



The CLEAN PAPER

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

Kvantna memorija napokon se poboljšavala kako je rasla — stvarna prekretnica, ali ne i gotov računar

Google Quantum AI prvi je put nedvosmisleno pokazao da kvantna memorija sa površinskim kodom može raditi ispod praga greške: kako su povećavali udaljenost koda sa 3 na 5 pa na 7, stopa logičkih grešaka opadala je eksponencijalno, oko 2,14 puta za svaka dva koraka, a najveća memorija udaljenosti 7 sa 101 kubitom trajala je duže od svog najboljeg fizičkog kubita — prešla je tačku izjednačenja — uz ispravljanje grešaka u stvarnom vremenu. To je dugo očekivana inženjerska prekretnica. Ali riječ je o jednom logičkom kubitu koji služi kao memorija, sa stopom grešaka još daleko od one potrebne za stvarne algoritme, bez izvedenih logičkih operacija, uz neobjašnjivu donju granicu stope grešaka koju autori ističu i uz još nekoliko redova veličine potrebnog skaliranja. Pređen je prag, nije isporučen računar.

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

Trenutak kada se krivulja napokon savila u pravom smjeru

Trideset godina kvantna korekcija grešaka počivala je na obećanju koje nikada nije bilo čisto ispunjeno. Teorija kaže da, ako jednu jedinicu kvantne informacije — jedan *logički* kubit — rasporedite preko mnogo bučnih fizičkih kubita i ako su ti fizički kubiti dovoljno dobri, dodavanje *još* kubita treba logički kubit učiniti *boljim*, tako da greške eksponencijalno padaju kako kod raste. Kvaka je u tom „ako”: ispod kritičnog praga šuma više kubita pomaže; iznad njega više kubita samo dodaje više šuma. Svaki prethodni eksperiment bio je na pogrešnoj strani te granice ili nije čisto pokazao trend. Veći kod pogoršavao je stvari umjesto da ih poboljšava.

