



The CLEAN PAPER

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

Egy kvantummemória végre jobb lett, ahogy nagyobbá vált — a valódi mérföldkő és a gép, ami még nincs

A Google Quantum AI először mutatta meg tisztán, hogy egy surface-code kvantummemória küszöb alatt működhet: ahogy a kódot 3-asról 5-ösre, majd 7-es távolságra növelték, a logikai hibaarány exponenciálisan csökkent (két lépésenként körülbelül $2,14\times$), a legnagyobb, 101 qubites 7-es távolságú memória pedig tovább élt, mint legjobb fizikai qubitje — breakeven fölé jutott —, valós időben futó hibajavítással. Ez régóta várt mérnöki mérföldkő. Ugyanakkor egyetlen, memóriaként működő logikai qubitról van szó, még mindig messze a valós algoritmusok hibaarány-igényétől, logikai műveletek nélkül, a szerzők által jelzett megmagyarázatlan hibapadlóval és több nagyságrendnyi skálázási úttal. Egy küszöböt átléptek, nem számítógépet szállítottak.

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

A pillanat, amikor a görbe végre jó irányba hajlott

Harminc éve a kvantum-hibajavítás egy olyan ígéretre épül, amelyet eddig soha nem sikerült tisztán teljesíteni. Az elmélet szerint ha egy egységnyi kvantuminformációt — egy *logikai* qubitet — sok zajos fizikai qubit között osztunk szét, és ezek a fizikai qubitek elég jók, akkor *több* qubit hozzáadásával a logikai qubitnek *jobbnak* kell lennie: a hibáknak exponenciálisan csökkenniük kell, ahogy a kód nő. A bökkenő az „ha”: egy kritikus zajküszöb alatt több qubit segít; fölötté több qubit csak még több zajt ad. Minden korábbi kísérlet a vonal rossz oldalán maradt, vagy nem mutatta tisztán a trendet. A kód növelése rontott a helyzeten, nem javított.

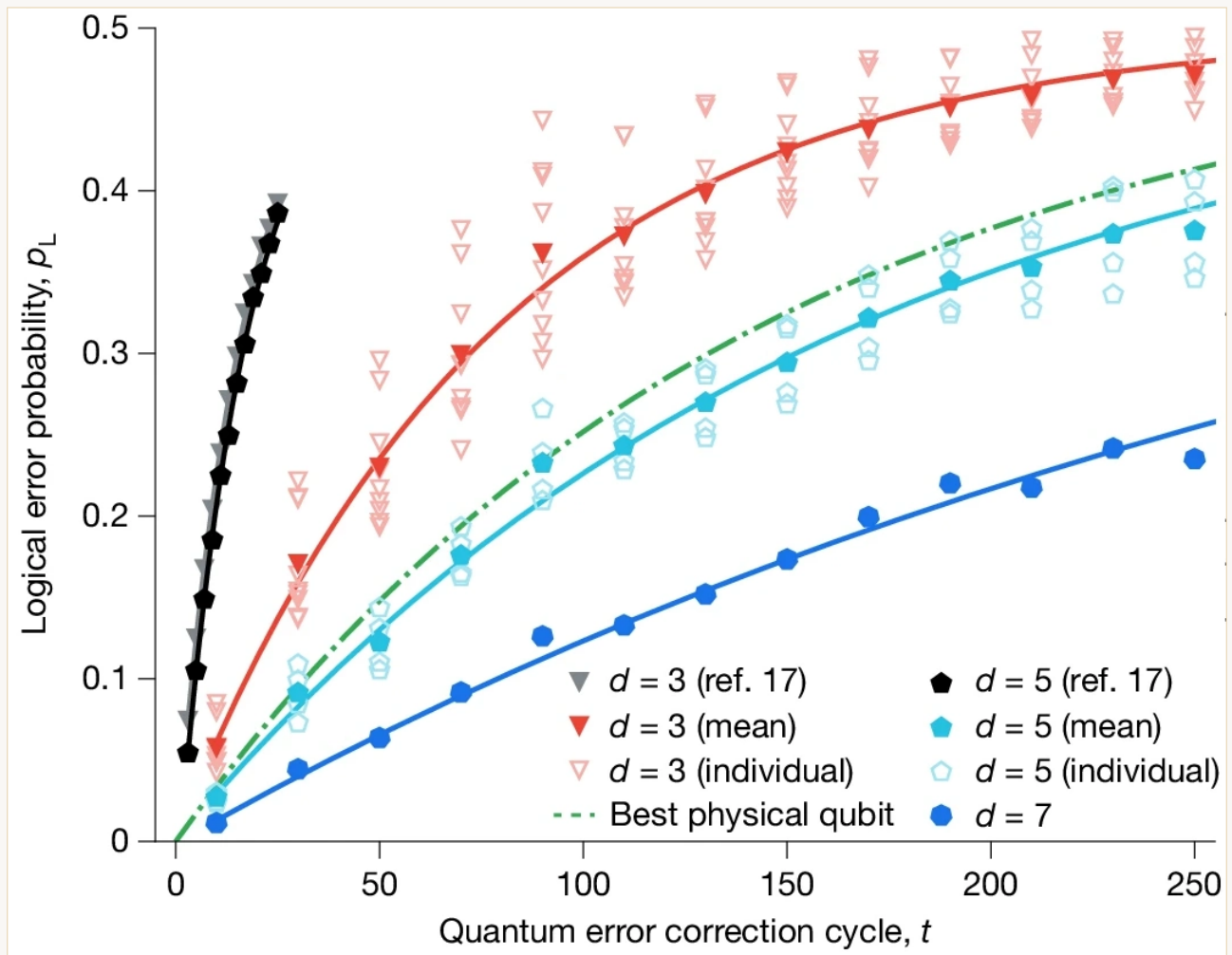
2024 decemberében a Google Quantum AI bejelentette ennek az ellenkező tartománynak az első világos demonstrációját. A Willow, legújabb szupravezető processzorgenerációjuk segítségével 3-as, 5-ös és 7-es kódtávolságú surface-code memóriákat építettek, és azt látták, hogy a logikai hibaarány *minden egyes méretnöveléskor csökkent* — két kódtávolság-lépésenként $\Lambda = 2,14 \pm 0,02$ faktorrallyal. A legnagyobb, 101 qubites, 7-es távolságú memória korrekciós ciklusonként $0,143\% \pm 0,003\%$ logikai hibával tárolt egy logikai qubitet, és — a főcímen belüli főcím — **2,4 $\pm$ 0,3-szor tovább élt, mint a chip saját legjobb fizikai qubitje**. Ezt „breakeven fölöttinek” nevezzük, és ezen a hardveren először fordult elő, hogy a teljes hibajavító apparátus valóban megtérült.

