



The CLEAN PAPER

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

Kvanttimuisti parani vihdoin kasvaessaan — todellinen virstanpylväs ja kone, jota se ei vielä ole

Google Quantum AI osoitti ensimmäistä kertaa selkeästi, että pintakoodiin perustuva kvanttimuisti voi toimia virhekyynnyksen alapuolella: kun koodietäisyys kasvatettiin 3:sta 5:een ja 7:ään, looginen virhetiheys laski eksponentiaalisesti (noin 2.14-kertaisesti jokaista kahden etäisyysyksikön kasvua kohti), ja suurin, 101 kubitin etäisyyden 7 muisti, eli pidempään kuin sen paras fyysinen kubitti — breakeven-rajan yli — samalla kun virheenkorjaus toimi reaaliajassa. Se on pitkään tavoiteltu tekninen virstanpylväs. Samalla kyse on yhdestä muistina toimivasta loogisesta kubitista, jonka virhetiheys on yhä kaukana oikeiden algoritmien vaatimasta, ilman loogisia operaatioita, kirjoittajien itse esiin nostaman selittämättömän virhepohjan kanssa ja valtava skaalausmatka edessä. Kynnys ylitettiin; tietokonetta ei vielä toimitettu.

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

Hetki, jolloin käyrä taipui oikeaan suuntaan

Kolmenkymmenen vuoden ajan kvanttivirheenkorjaus on nojannut lupaukseen, jota ei ollut koskaan osoitettu puhtaasti käytännössä. Teorian mukaan yksi kvantti-informaation yksikkö — yksi *looginen* kubitti — voidaan jakaa monen kohinaisen fyysisen kubitin kesken, ja jos fyysiset kubitit ovat riittävän hyviä, *useampien* kubittien lisäämisen pitäisi tehdä loogisesta kubitista *parempi*: virheiden pitäisi vähentyä eksponentiaalisesti koodin kasvaessa. Ratkaiseva kohta on ”jos”. Kriittisen kohinakynnyksen alapuolella useampi kubitti auttaa; sen yläpuolella useampi kubitti tuo vain lisää kohinaa. Kaikki aiemmat kokeet olivat olleet tämän rajan väärällä puolella tai eivät olleet näyttäneet kehitystä selkeästi. Koodin kasvattaminen teki asioista huonompia, ei parempia.

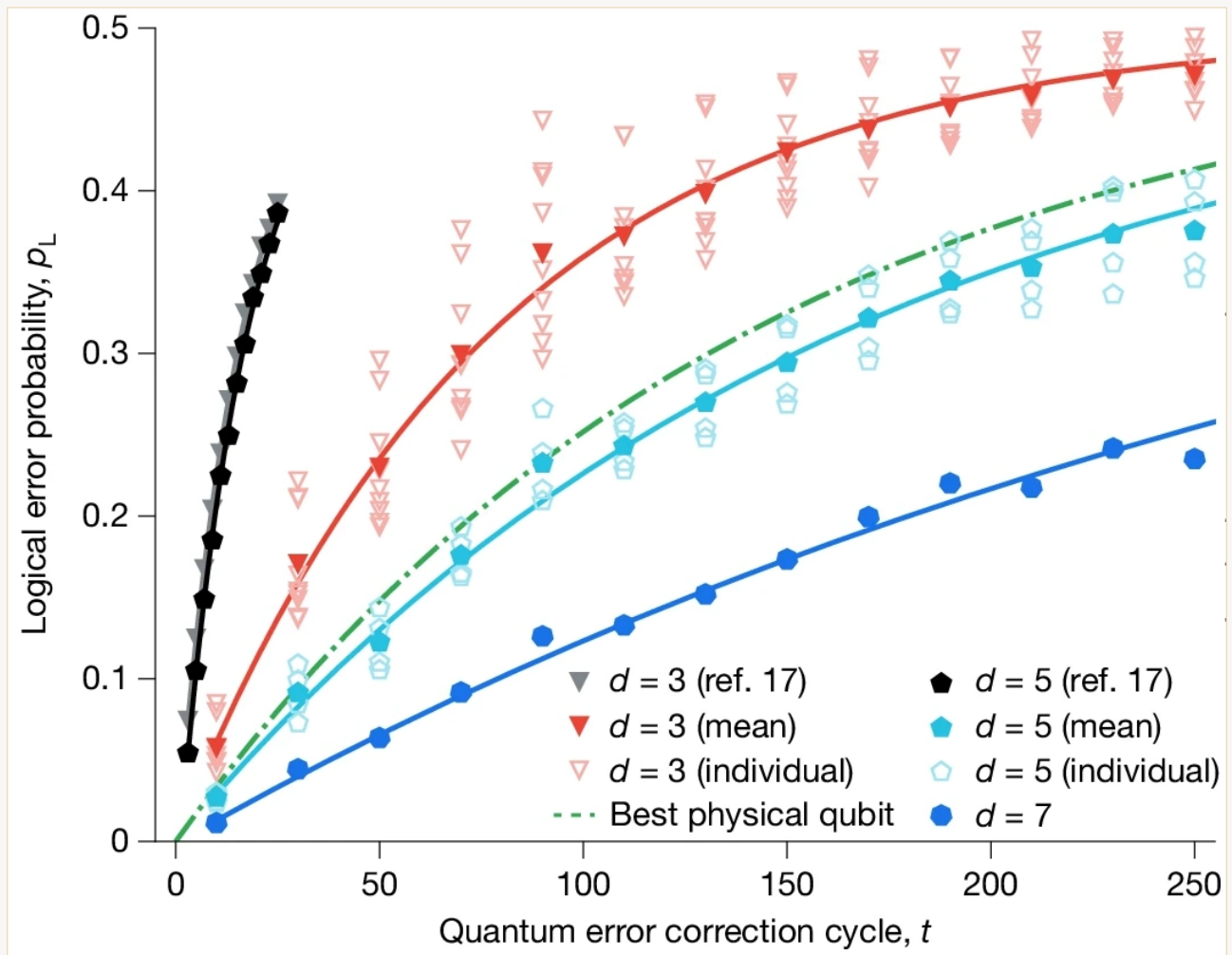
Joulukuussa 2024 Google Quantum AI raportoi ensimmäisen selvän näytön kynnyksen toiselta puolelta. Uusimman sukupolven Willow-suprajohtavilla prosessoreilla ryhmä rakensi pintakoodimuistit koodietäisyyksillä 3, 5 ja 7 ja näki loogisen virhetiheuden *laskevan* aina koodin kasvaessa — vaimennuskertoimella $\Lambda = 2.14 \pm 0.02$ jokaista kahden etäisyyksikön kasvua kohti. Suurin, 101 kubitin etäisyyden 7 muisti, säilytti loogisen kubitin virheellä $0.143\% \pm 0.003\%$ korjaussykliä kohden ja — vielä tärkeämpää — selviytyi *pidempään kuin oma paras fyysinen kubittinsa*, kertoimella 2.4 ± 0.3 . Tätä kutsutaan breakeven-rajan ylittämiseksi, ja tällä laitteistolla virheenkorjauksen koko koneisto maksoi ensimmäistä kertaa itsensä takaisin.

