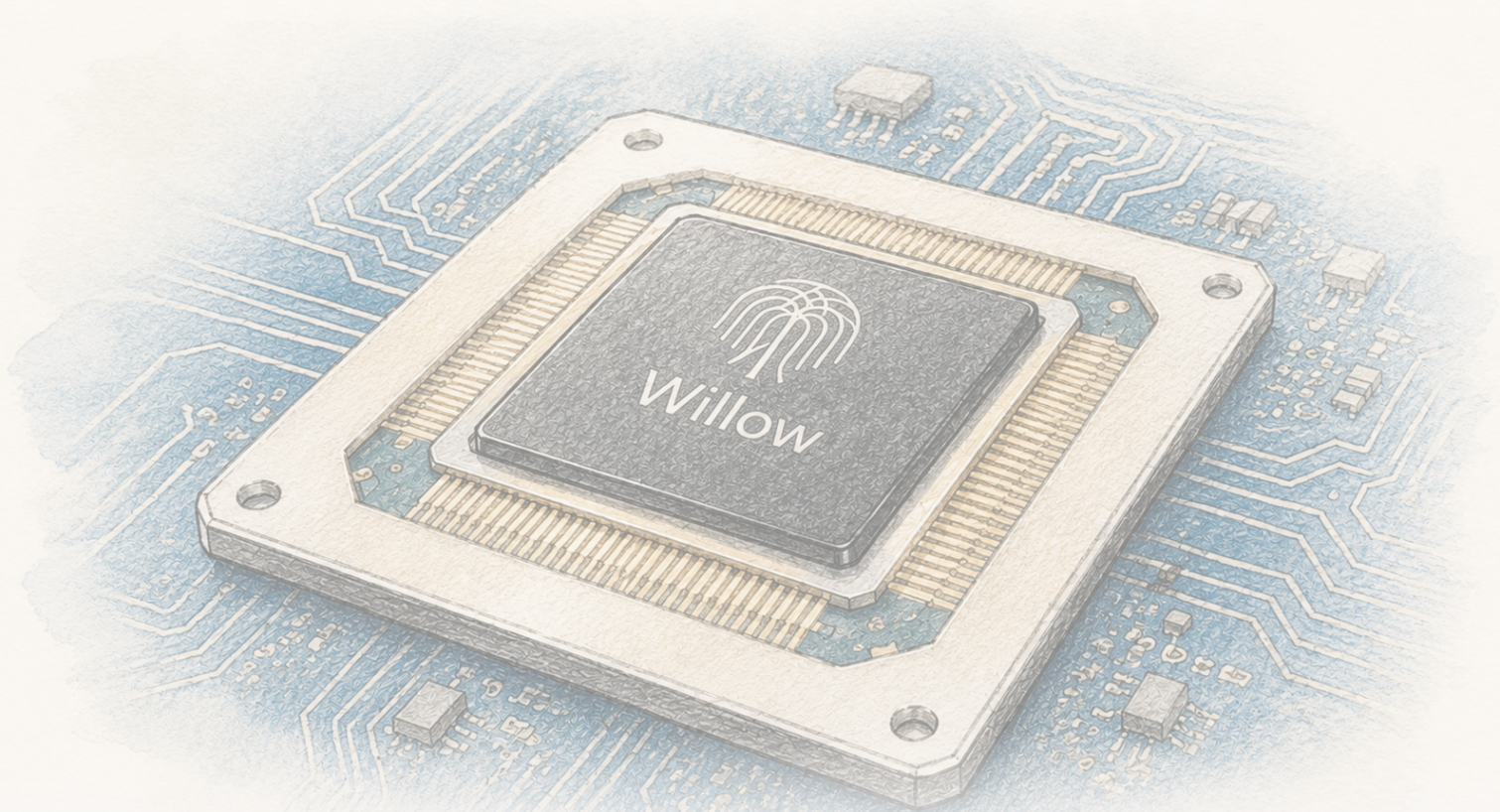




The CLEAN PAPER

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

En kvantehukommelse blev endelig bedre, da den voksede — den reelle milepæl og den maskine, den ikke er