Ez valódi mérföldkő, és érdemes pontosan megmondani, milyen. Bizonyíték arra, hogy a skálázás iránya most már helyes. Nem működő kvantumszámítógép, és a tanulmány nem is állítja ezt.

Mit jelent a „surface code”, a „distance” és a „küszöb alatt”?

Egy **logikai qubit** egyetlen védett kvantuminformációs egység, amely sok fizikai qubit között van kódolva. A **surface code** ennek egy konkrét módja kétdimenziós rácson, ahol további „mérő” qubitek folyamatosan ellenőrzik a hibákat anélkül, hogy megzavarnák a tárolt információt. A kód **távolsága** (d) nem fizikai távolság: az a legkisebb számú, megfelelő helyen bekövetkező hiba, amely úgy ronthatja el a logikai qubitet, hogy a kód nem veszi észre. A nagyobb d nagyobb, robusztusabb foltot jelent — több fizikai qubitet használ (nagyjából $2d^2 - 1$), és egyszerre több hibát korrigál, legfeljebb $(d - 1)/2$ -t. Az itt tesztelt 3-as, 5-ös és 7-es távolságú kódok tehát 1, 2 és 3 egyidejű hibát korrigálnak, és elméleti minimumként körülbelül 17, 49 és 97 fizikai qubitet használnak — a Google 7-es távolságú memóriája 101-et használt, valamivel e tankönyvi minimum fölött.

A **„küszöb alatt”** kifejezés döntő. A hibajavítás csak akkor segít, ha a fizikai hibaarány egy kritikus érték alatt van; ott minden távolságnövelés exponenciálisan csökkenti a logikai hibaarányt. A Λ elnyomási faktor ezt méri — $\Lambda > 1$ azt jelenti, hogy a kód növelése segít, és minél nagyobb Λ , annál jobb. A Google $\Lambda \approx 2,14$ -et közöl, vagyis minden kétegységnyi távolságnövelés nagyjából felére csökkentette a logikai hibaarányt. Az, hogy Λ kényelmesen 1 fölött van, maga az eredmény.



