



The CLEAN PAPER

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

'n Kwantumgeheue het uiteindelik beter geword soos dit gegroei het — die egte mylpaal, en die masjien wat dit nie is nie

Google Quantum AI het vir die eerste keer skoon gewys dat 'n surface-code-quantumgeheue onder die drempel kan loop: soos hulle die kode van afstand 3 deur 5 na 7 laat groei het, het die logiese foutkoers eksponensieel gedaal (omtrent $2,14 \times$ per twee treë), en die grootste, 'n afstand-7-geheue met 101 kubits, het sy beste fisiese kubit oorleef — beyond breakeven — met foutkorreksie wat intyds geloop het. Dit is 'n lank gesoekte ingenieursmylpaal. Dit is óók 'n enkele logiese kubit wat as geheue dien, teen 'n foutkoers steeds ver van wat werklike algoritmes nodig het, met geen uitgevoerde logiese bewerkings nie, 'n onverklaarde foutvloer wat die outeurs self aanmerk, en ordes van grootte se skalering wat voorlê. 'n Drempel oorgesteek, nie 'n rekenaar afgelewer nie.

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

Die oomblik toe die kromme die regte kant toe gebuig het

Dertig jaar lank het kwantumfoutkorreksie op 'n belofte gerus wat nooit netjies nagekom is nie. Die teorie sê: versprei een eenheid kwantuminligting — een *logiese* qubit — oor baie ruiserige fisiese qubits, en as daardie fisiese qubits goed genoeg is, behoort die byvoeg van *meer* daarvan die logiese qubit *beter* te maak, met foute wat eksponensieel daal namate die kode groei. Die vangplek is die “as”: onder 'n kritieke ruisdrempel help meer qubits; daarbo voeg meer qubits net meer ruis by. Elke vorige eksperiment het aan die verkeerde kant van daardie lyn geleef, of kon die tendens nie skoon wys nie. Om die kode te laat groei het sake vererger, nie verbeter nie.

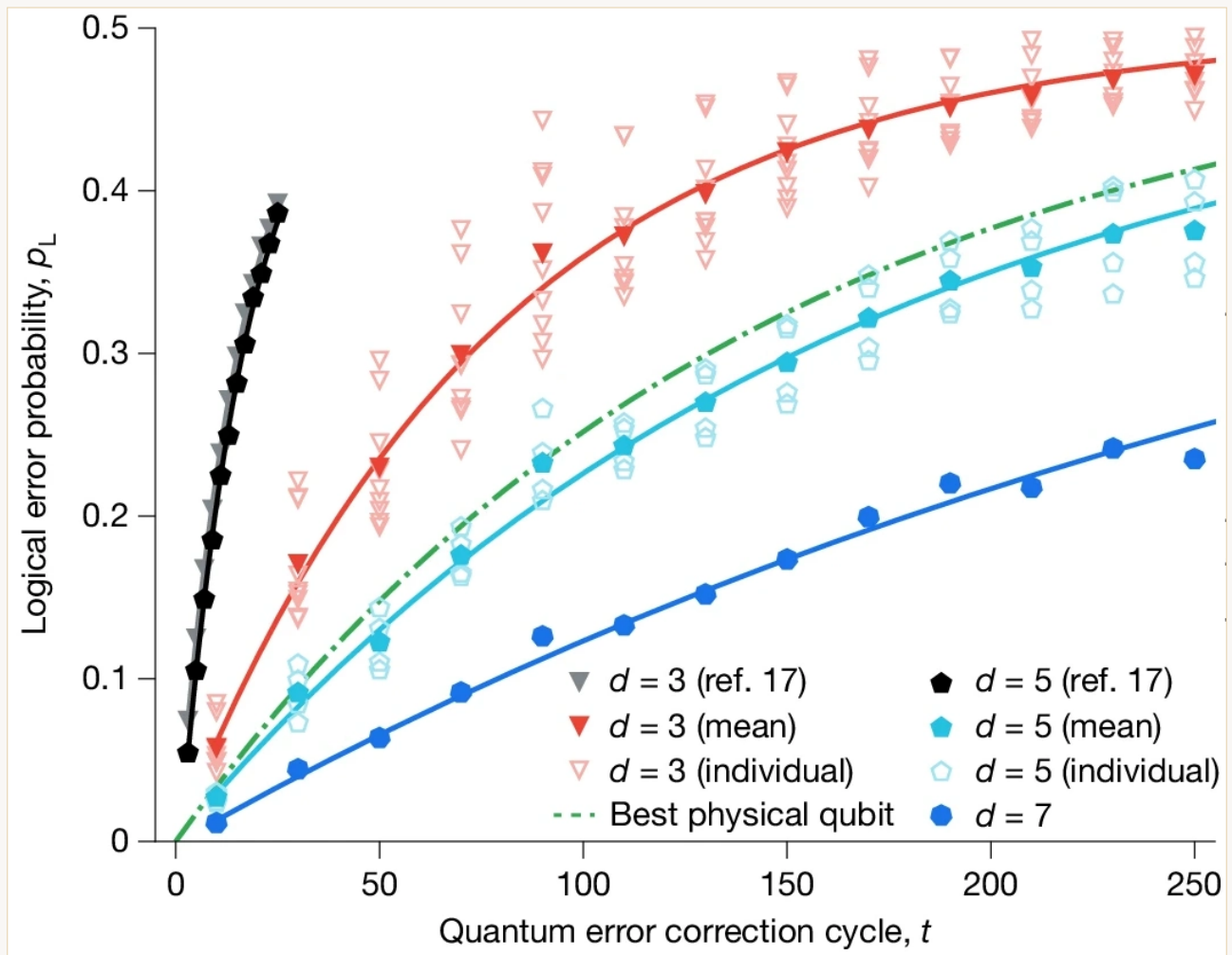
In Desember 2024 het Google Quantum AI die eerste duidelike demonstrasie van die ander bestel gerapporteer. Op Willow, hulle nuutste generasie supergeleidende verwerkers, het hulle surface-code-geheues gebou op kode-afstande 3, 5 en 7, en gesien hoe die logiese foutkoers *daal* elke keer as die kode groter word — met 'n faktor $\Lambda = 2,14 \pm 0,02$ vir elke twee treë op in afstand. Hulle grootste, 'n afstand-7-geheue met 101 qubits, het 'n logiese qubit gehou met 'n fout van $0,143\% \pm 0,003\%$ per korreksiesiklus, en — die opskrif binne die opskrif — dit het *langer oorleef as sy eie beste fisiese qubit*, met 'n faktor van $2,4 \pm 0,3$. Dit word “beyond breakeven” genoem, en dit is die eerste keer dat die hele apparaat van foutkorreksie homself op hierdie hardware betaal het.

