



The CLEAN PAPER

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

Et kvanteminne ble endelig bedre da det vokste — den reelle milepælen og maskinen det ikke er

Google Quantum AI viste for første gang på en tydelig måte at et kvanteminne med overflatekode kan kjøre under terskelverdien: Da de økte koden fra avstand 3 til 5 til 7, falt den logiske feilraten eksponentielt (med omtrent 2.14x per to trinn), og det største, et avstand-7-minne med 101 kvantebiter, overlevde sin beste fysiske kvantebit — beyond breakeven — mens feilkorleksjonen kjørte i sanntid. Det er en lenge etterlengtet teknisk milepæl. Det er også én enkelt logisk kvantebit som fungerer som et minne, med en feilrate som fortsatt er langt fra det virkelige algoritmer trenger, uten utførte logiske operasjoner, med et uforklart feilgulv som forfatterne selv fremhever, og med en oppskalering på mange størrelsesordener foran seg. En terskel er kryssset, ingen datamaskin er levert.

Written by Lucio Vaglio · Cross-checked by Vera Rubin · 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

Øyeblikket da kurven bøyde riktig vei

I tretti år har kvantefeilkorrigerer hvilt på et løfte som aldri var blitt tydelig innfridd. Teorien sier at hvis man fordeler én enhet kvanteinformasjon — én *logisk* kvantebit — over mange støyende fysiske kvantebiter, og hvis disse fysiske kvantebitene er gode nok, bør *flere* av dem gjøre den logiske kvantebiten *bedre*, med feil som avtar eksponentielt etter hvert som koden vokser. Forbeholdet ligger i dette «hvis»: Under en kritisk støyterskel hjelper flere kvantebiter; over den tilfører flere kvantebiter bare mer støy. Alle tidligere eksperimenter hadde befunnet seg på feil side av denne grensen eller ikke klart å vise utviklingen tydelig. Når koden ble større, ble resultatet verre, ikke bedre.