Így halmozódik a logikai hiba a korrekciós ciklusok során a 3-as, 5-ös és 7-es távolságú memóriáknál (fentről lefelé). A figyelendő vonal a zöld szaggatott — a chip *legjobb egyedi fizikai qubitje*. A 7-es távolságú memória (kék, legalul) lassabban gyűjti a hibát, mint ez a vonal, így a kódolt qubit tovább él, mint a legjobb fizikai qubit, amelyből felépül — „breakeven fölött”, $2,4\times$ faktossal. Ez az élettartam-eredmény; a küszöb alatti elnyomás maga ($\Lambda = 2,14$, ahogy a kód 3-ról 5-re, majd 7-re nő) a szöveg számaiban szerepel.

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

Mit csináltak a szerzők

- Surface-code memóriákat építettek két Willow chipen: egy 105 qubites processzor futtatta a skálázási teszt 3-as, 5-ös és 7-es távolságú kódjait (legnagyobbként a 101 qubites, 49 adatqubites 7-es távolságú memóriát), egy 72 qubites processzor pedig 5-ös távolságú memóriát valós idejű dekóderrel, valamint nagy távolságú repetition code-okat.
- Megmérték, hogyan változik a ciklusonkénti logikai hiba, amikor a kódtávolságot 3-ról 5-re, majd 7-re növelik, ebből meghatározva az Λ elnyomási faktort.
- A „breakeven” tesztjéhez összehasonlították a logikai qubit élettartamát ugyanazon chip legjobb egyedi fizikai qubitjével.
- A 5-ös távolságú kódot **valós idejű dekóderrel** futtatták — klasszikus hardverrel, amely olyan

gyorsan értelmezi a hibajelző méréseket, ahogy azok keletkeznek — akár egymillió cikluson át, megmutatva, hogy a hibajavítás képes lépést tartani a géppel.

- Egyszerűbb **repetition code-okat** 29-es távolságig növeltek, hogy felkutassák azokat a ritka, mély hibaforrásokat, amelyek alsó korlátot szabnak a teljesítménynek.

Mit találtak

- **A kód küszöb alatt működik.** A ciklusonkénti logikai hiba két távolságlépésenként $\Lambda = 2,14 \pm 0,02$ faktorral csökkent — tiszta exponenciális elnyomás, az elmélet által ígért viselkedés, amelyet korábban egyetlen processzor sem mutatott meg egyértelműen.
- **A 7-es távolságú memória ciklusonként 0,143% $\pm$ 0,003% hibát ért el, és 2,4 $\pm$ 0,3-szor tovább élt,** mint legjobb fizikai qubitje — breakeven fölött.
- **A valós idejű dekódolás lépést tartott.** A dekóder átlagos késleltetése 5-ös távolságnál 63 mikroszekundum volt az 1,1 mikroszekundumos ciklusidő mellett, egymillió cikluson át — a hibajavítás előben futott, nem csak utólagos elemzésként.
- **Egy ritka, mély hibaforrás megmaradt.** A repetition-code tesztekben a teljesítményt végül korrelált hibakitörések korlátozták, amelyek nagyjából **óránként egyszer** fordultak elő (mintegy egy 3×10^9 ciklusból), körülbelül 10^{-10} -es hibapadlót hozva létre. Ennek eredetét a szerzők szerint **még nem értik.**

Mit nem bizonyít

- **Nem kvantumszámítógép, amely számítást végez.** Ez kvantum-*memória*: egyetlen logikai qubitet tárol és véd. Nem hajt végre logikai műveleteket (kapukat) logikai qubitek között, és nem futtat algoritmust.
- **Nem egyetlen qubit választja el a hasznos gépektől.** Egy 7-es távolságú logikai qubit körülbelül 101 fizikai qubitet használ; a $\sim 0,1\%$ -os ciklusonkénti hibaarány még mindig messze van a valódi algoritmusok által igényelt nagyjából 10^{-6} – 10^{-10} tartománytól. Ennek áthidalásához jóval nagyobb kódtávolságok kellenek — sokkal több fizikai qubit logikai qubitenként —, a hasznos algoritmusok pedig egyszerre *több ezer* logikai qubitet igényelnek. Ehhez milliós nagyságrendű fizikaiqubit-költségvetés adódik.
- **Az „ha skálázódik” valódi munkát végez.** A tanulmány saját következtetése az, hogy a berendezés teljesítménye, *ha skálázódik*, megfelelhet nagy algoritmusok igényeinek. Egy logikai qubiten megmutatni, hogy jó irányú a trend, nem ugyanaz, mint megépíteni a skálázott gépet, és semmi nem garantálja, hogy a trend sokkal nagyobb méreteken is megmarad.
- **A megmagyarázatlan hibapadló élő probléma.** A repetition-code teljesítményt korlátozó korrelált kitörések a szerzők szavaival nagyságrendekkel nagyobbak a vártnál, és amíg nem értik meg őket, megakadályoznák a nagyobb hibatűrő alkalmazásokat — nyílt hiba, nem megoldott részlet.
- **Semmit nem mond titkosítás feltöréséről vagy hasznos feladatokban elért „quantum supremacy”-ről.** Ezekhez a teljes hibatűrő gép kell, amelyhez ez alapkövet helyez le, nem demonstrációja annak.

