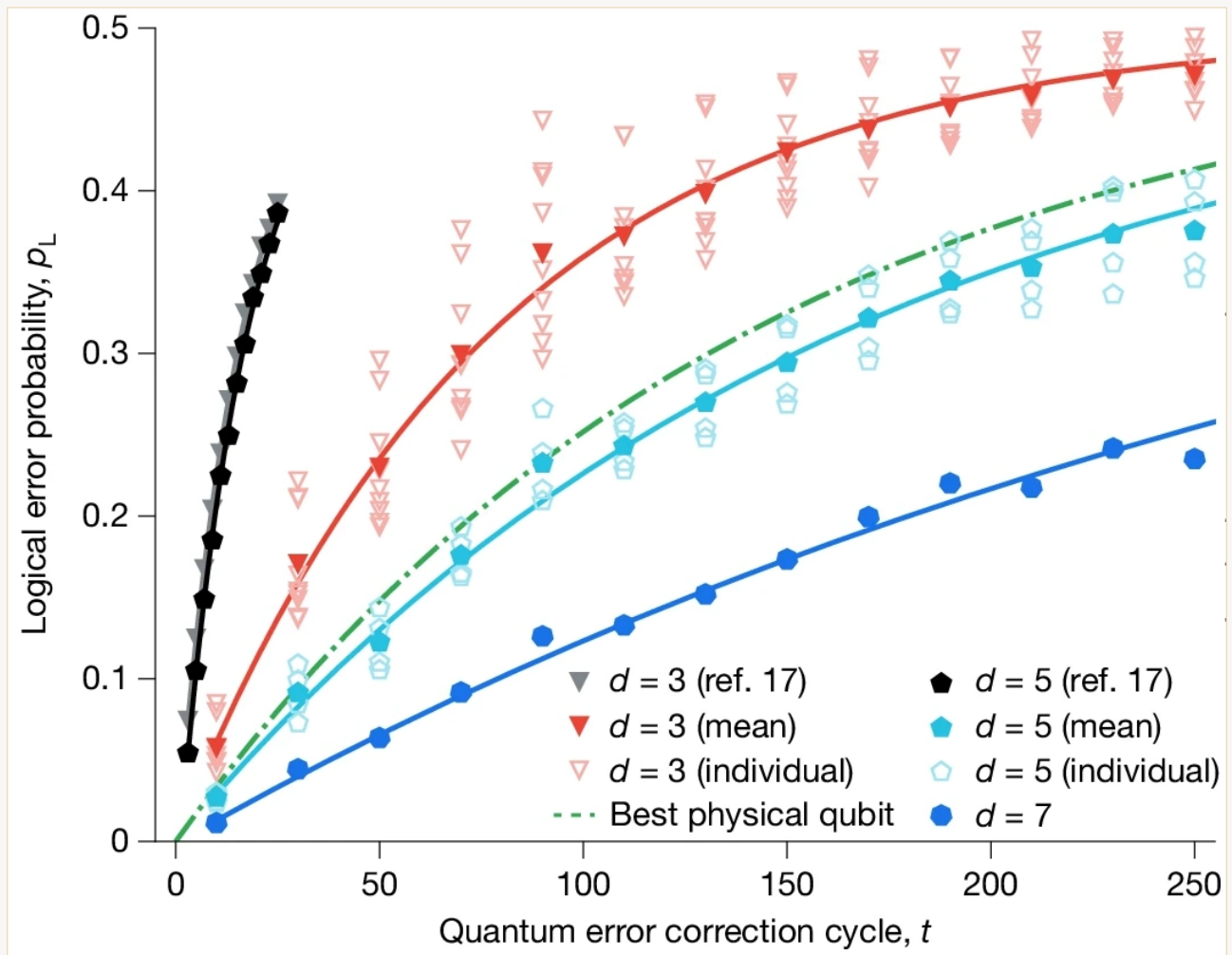
Për tridhjetë vjet, korrigjimi i gabimeve kuantike është mbështetur mbi një premtim që nuk ishte realizuar kurrë në mënyrë të pastër. Teoria thotë se, nëse shpërndan një njësi informacioni kuantik — një kubit *logjik* — në shumë kubitë fizikë me zhurmë, dhe nëse këta kubitë fizikë janë mjaftueshëm të mirë, atëherë shtimi i *më shumë* kubitëve duhet ta bëjë kubitin logjik *më të mirë*, me gabime që bien eksponencialisht ndërsa kodi rritet. Problemi është “nëse”: nën një prag kritik zhurme, më shumë kubitë ndihmojnë; mbi prag, më shumë kubitë vetëm shtojnë më shumë zhurmë. Çdo eksperiment i mëparshëm kishte mbetur në anën e gabuar të kësaj vije, ose nuk e kishte treguar prirjen në mënyrë të pastër. Rritja e kodit i bënte gjërat më keq, jo më mirë.

