



The CLEAN PAPER

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

Kvantová paměť se konečně zlepšovala s růstem — skutečný milník a stroj, kterým není

Google Quantum AI poprvé jasně ukázal, že kvantová paměť s povrchovým kódem může pracovat pod prahem: když se kód zvětšoval ze vzdálenosti 3 přes 5 na 7, logická chybovost exponenciálně klesala (asi 2,14krát na dva stupně) a největší, 101qubitová paměť se vzdáleností 7 přežila svůj nejlepší fyzický qubit — překročila bod zvratu — s korekcí chyb v reálném čase. Je to dlouho očekávaný technický milník. Je to také jediný logický qubit jako paměť, s chybovostí daleko od potřeb skutečných algoritmů, bez logických operací, s nevysvětleným chybovým dnem a mnoha řády škálování před sebou. Překročený práh, ne dodaný počítač.

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

Okamžik, kdy se křivka konečně ohnula správným směrem

Třicet let spočívala kvantová korekce chyb na slibu, který se nikdy nepodařilo jasně splnit. Teorie říká, že rozložíte-li jednu jednotku kvantové informace — jeden *logický* qubit — mezi mnoho hlučných fyzických qubitů a jsou-li tyto fyzické qubity dostatečně kvalitní, pak by přidávání *dalších* mělo logický qubit zlepšovat a chyby by s růstem kódu měly exponenciálně klesat. Háček je v onom „jsou-li“: pod kritickým prahem šumu více qubitů pomáhá; nad ním jen přidává další šum. Každý předchozí experiment se nacházel na špatné straně této hranice, nebo trend neukázal dostatečně čistě. Zvětšování kódu situaci zhoršovalo, ne zlepšovalo.

V prosinci 2024 Google Quantum AI oznámil první jasnou ukázkou opačného režimu. Na Willow, nejnovější generaci svých supravodivých procesorů, vytvořil paměť s povrchovým kódem o vzdálenostech 3, 5 a 7 a pozoroval, jak logická chybovost při každém zvětšení kódu *klesá* — o faktor $\Lambda = 2,14 \pm 0,02$ při každém zvýšení vzdálenosti o dva. Největší, 101qubitová paměť se vzdáleností 7 uchovávala logický qubit s chybou $0,143\% \pm 0,003\%$ na cyklus korekce a — hlavní zpráva uvnitř titulku — *vydržela déle než její vlastní nejlepší fyzický qubit*, o faktor $2,4 \pm 0,3$. Tomu se říká překročení bodu zvratu (beyond breakeven) a přínos korekce chyb na tomto hardwaru poprvé převážil její režii.

