



The CLEAN PAPER

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

Kvantmälu muutus lõpuks kasvades paremaks — tegelik verstapost, mitte valmis masin

Google Quantum AI näitas esimest korda puhtalt, et pinnakoodiga kvantmälu saab töötada allpool läve: koodi kasvatamisel kauguselt 3 kaugusele 5 ja 7 langes loogiline veamäär eksponentsiaalselt (umbes $2,14\times$ iga kahe sammu kohta) ning suurim, 101-kubitine kaugus-7 mälu, elas kauem kui selle parim füüsiline kubitt — üle breakeven'i — reaalses töötava veaparandusega. See on kaua oodatud insenertehniline verstapost. Samas on see üks loogiline kubitt mäluna, veamääraga, mis on päris algoritmide vajadusest veel kaugel, ilma loogiliste operatsioonideta, autorite enda märgitud seletamata veapõrandaga ja ees seisva mitme suurusjärgu skaleerimisega. Lävi ületatud, mitte arvuti tarnitud.

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

Hetk, mil kõver lõpuks õigesse suunda pöördus

Kolmkümmend aastat on kvantvigade parandus toetunud lubadusele, mida polnud kunagi puhtalt täidetud. Teooria ütleb, et kui jaotada üks kvantinformatsiooni ühik — üks *loogiline* kubitt — paljude müra- fütüsiliste kubittide vahel ning need fütüsilised kubitid on piisavalt head, siis peaks *rohkemate* kubittide lisamine muutma loogilise kubiti *paremaks*: koodi kasvades peaks vigade arv eksponentsiaalselt vähenema. Konks on selles „kui“-s: allpool kriitilist müraläve aitavad lisakubitid; üle läve lisavad need vaid rohkem müra. Kõik varasemad katsed olid kas selle piiri vaele poolel või ei suutnud trendi puhtalt näidata. Koodi kasvatamine tegi asja halvemaks, mitte paremaks.

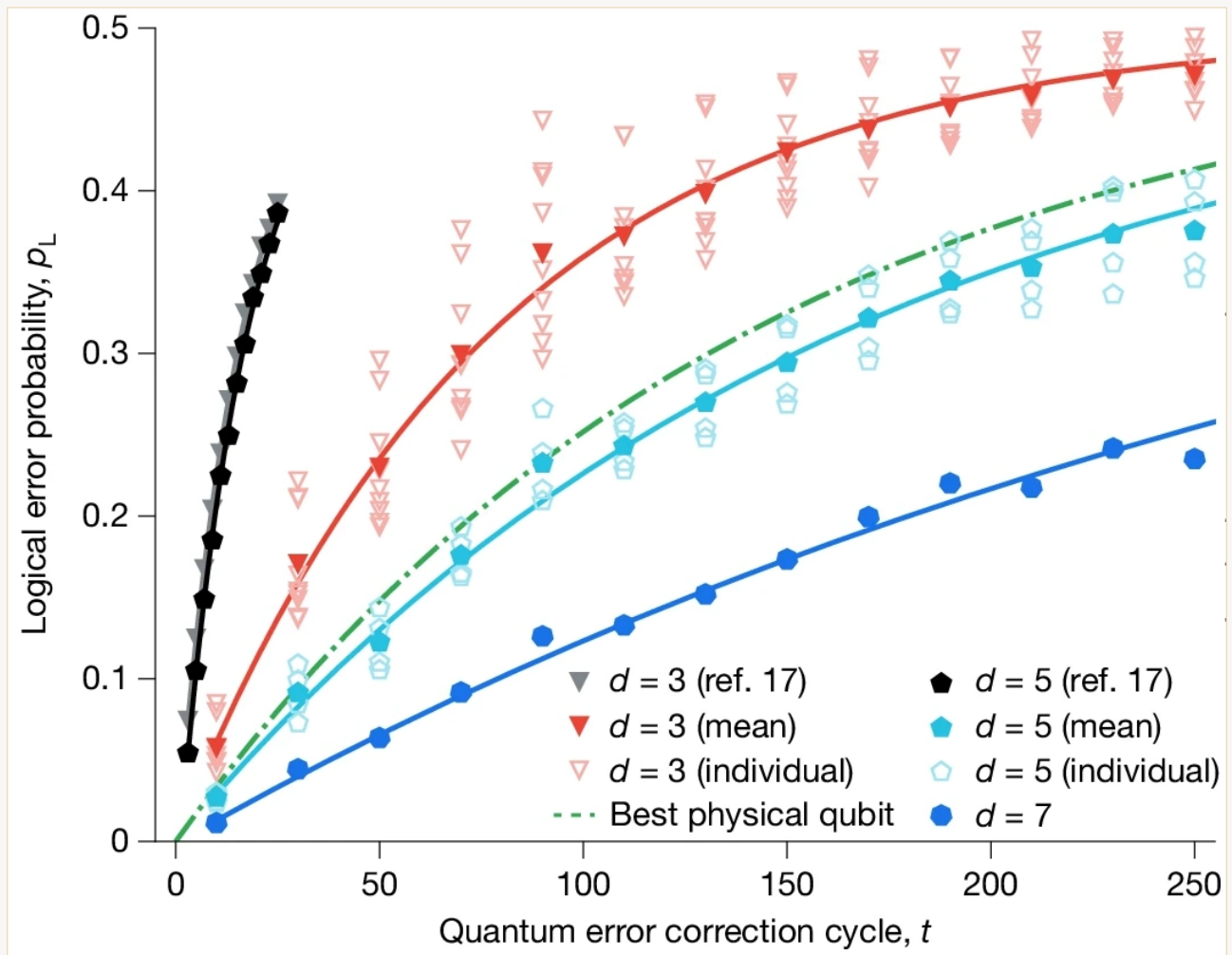
2024. aasta detsembris teatas Google Quantum AI esimesest selgest demonstratsioonist teises režiimis. Nende uusima põlvkonna ülijuhtprotsessoril Willow ehitati pinnakoodiga mälu koodikaugustega 3, 5 ja 7 ning jälgiti, kuidas loogiline veamäär *langes* iga kord, kui kood suuremaks tehti — teguriga $\Lambda = 2.14 \pm 0.02$ iga kahe kaugusastme kohta. Suurim, 101 kubitiga kaugus-7 mälu, hoidis loogilist kubitti veamääraga $0.143\% \pm 0.003\%$ parandustsükli kohta ning — pealkirja sees olev pealkiri — püsis elus *kauem kui sama kiibi parim fütüsiline kubitt*, teguriga 2.4 ± 0.3 . Seda nimetatakse „tasuvuspunktist kaugemale“ (*beyond breakeven*) jõudmiseks ning see on esimene kord, kui kogu veaparanduse aparaat selle riistvara puhul end ära tasus.