Në dhjetor 2024, Google Quantum AI raportoi demonstrimin e parë të qartë të regjimit tjetër. Në Willow, gjeneratën e tyre më të re të procesorëve superpërçues, ata ndërtuan memorie me kod sipërfaqësor me distanca kodi 3, 5 dhe 7 dhe panë që norma e gabimit logjik *binte* sa herë që kodi rritej — me faktor $\Lambda = 2.14 \pm 0.02$ për çdo rritje prej dy njësisht në distancë. Memoria më e madhe, me distancë 7 dhe 101 kubitë, mbajti një kubit logjik me gabim $0.143\% \pm 0.003\%$ për cikël korrigjimi dhe — titulli brenda titullit — mbijetoi *më gjatë se kubit fizik më i mirë* i vetë çipit, me faktor 2.4 ± 0.3 . Kjo quhet “beyond breakeven”, pra përtej pikës së barazimit, dhe është hera e parë që i gjithë aparati i korrigjimit të gabimeve e kompenson koston e vet në këtë hardware.

Ky është një piketë e vërtetë dhe ia vlen të jemi të saktë se çfarë lloj pikete është. Është provë se shkallëzimi tani shkon në drejtimin e duhur. Nuk është kompjuter kuantik funksional, dhe punimi nuk pretendon të jetë i tillë.

Çfarë nënkuptojnë “surface code”, “distance” dhe “below threshold”

Një **kubit logjik** është një njësi e mbrojtur informacioni kuantik e koduar në shumë kubitë fizikë. **Surface code**, ose kodi sipërfaqësor, është një mënyrë e veçantë për ta bërë këtë kodim në një rrjet 2D, ku kubitë shtesë “matës” kontrollojnë vazhdimisht gabimet pa prishur informacionin e ruajtur. **Distanca** d e kodit nuk është largësi fizike: është numri më i vogël i gabimeve të vendosura në mënyrën e duhur që mund ta korruptojnë kubitin logjik pa u dalluar nga kodi. Një d më e madhe do të thotë një zonë kodi më e madhe dhe më robuste — shpenzon më shumë kubitë fizikë (afërsisht $2d^2 - 1$) dhe korrigjon më shumë gabime të njëkohshme, deri në $(d - 1)/2$ prej tyre. Kështu, madhësitë e testuara këtu, distancat 3, 5 dhe 7, korrigjojnë përkatësisht 1, 2 dhe 3 gabime të njëkohshme dhe kërkojnë afërsisht 17, 49 dhe 97 kubitë fizikë — memoria me distancë 7 e Google përdori 101, pak më shumë se minimumi i tekstit shkollor. “**Nën prag**” është shprehja vendimtare. Korrigjimi i gabimeve ndihmon vetëm nëse norma e gabimit fizik është nën një vlerë kritike; aty, çdo rritje e distancës ul eksponencialisht normën e gabimit logjik. Faktori i shtypjes Λ e mat këtë — $\Lambda > 1$ do të thotë se rritja e kodit ndihmon, dhe sa më i madh Λ , aq më mirë. Google raporton $\Lambda \approx 2.14$, që do të thotë se çdo rritje prej dy njësisht në distancë e përgjysmoi afërsisht normën e gabimit logjik. Fakti që Λ është qartë mbi 1 është thelbi i rezultatit.



Si grumbullohet gabimi logjik gjatë cikleve të korrigjimit për memoriet me distancë 3, 5 dhe 7 (nga lart poshtë). Vija që duhet parë është ajo e gjelbër me ndërprerje — *kubiti fizik i vetëm më i mirë* në çip. Memoria me distancë 7 (blu, më e ulëta) e grumbullon gabimin më ngadalë se kjo vijë, kështu që kubiti i koduar jeton më gjatë se kubiti fizik më i mirë nga i cili është ndërtuar — “beyond breakeven”, me faktor $2.4\times$. Ky është rezultati i jetëgjatësisë; shtypja nën prag ($\Lambda = 2.14$ ndërsa kodi rritet nga distanca 3 në 5 në 7) gjendet në shifrat e tekstit.

