



The CLEAN PAPER

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

Kvantni pomnilnik je prvič postal boljši, ko je zrasel — resnični mejnik in računalnik, ki to še ni

Google Quantum AI je prvič jasno pokazal, da lahko kvantni pomnilnik s površinsko kodo deluje pod pragom: ko so razdaljo kode povečali s 3 na 5 in nato na 7, je logična stopnja napak eksponentno padala (približno 2,14-krat na vsaka dva koraka), največji, 101-kubitni pomnilnik razdalje 7, pa je preživel dlje kot njegov najboljši fizični kubit — presegel je točko izenačenja — ob sprotnem popravljanju napak. To je dolgo pričakovan inženirski mejnik. Hkrati gre za en sam logični kubit, ki deluje kot pomnilnik, z napako še daleč nad tem, kar potrebujejo resnični algoritmi, brez izvedenih logičnih operacij, z nepojasnenim pragom preostalih napak in z več redi velikosti skaliranja še pred nami. Presežen prag, ne dostavljen računalnik.

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

Trenutek, ko se je krivulja končno obrnila v pravo smer

Kvantno popravljjanje napak že trideset let temelji na obljubi, ki je eksperiment še nikoli ni tako jasno izpolnil. Teorija pravi: če eno enoto kvantne informacije — en *logični* kubit — razporedimo po številnih šumnih fizičnih kubitih in so ti fizični kubiti dovolj dobri, bi moral dodaten kubit logični kubit narediti *boljši*, stopnja napak pa bi morala z rastjo kode eksponentno padati. Težava je v pogoju »če«: pod kritičnim pragom šuma več kubitov pomaga; nad njim dodatni kubiti dodajo predvsem še več šuma. Vsi prejšnji poskusi so bili na napačni strani te meje ali pa trenda niso pokazali dovolj čisto. Večja koda je pomenila slabši rezultat, ne boljšega.

Decembra 2024 je Google Quantum AI poročal o prvi jasni demonstraciji drugega režima. Na Willowu, najnovejši generaciji njihovih superprevodnih procesorjev, so zgradili pomnilnike s površinsko kodo pri razdaljah 3, 5 in 7 ter opazili, da je logična stopnja napak ob vsakem povečanju kode *padla* — za faktor $\Lambda = 2,14 \pm 0,02$ na vsaka dva koraka razdalje. Največji, 101-kubitni pomnilnik razdalje 7, je hranil logični kubit z napako $0,143\% \pm 0,003\%$ na cikel popravljjanja in — ključni rezultat znotraj ključnega rezultata — je preživel *dlje od najboljšega fizičnega kubita*, iz katerega je bil zgrajen, za faktor $2,4 \pm 0,3$. Temu pravijo delovanje »onkraj točke izenačenja« (*beyond breakeven*): prvič je celotna dodatna infrastruktura popravljjanja napak na tej strojni opremi prinesla neto korist.

To je resničen mejnik, pomembno pa je natančno povedati, kakšne vrste. Dokazuje, da gre skaliranje zdaj v pravo smer. Ni delujoč kvantni računalnik in članek tega niti ne trdi.