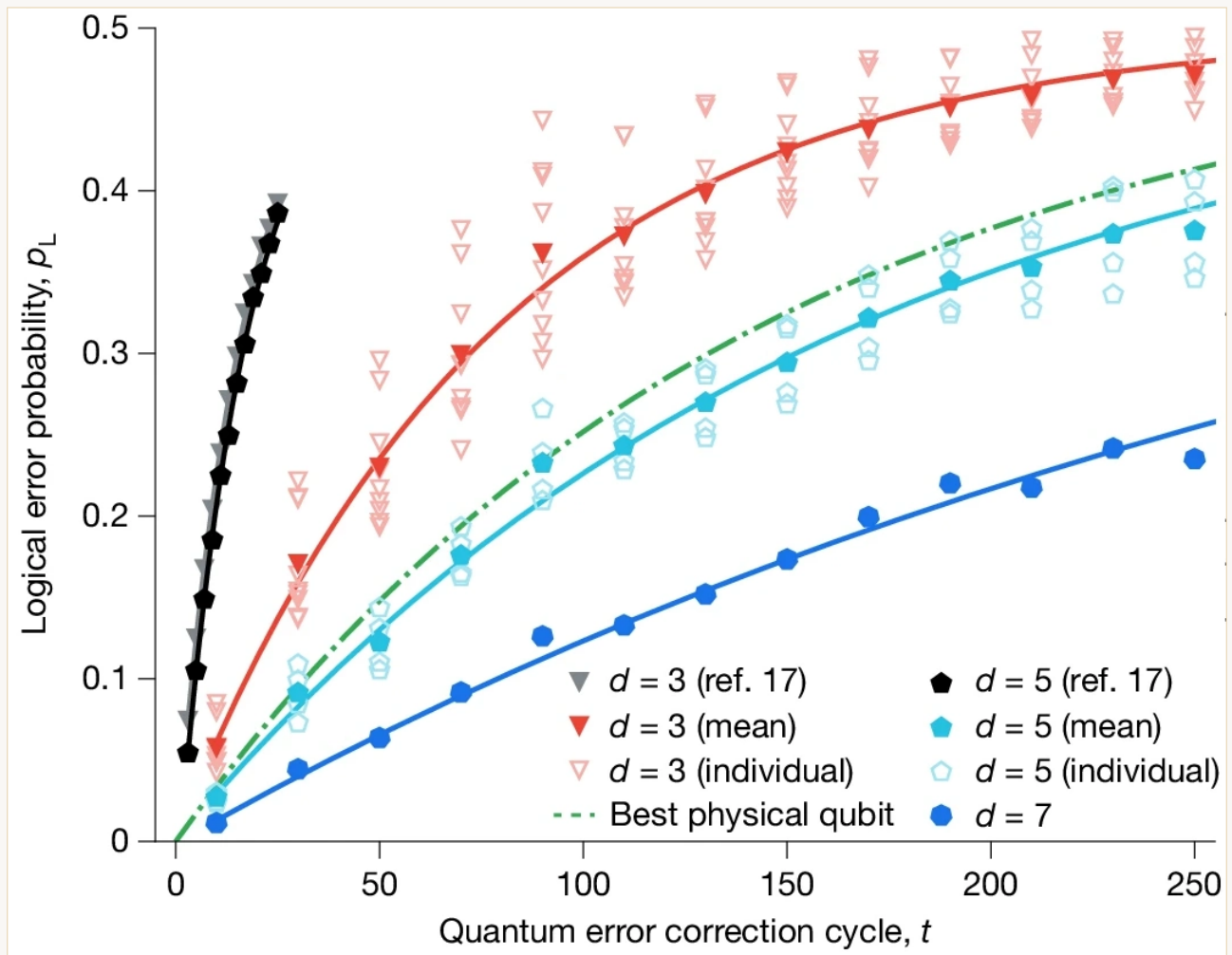
Je to skutečný milník a stojí za to přesně říci, jakého druhu. Dokazuje, že škálování nyní míří správným směrem. Není to fungující kvantový počítač a práce to ani netvrdí.

Co znamenají „povrchový kód“, „vzdálenost“ a „pod prahem“

Logický qubit je jedna chráněná jednotka kvantové informace zakódovaná do mnoha fyzických qubitů. **Povrchový kód** je konkrétní způsob tohoto kódování na dvourozměrné mřížce, kde další „měřicí“ qubity neustále kontrolují chyby, aniž by narušovaly uloženou informaci.

Vzdálenost kódu d není fyzická vzdálenost: je to nejmenší počet vhodně umístěných chyb, které mohou logický qubit poškodit, aniž si toho kód všimne. Větší d znamená větší a odolnější oblast — spotřebuje více fyzických qubitů (přibližně $2d^2 - 1$) a opraví více současných chyb, až $(d - 1)/2$. Tři zde testované velikosti, vzdálenosti 3, 5 a 7, tedy opravují 1, 2 a 3 současné chyby a spotřebují přibližně 17, 49 a 97 fyzických qubitů — paměť se vzdáleností 7 od Googlu použila 101, o něco více než učebnicové minimum.

„**Pod prahem**“ je klíčový výraz. Korekce chyb pomáhá pouze tehdy, když fyzická chybovost leží pod kritickou hodnotou; tam každé zvýšení vzdálenosti exponenciálně potlačuje logickou chybovost. Faktor potlačení Λ tento efekt měří — $\Lambda > 1$ znamená, že zvětšování kódu pomáhá, a čím vyšší Λ , tím lépe. Google uvádí $\Lambda \approx 2,14$, tedy že každé zvýšení vzdálenosti o dva snížilo logickou chybovost přibližně na polovinu. Podstatou výsledku je, že Λ pohodlně převyšuje 1.