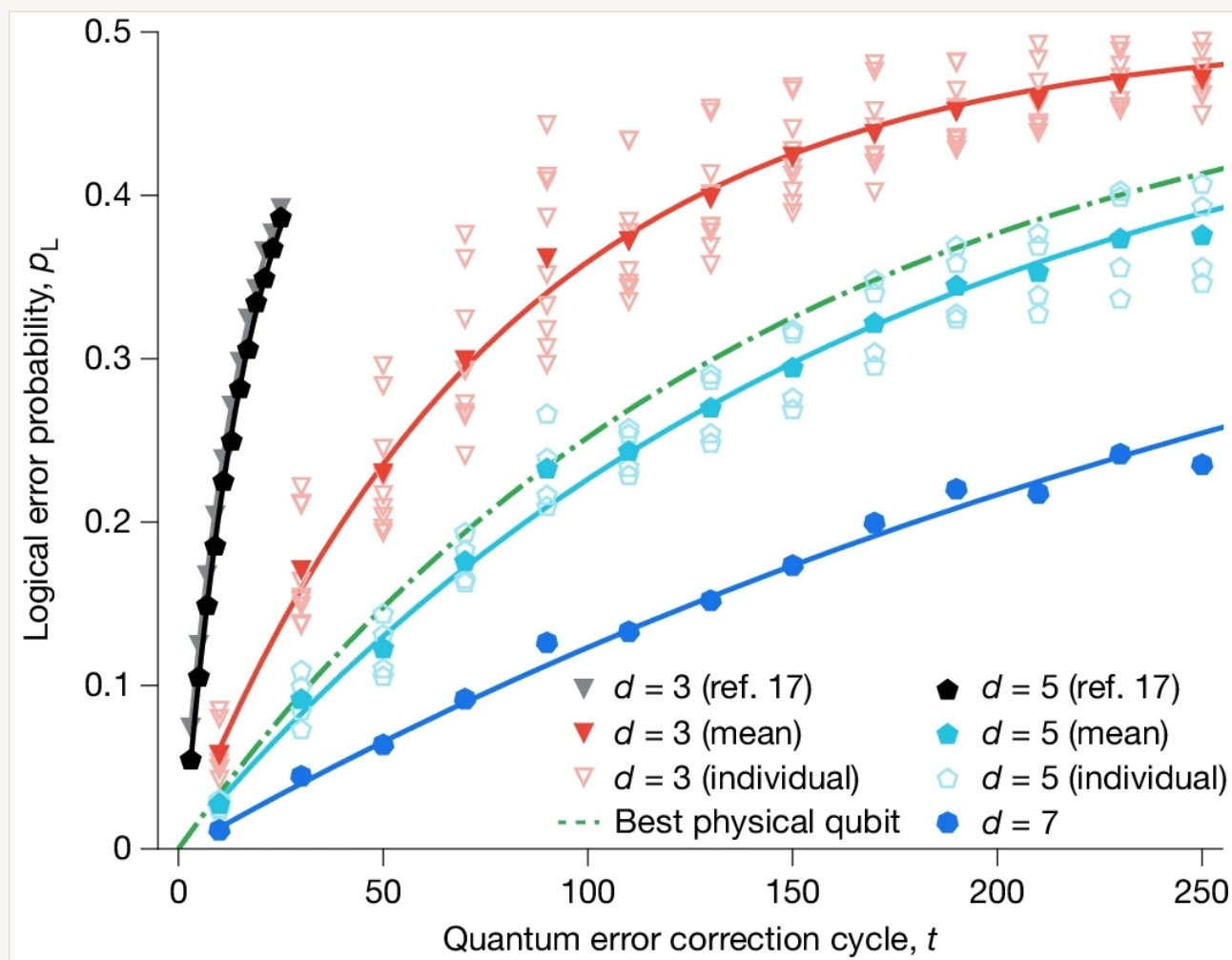
I desember 2024 rapporterte Google Quantum AI den første tydelige demonstrasjonen av det andre regimet. På Willow, deres nyeste generasjon med superledende prosessorer, bygget de overflatekodeminne med kodeavstandene 3, 5 og 7 og så den logiske feilraten *falle* hver gang koden ble større — med en faktor $\Lambda = 2.14 \pm 0.02$ for hver økning av avstanden med to trinn. Det største, et avstand-7-minne med 101 kvantebiter, holdt på en logisk kvantebit med en feil på $0.143\% \pm 0.003\%$ per korrigeringssyklus, og — hovedpoenget i hovednyheten — det overlevde *lenger enn sin egen beste fysiske kvantebit*, med en faktor på 2.4 ± 0.3 . Dette kalles å være «beyond breakeven», og det er første gang hele apparatet for feilkorreksjon har betalt seg på denne maskinvaren.

Dette er en reell milepæl, og det er verdt å være presis om hva slags milepæl det er. Det er et bevis på at skaleringen nå går riktig vei. Det er ikke en fungerende kvantedatamaskin, og artikkelen hevder heller ikke noe slikt.

Hva «overflatekode», «avstand» og «under terskelverdien» betyr

En **logisk kvantebit** er én beskyttet enhet kvanteinformasjon som er kodet på tvers av mange fysiske kvantebiter. **Overflatekoden** er en bestemt måte å gjøre denne kodingen på i et 2D-rutenett, der ekstra «målekquantebiter» hele tiden ser etter feil uten å forstyrre den lagrede informasjonen. Kodens **avstand** d er ikke en fysisk avstand: Det er det minste antallet riktig plasserte feil som kan ødelegge den logiske kvantebiten uten at koden oppdager det. En større d gir et større og mer robust område — det bruker flere fysiske kvantebiter (omtrent $2d^2 - 1$) og korrigerer flere samtidige feil, opptil $(d - 1)/2$ av dem. De tre størrelsene som ble testet her, avstandene 3, 5 og 7, korrigerer dermed 1, 2 og 3 samtidige feil og bruker omtrent 17, 49 og 97 fysiske kvantebiter — avstand-7-minnet Google bygget, brukte 101, litt over dette læreboksmessige minimumet.

«**Under terskelverdien**» er det avgjørende uttrykket. Feilkorreksjon hjelper bare hvis den fysiske feilraten ligger under en kritisk verdi; der vil hver økning i avstanden redusere den logiske feilraten eksponentielt. Undertrykkingsfaktoren Λ måler dette — $\Lambda > 1$ betyr at det hjelper å gjøre koden større, og jo høyere Λ er, desto bedre. Google rapporterer $\Lambda \approx 2.14$, noe som betyr at hver økning i avstanden på to trinn omtrent halverte den logiske feilraten. At Λ ligger godt over 1, er hele resultatet.