Google Quantum AI viste for første gang klart, at en kvantehukommelse med overfladecode kan køre under tærskelværdien: Da de øgede koden fra afstand 3 til 5 til 7, faldt den logiske fejlråde eksponentielt (med omkring 2.14x per to trin), og den største, en hukommelse med 101 qubits og afstand 7, overlevede sin bedste fysiske qubit — beyond breakeven — mens fejlkorraktionen kørte i realtid. Det er en længe eftertragtet teknisk milepæl. Det er også én enkelt logisk qubit, der fungerer som hukommelse, med en fejlråde, der stadig er langt fra, hvad virkelige algoritmer har brug for, uden udførte logiske operationer, med et uforklaret fejlgulv, som forfatterne selv fremhæver, og med en opskalering på mange størrelsesordener foran sig. En tærskel er krydset, ingen computer er leveret.

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

Øjeblikket, hvor kurven bøjede den rigtige vej

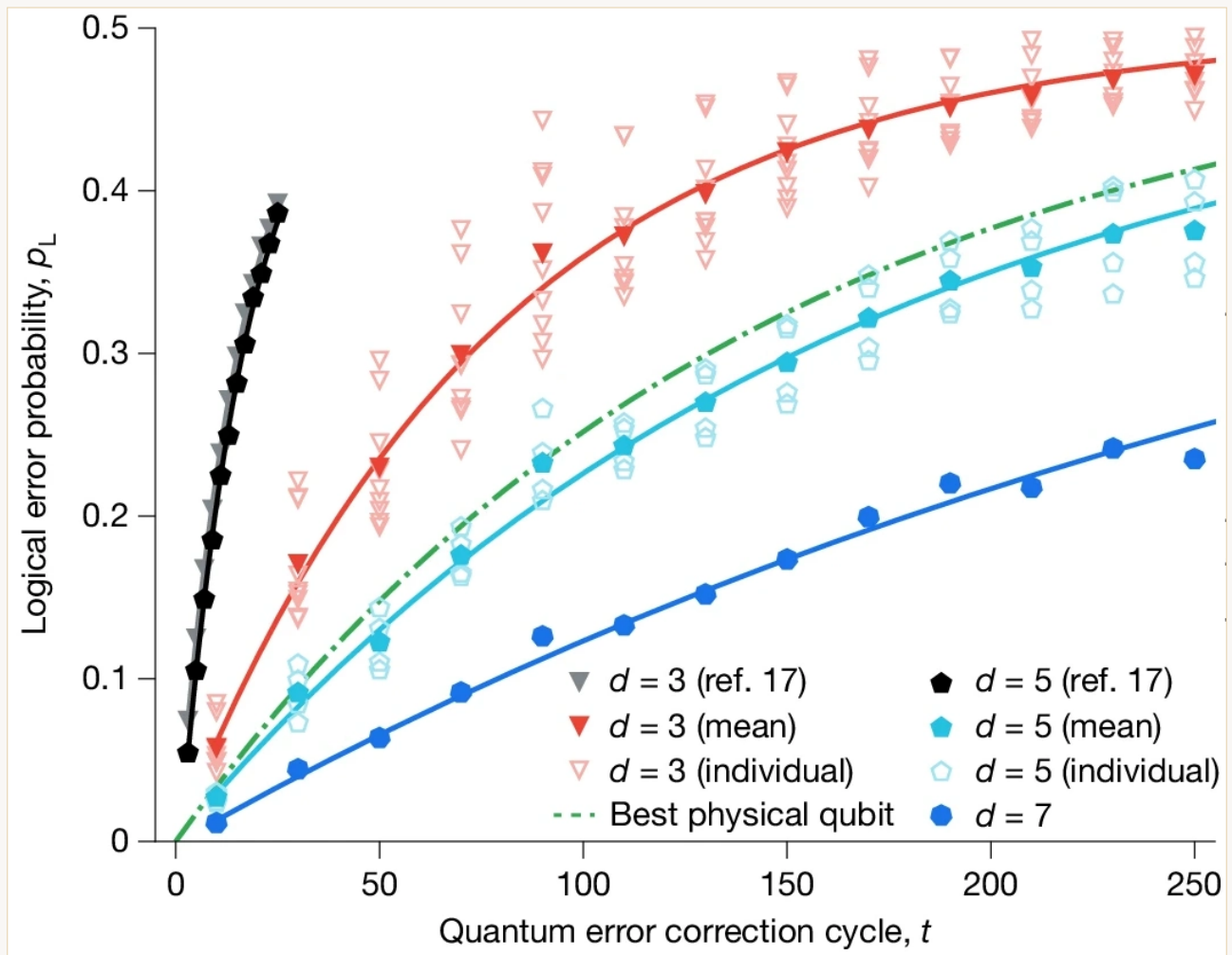
I tredive år har kvantefejlkorrektur hvilet på et løfte, som aldrig var blevet indfriet entydigt. Teorien siger, at hvis man fordeler én enhed kvanteinformation — én *logisk* qubit — over mange støjfyldte fysiske qubits, og hvis disse fysiske qubits er gode nok, så bør *flere* af dem gøre den logiske qubit *bedre*, idet fejlene falder eksponentielt, efterhånden som koden vokser. Hagen er dette ”hvis”: Under en kritisk støjtærskel hjælper flere qubits; over den tilføjer flere qubits blot mere støj. Alle tidligere eksperimenter havde befundet sig på den forkerte side af denne grænse eller havde ikke formået at vise udviklingen klart. Når koden voksede, blev resultatet værre, ikke bedre.