See on tõeline verstapost ja tasub täpselt öelda, milline. See tõestab, et skaleerumiskõver liigub nüüd õiges suunas. See ei ole töötav kvantarvuti ja artikkel ei väida seda olevat.

Mida tähendavad „pinnakood“, „kaugus“ ja „allpool läve“

Loogiline kubitt on üks kaitstud kvantinformatsiooni ühik, mis on kodeeritud paljude fütüsiliste kubittide vahel. **Pinnakood** on konkreetne viis sellist kodeerimist teha kahemõõtmelisel võrel, kus täiendavad „mõõtekubitid“ kontrollivad pidevalt vigu, salvestatud informatsiooni ennast häirimata. Koodi **kaugus** d ei ole fütüsiline vahemaa: see on väikseim õigesti paiknenud vigade arv, mis võib loogilise kubiti rikkuda nii, et kood seda ei märka. Suurem d tähendab suuremat ja robustsemat koodilaiku — see kulutab rohkem fütüsilisi kubitte (ligikaudu $2d^2 - 1$) ja parandab rohkem samaaegseid vigu, kuni $(d - 1)/2$. Seega parandavad siin testitud kaugused 3, 5 ja 7 vastavalt 1, 2 ja 3 samaaegset viga ning kasutavad ligikaudu 17, 49 ja 97 fütüsilist kubitti — Google'i ehitatud kaugus-7 mälu kasutas 101, veidi üle õpikunäite miinimumi.

„**Allpool läve**“ on võtmefraas. Veaparandus aitab ainult siis, kui fütüsiline veamäär on allpool kriitilist väärtust; seal vähendab iga kauguse kasv loogilist veamäära eksponentsiaalselt. Summutustegur Λ mõõdab seda — $\Lambda > 1$ tähendab, et koodi kasvatamine aitab, ning mida suurem Λ , seda parem. Google raporteerib $\Lambda \approx 2.14$, mis tähendab, et iga kahe astme võrra suurem kaugus vähendas loogilist veamäära ligikaudu poole võrra. Kogu tulemus seisneb selles, et Λ on mugavalt üle 1.