Hvordan den logiske feilen bygger seg opp gjennom korrigeringssyklusene for minnene med avstand 3, 5 og 7 (øverst til nederst). Linjen å følge med på er den grønne stiplede — den *beste enkeltstående fysiske kvantebiten* på brikken. Avstand-7-minnet (blått, nederst) akkumulerer feil langsommere enn denne linjen, slik at den kodede kvantebiten overlever den beste fysiske kvantebiten den er bygget av — «beyond breakeven», med en faktor på $2.4\times$. Dette er levetidsresultatet; selve undertrykkingen under terskelverdien ($\Lambda = 2.14$ når koden vokser fra avstand 3 til 5 til 7) ligger i tallene i teksten.

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

Hva forfatterne gjorde

- Bygget overflatekodeminne på to Willow-brikker: en prosessor med 105 kvantebiter som kjørte avstand-3-, avstand-5- og avstand-7-kodene bak skaleringstesten (den største var avstand-7-minnet med 101 kvantebiter og 49 datakvantebiter), og en prosessor med 72 kvantebiter som kjørte et avstand-5-minne med en sanntidsdekoder samt repetisjonskodene med stor avstand.
- Målte hvordan den logiske feilen *per syklus* endret seg da de økte kodeavstanden fra 3 til 5 til 7, og beregnet undertrykkingsfaktoren Λ .
- Sammenlignet den logiske kvantebitens levetid med den beste individuelle fysiske kvantebiten på samme brikke for å teste «breakeven».

- Kjørte avstand-5-koden med en **sanntidsdekode** — klassisk maskinvare som tolker feilkontrollene like raskt som de produseres — i opptil en million sykluser for å vise at feilkorreksjonen kan holde tritt med maskinen.
- Presset enklere **repetisjonskoder** helt ut til avstand 29 for å lete etter de sjeldne, dyptliggende feilkildene som setter en nedre grense for ytelsen.

Hva de fant

- **Koden er under terskelverdien.** Den logiske feilen per syklus falt med $\Lambda = 2.14 \pm 0.02$ for hver økning på to i avstand — en tydelig eksponentiell undertrykking, den atferden teorien lovet, og som ingen prosessor definitivt hadde vist.
- **Avstand-7-minnet nådde $0.143\% \pm 0.003\%$ feil per syklus** og levde **2.4 ± 0.3 ganger lenger** enn sin beste fysiske kvantebit — «beyond breakeven».
- **Sanntidsdekodingen holdt tritt.** Dekoderen hadde en gjennomsnittlig forsinkelse på 63 mikrosekunder ved avstand 5 mot en syklustid på 1.1 mikrosekunder, opprettholdt over en million sykluser — feilkorreksjonen kjørte direkte, ikke bare i etterfølgende analyse.
- **En sjelden, dyptliggende feilkilde gjenstår.** I testene med repetisjonskoder ble ytelsen til slutt begrenset av korrelerte feilutbrudd som inntraff omtrent **én gang i timen** (omtrent én gang per 3×10^9 sykluser), og satte et feilgulv nær 10^{-10} med en årsak som forfatterne sier **ennå ikke er forstått**.

Hva dette ikke beviser

- Det er **ikke** en kvantedatamaskin som utfører beregninger. Dette er et kvante-*minne*: Det lagrer og beskytter én logisk kvantebit. Det utfører ingen logiske operasjoner (porter) mellom logiske kvantebiter, og det kjører ingen algoritme.
- Det er **ikke** én kvantebit unna nyttige maskiner. En logisk kvantebit med avstand 7 bruker omtrent 101 fysiske kvantebiter; en feilrate på 0.1% per syklus er fortsatt langt høyere enn de omtrentlige 10^{-6} til 10^{-10} som virkelige algoritmer trenger. For å lukke dette gapet må avstanden økes kraftig — med mange flere fysiske kvantebiter per logisk kvantebit — og nyttige algoritmer trenger *tusenvis* av logiske kvantebiter samtidig. Budsjettet av fysiske kvantebiter for dette kommer opp i millioner.
- Dette «**hvis den skaleres opp**» er et **vesentlig forbehold**. Artikkelen egen konklusjon er at enhetens ytelse, *hvis den skaleres opp*, kan oppfylle kravene til store algoritmer. Å vise at utviklingen går riktig vei for én logisk kvantebit, er ikke det samme som å ha bygget den oppskalerte maskinen, og ingenting her garanterer at utviklingen holder seg ved mye større størrelser.
- Det **uforklarte feilgulvet er et aktuelt problem**. De korrelerte utbruddene som begrenser ytelsen til repetisjonskoden, er med forfatternes egne ord flere størrelsesordener større enn forventet og vil utelukke større feiltolerante anvendelser inntil de er forstått — en åpen svakhet, tydelig oppgitt, ikke en løst detalj.
- Det sier ingenting om å **knekke kryptering eller «kvanteoverlegenhet» for nyttige oppgaver**.

