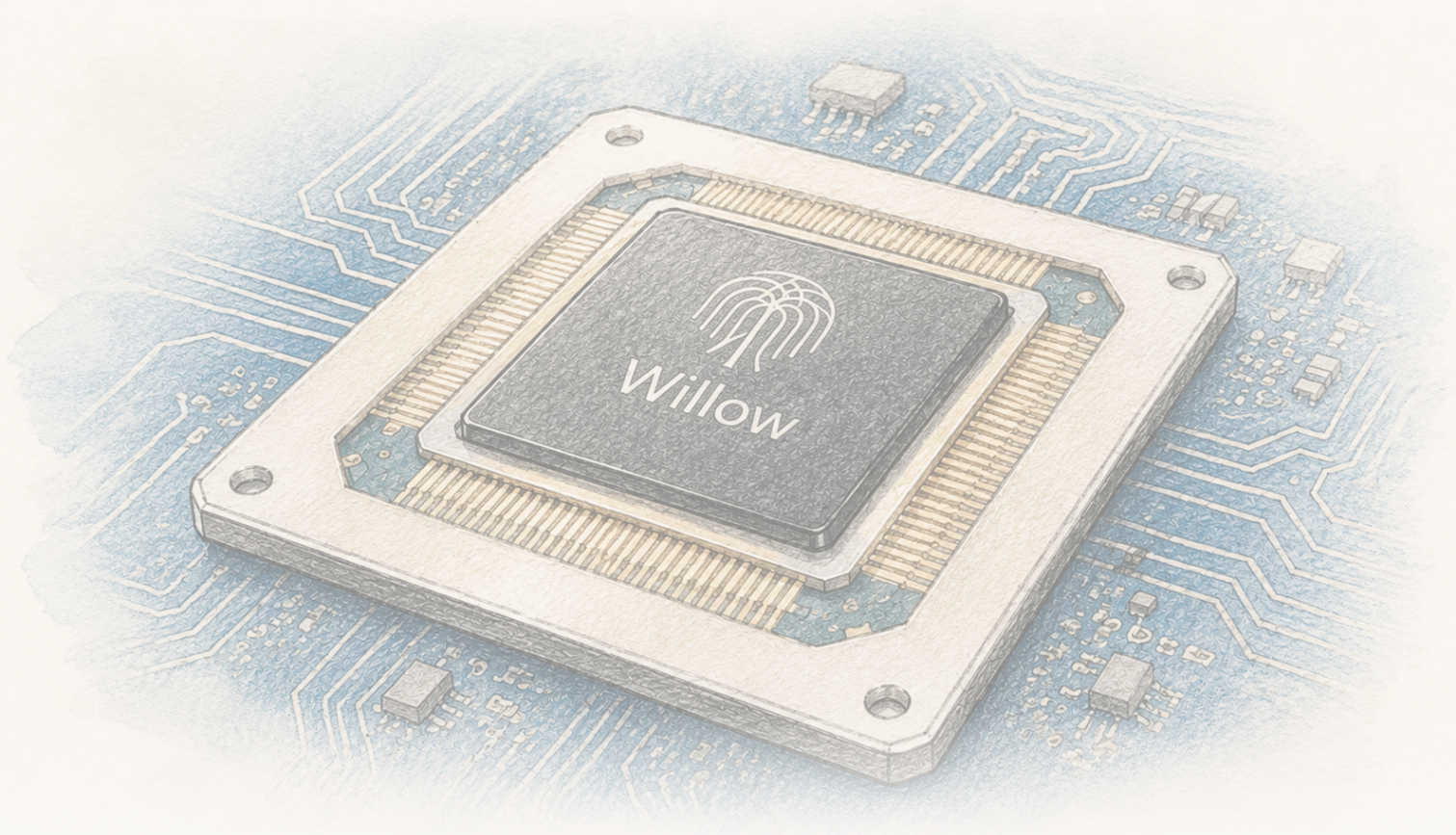




The CLEAN PAPER

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

Ett kvantminne blev till slut bättre när det växte — den verkliga milstolpen och den maskin det inte är