Kuidas loogiline viga parandustsüklite jooksul koguneb kaugus-3, -5 ja -7 mäludes (ülalt alla). Jälgitav joon on roheline katkendjoon — kiibi *parim üksik füüsiline kubitt*. Kaugus-7 mälu (sinine, kõige alumine) kogub viga aeglasemalt kui see joon, seega elab kodeeritud kubitt kauem kui parim füüsiline kubitt, millest see on ehitatud — „beyond breakeven“, teguriga 2,4×. See on eluea tulemus; allpool-läve summutus ise ($\Lambda = 2.14$, kui kood kasvab kauguselt 3 kaugusele 5 ja 7) on artikli tekstis olevates arvudes.

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

Mida autorid tegid

- Ehitasid pinnakoodiga mälud kahele Willow kiibile: 105-kubitilisele protsessorile, millel jooksid skaleerimistesti kaugus-3, -5 ja -7 koodid (suurim neist 101-kubitine, 49 andmekubitiga kaugus-7 mälu), ning 72-kubitilisele protsessorile, millel jooksis kaugus-5 mälu reaallajdekooderiga ja suure kaugusega korduskoodid.
- Mõõtsid, kuidas loogiline viga *tsükli kohta* muutus, kui koodikaugust kasvatati 3-lt 5-le ja 7-le, ning tuletasid summutusteguri Λ .
- Võrdlesid loogilise kubiti eluiga sama kiibi parima üksiku füüsilise kubitiga, et testida „breakeven’i“.
- Käitasid kaugus-5 koodi **reaallajdekooderiga** — klassikalise riistvaraga, mis tõlgendab veakontrolle sama kiiresti, kui neid toodetakse — kuni miljon tsükli, näitamaks, et veaparandus suudab

masinaga sammu pidada.

- Venitasid lihtsamad **korduskoodid** kuni kauguseni 29, et otsida haruldasi sügavaid veaallikaid, mis seavad jõudlusele põhja.

Mida nad leidsid

- **Kood töötab allpool läve.** Loogiline viga tsükli kohta langes teguriga $\Lambda = 2.14 \pm 0.02$ iga kahe kaugusastme kasvu kohta — puhas eksponentsiaalne summutus, käitumine, mida teooria lubas ja mida ükski protsessor polnud seni lõplikult näidanud.
- **Kaugus-7 mälu jõudis veamäärani $0.143\% \pm 0.003\%$ tsükli kohta** ning elas 2.4 ± 0.3 korda kauem kui selle parim füüsiline kubitt — üle breakeven'i.
- **Reaalajadekodeerimine pidas sammu.** Kaugusel 5 oli dekodeeri keskmine latentsus 63 mikrosekundit võrreldes 1,1-mikrosekundilise tsükliajaga ning süsteem töötas miljoni tsükli vältel — veaparandus toimis töö käigus, mitte ainult hilisemas analüüsis.
- **Alles on haruldane ja sügav veaallikas.** Korduskoodide katsetes piiras jõudlust lõpuks korreleeritud veapurse teke ligikaudu **kord tunnis** (umbes üks iga 3×10^9 tsükli kohta), seades veapõranda 10^{-10} lähedale; autorite sõnul on selle päritolu **endiselt teadmata**.

Mida see ei tõesta

- See **ei ole** kvantarvuti, mis teeb arvutusi. See on kvant-mälu: see salvestab ja kaitseb üht loogilist kubitti. See ei tee loogilisi operatsioone (väravaid) loogiliste kubittide vahel ega käita algoritmi.
- See **ei ole** kasulikust masinast ühe kubiti kaugusel. Kaugus-7 loogiline kubitt kulutab umbes 101 füüsilist kubitti; 0,1%-line veamäär tsükli kohta on endiselt väga kaugel ligikaudu 10^{-6} kuni 10^{-10} tasemest, mida päris algoritmid vajavad. Selle vahe sulgemine tähendab palju suuremaid kaugusi — palju rohkem füüsilisi kubitte ühe loogilise kubiti kohta — ning kasulikud algoritmid vajavad korraga tuhandeid loogilisi kubitte. Füüsiliste kubittide eelarve liigub seega miljonitesse.
- Fraasis „**kui skaleerida**“ teeb „kui“ päris tööd. Artikli enda järeldus on, et seadme jõudlus, *kui seda skaleerida*, võiks vastata suurte algoritmide nõuetele. Ühe loogilise kubiti peal õige trendi näitamine ei ole sama mis skaleeritud masina ehitamine ning miski siin ei garanteeri, et trend püsib palju suuremates süsteemides.
- **Seletamata veapõrand on elus probleem.** Korreleeritud pursked, mis piiravad korduskoodide jõudlust, on autorite sõnul suurusjärkude võrra oodatust suuremad ning muudaksid suuremad vea kindlad rakendused võimatuks, kuni põhjus on mõistetud — avatud puudus, mida öeldakse otse, mitte lahendatud detail.
- See ei ütle midagi **krüptograafia murdmise või kasulike ülesannete „kvantüleoleku“** kohta. Need nõuavad täielikku vea kindlat masinat, mille jaoks see tulemus on vundamendikivi, mitte demonstratsioon.