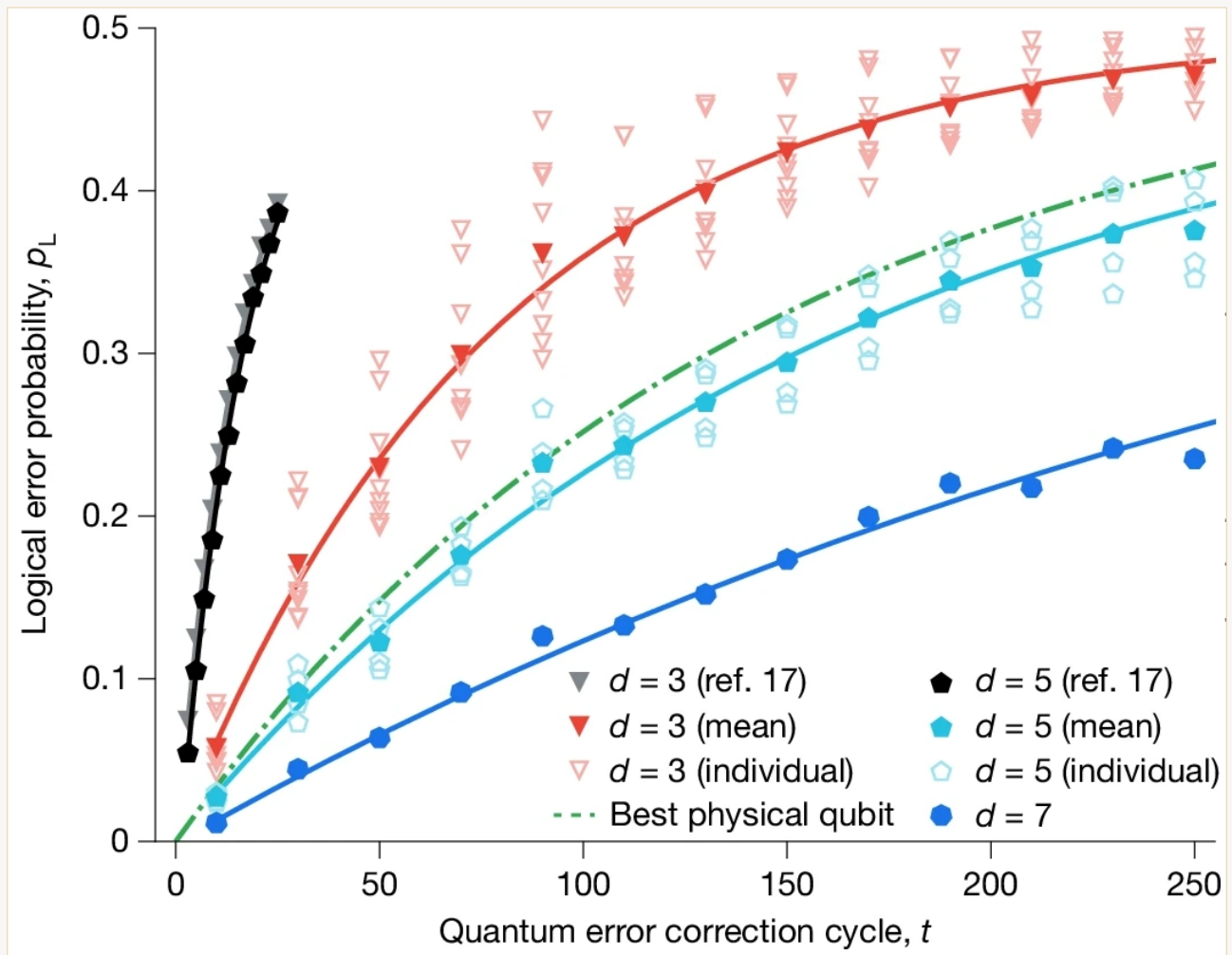
U decembru 2024. Google Quantum AI izvijestio je o prvoj jasnoj demonstraciji suprotnog režima. Na Willowu, njihovoj najnovijoj generaciji supravodljivih procesora, izgradili su memorije sa surface codeom udaljenosti 3, 5 i 7 te gledali kako stopa logičkih grešaka *pada* svaki put kada kod postane veći — za faktor $\Lambda = 2.14 \pm 0.02$ pri svakom povećanju udaljenosti za dva. Najveća memorija, s udaljenosti 7 i 101 kubitom, držala je logički kubit uz grešku od $0.143\% \pm 0.003\%$ po ciklusu korekcije i — naslov unutar naslova — trajala je *duže od najboljeg fizičkog kubita od kojeg je bila sastavljena*, za faktor 2.4 ± 0.3 . To se naziva „iznad breakevena” i prvi je put da je cijeli aparat korekcije grešaka na ovom hardveru nadoknadio vlastiti trošak.

To je stvarna prekretnica i vrijedi biti precizan kakva. To je dokaz da skaliranje sada ide u pravom smjeru. Nije funkcionalni kvantni računar i rad ne tvrdi da jest.

Šta znače „surface code”, „udaljenost” i „ispod praga”

Logički kubit jedna je zaštićena jedinica kvantne informacije kodirana kroz mnogo fizičkih kubita. **Surface code** je poseban način takvog kodiranja na 2D mreži, gdje dodatni „mjerni” kubiti stalno provjeravaju greške bez narušavanja pohranjene informacije. **Udaljenost koda** d nije fizička udaljenost: to je najmanji broj dobro raspoređenih grešaka koji može oštetiti logički kubit a da ih kod ne primijeti. Veći d znači veći i robusniji komad koda — troši više fizičkih kubita (otprilike $2d^2 - 1$) i ispravlja više istovremenih grešaka, do $(d - 1)/2$. Zato tri veličine testirane ovdje, udaljenosti 3, 5 i 7, ispravlja 1, 2 i 3 istovremene greške i troše otprilike 17, 49 i 97 fizičkih kubita — Googleova memorija udaljenosti 7 koristila je 101, malo više od tog udžbeničkog minimuma.

„**Ispod praga**” ključna je fraza. Korekcija grešaka pomaže samo ako je stopa fizičkih grešaka ispod kritične vrijednosti; tada svako povećanje udaljenosti eksponencijalno smanjuje stopu logičkih grešaka. Faktor potiskivanja Λ to mjeri — $\Lambda > 1$ znači da povećavanje koda pomaže, a što je Λ veći, to bolje. Google izvještava $\Lambda \approx 2.14$, što znači da je svako povećanje udaljenosti za dva otprilike prepolovilo stopu logičkih grešaka. Činjenica da je Λ komotno iznad 1 cijeli je rezultat.