Google Quantum AI visade för första gången på ett entydigt sätt att ett kvantminne med ytkod kan arbeta under tröskelvärdet: när de ökade koden från avstånd 3 till 5 till 7 sjönk den logiska felfrekvensen exponentiellt (med omkring 2.14x för vartannat steg), och det största, ett avstånd-7-minne med 101 kvantbiter, överlevde sin bästa fysiska kvantbit — beyond breakeven — samtidigt som felkorrigeringen kördes i realtid. Det är en sedan länge eftersträvad teknisk milstolpe. Det är också en enda logisk kvantbit som fungerar som minne, med en felfrekvens som fortfarande är långt från vad verkliga algoritmer behöver, utan utförda logiska operationer, med ett oförklarat felgolv som författarna själva uppmärksammar och med en uppskalning på många storleksordningar framför sig. Ett tröskelvärde har passerats, ingen dator har levererats.

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

Ögonblicket då kurvan böjde åt rätt håll

I trettio år har kvantfelskorrigering vilat på ett löfte som aldrig tidigare hade infriats på ett entydigt sätt. Teorin säger att om man fördelar en enhet kvantinformation — en *logisk* kvantbit — över många brusiga fysiska kvantbitar, och om dessa fysiska kvantbitar är tillräckligt bra, bör *fler* av dem göra den logiska kvantbiten *bättre*, med fel som minskar exponentiellt när koden växer. Haken ligger i detta ”om”: under en kritisk bruströskel hjälper fler kvantbitar; över den tillför fler kvantbitar bara mer brus. Alla tidigare experiment hade befunnit sig på fel sida om den gränsen eller inte kunnat visa utvecklingen tydligt. När koden gjordes större blev det sämre, inte bättre.

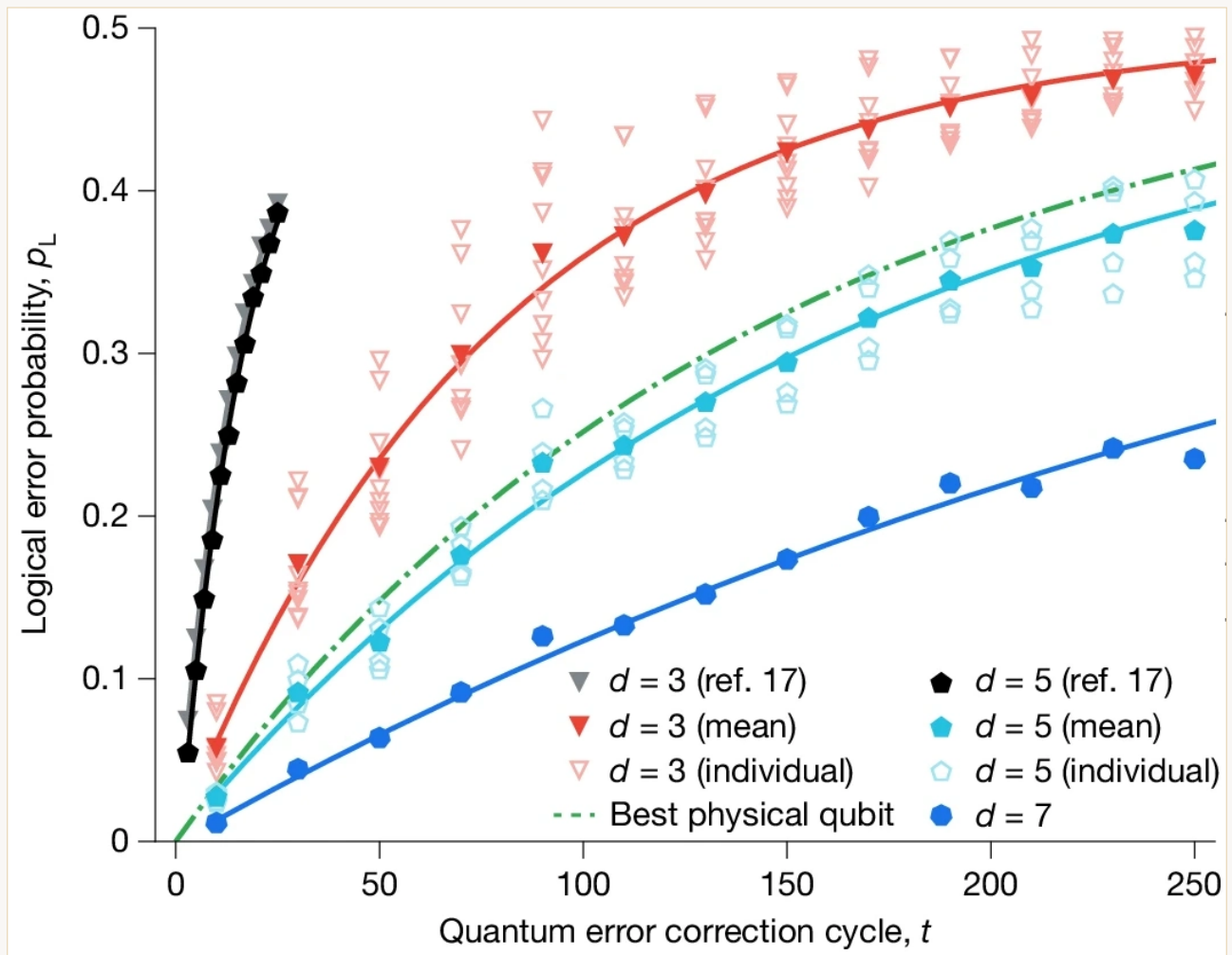
I december 2024 rapporterade Google Quantum AI den första tydliga demonstrationen av den andra regimen. På Willow, deras senaste generation supraledande processorer, byggde de ytkodsminnen med kodavstånden 3, 5 och 7 och såg den logiska felfrekvensen *sjunka* varje gång koden blev större — med en faktor $\Lambda = 2.14 \pm 0.02$ för vartannat steg som avståndet ökade. Det största, ett avstånd-7-minne med 101 kvantbitar, höll en logisk kvantbit med ett fel på $0.143\% \pm 0.003\%$ per korrigeringscykel och — huvudpoängen i huvudnyheten — den överlevde *längre än chipets egen bästa fysiska kvantbit*, med en faktor 2.4 ± 0.3 . Detta kallas att vara ”beyond breakeven”, och det är första gången hela felkorrigeringsapparaten har betalat sig på denna hårdvara.