Jak logická chyba narůstá s cykly korekce u paměti se vzdáleností 3, 5 a 7 (shora dolů). Sledujte zelenou přerušovanou čáru — *nejlepší jednotlivý fyzický qubit* na čipu. Paměť se vzdáleností 7 (modrá, nejnižší) hromadí chyby pomaleji než tato čára, takže zakódovaný qubit přežívá nejlepší fyzický qubit, z něhož je vytvořen — překračuje bod zvratu faktorem $2,4 \times$. Toto je výsledek pro životnost; samotné potlačení pod prahem ($\Lambda = 2,14$ při růstu kódu ze vzdálenosti 3 přes 5 na 7) vyjadřují čísla v textu.

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

Co autoři udělali

- Vytvořili paměti s povrchovým kódem na dvou čipech Willow: 105qubitovém procesoru, na kterém běžely kódy se vzdáleností 3, 5 a 7 použité ve zkoušce škálování (největší byla 101qubitová paměť se vzdáleností 7 a 49 datovými qubity), a 72qubitovém procesoru, na kterém běžela paměť se vzdáleností 5 s dekodérem v reálném čase a repetiční kódy vysokých vzdáleností.
- Měřili, jak se logická chyba *na cyklus* mění při zvyšování vzdálenosti kódu ze 3 na 5 a 7, a určili faktor potlačení Λ .
- Porovnali životnost logického qubitu s nejlepším jednotlivým fyzickým qubitem na stejném čipu, aby otestovali „bod zvratu“.
- Provozovali kód se vzdáleností 5 s **dekodérem v reálném čase** — klasickým hardwarem, který

vyhodnocuje kontroly chyb tak rychle, jak vznikají — až milion cyklů, aby ukázali, že korekce chyb dokáže držet krok se strojem.

- Rozšířili jednodušší **repetiční kódy** až na vzdálenost 29, aby hledali vzácné hluboké zdroje chyb, které vytvářejí výkonnostní dno.

Co zjistili

- **Kód pracuje pod prahem.** Logická chyba na cyklus klesala o faktor $\Lambda = 2,14 \pm 0,02$ při každém zvýšení vzdálenosti o dva — čisté exponenciální potlačení, které slibovala teorie a žádný procesor předtím přesvědčivě neukázal.
- **Paměť se vzdáleností 7 dosáhla chyby 0,143 % $\pm$ 0,003 % na cyklus** a žila **$2,4 \pm 0,3$ krát déle** než její nejlepší fyzický qubit — překročila bod zvratu.
- **Dekódování v reálném čase drželo krok.** Při vzdálenosti 5 měl dekodér průměrnou latenci 63 mikrosekund proti délce cyklu 1,1 mikrosekundy a vydržel milion cyklů — korekce chyb běžela živě, ne jen v následné analýze.
- **Zůstává vzácný hluboký zdroj chyb.** V testech repetičních kódů výkon nakonec omezovaly korelované výbuchy chyb přibližně **jednou za hodinu** (asi jednou za 3×10^9 cyklů), které vytvářejí chybové dno kolem 10^{-10} a jejichž původ podle autorů **dosud není znám**.