Google Quantum AI and Collaborators / Nature CC BY-NC-ND 4.0

Çfarë bënë autorët

- Ndërtuan memorie me kod sipërfaqësor në dy çipa Willow: një procesor me 105 kubitë që ekzekutoi kodet me distancë 3, 5 dhe 7 për testin e shkallëzimit (më i madhi ishte memoria me distancë 7, 101 kubitë gjithsej dhe 49 kubitë të dhënash), dhe një procesor me 72 kubitë që ekzekutoi një memorie me distancë 5 me dekoder në kohë reale, si dhe kode përsëritjeje me distanca të larta.
- Matën si ndryshonte gabimi logjik *për cikël* ndërsa distanca e kodit rritej nga 3 në 5 në 7, duke nxjerrë faktorin e shtypjes Λ .
- Krahasuan jetëgjatësinë e kubitit logjik me kubitin fizik individual më të mirë në të njëjtin çip, për të testuar “breakeven”.
- Ekzekutuan kodin me distancë 5 me një **dekoder në kohë reale** — hardware klasik që interpreton

kontrollet e gabimeve po aq shpejt sa prodhohen — deri në një milion cikle, për të treguar se korrigjimi mund të ecë bashkë me makinën.

- Zgjatën kode më të thjeshta **përsëritjeje** deri në distancën 29 për të kërkuar burimet e rralla dhe të thella të gabimeve që vendosin një dysheme performance.

Çfarë gjetën

- **Kodi punon nën prag.** Gabimi logjik për cikël ra me $\Lambda = 2.14 \pm 0.02$ për çdo rritje prej dy në distancë — shtypje e pastër eksponenciale, sjellja që teoria premtonte dhe që asnjë procesor nuk e kishte treguar qartë më parë.
- **Memoria me distancë 7 arriti $0.143\% \pm 0.003\%$ gabim për cikël dhe jetoi 2.4 ± 0.3 herë më gjatë** se kubiti i saj fizik më i mirë — përtej breakeven-it.
- **Dekodimi në kohë reale e përballoi ritmin.** Dekoderi kishte mesatarisht 63 mikrosekonda latencë në distancën 5 kundrejt cikleve 1.1-mikrosekondëshe, dhe u mbajt për një milion cikle — korrigjimi i gabimeve po funksiononte drejtpërdrejt, jo vetëm në analizë pas eksperimentit.
- **Mbetet një burim i rrallë dhe i thellë gabimi.** Në testet me kode përsëritjeje, performanca përfundimisht u kufizua nga shpërthime të korreluara gabimesh që ndodhnin afërsisht **një herë në orë** (rreth një në çdo 3×10^9 cikle), duke vendosur një dysheme gabimi pranë 10^{-10} , origjinën e së cilës autorët thonë se **ende nuk e kuptojnë**.

Çfarë nuk provon

- **Nuk** është kompjuter kuantik që po bën llogaritje. Kjo është një *memorie* kuantike: ruan dhe mbron një kubit logjik. Nuk kryen operacione logjike (porta) mes kubitëve logjikë dhe nuk ekzekuton algoritëm.
- **Nuk** është një kubit larg makinave të dobishme. Një kubit logjik me distancë 7 harxhon rreth 101 kubitë fizikë; një normë gabimi rreth 0.1% për cikël është ende shumë mbi afërsisht 10^{-6} deri 10^{-10} që kërkojnë algoritmet reale. Mbyllja e atij hendeku kërkon distanca shumë më të mëdha — pra shumë më tepër kubitë fizikë për kubit logjik — dhe algoritmet e dobishme kërkojnë *mijëra* kubitë logjikë njëherësh. Buxheti fizik shkon në miliona kubitë.
- Fjala “**nëse shkallëzohet**” është **thelbësore**. Vetë përfundimi i punimit thotë se performanca e pajisjes, *nëse shkallëzohet*, mund të plotësojë kërkesat e algoritmeve të mëdha. Të tregosh se prirja është e drejtë për një kubit logjik nuk është e njëjta gjë me të ndërtuar makinën e shkallëzuar, dhe asgjë këtu nuk garanton se prirja mbijeton në madhësi shumë më të mëdha.
- **Dyshemeja e pashpjeguar e gabimit është problem real.** Shpërthimet e korreluara që kufizojnë kodet e përsëritjes janë, sipas autorëve, disa rende madhësie më të mëdha se pritej dhe do t'i pengonin aplikimet më të mëdha fault-tolerant derisa të kuptohen — një mangësi e hapur, e thënë qartë, jo detaj i zgjidhur.
- Rezultati nuk thotë asgjë për **thyerjen e enkriptimit ose “supremacinë kuantike” në detyra të dobishme**. Këto kërkojnë makinën e plotë fault-tolerant, për të cilën ky rezultat është vetëm gur themeli.