I december 2024 rapporterede Google Quantum AI den første klare demonstration af det andet regime. På Willow, deres nyeste generation af superledende processorer, byggede de kvantehukommelser med overfladekode ved kodeafstandene 3, 5 og 7 og så den logiske fejlrate *falde*, hver gang koden blev større — med en faktor $\Lambda = 2.14 \pm 0.02$ for hver forøgelse af afstanden med to trin. Den største, en hukommelse med 101 qubits og afstand 7, lagrede en logisk qubit med en fejl på $0.143\% \pm 0.003\%$ per korrektionscyklus, og — hovedpointen i hovednyheden — den overlevede *længere end sin egen bedste fysiske qubit*, med en faktor på 2.4 ± 0.3 . Dette kaldes at være ”beyond breakeven”, og det er første gang, hele apparatet til fejlkorrektur har tjent sig hjem på denne hardware.

Dette er en reel milepæl, og det er værd at være præcis om, hvilken slags milepæl det er. Det er et bevis på, at skaleringen nu går den rigtige vej. Det er ikke en fungerende kvantecomputer, og artiklen fremsætter heller ikke en sådan påstand.

Hvad ”overfladekode”, ”afstand” og ”under tærskelværdien” betyder

En **logisk qubit** er én beskyttet enhed kvanteinformation, som er kodet på tværs af mange fysiske qubits. **Overfladekoden** er en bestemt måde at udføre denne kodning på i et 2D-gitter, hvor ekstra ”målequbits” konstant leder efter fejl uden at forstyrre den lagrede information. Kodens **afstand** d er ikke en fysisk afstand: Det er det mindste antal velplacerede fejl, som kan ødelægge den logiske qubit, uden at koden opdager det. En større d giver et større og mere robust område — det kræver flere fysiske qubits (omtrent $2d^2 - 1$) og korrigerer flere samtidige fejl, op til $(d - 1)/2$ af dem. De tre størrelser, som blev testet her, afstandene 3, 5 og 7, korrigerer således 1, 2 og 3 samtidige fejl og kræver omtrent 17, 49 og 97 fysiske qubits — den hukommelse med afstand 7, som Google byggede, brugte 101, lidt over dette lærebogsmæssige minimum. ”**Under tærskelværdien**” er det afgørende udtryk. Fejlkorrektur hjælper kun, hvis den fysiske fejlrate ligger under en kritisk værdi; dér reducerer hver forøgelse af afstanden den logiske fejlrate eksponentielt. Undertrykkelsesfaktoren Λ måler dette — $\Lambda > 1$ betyder, at det hjælper at gøre koden større, og jo højere Λ er, desto bedre. Google rapporterer $\Lambda \approx 2.14$, hvilket betyder, at hver forøgelse af afstanden med to trin omtrent halverede den logiske fejlrate. At Λ ligger klart over 1, er hele resultatet.



