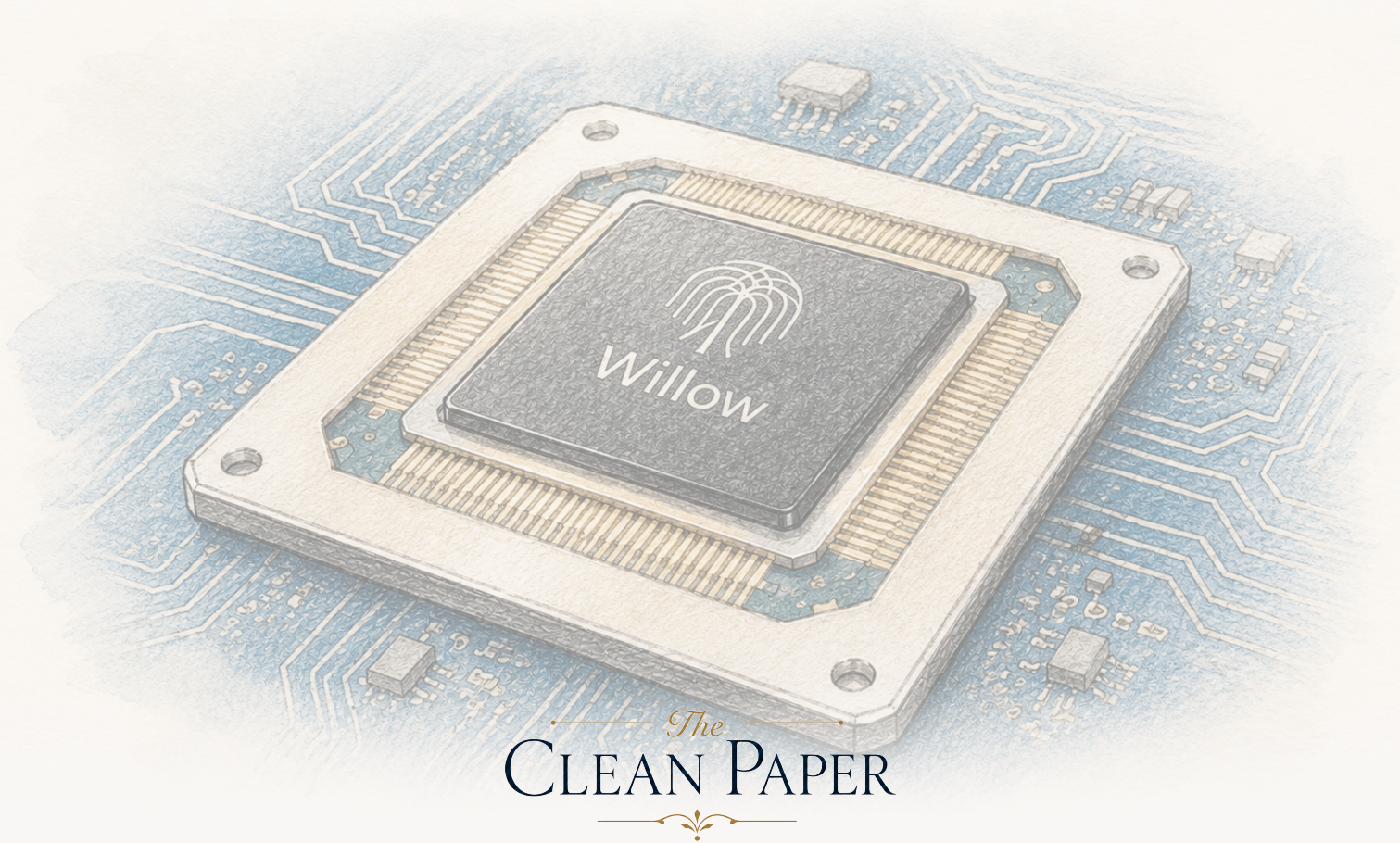




WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Քվանտային հիշողությունը վերջապես լավացավ՝ մեծանալով. իրական milestone-ը, բայց ոչ դեռ quantum computer