Detta är en verklig milstolpe, och det är viktigt att vara precis med vilken sorts milstolpe det är. Det är ett bevis för att skalningen nu går åt rätt håll. Det är inte en fungerande kvantdator, och artikeln gör inte heller något sådant påstående.

Vad ”ytkod”, ”avstånd” och ”under tröskelvärde” betyder

En **logisk kvantbit** är en skyddad enhet kvantinformation som kodats över många fysiska kvantbitar. **Ytkoden** är ett särskilt sätt att göra denna kodning i ett 2D-rutnät, där extra ”mätkvantbitar” ständigt söker efter fel utan att störa den lagrade informationen. Kodens **avstånd** d är inte ett fysiskt avstånd: det är det minsta antalet rätt placerade fel som kan förvanska den logiska kvantbiten utan att koden märker det. Ett större d innebär ett större, robustare område — det kräver fler fysiska kvantbitar (ungefär $2d^2 - 1$) och korrigerar fler samtidiga fel, upp till $(d - 1)/2$ av dem. De tre storlekar som testades här, avstånden 3, 5 och 7, korrigerar alltså 1, 2 respektive 3 samtidiga fel och kräver ungefär 17, 49 respektive 97 fysiska kvantbitar — det avstånd-7-minne som Google byggde använde 101, något över detta läroboksmässiga minimum.

”**Under tröskelvärde**” är det avgörande uttrycket. Felkorrigering hjälper bara om den fysiska felfrekvensen ligger under ett kritiskt värde; där minskar varje ökning av avståndet den logiska felfrekvensen exponentiellt. Undertryckningsfaktorn Λ mäter detta — $\Lambda > 1$ betyder att det hjälper att göra koden större, och ju högre Λ , desto bättre. Google rapporterar $\Lambda \approx 2.14$, vilket innebär att varje ökning av avståndet med två steg ungefär halverade den logiska felfrekvensen. Att Λ ligger klart över 1 är hela resultatet.