Kui tugev on tõendus

- **Põhiväide on tugev ja oluline.** Allpool-läve töö puhta eksponentsiaalse summutusega kolme koodikauguse lõikes, koos breakeven'ist parema eluea ja töötava reaalarajdekooderiga, on täpselt see kombinatsioon, milleni valdkond on püüdnud jõuda, ning seda näidatakse otseselt, mitte ei tuledata kaudselt. See ei ole hüpeartefakt; see on juhtiva meeskonna reaalne insenertehniline tulemus.
- **Autorid on ulatuse suhtes ettevaatlikud.** Nad raamivad seda allpool-läve *mäluna*, märgivad ise seletamata korreleeritud vigade pööranda ja seovad tuleviku silmatorkava „kui skaleerida“ tingimusega. Ülepaisutus tekib pigem ümbritsevas kajastuses, kus „veaparandusega mälu kubitt muutus kasvades paremaks“ ümardatakse väiteks „kvantarvutus on kohal“.
- **Aus staatus on vundamendisamm, mis on puhtalt tehtud.** Üks loogiline kubitt, kaitstud piisavalt hästi, et redundantsi lisamine lõpuks aitab — kuid ees on pikk, raske ja mitte garanteeritud tee skaleerimise, loogiliste värvate ja seletamata vigade lahendamiseni.

Miks see oluline on

Veakindlal kvantarvutusel on alati olnud kana-ja-muna probleem: kasulikud masinad vajavad veamäärasid, milleni ükski füüsiline kubitt ei küüni, ning lahendus — veaparandus — töötab ainult siis, kui riistvara on juba piisavalt hea, et olla allpool läve. Selle piiri ületamine kasvõi üks kord, kasvõi ühe loogilise kubiti peal, muudab küsimuse „kas see on üldse võimalik?“ küsimuseks „kui kaugelt seda saab skaleerida ja kui kiiresti?“. See on tõeline ja tähenduslik nihe ning seepärast väärib tulemus tähelepanu.

Kuid sama hoolikus, mis teeb tulemuse usaldusväärseks, peaks ka seda ümbritsevat lugu tagasi hoidma. See on vundamendi esimene hästi laotud tellis. See ei ole hoone, ning selle ladunud inimesed ütlevad seda esimesena. Õige viis kvantarvutuse arengut järgmistel aastatel jälgida ongi vaadata seda ehamugavalt igavat kõverat: kas summutustegur püsib koodide kasvades, kas loogilisi värvaid saab teha sama puhtalt kui loogilist mälu ning kas see salapärase kord-tunnis-viga saab kunagi seletuse.

Lühidalt

Google Quantum AI näitas esimest korda puhtalt, et pinnakoodiga kvantmälu saab töötada *allpool läve*: kui koodi kasvatati kauguselt 3 kaugusele 5 ja 7, langes loogiline veamäär eksponentsiaalselt (umbes $2,14 \times$ iga kahe sammu kohta) ning suurim, 101-kubitine kaugus-7 mälu, elas kauem kui selle parim füüsiline kubitt — üle breakeven'i — samal ajal kui veaparandus töötas reaalarajas. See on tõeline ja kaua oodatud verstapost kvantarvutite inseneritöös. Samas on see üksainus loogiline kubitt, mis toimib mäluna, veamääraga, mis on päris algoritmide vajadusest veel kaugel; loogilisi operatsioone ei tehtud, autorid ise märgivad seletamata veapöörandat ning ees on mitu suurusjärku skaleerimist. Tõeline lävi ületatud — mitte valmis kvantarvuti.

Allikad

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>