Tämä on aito virstanpylväs, ja on syytä olla tarkka siitä, millainen. Se osoittaa, että skaalaus menee nyt oikeaan suuntaan. Se ei ole toimiva kvanttietokone, eikä artikkeli väitä olevansa sellainen.

Mitä ”pintakoodi”, ”etäisyys” ja ”kynnyksen alapuolella” tarkoittavat

Looginen kubitti on yksi suojattu kvantti-informaation yksikkö, joka on koodattu monen fyysisen kubitin yli. **Pintakoodi** on tietty tapa tehdä tämä koodaus 2D-hilassa, jossa ylimääräiset ”mittauskubitit” tarkistavat jatkuvasti virheitä häiritsemättä tallennettua informaatiota. Koodin **etäisyys** d ei ole fyysinen etäisyys: se on pienin määrä sopivasti sijoittuneita virheitä, jotka voivat korruptoida loogisen kubitin koodin huomaamatta. Suurempi d tarkoittaa suurempaa ja kestävämpää aluetta — siihen käytetään enemmän fyysisiä kubitteja (suunnilleen $2d^2 - 1$) ja se korjaa enemmän samanaikaisia virheitä, korkeintaan $(d - 1)/2$. Tässä testatut etäisyydet 3, 5 ja 7 korjaavat siis 1, 2 ja 3 samanaikaista virhettä ja käyttävät karkeasti 17, 49 ja 97 fyysistä kubittia; Googlen etäisyyden 7 muisti käytti 101 kubittia, hieman oppikirjan minimiä enemmän.

”**Kynnyksen alapuolella**” on ratkaiseva ilmaus. Virheenkorjaus auttaa vain, jos fyysinen virhetiheys on kriittisen arvon alapuolella; siellä jokainen etäisyyden kasvatus vaimentaa loogista virhetiheyttä eksponentiaalisesti. Vaimennuskerroin Λ mittaa tätä — $\Lambda > 1$ tarkoittaa, että koodin kasvattaminen auttaa, ja mitä suurempi Λ , sitä parempi. Google raportoi $\Lambda \approx 2.14$, eli jokainen kahden yksikön kasvu etäisyydessä suunnilleen puolitti loogisen virhetiheuden. Se, että Λ on selvästi yli 1:n, on koko tuloksen ydin.