Google Quantum AI-ը առաջին անգամ մաքուր ցույց տվեց, որ surface-code quantum memory-ն կարող է աշխատել below threshold. distance-ը 3-ից 5, ապա 7 բարձրացնելիս logical error rate-ը exponential նվազեց՝ մոտ $2.14 \times$ յուրաքանչյուր երկու step-ի համար, իսկ 101-qubit distance-7 memory-ն իր լավագույն physical qubit-ից ավելի երկար ապրեց՝ beyond breakeven, real-time error correction-ով: Սա երկար սպասված engineering milestone է, բայց դեռ մեկ logical qubit է memory-ի դերում՝ իրական algorithm-ներին անհրաժեշտից շատ ավելի բարձր error rate-ով, առանց logical operation-ների և unexplained error floor-ով:

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

Այն պահը, երբ կորն առաջին անգամ ճիշտ ուղղությամբ թեքվեց

Երեսուն տարի քվանտային սխալ ուղղումը հենվում էր մի խոստման վրա, որը դեռ երբեք մաքուր ձևով չէր իրականացվել: Տեսությունն ասում է. եթե քվանտային տեղեկությունի մեկ միավորը՝ մեկ տրամաբանական քուրբիթ, տարածեք բազմաթիվ noisy ֆիզիկական քուրբիթների վրա, և եթե այդ ֆիզիկական քուրբիթները բավական լավն են, ապա դրանցից *ավելի շատ* ավելացնելը պետք է տրամաբանական քուրբիթը *ավելի լավ* դարձնի՝ կողք մեծանալիս սխալները էքսպոնենցիալ նվազեցնելով: Խնդիրը այդ «եթե»-ն է. critical աղմուկ շեմից ցածր ավելի շատ քուրբիթը օգնում է, իսկ դրանից վեր՝ պարզապես ավելի շատ աղմուկ է ավելացնում: Նախորդ փորձերը կամ շեմի սխալ կողմում էին, կամ միտում-ը մաքուր չէին ցույց տվել: կողք մեծացնելը արդյունքը վատացնում էր, ոչ լավացնում:

2024 թվականի դեկտեմբերին Google քվանտային ԱԲը հաղորդեց առաջին հստակ ցուցադրումը հակառակ regime-ի: Willow-ի՝ իրենց նորագույն superconducting պրոցեսոր-ների վրա նրանք մակերեսային հիշողություններ կառուցեցին հեռավորություն 3, 5 և 7-ով և տեսան, որ տրամաբանական սխալի հաճախականությունը կողք մեծացնելիս ամեն անգամ *նվազում* է՝ յուրաքանչյուր երկու հեռավորություն քայլի համար $\Lambda = 2.14 \pm 0.02$ ճնշում գործոնով: Ամենամեծը՝ 101-քուրբիթ հեռավորություն-7 հիշողությունն, մեկ տրամաբանական քուրբիթ պահեց $0.143\% \pm 0.003\%$ սխալ per ուղղում ցիկլով և՛ վերնագրի ամենակարևոր մասով, ապրեց *ավելի երկար, քան նույն չիպի լավագույն ֆիզիկական քուրբիթը՝ 2.4 ± 0.3 անգամ:* Սա կոչվում է **սահմանից այն կողմ շահավետության շեմ**: Առաջին անգամ այս սարքավորում-ի վրա սխալ ուղղումի ամբողջ overhead-ը իր արժեքը վերադարձրեց:

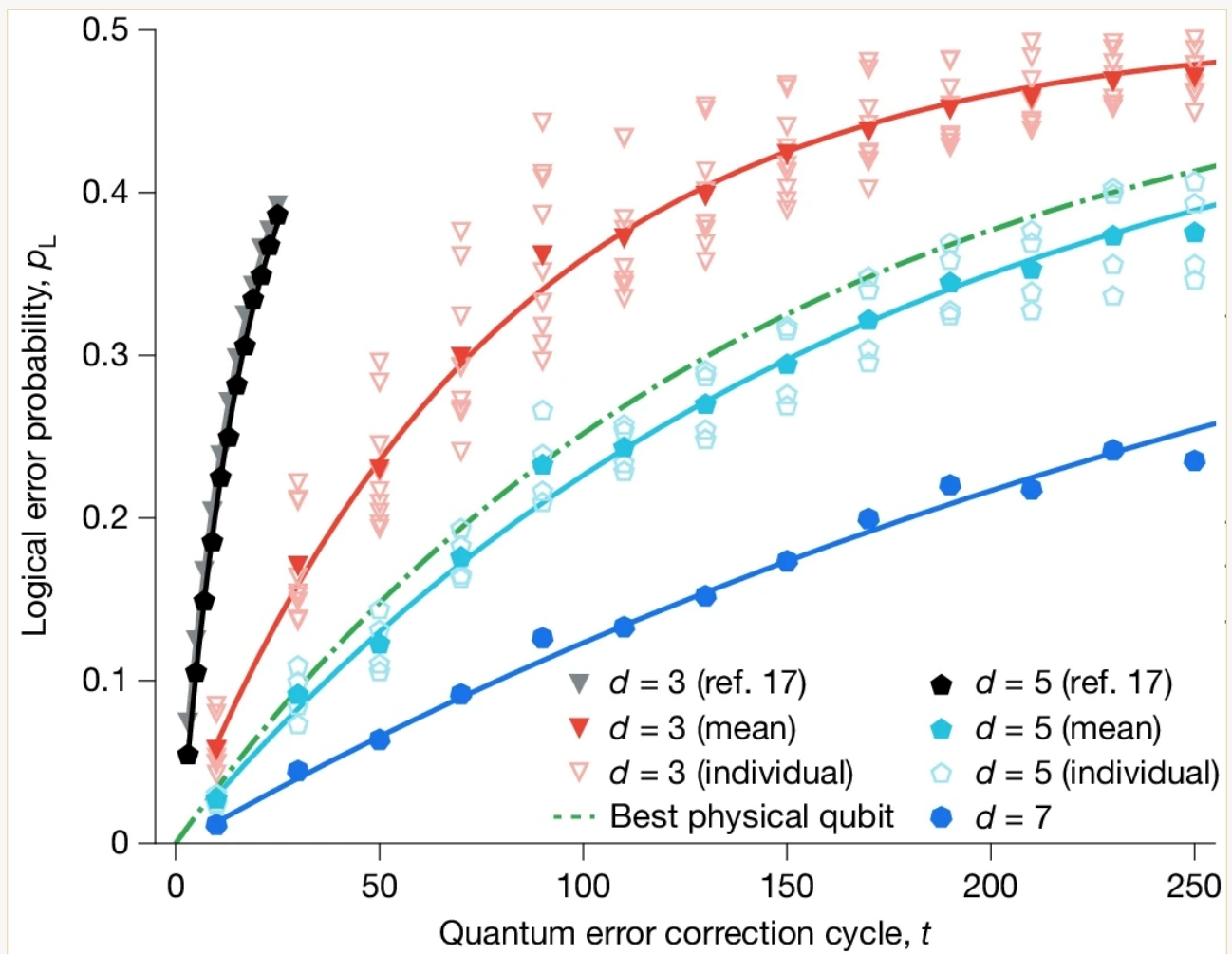
Սա իրական հանգրվան է, և կարևոր է ճշգրտել՝ ինչ հանգրվան: Արդյունքը ցույց է տալիս, որ մասշտաբավորումը հիմա ճիշտ ուղղությամբ է գնում: Սա աշխատող քվանտային համակարգիչ չէ, և հոգվածը դա չի էլ պնդում:

Ինչ են նշանակում «մակերես code»-ը, «հեռավորություն»-ը և «below շեմ»-ը

տրամաբանական քուրբիթ-ը քվանտային տեղեկությունի մեկ պաշտպանված միավոր է, կոդավորված բազմաթիվ ֆիզիկական քուրբիթների վրա: **Surface կոդ**-ը այդ encoding-ի կոնկրետ մեթոդ է 2D ցանցի վրա, որտեղ լրացուցիչ չափել քուրբիթներ անընդհատ ստուգում են սխալները՝ առանց պահված տեղեկությունը խանգարելու: կոդ **հեռավորություն d** -ն ֆիզիկական հեռավորությունն է: Դա ճիշտ տեղերում առաջացած սխալների ամենափոքր թիվն է, որը կարող է տրամաբանական քուրբիթը corrupt անել՝ առանց կոդի կողմից հայտնաբերվելու: Ավելի մեծ d -ն ավելի մեծ և կայուն patch է. այն ծախսում է ավելի շատ ֆիզիկական քուրբիթ (մոտ $2d^2 - 1$) և միաժամանակ ավելի շատ սխալ է ուղղում՝ մինչև $(d - 1)/2$: Այսպիսով այստեղ փորձարկված distances 3, 5 և 7-ը ուղղում են համապատասխանաբար 1, 2 և 3 simultaneous