Dit is 'n egte mylpaal, en dit is die moeite werd om presies te wees oor watter soort. Dit is 'n bewys dat die skalering nou die regte kant toe loop. Dit is nie 'n werkende kwantumrekenaar nie, en die artikel maak ook nie daarop aanspraak nie.

Wat “surface code”, “afstand” en “onder die drempel” beteken

'n **Logiese qubit** is een beskermde eenheid kwantuminligting, gekodeer oor baie fisiese qubits. Die **surface code** is 'n bepaalde manier om daardie kodering op 'n 2D-rooster te doen, waar ekstra “meet”-qubits voortdurend vir foute kyk sonder om die gestoorde inligting te versteur. Die kode-**afstand** d is nie 'n fisiese afstand nie: dit is die kleinste aantal goed geplaaste foute wat die logiese qubit kan bederf sonder dat die kode dit agterkom. 'n Groter d is 'n groter, stewiger lap — dit bestee meer fisiese qubits (ruweg $2d^2 - 1$) en korrigeer meer gelyktydige foute, tot $(d - 1)/2$ daarvan. Die drie groottes wat hier getoets is, afstande 3, 5 en 7, korrigeer dus 1, 2 en 3 gelyktydige foute en bestee ruweg 17, 49 en 97 fisiese qubits — die afstand-7-geheue wat Google gebou het, het 101 gebruik, effens bo hierdie handboekminimum.

“**Onder die drempel**” (“below threshold”) is die deurslaggewende frase. Foutkorreksie help net as jou fisiese foutkoers onder 'n kritieke waarde lê; daar onderdruk elke toename in afstand die logiese foutkoers eksponensieel. Die onderdrukkingsfaktor Λ meet dit — $\Lambda > 1$ beteken dat dit help om die kode te laat groei, en hoe hoër Λ , hoe beter. Google rapporteer $\Lambda \approx 2,14$: elke toename van twee in afstand het die logiese foutkoers ruweg gehalveer. Dat Λ gemaklik bo 1 lê, is die hele resultaat.