Loogisen virheen kasvu korjaussykliden myötä etäisyyden 3, 5 ja 7 muisteissa (ylhäältä alas). Oleellinen viiva on vihreä katkoviiva — sirun *paras yksittäinen fyysinen kubitti*. Etäisyyden 7 muisti (sininen, alin) kerää virhettä hitaammin kuin tämä viiva, joten koodattu kubitti elää pidempään kuin paras fyysinen kubitti, josta se on rakennettu — ”beyond breakeven”, kertoimella $2.4\times$. Tämä on elinikäätulos; itse kynnyksen alapuolinen vaimennus ($\Lambda = 2.14$, kun koodi kasvaa etäisyydestä 3 etäisyyksiin 5 ja 7) näkyy tekstin luvuissa.

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

Mitä tutkijat tekivät

- Rakensivat pintakoodimuistit kahdelle Willow-sirulle: 105 kubitin prosessorille, jolla ajettiin skaalaustestin etäisyyksien 3, 5 ja 7 koodit (suurimpana 101 kubitin, 49 datakubitin etäisyyden 7 muisti), sekä 72 kubitin prosessorille, jolla ajettiin etäisyyden 5 muisti reaaliaikaisen dekooderin kanssa ja korkean etäisyyden toistokodeja.
- Mittasivat loogisen virheen sykliä *kohden* koodietäisyyden kasvaessa 3:sta 5:een ja 7:ään ja määrittivät vaimennuskertoimen Λ .
- Vertailivat loogisen kubitin elinikää saman sirun parhaaseen yksittäiseen fyysiseen kubittiin testatakseen ”breakeveniä”.
- Ajoivat etäisyyden 5 koodia **reaaliaikaisella dekooderilla** — klassisella laitteistolla, joka tulkit-

see virhetarkistukset samaa tahtia kuin niitä syntyy — jopa miljoonan syklin ajan osoittaakseen virheenkorjauksen pysyvän laitteen tahdissa.

- Veivät yksinkertaisemmat **toistokoodit** etäisyyteen 29 etsiäkseen harvinaisia syviä virhelähteitä, jotka asettavat suorituskyyvylle pohjan.

Mitä he löysivät

- **Koodi toimii kynnyksen alapuolella.** Looginen virhe sykliä kohden pieneni kertoimella $\Lambda = 2.14 \pm 0.02$ jokaista kahden etäisyysyksikön kasvua kohti — puhdas eksponentiaalinen vaimennus, jonka teoria lupasi mutta jota yksikään prosessori ei ollut kiistattomasti osoittanut.
- **Etäisyyden 7 muisti saavutti $0.143 \% \pm 0.003 \%$ virheen sykliä kohden** ja eli 2.4 ± 0.3 kertaa **pidempään** kuin paras fyysinen kubittinsa — breakeven-ajan yli.
- **Reaaliaikainen dekodaus pysyi mukana.** Dekooderin keskimääräinen latenssi oli etäisyydellä 5 63 mikrosekuntia 1.1 mikrosekunnin sykliajalla, ja toiminta jatkui miljoonan syklin ajan — virheenkorjaus tapahtui livenä, ei vain jälkianalyysissä.
- **Harvinainen syvä virhelähde on yhä olemassa.** Toistokooditesteissä suorituskyykyä rajoittivat lopulta korreloituneet virhepurskeet, joita tapahtui noin **kerran tunnissa** (noin yksi jokaista 3×10^9 sykliä kohti). Ne asettivat virhepohjan lähelle 10^{-10} :tä, eikä kirjoittajien mukaan niiden alkuperää vielä ymmärretä.

Mitä tämä ei todista

- Kyse **ei** ole laskentaa tekevästä kvanttietokoneesta. Tämä on kvanttimuisti: se tallentaa ja suojaa yhtä loogista kubittia. Se ei tee loogisia operaatioita (portteja) loogisten kubittien välillä eikä aja algoritmia.
- Käyttökelpoiseen koneeseen **ei** ole enää vain yhtä kubittia matkaa. Etäisyyden 7 looginen kubitti käyttää noin 101 fyysistä kubittia; noin 0.1 %:n virhe sykliä kohden on yhä kaukana oikeiden algoritmien tarvitsemasta noin 10^{-6} – 10^{-10} tasosta. Kuilun sulkeminen edellyttää paljon suurempia etäisyyksiä — paljon enemmän fyysisiä kubitteja yhtä loogista kohden — ja hyödylliset algoritmit tarvitsevat *tuhansia* loogisia kubitteja samanaikaisesti. Fyysisten kubittien budjetti nousee miljooniin.
- **”Jos skaalataan” on olennainen ehto.** Artikkelin oma johtopäätös on, että laitteen suorituskyyky *skaalattuna* voisi täyttää suurten algoritmien vaatimukset. Oikeansuuntaisen trendin osoittaminen yhdellä loogisella kubitilla ei ole sama asia kuin skaalatun koneen rakentaminen, eikä mikään tässä takaa trendin säilyvän paljon suuremmissa ko’oissa.
- **Selittämätön virhepohja on avoin ongelma.** Korreloituneet purskeet, jotka rajoittavat toistokoodin suorituskyykyä, ovat kirjoittajien mukaan kertaluokkia odotettua suurempia ja estäisivät suuremmat vikasietoiset sovellukset, kunnes ne ymmärretään — avoin vika, joka sanotaan suoraan, ei ratkaistu yksityiskohta.
- Tulos ei sano mitään **salauksen murtamisesta tai hyödyllisten tehtävien ”kvanttiylivallasta”**. Ne tarvitsevat täyden vikasietoisin koneen, jolle tämä on vain yksi peruskivi.