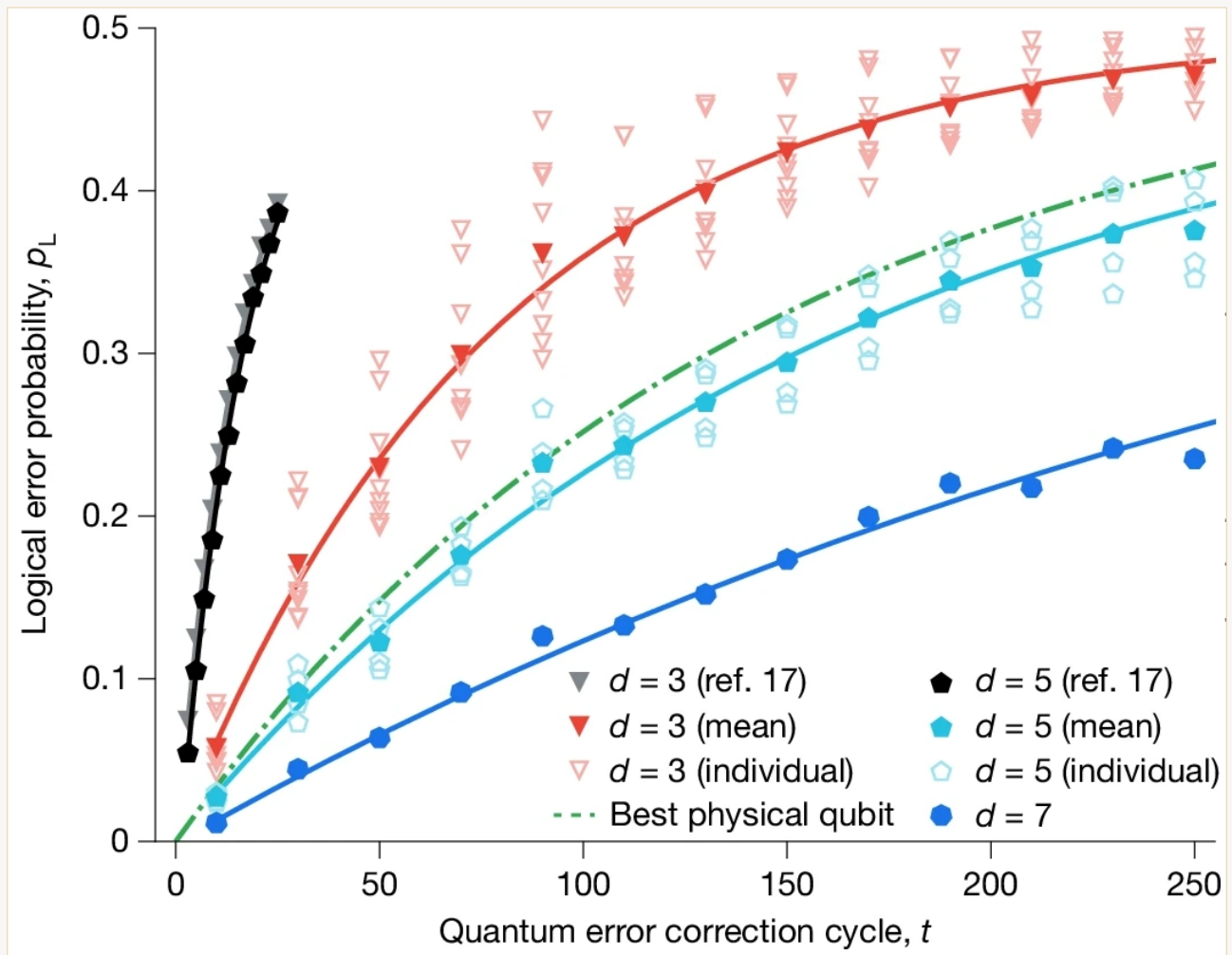
Kaj pomenijo »površinska koda«, »razdalja« in »pod pragom«

Logični kubit je zaščitena enota kvantne informacije, kodirana v številnih fizičnih kubitih.

Površinska koda je določen način takega kodiranja na dvodimenzionalni mreži, kjer dodatni »merilni« kubiti neprestano preverjajo napake, ne da bi pri tem uničili shranjeno informacijo.

Razdalja kode d ni fizična razdalja: je najmanjše število primerno razporejenih napak, ki lahko pokvari logični kubit, ne da bi jih koda zaznala. Večji d pomeni večjo in robustnejšo zaplato kode — porabi več fizičnih kubitov (približno $2d^2 - 1$) in lahko popravi več sočasnih napak, do $(d - 1)/2$. Tri velikosti, preizkušene tukaj, z razdaljami 3, 5 in 7, tako popravijo 1, 2 oziroma 3 sočasne napake ter porabijo približno 17, 49 oziroma 97 fizičnih kubitov — Googlov pomnilnik razdalje 7 jih je uporabljal 101, nekoliko nad tem učbeniškim minimumom.

Izraz »**pod pragom**« je ključen. Popravljjanje napak pomaga le, če je fizična stopnja napak pod kritično vrednostjo; tam vsako povečanje razdalje eksponentno zatre logično stopnjo napak. To meri faktor zatiranja Λ — $\Lambda > 1$ pomeni, da rast kode pomaga, večji Λ pa je boljši. Google poroča $\Lambda \approx 2,14$, kar pomeni, da je vsako povečanje razdalje za dve stopnji logično stopnjo napak približno prepolovilo. Bistvo rezultata je prav to, da je Λ udobno nad 1.