Hoe die logiese fout oor die korreksiesiklusse opbou, vir die afstand-3-, -5- en -7-geheues (van bo na onder). Die lyn om dop te hou is die groen stippellyn — die *beste enkele fisiese qubit* op die skyfie. Die afstand-7-geheue (blou, laagste) versamel foute stadiger as daardie lyn, so die gekodeerde qubit oorleef die beste fisiese qubit waaruit dit gebou is — “beyond breakeven”, met ’n faktor van 2,4×. Dit is die leeftydresultaat; die onderdrukking onder die drempel self ($\Lambda = 2,14$, soos die kode groei van afstand 3 na 5 na 7) lê in die syfers in die teks.

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

Wat die outeurs gedoen het

- Hulle het surface-code-geheues op twee Willow-skyfies gebou: ’n 105-qubit-verwerker waarop die afstand-3-, -5- en -7-kodes agter die skaleringstoets geloop het (die grootste daarvan die afstand-7-geheue met 101 qubits en 49 data-qubits), en ’n 72-qubit-verwerker wat ’n afstand-5-geheue met ’n intydse dekodeerder plus die hoë-afstand-herhalingskodes uitgevoer het.
- Hulle het gemeet hoe die logiese fout *per siklus* verander het toe hulle die kode-afstand van 3 deur 5 na 7 verhoog het, en die onderdrukkingsfaktor Λ daaruit onttrek.
- Hulle het die logiese qubit se leeftyd vergelyk met die beste individuele fisiese qubit op dieselfde skyfie — die toets vir “breakeven”.

- Hulle het die afstand-5-kode met 'n **intydse dekodeerder** laat loop — klassieke hardeware wat die foutkontroles interpreteer so vinnig as wat hulle inkom — vir tot 'n miljoen siklusse, om te wys die foutkorreksie kan by die masjien byhou.
- Hulle het eenvoudiger **herhalingskodes** tot by afstand 29 gedryf om die seldsame, diep foutbronne op te spoor wat 'n vloer onder die werkverrigting lê.

Wat hulle gevind het

- **Die kode is onder die drempel.** Die logiese fout per siklus het met $\Lambda = 2,14 \pm 0,02$ gedaal vir elke toename van twee in afstand — duidelike eksponensiële onderdrukking, die gedrag wat die teorie belooft en wat geen verwerker beslissend gewys het nie.
- **Die afstand-7-geheue het $0,143\% \pm 0,003\%$ fout per siklus bereik en het $2,4 \pm 0,3$ keer langer** gelewe as sy beste fisiese qubit — beyond breakeven.
- **Intydse dekodering het bygehou.** Die dekodeerder het by afstand 5 gemiddeld 63 mikrosekondes latensie gehad teenoor 'n siklustyd van 1,1 mikrosekondes, volgehou oor 'n miljoen siklusse — die foutkorreksie het lewendig geloop, nie net in agterna-ontleding nie.
- **'n Seldsame, diep foutbron bly oor.** In die herhalingskode-toetse is die werkverrigting uiteindelik beperk deur gekorreleerde foutuitbarstings wat ruweg **een keer per uur** gebeur (omtrent een uit elke 3×10^9 siklusse), wat 'n foutvloer naby 10^{-10} lê waarvan die oorsprong volgens die outeurs **nog nie verstaan word nie**.

Wat dit nie bewys nie

- Dit is **nie** 'n kwantumrekenaar wat berekeninge doen nie. Dit is 'n kwantum-geheue: dit stoor en beskerm een logiese qubit. Dit voer geen logiese bewerkinge (hekke) tussen logiese qubits uit nie en laat geen algoritme loop nie.
- Dit is **nie** een qubit weg van nuttige masjiene nie. 'n Logiese afstand-7-qubit bestee omtrent 101 fisiese qubits; 'n foutkoers van 0,1% per siklus is steeds ver bo die ruweg 10^{-6} tot 10^{-10} wat werklike algoritmes nodig het. Om daardie gaping toe te maak beteken om na veel groter afstande te stoot — baie meer fisiese qubits per logiese qubit — en nuttige algoritmes het *duisende* logiese qubits tegelyk nodig. Die begroting aan fisiese qubits daarvoor loop in die miljoene.
- Die **“as dit geskaleer word” doen werklike werk.** Die artikel se eie slotsom is dat die toestel se werkverrigting, *as dit geskaleer word*, aan die vereistes van groot algoritmes kan voldoen. Om te wys die tendens is reg op een logiese qubit is nie dieselfde as om die geskaleerde masjien gebou te het nie, en niks hier waarborg dat die tendens tot by veel groter groottes oorleef nie.
- Die **onverklaarde foutvloer is 'n lewende probleem.** Die gekorreleerde uitbarstings wat die herhalingskode-werkverrigting begrens, is, in die outeurs se woorde, ordes van grootte groter as verwag en sou groter fouttolerante toepassings uitsluit totdat dit verstaan word — 'n oop gebrek, reguit gestel, nie 'n afgehandelde besonderheid nie.
- Dit sê niks oor **die breek van enkripsie of “kwantumoppergesag” vir nuttige take nie.** Daarvoor is die volle fouttolerante masjien nodig waarvoor dit 'n hoeksteen is — nie 'n demonstrasie daarvan