Hur det logiska felet byggs upp under korrigeringscyklerna för minnena med avstånd 3, 5 och 7 (uppfifrån och ned). Linjen att hålla ögonen på är den gröna streckade — chipets *bästa enskilda fysiska kvantbit*.

Avstånd-7-minnet (blått, längst ned) samlar på sig fel långsammare än den linjen, så den kodade kvantbiten överlever den bästa fysiska kvantbit som den är uppbyggd av — "beyond breakeven", med en faktor $2.4\times$. Detta är livslängdsresultatet; själva undertryckningen under tröskelvärdet ($\Lambda = 2.14$, när koden växer från avstånd 3 till 5 till 7) finns i siffrorna i texten.

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

Vad författarna gjorde

- Byggede ytkodsminnen på två Willow-chip: en processor med 105 kvantbitar som körde avstånd-3-, avstånd-5- och avstånd-7-koderna bakom skalningstestet (den största var avstånd-7-minnet med 101 kvantbitar och 49 datakvantbitar), samt en processor med 72 kvantbitar som körde ett avstånd-5-minne med en realtidsavkodare plus repetitionskoderna med stora avstånd.
- Mätte hur det logiska felet *per cykel* förändrades när de ökade kodavståndet från 3 till 5 till 7 och tog fram undertryckningsfaktorn Λ .
- Jämförde den logiska kvantbitens livslängd med den bästa enskilda fysiska kvantbiten på samma chip för att pröva "breakeven".

- Körde avstånd-5-koden med en **realtidsavkodare** — klassisk hårdvara som tolkar felkontrollerna lika snabbt som de produceras — i upp till en miljon cykler, för att visa att felkorrigeringen kan hålla jämna steg med maskinen.
- Pressade enklare **repetitions-koder** ända till avstånd 29 för att leta efter de sällsynta, djupt ligande felkällor som sätter en nedre gräns för prestandan.

Vad de fann

- **Koden är under tröskelvärdet.** Det logiska felet per cykel minskade med $\Lambda = 2.14 \pm 0.02$ för varje ökning av avståndet med två — en tydlig exponentiell undertryckning, det beteende som teorin utlovade och som ingen processor entydigt hade uppvisat.
- **Avstånd-7-minnet nådde $0.143\% \pm 0.003\%$ fel per cykel** och levde **2.4 ± 0.3 gånger längre** än sin bästa fysiska kvantbit — ”beyond breakeven”.
- **Realtidsavkodningen höll jämna steg.** Avkodarens genomsnittliga latens var 63 mikrosekunder vid avstånd 5, jämfört med en cykeltid på 1.1 mikrosekunder, och detta upprätthölls under en miljon cykler — felkorrigeringen kördes direkt, inte bara i efterhandsanalys.
- **En sällsynt, djupt liggande felkälla återstår.** I testerna med repetitions-koder begränsades prestandan till slut av korrelerade felutbrott som inträffade ungefär **en gång i timmen** (omkring en gång per 3×10^9 cykler), vilket satte ett felgolv nära 10^{-10} vars ursprung enligt författarna **ännu inte är klarlagt**.

Vad detta inte bevisar

- Det är **inte** en kvantdator som utför beräkningar. Detta är ett kvant-*minne*: det lagrar och skyddar en logisk kvantbit. Det utför inga logiska operationer (grindar) mellan logiska kvantbitar och kör ingen algoritm.
- Det är **inte** en enda kvantbit från användbara maskiner. En logisk kvantbit med avstånd 7 kräver omkring 101 fysiska kvantbitar; en felfrekvens på 0.1% per cykel ligger fortfarande långt över de ungefärliga 10^{-6} till 10^{-10} som verkliga algoritmer behöver. För att sluta det gapet krävs mycket större avstånd — många fler fysiska kvantbitar per logisk kvantbit — och användbara algoritmer behöver *tusentals* logiska kvantbitar samtidigt. Budgeten av fysiska kvantbitar för detta uppgår till miljontals.
- Detta **”om det skalas upp” är ett avgörande förbehåll**. Artikelns egen slutsats är att enhetens prestanda, *om den skalas upp*, skulle kunna uppfylla kraven för stora algoritmer. Att visa att utvecklingen går åt rätt håll för en logisk kvantbit är inte detsamma som att ha byggt den uppskalade maskinen, och ingenting här garanterar att utvecklingen håller i sig vid mycket större storlekar.
- Det **oförklarade felgolvet är ett aktuellt problem**. De korrelerade utbrott som begränsar repetitions-kodens prestanda är, med författarnas ord, storleksordningar större än väntat och skulle omöjliggöra större feltoleranta tillämpningar tills de har klarlagts — en öppen brist som redovisas tydligt, inte en löst detalj.
- Det säger ingenting om att **knäcka kryptering eller ”kvantöverlägsenhet” för användbara**