Kako se logična napaka kopiči skozi cikle popravljanja pri pomnilnikih razdalje 3, 5 in 7 (od zgoraj navzdol). Črta, ki jo velja spremljati, je zelena črtkana — *najboljši posamezni fizični kubit* na čipu. Pomnilnik razdalje 7 (modra, najnižja krivulja) napako kopiči počasneje kot ta črta, zato kodirani kubit preživi najboljšega fizičnega kubita, iz katerega je sestavljen — deluje »onkraj točke izenačenja«, za faktor 2,4×. To je rezultat življenjske dobe; samo zatiranje pod pragom ($\Lambda = 2,14$ ob rasti kode z razdalje 3 na 5 in 7) je podano v številkah v besedilu.

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

Kaj so naredili avtorji

- Zgradili so pomnilnike s površinsko kodo na dveh čipih Willow: 105-kubitnem procesorju, na katerem so za test skaliranja poganjali kode razdalje 3, 5 in 7 (največji je bil 101-kubitni pomnilnik razdalje 7 z 49 podatkovnimi kubitami), ter 72-kubitnem procesorju, na katerem so poganjali pomnilnik razdalje 5 s sprotnim dekoderjem in ponovitvene kode z velikimi razdaljami.
- Merili so, kako se logična napaka *na cikel* spreminja, ko razdaljo kode povečajo s 3 na 5 in 7, ter iz tega izračunali faktor zatiranja Λ .
- Življenjsko dobo logičnega kubita so primerjali z najboljšim posameznim fizičnim kubitom na istem čipu, da bi preverili, ali dosežejo »izenačenje«.
- Kodo razdalje 5 so poganjali s **sprotnim dekoderjem** — klasično strojno opremo, ki razlaga mer-

itve napak dovolj hitro, da sledi njihovi produkciji — do milijona ciklov, s čimer so pokazali, da popravljanje napak lahko teče med delovanjem naprave.

- Preprostejše **ponovitvene kode** so razširili do razdalje 29, da bi iskali redke globoke vire napak, ki določijo končno dno zmogljivosti.

Kaj so ugotovili

- **Koda deluje pod pragom.** Logična napaka na cikel je padla za $\Lambda = 2,14 \pm 0,02$ ob vsakem povečanju razdalje za dve — jasno eksponentno zatiranje, vedenje, ki ga je napovedovala teorija in ga noben prejšnji procesor ni dokončno pokazal.
- **Pomnilnik razdalje 7 je dosegel 0,143 % $\pm$ 0,003 % napake na cikel** in živel **2,4 $\pm$ 0,3-krat dlje** od najboljšega fizičnega kubita — onkraj točke izenačenja.
- **Sprotno dekodiranje je zmoglo slediti.** Dekoder je pri razdalji 5 imel povprečno latenco 63 mikrosekund ob 1,1-mikrosekundnem ciklu ter je deloval skozi milijon ciklov — popravljanje napak je teklo v živo, ne le v naknadni analizi.
- **Ostaja redek in globok vir napak.** Pri preizkusih ponovitvene kode so zmogljivost nazadnje omejili korelirani izbruhi napak, ki so se zgodili približno **enkrat na uro** (okoli enkrat na 3×10^9 ciklov) ter povzročili dno napak blizu 10^{-10} . Avtorji pravijo, da njihov izvor **še ni razumljen**.