nie.

Hoe sterk is die getuienis

- **Die kernbewering is stewig en belangrik.** Werking onder die drempel met duidelike eksponensiële onderdrukking oor drie kode-afstande, plus 'n leeftyd verby breakeven en 'n werkende in-tydse dekodeerder — presies die kombinasie wat die veld probeer bereik het, en dit word direk gedemonstreer eerder as afgelei. Dit is nie 'n hype-artefak nie; dit is 'n egte ingenieursresultaat van 'n voorste groep.
- **Die outeurs is versigtig oor die omvang.** Hulle raam dit as *geheue* onder die drempel, merk self die onverklaarde gekorreleerde foutvloer aan, en hang die toekoms aan daardie opvallende “as dit geskaleer word”. Die oordrywing, waar dit verskyn, sit in die omringende dekking wat “'n foutgecorrigeerde geheue-qubit het beter geword soos dit gegroei het” afrond na “kwantumrekenaars is hier”.
- **Die eerlike status is 'n grondslagtree, netjies gegee.** Een logiese qubit, goed genoeg beskerm dat oortolligheid uiteindelik help — met 'n lang, harde en nog-nie-gewaarborgde pad van skalering, logiese hekke en onverklaarde foute wat voorlê.

Waarom dit saak maak

Fouttolerante kwantumrekenaars het nog altyd 'n hoender-en-eier-gevoel gehad: die masjiene wat nuttig sou wees, het foutkoerse nodig wat geen fisiese qubit kan bereik nie, en die oplossing — foutkorreksie — werk net as die hardeware reeds goed genoeg is om onder die drempel te wees. Om daardie lyn oor te steek, al is dit een keer, al is dit met 'n enkele logiese qubit, verander die vraag van “is dit hoegenaamd moontlik?” na “hoe ver kan dit geskaleer word, en hoe vinnig?”. Dit is 'n werklike, betekenisvolle verskuiwing, en daarom verdien die resultaat aandag.

Maar dieselfde sorg wat die resultaat vertrouenswaardig maak, behoort die storie daaromheen te temper. Dit is die eerste steen van 'n fondament, goed gelê. Dit is nie die gebou nie, en die mense wat dit gelê het, is die eerstes om dit te sê. Die regte manier om kwantumrekenaars oor die volgende paar jaar te volg, is presies hierdie onglansryke kromme: of die onderdrukkingsfaktor hou soos die kodes groei, of logiese hekke so betroubaar uitgevoer kan word soos logiese geheue, en of daardie geheimsinnige een-keer-per-uur-fout ooit verklaar word.

Kortliks

Google Quantum AI het vir die eerste keer duidelik gewys dat 'n surface-code-kwantumgeheue *onder die drempel* kan werk: soos hulle die kode van afstand 3 deur 5 na 7 laat groei het, het die logiese foutkoers eksponensieel gedaal (met omtrent $2,14 \times$ per twee treë), en die grootste, 'n afstand-7-geheue

met 101 qubits, het sy beste fisiese qubit oorleef — beyond breakeven — terwyl die foutkorreksie intyds geloop het. Dit is 'n egte, lank gesoekte mylpaal in die ingenieurswese van kwantumrekenaars. Dit is óók 'n enkele logiese qubit wat as geheue dien, met 'n foutkoers steeds ver van wat werklike algoritmes vereis, geen uitgevoerde logiese bewerkings nie, 'n onverklaarde foutvloer wat die outeurs self aanmerk, en 'n skaleringspad van baie ordes van grootte wat voorlê. 'n Egte drempel oorgesteek — nie 'n kwantumrekenaar afgelewer nie.

Bronne

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>