uppgifter. Det kräver den fullständiga feltoleranta maskin som detta är en grundsten för, inte en demonstration av.

Hur starka är beläggen

- **Kärnpåståendet är väl belagt och viktigt.** Drift under tröskelvärdet med tydlig exponentiell undertryckning över tre kodavstånd, plus en livslängd bortom breakeven och en fungerande realtid-savkodare, är exakt den kombination som forskningsfältet har försökt uppnå, och den demonstreras direkt i stället för att härledas. Detta är inte ett resultat skapat av hajp; det är ett verkligt ingenjörresultat från en ledande grupp.
- **Författarna är noga med resultatets räckvidd.** De beskriver det som ett *minne* under tröskelvärdet, uppmärksammar själva det oförklarade golvet av korrelerade fel och villkorar framtiden med detta tydliga ”om det skalas upp”. Överdriften, där den förekommer, finns i den omgivande rapportering som förvandlar ”en felkorrigerad minneskvantbit blev bättre när den växte” till ”kvantdatorn är här”.
- **Den ärliga lägesbilden är att ett grundläggande steg har tagits på ett entydigt sätt.** En logisk kvantbit, tillräckligt väl skyddad för att ytterligare redundans äntligen ska hjälpa — samtidigt återstår en lång, svår och ännu inte garanterad väg med uppskalning, logiska grindar och oförklarade fel.

Varför det spelar roll

Feltoleranta kvantdatorer har alltid varit något av ett hönan-och-ägget-problem: de maskiner som skulle vara användbara behöver felfrekvenser som ingen fysisk kvantbit kan uppnå, och lösningen — felkorrigering — fungerar bara om hårdvaran redan är tillräckligt bra för att ligga under tröskelvärdet. Att korsa den gränsen, om så bara en gång och för en enda logisk kvantbit, ändrar frågan från ”är detta över huvud taget möjligt?” till ”hur långt kan det skalas upp, och hur snabbt?” Det är en verklig och betydelsefull förändring, och därför förtjänar resultatet uppmärksamhet.

Men samma noggrannhet som gör resultatet trovärdigt bör också dämpa berättelsen kring det. Detta är den första tegelstenen i en grund, väl lagd. Det är inte byggnaden, och de som lade stenen är de första att säga det. Det rätta sättet att följa kvantdatorernas utveckling de närmaste åren är just denna föga glamorösa kurva: om undertryckningsfaktorn håller när koderna växer, om logiska grindar kan utföras lika tillförlitligt som logiskt minne och om det där gåtfulla felet en gång i timmen någonsin får en förklaring.

Kort sammanfattning

Google Quantum AI visade för första gången på ett entydigt sätt att ett kvantminne med ytkod kan arbeta *under tröskelvärde*: när de ökade koden från avstånd 3 till 5 till 7 sjönk den logiska felfrekvensen exponentiellt (med omkring $2.14\times$ för vartannat steg), och det största, ett avstånd-7-minne med 101 kvantbitar, överlevde sin bästa fysiska kvantbit — ”beyond breakeven” — samtidigt som dess felkorrigering kördes i realtid. Detta är en verklig och sedan länge eftersträvad milstolpe inom utvecklingen av kvantdatorer. Det är också en enda logisk kvantbit som fungerar som minne, med en felfrekvens som fortfarande är långt från vad verkliga algoritmer kräver, utan utförda logiska operationer, med ett oförklarat felgolv som författarna själva uppmärksammar och med en uppskalningsväg på många storleksordningar framför sig. Ett verkligt tröskelvärde har passerats — ingen kvantdator har levererats.

Källor

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>