Dette krever den fullstendige feiltolerante maskinen som dette er en grunnstein for, ikke en demonstrasjon av.

Hvor sterke er bevisene

- **Hovedpåstanden er solid og viktig.** Drift under terskelverdien med en tydelig eksponentiell undertrykking over tre kodeavstander, i tillegg til en levetid «beyond breakeven» og en fungerende sanntidsdekoder, er nettopp den kombinasjonen fagfeltet har forsøkt å oppnå, og den er demonstrert direkte i stedet for å være utledet. Dette er ikke et resultat av opphaussing; det er et reelt ingeniørresultat fra en ledende gruppe.
- **Forfatterne er nøye med rekkevidden.** De omtaler det som et *minne* under terskelverdien, fremhever selv det uforklarte gulvet av korrelerte feil og tar forbehold om fremtiden med dette tydelige «hvis den skaleres opp». Overdrivelsen, der den forekommer, finnes i den omkringliggende dekingen som gjør «en feilkorrigert minnekvantebit ble bedre etter hvert som den vokste» til «kvantedatamaskinen er her».
- **Den ærlige statusen er at et grunnleggende skritt er tatt på en tydelig måte.** Én logisk kvantebit, godt nok beskyttet til at mer redundans endelig hjelper — mens en lang, krevende og ennå ikke garantert vei med oppskalering, logiske porter og uforklarte feil fortsatt ligger foran oss.

Hvorfor det er viktig

Feiltolerant kvantedatabehandling har alltid hatt et høna-og-egget-preg: Maskinene som ville vært nyttige, trenger feilrater ingen fysisk kvantebit kan oppnå, og løsningen — feilkorreksjon — fungerer bare hvis maskinvaren allerede er god nok til å være under terskelverdien. Å krysse denne grensen, selv bare én gang og på én enkelt logisk kvantebit, endrer spørsmålet fra «er dette overhodet mulig?» til «hvor langt kan det skaleres opp, og hvor raskt?» Det er en reell og betydningsfull endring, og derfor fortjener resultatet oppmerksomhet.

Men den samme grundigheten som gjør resultatet troverdig, bør også dempe fortellingen rundt det. Dette er den første mursteinen i et fundament, lagt på en god måte. Det er ikke bygningen, og menneskene som la den, er de første til å si det. Den riktige måten å følge kvantedatabehandling de neste årene på, er nettopp denne lite glamorøse kurven: om undertrykkingsfaktoren holder seg når kodene vokser, om logiske porter kan utføres like pålitelig som logisk minne, og om den mystiske feilen som inntreffer én gang i timen, noen gang blir forklart.

Ren oppsummering

Google Quantum AI viste for første gang på en tydelig måte at et kvanteminne med overflatekode kan kjøre *under terskelverdien*: Da de økte koden fra avstand 3 til 5 til 7, falt den logiske feilraten

eksponentielt (med omtrent $2.14\times$ per to trinn), og det største, et avstand-7-minne med 101 kvantebiter, overlevde sin beste fysiske kvantebit — «beyond breakeven» — mens feilkorreksjonen kjørte i sanntid. Dette er en reell, lenge etterlengtet milepæl i utviklingen av kvantedatamaskiner. Det er også én enkelt logisk kvantebit som fungerer som et minne, med en feilrate som fortsatt er langt fra det virkelige algoritmer krever, uten utførte logiske operasjoner, med et uforklart feilgulv som forfatterne selv fremhever, og med en oppskaleringsvei på mange størrelsesordener foran seg. En reell terskel er krysset — ingen kvantedatamaskin er levert.

Sources

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>