Kako se logička greška nakuplja tokom ciklusa korekcije za memorije udaljenosti 3, 5 i 7 (odozgo prema dolje). Linija koju treba pratiti zelena je isprekidana — *najbolji pojedinačni fizički kubit* na čipu. Memorija udaljenosti 7 (plava, najniža) nakuplja grešku sporije od te linije, pa kodirani kubit nadživi najbolji fizički kubit od kojeg je sastavljen — „iznad breakevena“, za faktor $2.4\times$. To je rezultat životnog vijeka; samo potiskivanje ispod praga ($\Lambda = 2.14$ dok kod raste s udaljenosti 3 na 5 pa 7) nalazi se u brojkama u tekstu.

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

Šta su autori uradili

- Izgradili su surface-code memorije na dva Willow čipa: procesoru sa 105 kubita koji je pokretao kodove udaljenosti 3, 5 i 7 u testu skaliranja (najveći je bio kod udaljenosti 7 sa 101 kubitom i 49 podatkovnih kubita) te procesoru sa 72 kubita koji je pokretao memoriju udaljenosti 5 s dekoderom u stvarnom vremenu i kodove ponavljanja velike udaljenosti.
- Mjerili su kako se logička greška *po ciklusu* mijenja kada udaljenost koda raste s 3 na 5 pa 7, izvodeći faktor potiskivanja Λ .
- Uporedili su životni vijek logičkog kubita s najboljim pojedinačnim fizičkim kubitom na istom čipu kako bi testirali „breakeven“.
- Pokretali su kod udaljenosti 5 s **dekoderom u stvarnom vremenu** — klasičnim hardverom koji

tumači provjere grešaka brzinom kojom nastaju — do milion ciklusa, kako bi pokazali da korekcija može pratiti mašina.

- Jednostavnije **kodove ponavljanja** rastegnuli su do udaljenosti 29 kako bi tražili rijetke, duboke izvore grešaka koji postavljaju donju granicu performansi.

Šta su pronašli

- **Kod je ispod praga.** Logička greška po ciklusu smanjivala se za faktor $\Lambda = 2.14 \pm 0.02$ pri svakom povećanju udaljenosti za dva — čisto eksponencijalno potiskivanje, ponašanje koje je teorija obećavala a nijedan procesor nije definitivno pokazao.
- **Memorija udaljenosti 7 dosegula je grešku od $0.143\% \pm 0.003\%$ po ciklusu** i živjela 2.4 ± 0.3 puta duže od najboljeg fizičkog kubita — iznad breakevena.
- **Dekodiranje u stvarnom vremenu držalo je korak.** Dekoder je na udaljenosti 5 imao prosječnu latenciju od 63 mikrosekunde nasuprot ciklusu od 1.1 mikrosekunde, održano kroz milion ciklusa — korekcija grešaka radila je uživo, a ne samo u naknadnoj analizi.
- **Ostaje rijedak, dubok izvor grešaka.** U testovima s kodovima ponavljanja performanse su na kraju ograničavali korelirani naleti grešaka koji se događaju približno **jednom na sat** (oko jednom u svakih 3×10^9 ciklusa), postavljajući pod greške blizu 10^{-10} , čiji uzrok autori kažu da **još nije shvaćen**.