Hvordan den logiske fejl bygger sig op gennem korrektionscyklusserne for hukommelserne med afstand 3, 5 og 7 (oppefra og ned). Linjen, man skal holde øje med, er den grønne stiplede — chippens *bedste enkelte fysiske qubit*. Hukommelsen med afstand 7 (blå, nederst) akkumulerer fejl langsommere end denne linje, så den kodede qubit overlever den bedste fysiske qubit, den er bygget af — "beyond breakeven", med en faktor på $2.4\times$. Dette er levetidsresultatet; selve undertrykkelsen under tærskelværdien ($\Lambda = 2.14$, når koden vokser fra afstand 3 til 5 til 7) findes i tallene i teksten.

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

Hvad forfatterne gjorde

- Byggede kvantehukommelser med overfladekode på to Willow-chips: en processor med 105 qubits, som kørte koderne med afstand 3, 5 og 7 bag skaleringstesten (den største var hukommelsen med afstand 7, 101 qubits og 49 dataqubits), og en processor med 72 qubits, som kørte en hukommelse med afstand 5 med en realtidsdekoder samt repetitionskoderne med stor afstand.
- Målte, hvordan den logiske fejl *per cyklus* ændrede sig, da de øgede kodeafstanden fra 3 til 5 til 7, og udledte undertrykkelsesfaktoren Λ .
- Sammenlignede den logiske qubits levetid med den bedste individuelle fysiske qubit på samme

chip for at teste ”breakeven”.

- Kørte koden med afstand 5 med en **realidsdekoder** — klassisk hardware, som fortolker fejlkorrigerne lige så hurtigt, som de produceres — i op til en million cyklusser for at vise, at fejlkorrektionen kan holde trit med maskinen.
- Pressede enklere **repetitions-koder** helt ud til afstand 29 for at lede efter de sjældne, dybtliggende fejlkilder, som sætter en nedre grænse for ydeevnen.

Hvad de fandt

- **Koden er under tærskelværdien.** Den logiske fejl per cyklus faldt med $\Lambda = 2.14 \pm 0.02$ for hver forøgelse af afstanden med to — en klar eksponentiel undertrykkelse, den adfærd, teorien lovede, og som ingen processor definitivt havde vist.
- **Hukommelsen med afstand 7 nåede $0.143\% \pm 0.003\%$ fejl per cyklus og levede 2.4 ± 0.3 gange længere** end sin bedste fysiske qubit — ”beyond breakeven”.
- **Realidsafkodningen holdt trit.** Dekoderen havde en gennemsnitlig latenstid på 63 mikrosekunder ved afstand 5 mod en cyklostid på 1.1 mikrosekunder, opretholdt over en million cyklusser — fejlkorrektionen kørte direkte, ikke blot i efterfølgende analyse.
- **En sjælden, dybtliggende fejlkilde er stadig tilbage.** I testene med repetitionskoder blev ydeevnen til sidst begrænset af korrelerede fejludbrud, som indtraf omtrent **én gang i timen** (omkring én gang per 3×10^9 cyklusser), og satte et fejlgulv nær 10^{-10} med en årsag, som forfatterne siger **endnu ikke er forstået**.

Hvad dette ikke beviser

- Det er **ikke** en kvantecomputer, der udfører beregninger. Dette er en kvante-*hukommelse*: Den lagrer og beskytter én logisk qubit. Den udfører ikke logiske operationer (gates) mellem logiske qubits og kører ingen algoritme.
- Den er **ikke** én qubit fra nyttige maskiner. En logisk qubit med afstand 7 kræver omtrent 101 fysiske qubits; en fejlrate på 0.1% per cyklus ligger stadig langt over de omtrentlige 10^{-6} til 10^{-10} , som virkelige algoritmer har brug for. At lukke dette gab kræver langt større afstande — mange flere fysiske qubits per logisk qubit — og nyttige algoritmer kræver *tusindvis* af logiske qubits på én gang. Budgettet af fysiske qubits til dette løber op i millioner.
- Dette **”hvis den skaleres op” er et væsentligt forbehold**. Artiklens egen konklusion er, at enhedens ydeevne, *hvis den skaleres op*, vil kunne opfylde kravene til store algoritmer. At vise, at udviklingen går den rigtige vej for én logisk qubit, er ikke det samme som at have bygget den opskalerede maskine, og intet her garanterer, at udviklingen holder ved langt større størrelser.
- Det **uforklarede fejlgulv er et aktuelt problem**. De korrelerede udbrud, som begrænser repetitionskodens ydeevne, er med forfatternes egne ord flere størrelsesordener større end forventet og vil umuliggøre større fejltolerante anvendelser, indtil de er forstået — en åben svaghed, der oplyses klart, ikke en løst detalje.
- Det siger intet om **at bryde kryptering eller ”kvanteoverlegenhed” til nyttige opgaver**. Det