Co to nedokazuje

- **Není to kvantový počítač provádějící výpočty.** Jde o kvantovou *paměť*: ukládá a chrání jeden logický qubit. Neprovádí logické operace (hradla) mezi logickými qubity ani nespouští žádný algoritmus.
- Od užitečných strojů nás **nedělí jediný qubit**. Logický qubit se vzdáleností 7 spotřebuje asi 101 fyzických qubitů; chyba 0,1 % na cyklus je stále daleko nad přibližně 10^{-6} až 10^{-10} , které vyžadují skutečné algoritmy. Uzavření této mezery znamená mnohem větší vzdálenosti — mnohem více fyzických qubitů na každý logický — a užitečné algoritmy potřebují *tisíce* logických qubitů současně. Rozpočet fyzických qubitů pro takový stroj sahá do milionů.
- Výraz „**pokud se škáluje**“ **odvádí skutečnou práci**. Vlastní závěr práce říká, že výkon zařízení by *při škálování* mohl splnit požadavky velkých algoritmů. Ukázat správný trend na jednom logickém qubitu není totéž jako postavit stroj v cílovém měřítku a nic nezaručuje, že trend přežije při mnohem větších velikostech.
- **Nevysvětlené chybové dno je živý problém.** Korelované výbuchy omezující výkon repetičního kódu jsou podle autorů o řády větší, než se očekávalo, a dokud nebudou pochopeny, znemožní větší aplikace odolné proti chybám — otevřená vada popsaná jasně, nikoli vyřešený detail.
- Výsledek neříká nic o **prolamování šifer ani „kvantové nadvládě“ v užitečných úlohách**. Ty vyžadují úplný stroj odolný proti chybám, pro který je toto základním kamenem, ne předvedením.

Jak silné jsou důkazy

- **Základní tvrzení je pevné a důležité.** Provoz pod prahem s čistým exponenciálním potlačením přes tři vzdálenosti kódu, životností za bodem zvratu a fungujícím dekodérem v reálném čase je přesně kombinace, o kterou obor usiloval, a je prokázána přímo, nikoli odvozena. Nejde o výsledek vytvořený přehnaným vyprávěním; je to skutečný technický výsledek přední skupiny.
- **Autoři pečlivě vymezují rozsah.** Popisují jej jako *paměť* pod prahem, sami upozorňují na nevysvětlené dno korelovaných chyb a budoucnost podmiňují nápadným „pokud se škáluje“. Přehánění vzniká v okolním zpravodajství, které „paměťový qubit s korekcí chyb se při růstu zlepšoval“ zaokrouhluje na „kvantové počítače jsou tady“.
- **Poctivý stav je čistě provedený základní krok.** Jeden logický qubit chráněný natolik, že přidání redundance konečně pomáhá — s dlouhou, těžkou a nezaručenou cestou škálování, logických hradel a nevysvětlených chyb stále před sebou.

Proč na tom záleží

Kvantové výpočty odolné proti chybám vždy připomínaly problém slepice a vejce: užitečné stroje potřebují chybovost, které žádný fyzický qubit nedosáhne, a řešení — korekce chyb — funguje jen tehdy, když je hardware už dost dobrý na provoz pod prahem. Překročení této hranice byť jednou, byť na jediném logickém qubit, mění otázku z „je to vůbec možné?“ na „jak daleko a jak rychle se to dá škálovat?“. To je skutečný a významný posun a důvod, proč výsledek zasluhuje pozornost.

Stejná pečlivost, díky níž je výsledek důvěryhodný, však má tlumit příběh kolem něj. Je to první cihla základů, položená dobře. Není to budova a lidé, kteří ji položili, to říkají jako první. Správným způsobem, jak v příštích letech sledovat kvantové výpočty, je právě tato nepřitažlivá křivka: zda faktor potlačení vydrží při růstu kódů, zda lze logická hradla provádět stejně čistě jako logickou paměť a zda se někdy vysvětlí záhadná chyba jednou za hodinu.

Čisté shrnutí

Google Quantum AI poprvé jasně ukázal, že kvantová paměť s povrchovým kódem může pracovat *pod prahem*: když se kód zvětšoval ze vzdálenosti 3 přes 5 na 7, logická chybovost exponenciálně klesala (asi 2,14krát na každé dva stupně) a největší, 101qubitová paměť se vzdáleností 7 přežila svůj nejlepší fyzický qubit — překročila bod zvratu — zatímco korekce chyb běžela v reálném čase. Je to skutečný, dlouho očekávaný milník v technice kvantových počítačů. Je to také jediný logický qubit fungující jako paměť, s chybovostí stále daleko od požadavků skutečných algoritmů, bez provedených logických operací, s nevysvětleným chybovým dnem, na které upozorňují sami autoři, a s mnoha řády škálování před sebou. Překročen skutečný práh — nikoli dodán kvantový počítač.

Zdroje

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>