Česa to ne dokazuje

- To **ni** kvantni računalnik, ki računa. Gre za kvantni *pomnilnik*: hrani in ščiti en logični kubit. Ne izvaja logičnih operacij (vrat) med logičnimi kubitami in ne poganja algoritma.
- Do uporabnega stroja **ne manjka le še en kubit**. Logični kubit razdalje 7 porabi približno 101 fizični kubit; napaka približno 0,1 % na cikel je še vedno daleč nad približno 10^{-6} do 10^{-10} , ki jih potrebujejo resnični algoritmi. Zapiranje te vrzeli pomeni precej večje razdalje — mnogo več fizičnih kubitov na logični kubit — uporabni algoritmi pa potrebujejo *tisoče* logičnih kubitov hkrati. Skupni fizični proračun zato hitro sega v milijone kubitov.
- Izraz **»če se skalira«** je bistven. Lastni zaključek članka je, da bi zmogljivost naprave, *če se skalira*, lahko zadovoljila zahteve velikih algoritmov. Pravilna smer trenda pri enem logičnem kubitni ni isto kot zgrajen skaliran stroj, nič pa ne zagotavlja, da bo trend preživel pri mnogo večjih velikostih.
- **Nepojasnjeno dno napak je dejanski odprt problem.** Korelirani izbruhi, ki omejujejo ponovitvene kode, so po besedah avtorjev za več redov velikosti večji od pričakovanih in bi, dokler jih ne razumejo, onemogočili večje odporne aplikacije. To je odkrito navedena pomanjkljivost, ne rešena podrobnost.
- Rezultat ne pove ničesar o **razbijanju šifriranja ali »kvantni premoči« pri uporabnih nalogah**. Za to je potreben popoln odporen kvantni računalnik, za katerega je ta rezultat temeljni kamen, ne pa demonstracija takega stroja.

Kako močni so dokazi

- **Osrednja trditev je trdna in pomembna.** Delovanje pod pragom z jasnim eksponentnim zatiranjem pri treh razdaljah kode, življenjska doba onkraj točke izenačenja in delujoč sprotni dekodeer so prav kombinacija, ki jo je področje skušalo doseči, rezultat pa je neposredno prikazan in ne le sklepan. To ni artefakt pretiranega poročanja; gre za resničen inženirski rezultat vodilne skupine.
- **Avtorji so previdni glede obsega.** Govorijo o *pomnilniku* pod pragom, sami izpostavijo nepojasnjeno dno koreliranih napak in prihodnost pogojijo z očitnim »če se skalira«. Pretiravanje se pojavi predvsem v poročanju, ki »pomnilniški kubit z odpravljanjem napak je ob povečanju postal boljši« zaokroži v »kvantno računalništvo je tu«.
- **Pošten status je: temeljni korak, opravljen jasno.** En logični kubit, zaščiten dovolj dobro, da dodatna redundanca končno pomaga — pred njim pa dolga, zahtevna in nezagotovljena pot skaliranja, logičnih vrat ter še nepojasnjenih napak.

Zakaj je pomembno

Odporno kvantno računalništvo je vedno imelo občutek začaranega kroga: uporabni stroji potrebujejo stopnje napak, ki jih noben fizični kubit ne doseže, rešitev — popravljanje napak — pa deluje le, če je strojna oprema že dovolj dobra, da je pod pragom. Prestop te meje, četudi samo enkrat in na enem logičnem kubit, spremeni vprašanje iz »ali je to sploh mogoče?« v »kako daleč in kako hitro lahko to skaliramo?«. To je resnična in pomembna sprememba, zato rezultat zasluži pozornost.

Toda ista previdnost, zaradi katere je rezultat zaupanja vreden, mora omejiti zgodbo, ki jo zgradimo okoli njega. To je prva opeka temelja, dobro položena. Ni stavba, in ljudje, ki so jo položili, to sami povedo. Kvantno računalništvo je v naslednjih letih smiselno spremljati prav prek te nebleščeče krivulje: ali faktor zatiranja ostane ugoden pri večjih kodah, ali je mogoče logična vrata izvajati tako čisto kot logični pomnilnik in ali bodo tisto skrivnostno napako enkrat na uro kdaj pojasnili.

Kratek povzetek

Google Quantum AI je prvič jasno pokazal, da lahko kvantni pomnilnik s površinsko kodo deluje *pod pragom*: ko so kodo povečali z razdalje 3 na 5 in nato 7, je logična stopnja napak eksponentno padala (približno 2,14-krat na vsaka dva koraka), največji, 101-kubitni pomnilnik razdalje 7, pa je preživel svoj najboljši fizični kubit — onkraj točke izenačenja — medtem ko je popravljanje napak potekalo sproti. To je resničen in dolgo pričakovan mejnik inženirstva kvantnih računalnikov. Hkrati gre za en sam logični kubit, ki služi kot pomnilnik, z napako še daleč od zahtev resničnih algoritmov, brez izvedenih logičnih operacij, z nepojasnjениm dnom napak, na katerega opozarjajo avtorji sami, in s potjo skaliranja čez mnoge rede velikosti. Resnično presežen prag — ne dostavljen kvantni računalnik.

Viri

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>