kræver den fulde fejltolerante maskine, som dette er en grundsten for, ikke en demonstration af.

Hvor stærk er dokumentationen

- **Hovedpåstanden er solid og vigtig.** Drift under tærskelværdien med en klar eksponentiel undertrykkelse på tværs af tre kodeafstande, plus en levetid "beyond breakeven" og en fungerende realtidsdekoder, er præcis den kombination, feltet har forsøgt at nå, og den demonstreres direkte frem for at blive udledt. Dette er ikke et produkt af hype; det er et reelt ingeniørresultat fra en førende gruppe.
- **Forfatterne er omhyggelige med resultatets rækkevidde.** De beskriver det som en *hukommelse* under tærskelværdien, fremhæver selv det uforklarede gulv af korrelerede fejl og tager forbehold om fremtiden med dette tydelige "hvis den skales op". Overdrivelsen, hvor den forekommer, findes i den omgivende dækning, som gør "en fejlkorrigeret hukommelsesqubit blev bedre, da den voksede" til "kvantecomputeren er her".
- **Den ærlige status er, at et grundlæggende skridt er taget på overbevisende vis.** Én logisk qubit, beskyttet godt nok til, at yderligere redundans endelig hjælper — mens en lang, vanskelig og endnu ikke garanteret vej med opskalering, logiske gates og uforklarede fejl stadig ligger foran os.

Hvorfor det er vigtigt

Fejltolerant kvanteberegning har altid haft karakter af et hønen-og-ægget-problem: De maskiner, der ville være nyttige, har brug for fejlrater, som ingen fysisk qubit kan nå, og løsningen — fejlkorrektion — virker kun, hvis hardwaren allerede er god nok til at ligge under tærskelværdien. At krydse denne grænse, selv blot én gang og på blot én logisk qubit, ændrer spørgsmålet fra "er dette overhovedet muligt?" til "hvor langt kan det skales op, og hvor hurtigt?" Det er en reel og betydningsfuld ændring, og derfor fortjener resultatet opmærksomhed.

Men den samme omhu, som gør resultatet troværdigt, bør også dæmpe fortællingen omkring det. Dette er den første mursten i et fundament, lagt ordentligt. Det er ikke bygningen, og de mennesker, der lagde den, er de første til at sige det. Den rigtige måde at følge kvanteberegning på i de næste par år er netop denne uglamourøse kurve: om undertrykkelsesfaktoren holder, når koderne vokser, om logiske gates kan udføres lige så pålideligt som logisk hukommelse, og om den mystiske fejl én gang i timen nogensinde bliver forklaret.

Kort fortalt

Google Quantum AI viste for første gang klart, at en kvantehukommelse med overfladecode kan køre *under tærskelværdien*: Da de øgede koden fra afstand 3 til 5 til 7, faldt den logiske fejlrater eksponen-

tielt (med omkring $2.14\times$ per to trin), og den største, en hukommelse med 101 qubits og afstand 7, overlevede sin bedste fysiske qubit — ”beyond breakeven” — mens dens fejlkorrektion kørte i realtid. Dette er en reel og længe eftertragtet milepæl i udviklingen af kvantecomputere. Det er også én enkelt logisk qubit, der fungerer som hukommelse, med en fejlrate, der stadig er langt fra, hvad virkelige algoritmer kræver, uden udførte logiske operationer, med et uforklaret fejlgulv, som forfatterne selv fremhæver, og med en opskaleringsvej på mange størrelsesordener foran sig. En reel tærskel er krydset — ingen kvantecomputer er leveret.

Kilder

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>