սխալ և տեսականորեն պահանջում մոտ 17, 49 և 97 ֆիզիկական քուրիթ: Google-ի հեռավորություն-7 հիշողությունն օգտագործել է 101` մի փոքր ավելի, քան textbook minimum-ը:

Below շեմ-ը կարևոր արտահայտությունն է: սխալ ուղղումը օգնում է միայն այն դեպքում, երբ ֆիզիկական սխալի հաճախականությունը կրիտիկական արժեքից ցածր է: Այդ ռեժիմում հեռավորությունի յուրաքանչյուր աճ տրամաբանական սխալի հաճախականությունը էքսպոնենցիալ ճնշում է: ճնշում գործոն Λ -ն չափում է դա. $\Lambda > 1$ նշանակում է կոդը մեծացնելն օգնում է, և որքան բարձր է Λ -ն, այնքան լավ: Google-ը հայտնում է $\Lambda \approx 2.14$, այսինքն` հեռավորությունի յուրաքանչյուր երկուքայլ աճը տրամաբանական սխալի հաճախականությունը մոտ կեսով նվազեցրել է: Հենց այն, որ Λ -ն հստակ 1-ից մեծ է, ամբողջ արդյունքի սիրտն է:



Ինչպես է տրամաբանական սխալը կուտակվում ուղղում ցիկլների ընթացքում հեռավորություն-3, -5 և -7 հիշողությունների համար` վերևից ներքև: Դիտելու հիմնական գիծը կանաչ dashed-ն է` չիպի լավագույն մեկ ֆիզիկական քուրիթը: հեռավորություն-7 հիշողությունն` կապույտ, ամենացածր curve-ը, այդ գծից ավելի դանդաղ է սխալ կուտակում, այսինքն կոդավորված քուրիթը ապրում է ավելի երկար, քան այն լավագույն ֆիզիկական քուրիթը, որոնցից կառուցված է` $2.4\times$ սահմանից այն կոդմ շահավետության

շեմ: Մա lifetime արդյունքն է. below-շեմ ճնշումը ինքնին ($\Lambda = 2.14$, երբ հեռավորությունը 3-ից 5, ապա 7 է աճում) տեքստի թվերում է:

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

Ինչ արեցին հեղինակները

- Երկու Willow չիպի վրա կառուցեցին մակերեսկող հիշողություններ: 105-քուբիթ պրոցեսոր-ը գործարկեց մասշտաբավորում թեստի հեռավորություն-3, -5 և -7 կոդերը՝ ամենամեծը 101-քուբիթ, 49-տվյալներքուբիթ հեռավորություն-7 հիշողություն: 72-քուբիթ պրոցեսոր-ը գործարկեց հեռավորություն-5 հիշողություն՝ իրականժամանակ decoder-ով, ինչպես նաև բարձրհեռավորություն կրկնության կոդեր:
- Չափեցին, թե տրամաբանական սխալը *per qubit* ինչպես է փոխվում կող հեռավորությունը 3-ից 5, ապա 7 բարձրացնելիս, և դրանից հանեցին Λ ճնշում գործոնը:
- տրամաբանական քուբիթի lifetime-ը համեմատեցին նույն չիպի լավագույն անհատ ֆիզիկական քուբիթի հետ՝ շահավետության շեմը փորձարկելու համար:
- հեռավորություն-5 կոդը գործարկեցին **իրականժամանակ decoder**-ով՝ դասական սարքավորում, որը սխալ ստուգումները մեկնաբանում է գրեթե նույն արագությամբ, որով դրանք ստեղծվում են, մինչև մեկ միլիոն ցիկլ՝ ցույց տալու, որ սխալ ուղղումը կարող է live աշխատել մեքենայի հետ:
- Ավելի պարզ **կրկնության կոդ**-երը հասցրին մինչև հեռավորություն 29՝ փնտրելով հազվադեպ, խոր սխալ աղբյուրներ, որոնք արդյունավետության համար floor են սահմանում:

Ինչ գտան

- **կոդը below շեմ է:** տրամաբանական սխալ *per qubit* հեռավորությունի յուրաքանչյուր երկուքայլ աճի համար նվազեց $\Lambda = 2.14 \pm 0.02$ գործոնով: Մա մաքուր էքսպոնենցիալ ճնշում է՝ այն վարքը, որը տեսությունն վաղուց խոստանում էր և ոչ մի պրոցեսոր դեռ համոզիչ կերպով չէր ցույց տվել:
- **հեռավորություն-7 հիշողությունն հասավ $0.143\% \pm 0.003\%$ սխալ *per qubit* և իր լավագույն ֆիզիկական քուբիթից 2.4 ± 0.3 անգամ ավելի երկար ապրեց՝** սահմանից այն կողմ շահավետության շեմ:
- **իրականժամանակ decoding-ը հասցրեց մեքենային:** հեռավորություն 5-ի դեպքում ապակոդավորիչի միջին ուշացումը 63 միկրովայրկյան էր՝ 1.1-միկրովայրկյան ցիկլ ժամանակի պայմաններում, և դա պահպանվեց մեկ միլիոն ցիկլի ընթացքում: սխալ ուղղումը աշխատում էր իրական ժամանակում, ոչ միայն հետագա վերլուծությունում:
- **Մնում է հազվադեպ, խոր սխալ աղբյուր:** Repetition-կոդ թեստերում արդյունավետություն վերջում սահմանափակվեց correlated սխալ burst-երով, որոնք լինում էին մոտ

Ժամը մեկ անգամ`** մոտ  $3 \times 10^9$  ցիկլից մեկում, առաջացնելով մոտ  $10^{-10}$  սխալ floor: Հեղինակներն ասում են` դրա պատճառը **դեռ հասկանալի չէ:

Ինչ սա չի ապացուցում

- Սա **քվանտային համակարգիչ չէ, որը computation է անում**: Սա քվանտային *հիշողություն* է` պահում և պաշտպանում է մեկ տրամաբանական քուրիթ: Այն տրամաբանական քուրիթների միջև տրամաբանական գործողություն կամ դարպաս չի կատարում և այգորիթմ չի աշխատեցնում:
- Սա **օգտակար մեքենայից մեկ քուրիթ հեռու լինել չէ**: հեռավորություն-7 տրամաբանական քուրիթը ծախսում է մոտ 101 ֆիզիկական քուրիթ: 0.1% սխալ per ցիկլը դեռ շատ բարձր է այն մոտ 10^{-6} - 10^{-10} միջակայքից, որը իրական այգորիթմներին է պետք: Այդ տարբերությունը փակելու համար պետք են շատ ավելի մեծ հեռավորություններ` յուրաքանչյուր տրամաբանական քուրիթի համար շատ ավելի շատ ֆիզիկական քուրիթ, իսկ օգտակար այգորիթմները միաժամանակ *հազարավոր* տրամաբանական քուրիթ են պահանջում: ֆիզիկական քուրիթ budget-ը կարող է հասնել միլիոնների:
- Հոդվածի «**եթե մասշտաբավորվի**»-ը կարևոր սահմանափակում է: Իր եզրակացության մեջ հոդվածը ասում է, որ սարք արդյունավետությունը *եթե սանդղակ արվի*, կարող է բավարարել մեծ այգորիթմների պահանջները: Մեկ տրամաբանական քուրիթի վրա ճիշտ միտում ցույց տալը նույնը չէ, ինչ մասշտաբավորված մեքենա կառուցելը, և այստեղ ոչինչ չի երաշխավորում, որ միտում-ը շատ ավելի մեծ չափերի վրա նույնը կմնա:
- **Անբացատրելի սխալի ստորին սահմանը իրական բաց խնդիր է**: Կապակցված պոռթկումները, որոնք սահմանափակում են repetition-կող արդյունավետությունը, հեղինակների խոսքով մի քանի կարգով ավելի մեծ են սպասվածից և, մինչև հասկացվելը, կարող են խանգարել ավելի մեծ խափանումադիմացկուն կիրառումներին: Սա բաց խնդիր է, ոչ փակված մանրուք:
- Արդյունքը ոչինչ չի ասում **գաղտնագրում կոտրելու կամ օգտակար առաջադրանքերի համար «քվանտային գերակայություն»-ի** մասին: Դրա համար պետք է ամբողջական խափանումադիմացկուն մեքենա, որի հիմք stone-երից մեկն է այս արդյունքը, ոչ ինքնին ցուցադրումը:

Որքա՞ն ուժեղ է ապացույցը

- **Հիմնական պնդումը ուժեղ և կարևոր է**: Երեք կող հեռավորությունի վրա հստակ էքսպոնենցիալ ճնշումով շեմից ցածր աշխատանք, և սահմանից այն կողմնահավետության շեմ lifetime և աշխատող իրականժամանակ decoder` հենց այն համադրությունն է, որ դաշտը փորձում էր հասնել: Այն ցուցադրված է ուղիղ, ոչ անուղղակի եզրակացությունով: Սա գովազդային չափազանցություն չէ, այլ իրական ինժեներական արդյունք առաջատար խմբից:

- **Հեղինակները սահմանների հարցում զգույշ են:** Նրանք արդյունքը նկարագրում են որպես below-շեմ *հիշողություն*, իրենք են մատնանշում unexplained correlated-սխալի ստորին սահմանը և ապագայի մասին խոսելիս հատուկ պահում «եթե մասշտաբավորվի» վերապահումը: Չափագանցումը հիմնականում surrounding ծածկություն է, երբ «սխալուղղված հիշողություն քուրիթը մեծանալիս լավացավ» դարձնում են «քվանտային հաշվարկը արդեն այստեղ է»:
- **Ազնիվ կարգավիճակը հիմքային քայլ է:** Մեկ տրամաբանական քուրիթ, այնքան լավ պաշտպանված, որ redundancy-ն վերջապես օգնում է, բայց առջևում երկար, դժվար և ոչ երաշխավորված ճանապարհ է՝ մասշտաբավորում, տրամաբանական դարպասներ և unexplained սխալներ:

Ինչու է սա կարևոր

Խափանումադիմացկուն քվանտային computing-ը միշտ մի փոքր chicken-and-egg խնդիր էր հիշեցնում: Օգտակար մեքենաներին պետք են սխալի հաճախականություններ, որոնց ոչ մի ֆիզիկական քուրիթ չի հասնում, իսկ լուծումը՝ սխալ ուղղումը, աշխատում է միայն այն ժամանակ, երբ սարքավորում-ն արդեն բավական լավ է շեմից ցածր լինելու համար: Այդ գիծը թեկուզ մեկ անգամ, մեկ տրամաբանական քուրիթի վրա անցնելը հարցը փոխում է «ընդհանրապես հնարավո՞ր է» ձևից դեպի «որքա՞ն հեռու և որքա՞ն արագ կարելի է սանդղակ անել»: Սա իրական և կարևոր փոփոխություն է:

Բայց հենց այն նույն զգուշությունը, որը արդյունքը վստահելի է դարձնում, պետք է զսպի դրա շուրջ պատմությունը: Սա հիմքի առաջին լավ դրված աղյուսն է: Սա building-ը չէ, և այն դրած մարդիկ առաջիններից են դա ասելու: քվանտային computing-ին հաջորդ տարիներին հետևելու լավագույն ձևը հենց այս ոչ glamorous curve-ն է՝ ճնշում գործոնը կապահպանվի՞ կողը մեծացնելիս, տրամաբանական դարպասները կաշխատե՞ն այնքան հստակ, որքան տրամաբանական հիշողությունն, և արդյոք ժամում մեկ անգամ հայտնվող այդ տարօրինակ սխալը երբևէ կրացատրվի:

Կարճ ամփոփում

Google քվանտային ԱԲը առաջին անգամ մաքուր ցույց տվեց, որ մակերեսկող քվանտային հիշողությունն կարող է աշխատել *շեմից ցածր*: կող հեռավորությունը 3-ից 5, ապա 7 բարձրացնելիս տրամաբանական սխալի հաճախականությունը էքսպոնենցիալ նվազեց՝ յուրաքանչյուր երկու քայլի համար մոտ $2.14\times$, իսկ ամենամեծ՝ 101-քուրիթ հեռավորություն-7 հիշողությունն իր լավագույն ֆիզիկական քուրիթից ավելի երկար ապրեց՝ սահմանից այն կողմ շահավետության շեմ, իրական ժամանակ աշխատող սխալ ուղղումով: Սա քվանտային համակարգչային ինժեներիայի երկար սպասված իրական հանգրվան է: Բայց այն նաև ընդամենը մեկ տրամաբանական քուրիթ է՝ հիշողությունի դերում, սխալի հաճախականությունով, որը դեռ շատ հեռու

Է իրական ալգորիթմների պահանջներից, առանց տրամաբանական գործողությունների, հեղինակների իսկ մատնանշած unexplained սխալ floor-ով և մասշտաբավորումի մի քանի կարգի մասշտաբավորման ճանապարհով: Իրական շեմ է անցել, ոչ քվանտային համակարգիչ է հանձնվել:

Աղբյուրներ

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>