Kuinka vahvaa näyttö on

- **Ydinväite on vankka ja tärkeä.** Kynnyksen alapuolinen toiminta puhtaalla eksponentiaalisella vaimennuksella kolmen koodietäisyyden yli, breakeven-ajan ylittävä elinikä ja toimiva reaaliaikainen dekooderi ovat täsmälleen se yhdistelmä, jota ala on tavoitellut. Se osoitetaan suoraan eikä päätellä epäsuorasti. Tämä ei ole hypeartefakti vaan johtavan ryhmän todellinen tekninen tulos.
- **Kirjoittajat rajaavat tuloksen huolellisesti.** He puhuvat kynnyksen alapuolisesta *muistista*, nostavat itse esiin selittämättömän korreloituneen virhepohjan ja sitovat tulevaisuuden näkyvästi ehtoihin ”jos skaalataan”. Ylilyönti syntyy ympäröivässä uutisoinnissa, jossa ”virheenkorjattu muistkubitti parani kasvaessaan” pyöristyy väitteeksi ”kvanttilaskenta on täällä”.
- **Rehellinen tila on perustavanlaatuinen askel, joka otettiin puhtaasti.** Yksi looginen kubitti, joka on suojattu tarpeeksi hyvin, että redundanssin lisääminen vihdoinkin auttaa — edessä silti pitkä, vaikea ja takaamaton tie skaalaukseen, loogisiin portteihin ja selittämättömien virheiden ratkaisemiseen.

Miksi tällä on merkitystä

Vikasietoisessa kvanttilaskennassa on aina ollut muna–kana-ongelma: hyödylliset koneet tarvitsevat virhetiheysä, joihin mikään fyysinen kubitti ei yllä, ja korjaus — virheenkorjaus — toimii vain, jos laitteisto on jo tarpeeksi hyvä ollakseen kynnyksen alapuolella. Tämän rajan ylittäminen edes kerran, edes yhdellä loogisella kubitilla, muuttaa kysymyksen muodosta ”onko tämä ylipäättään mahdollista?” muotoon ”kuinka pitkälle ja kuinka nopeasti tämä voidaan skaalata?”. Se on todellinen ja merkityksellinen muutos, minkä vuoksi tulos ansaitsee huomiota.

Mutta juuri sama huolellisuus, joka tekee tuloksesta uskottavan, pitäisi ulottaa sen ympärille kerrottuun tarinaan. Tämä on perustuksen ensimmäinen tiili, hyvin asetettuna. Se ei ole rakennus, ja tiilen asettaneet sanovat sen ensimmäisinä. Oikea tapa seurata kvanttilaskentaa tulevina vuosina on juuri tämä epähohdokas käyrä: säilyykö vaimennuskerroin koodien kasvaessa, saadaanko loogiset portit yhtä puhtaiksi kuin looginen muisti ja selviääkö mystisen kerran tunnissa tapahtuvan virheen syy.

Tiivistelmä

Google Quantum AI osoitti ensimmäistä kertaa selkeästi, että pintakoodiin perustuva kvanttimuisti voi toimia *kynnyksen alapuolella*: koodin kasvaessa etäisyydestä 3 etäisyyksiin 5 ja 7 looginen virhetiheys pieneni eksponentiaalisesti (noin $2.14 \times$ jokaista kahta askelta kohti), ja suurin, 101 kubitin etäisyyden 7 muisti eli parasta fyysistä kubitistaan pidempään — breakeven-ajan yli — samalla kun virheenkorjaus toimi reaaliajassa. Tämä on aito ja pitkään tavoiteltu kvanttietokoneiden rakentamisen virstanpylväs. Samalla kyse on yhdestä muistina toimivasta loogisesta kubitista, jonka virhetiheys on yhä kaukana oikeiden algoritmien vaatimasta, ilman loogisia operaatioita, kirjoittajien itse esiin nostaman selittämättömän virhepohjan kanssa ja monen kertaluokan

skaalausmatka edessä. Todellinen kynnys ylitettiin — kvanttietokonetta ei vielä toimitettu.

Lähteet

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>