Sa e fortë është prova

- **Pretendimi qendror është i fortë dhe i rëndësishëm.** Funksionimi nën prag me shtypje të pastër eksponenciale në tri distanca kodi, bashkë me jetëgjatësi përtej breakeven-it dhe dekoder në kohë reale, është pikërisht kombinimi që fusha përpigëj të arrinte, dhe është demonstruar drejtpërdrejt, jo nxjerrë me hamendje. Ky nuk është artefakt marketingu; është rezultat real inxhinierik nga një grup udhëheqës.
- **Autorët janë të kujdesshëm për shtrirjen.** E paraqesin si *memorie* nën prag, e sinjalizojnë vetë dyshtetësi e pashpjeguar të gabimeve të korreluara dhe e kushtëzojnë të ardhmen me atë “nëse shkallëzohet”. Teprimi, ku shfaqet, është në mbulimin që e rrumbullakos “një kubit memorie i korrigjuar për gabime u përmirësua kur u rrit” në “kompjuterët kuantikë mbërritën”.
- **Statusi i ndershëm është një hap themelor, i bërë qartë.** Një kubit logjik, i mbrojtur mjaftueshëm që shtimi i redundancës më në fund ta ndihmojë — me një rrugë të gjatë, të vështirë dhe jo të garantuar për shkallëzim, porta logjike dhe gabime të pashpjeguara ende përpara.

Pse ka rëndësi

Kompjutimi kuantik fault-tolerant gjithmonë ka pasur një problem “pulë apo vezë”: makinat që do të ishin të dobishme kërkojnë norma gabimi që asnjë kubit fizik nuk mund t’i arrijë, ndërsa zgjidhja — korrigjimi i gabimeve — funksionon vetëm nëse hardware-i është tashmë mjaft i mirë për të qenë nën prag. Kalimi i asaj vije, qoftë edhe një herë dhe me vetëm një kubit logjik, e ndryshon pyetjen nga “a është e mundur fare?” në “sa larg mund të shkallëzohet dhe sa shpejt?”. Ky është ndryshim real e me kuptim, dhe arsyeja pse rezultati meriton vëmendje.

Por i njëjti kujdes që e bën rezultatin të besueshëm duhet edhe ta zbusë historinë rreth tij. Ky është tulla e parë e një themeli, e vendosur mirë. Nuk është ndërtesa, dhe njerëzit që e vendosën janë të parët që e thonë. Mënyra e duhur për të ndjekur kompjutimin kuantik në vitet e ardhshme është pikërisht kjo kurbë jo glamuroze: a mbahet faktori i shtypjes ndërsa kodet rriten, a mund të kryhen portat logjike po aq pastër sa memoria logjike, dhe a do të shpjegohet ai gabim misterioz që ndodh rreth një herë në orë.

Përmbledhje e shkurtër

Google Quantum AI tregoi për herë të parë në mënyrë të pastër se një memorie kuantike me kod sipërfaqësor mund të punojë *nën prag*: ndërsa kodi u rrit nga distanca 3 në 5 në 7, norma e gabimit logjik ra eksponencialisht (rreth $2.14 \times$ për çdo dy hapa), dhe memoria më e madhe, me distancë 7 dhe 101 kubitë, jetoi më gjatë se kubiti i saj fizik më i mirë — përtej breakeven-it — ndërsa korrigjimi i gabimeve funksiononte në kohë reale. Kjo është një piketë reale dhe e kërkuar prej kohësh në inxhinierinë e kompjuerëve kuantikë. Por është gjithashtu vetëm një kubit logjik që vepron si memorie, me normë gabimi ende larg asaj që kërkojnë algoritmet reale, pa operacione logjike, me një dyshtetësi gabimi të pashpjeguar që vetë autorët e theksojnë dhe me shumë rende madhësie shkallëzim përpara.

Një prag real u kalua — jo se u dorëzua një kompjuter kuantik.

Burimet

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>