Mennyire erős a bizonyíték?

- **A központi állítás szilárd és fontos.** A három kódtávolságon át mutatott tiszta exponenciális elnyomással végzett küszöb alatti működés, a breakeven fölötti élettartam és a működő valós idejű dekóder együtt pontosan az a kombináció, amelyet a terület régóta próbált elérni. Közvetlenül demonstrálták, nem következették. Ez nem hype-műtermék, hanem egy vezető csoport valódi mérnöki eredménye.
- **A szerzők óvatosa a hatókörrel.** Küszöb alatti *memóriaként* keretezik, maguk emelik ki a megmagyarázatlan korrelált hibapadlót, és a jövőt a feltűnő „ha skálázódik” feltételhez kötik. A túlzás, ahol megjelenik, inkább a környező tudósításban van, amely „egy hibajavított memóriaqubit javult, ahogy nőtt” mondatból „megérkezett a kvantumszámítástechnika” kijelentést csinál.
- **Az őszinte állapot: egy alapvető lépést tisztán megtettek.** Egy logikai qubit, amelyet már elég jól védenek ahhoz, hogy a redundancia növelése valóban segítsen — hosszú, nehéz és még nem garantált skálázási úttal, logikai kapukkal és megmagyarázatlan hibákkal előtte.

Miért fontos

A hibátűrő kvantumszámítástechnika mindig tyúk-tojás problémának tűnt: a hasznos gépek olyan hibaarányt igényelnek, amelyet egyetlen fizikai qubit sem ér el, a megoldás — a hibajavítás — pedig csak akkor működik, ha a hardver már eleve elég jó ahhoz, hogy a küszöb alatt legyen. E vonal átlépése akár egyszer, akár egyetlen logikai qubiten is megváltoztatja a kérdést „lehetséges-e egyáltalán?”-ről „meddig és milyen gyorsan skálázható?”-ra. Ez valódi és jelentős változás, ezért érdemes figyelmet.

Ugyanakkor éppen az eredményt hitelessé tevő óvatosságnak kell lehűtenie a köré épülő történetet. Ez egy alap első téglája, jól lerakva. Nem az épület, és akik lerakták, elsőként mondják ezt. A kvantumszámítástechnikát a következő években éppen ezen a nem látványos görbén érdemes követni: megmarad-e az elnyomási faktor, ahogy nőnek a kódok; végrehajthatók-e a logikai kapuk olyan tisztán, mint a logikai memória; és kiderül-e valaha, mi okozza azt a rejtélyes óránkénti hibát.

Röviden

A Google Quantum AI először mutatta meg tisztán, hogy egy surface-code kvantummemória *küszöb alatt* működhet: ahogy a kódot 3-asról 5-ösre, majd 7-es távolságra növelték, a logikai hibaarány exponenciálisan csökkent (két lépésenként körülbelül $2,14\times$), a legnagyobb, 101 qubites 7-es távolságú memória pedig tovább élt, mint legjobb fizikai qubitje — breakeven fölé jutott —, miközben a hibajavítás valós időben futott. Ez valódi, régóta keresett mérföldkő a kvantumszámítógépek mérnöki fejlesztésében. Ugyanakkor egyetlen, memóriaként működő logikai qubitról van szó, még mindig jóval a valós algoritmusokhoz szükséges hibaarány fölött, logikai műveletek nélkül, a szerzők által is kiemelt megmagyarázatlan hibapadlóval, és több nagyságrendnyi skálázási úttal előtte. Egy valódi küszöböt átléptek — kvantumszámítógépet még nem adtak át.

Források

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>