Šta ovo ne dokazuje

- Ovo **nije** kvantni računar koji računa. Ovo je kvantna *memorija*: pohranjuje i štiti jedan logički kubit. Ne izvodi logičke operacije (vrata) između logičkih kubita i ne pokreće algoritam.
- **Nije** jedan kubit udaljeno od korisnih mašina. Logički kubit udaljenosti 7 troši oko 101 fizički kubit; stopa greške od 0.1% po ciklusu još je daleko iznad približno 10^{-6} do 10^{-10} koliko trebaju stvarni algoritmi. Zatvaranje tog jaza znači prelazak na mnogo veće udaljenosti — mnogo više fizičkih kubita po logičkom kubit — a korisni algoritmi trebaju *hiljade* logičkih kubita odjednom. Budžet fizičkih kubita tada ide u milione.
- Izraz „**ako se skalira**” **nosi stvarnu težinu**. Vlastiti zaključak rada glasi da bi performanse uređaja, *ako se skaliraju*, mogle zadovoljiti zahtjeve velikih algoritama. Pokazati pravi trend na jednom logičkom kubit nije isto što i izgraditi skalirani mašina i ništa ovdje ne jamči da će trend preživjeti na mnogo većim veličinama.
- **Neobjašnjeni pod grešaka živ je problem.** Korelirani naleti koji ograničavaju kodove ponavljanja, riječima autora, redovima su veličine veći od očekivanog i onemogućili bi veće fault-tolerant primjene dok se ne razumiju — otvorena mana, jasno navedena, a ne riješen detalj.
- Rad ne govori ništa o **razbijanju enkripcije ili „kvantnoj nadmoći” za korisne zadatke**. Za to je potreban puni fault-tolerant mašina kojem je ovo temeljni kamen, a ne njegova demonstracija.

Koliko su dokazi jaki

- **Centralna tvrdnja čvrsta je i važna.** Rad ispod praga s čistim eksponencijalnim potiskivanjem kroz tri udaljenosti koda, uz životni vijek iznad breakevena i funkcionalan dekodirer u stvarnom vremenu, upravo je kombinacija koju je područje pokušavalo postići i ovdje je direktno demonstrirana, a ne izvedena posredno. To nije artefakt hypea; stvaran je inženjerski rezultat vodeće grupe.
- **Autori su oprezni s opsegom.** Rezultat opisuju kao memoriju *ispod praga*, sami ističu neobjašnjeni pod koreliranih grešaka i budućnost uslovljavaju upadljivim „ako se skalira”. Pretjerivanje, gdje ga ima, dolazi iz okolnog izvještavanja koje „memorijski kubit s korekcijom grešaka poboljšavao se kako je rastao” zaokružuje u „kvantno računarstvo je stiglo”.
- **Pošten status je temeljni korak, uredno napravljen.** Jedan logički kubit, zaštićen dovoljno dobro da dodatna redundancija napokon pomaže — uz dug, težak i još nezajamčen put skaliranja, logičkih vrata i neobjašnjenih grešaka koji tek slijedi.

Zašto je to važno

Fault-tolerant kvantno računarstvo oduvijek je imalo problem kokoši i jajeta: korisnim strojevima trebaju stope greške koje nijedan fizički kubit ne može postići, a rješenje — korekcija grešaka — djeluje samo ako je hardver već dovoljno dobar da bude ispod praga. Prelazak te granice, makar jednom i na samo jednom logičkom kubit, mijenja pitanje iz „je li ovo uopće moguće?” u „koliko se daleko može skalirati i koliko brzo?” To je stvarna i važna promjena i razlog zašto rezultat zaslužuje pažnju.

Ali ista pažnja zbog koje je rezultat vjerodostojan treba umiriti priču oko njega. Ovo je prva cigla temelja, dobro položena. Nije zgrada, a ljudi koji su je položili prvi to kažu. Pravi način praćenja kvantnog računarstva sljedećih godina upravo je ova neglamurozna krivulja: drži li faktor potiskivanja dok kodovi rastu, mogu li se logička vrata izvoditi jednako čisto kao logička memorija i hoće li ona misteriozna greška jednom na sat ikada biti objašnjena.

Kratak sažetak

Google Quantum AI prvi je put čisto pokazao da kvantna memorija sa surface codeom može raditi *ispod praga*: dok su povećavali kod s udaljenosti 3 na 5 pa 7, stopa logičkih grešaka padala je eksponencijalno (oko $2.14 \times$ pri svakom povećanju za dva), a najveća memorija, udaljenosti 7 i sa 101 kubitom, nadživjela je svoj najbolji fizički kubit — iznad breakevena — uz korekciju grešaka u stvarnom vremenu. To je stvarna, dugo tražena prekretnica u inženjerstvu kvantnih računara. Ali to je i jedan logički kubit koji služi kao memorija, sa stopom greške još daleko od onoga što traže stvarni algoritmi, bez izvedenih logičkih operacija, s neobjašnjenim podom grešaka koji sami autori naglašavaju i uz skaliranje od mnogo redova veličine pred nama. Prijeden stvaran prag — nije isporučen kvantni

računar.

Izvori

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>