



The CLEAN PAPER

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

یک حافظه کوانتومی بالاخره با بزرگ تر شدن بهتر شد — نقطه عطف واقعی، و ماشینی که هنوز نیست

Google AI Quantum برای نخستین بار به طور روشن نشان داد یک حافظه کوانتومی *code surface* می تواند *threshold below* کار کند: با رشد که از فاصله ۳ به ۵ به ۷، نرخ خطای منطقی به صورت نمایی پایین آمد (حدود $14.2 \times$ به ازای هر دو واحد)، و بزرگ ترین نمونه، حافظه فاصله ۷-با ۱۰۱ کیوبیت، از بهترین کیوبیت فیزیکی خودش بیشتر عمر کرد — *beyond breakeven* — در حالی که تصحیح خطا بلا درنگ اجرا می شد. نقطه عطف مهندسی واقعی است. اما یک کیوبیت منطقی در نقش حافظه است، با خطایی هنوز بسیار بالاتر از نیاز الگوریتم ها، بدون *gate* منطقی، با کف خطای ناشناخته و چندین مرتبه بزرگی *scaling* در پیش. آستانه ای رد شد، نه رایانه ای تحویل داده شد.

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: y-08449-024-s41586/1038.10

لحظه‌ای که منحنی بالاخره در جهت درست خم شد

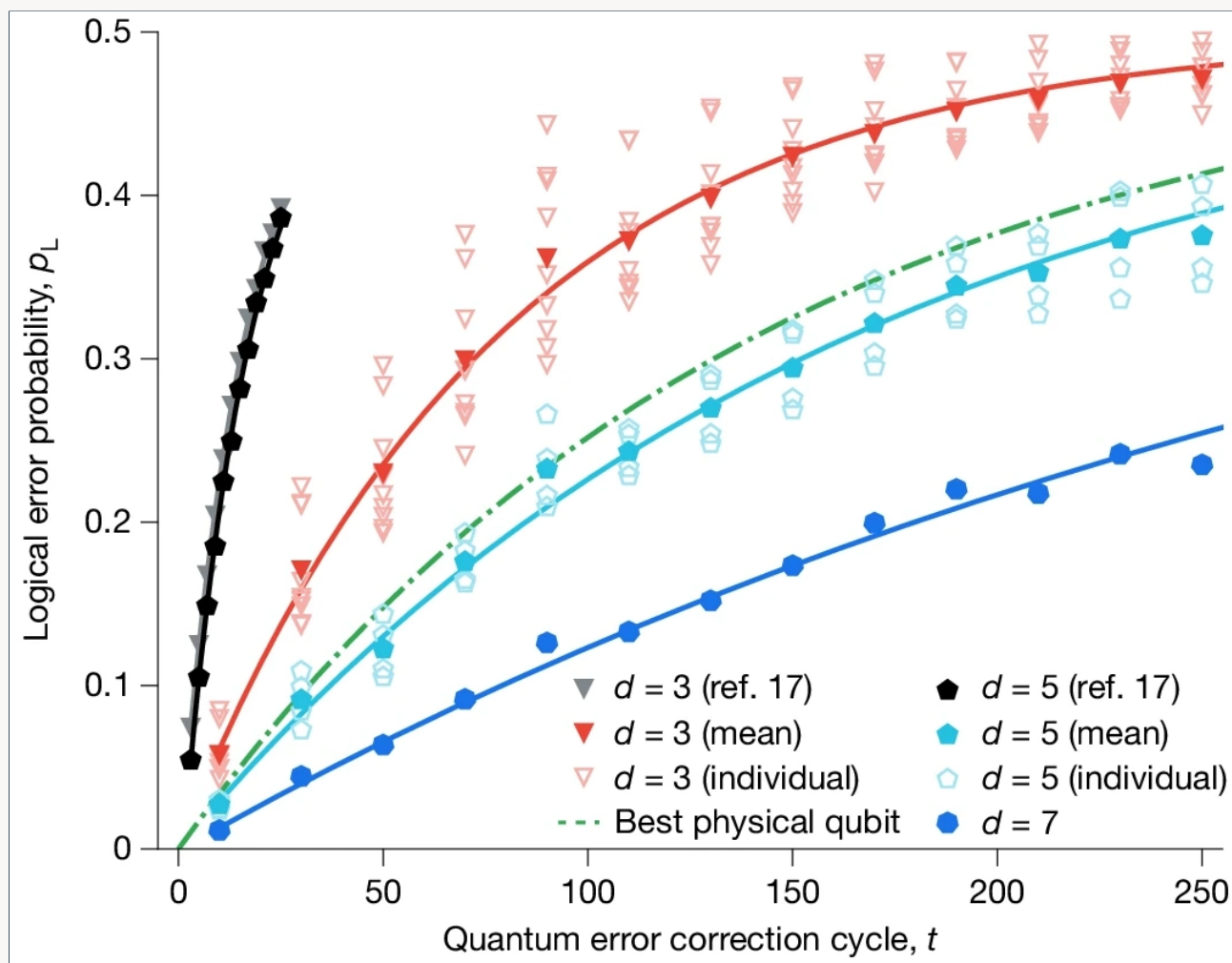
سی سال است تصحیح خطای کوانتومی بر وعده‌ای تکیه دارد که هیچ آزمایشی آن را تمیز و روشن محقق نکرده بود. نظریه می‌گوید اگر یک واحد اطلاعات کوانتومی – یک کیوبیت منطقی – را روی تعداد زیادی کیوبیت فیزیکی پرنویز پخش کنید و اگر آن کیوبیت‌های فیزیکی به اندازه کافی خوب باشند، افزودن تعداد بیشتر باید کیوبیت منطقی را بهتر کند و با بزرگ شدن کد، خطا به صورت نمایی پایین بیاید. گیر کار همان «اگر» است: پایین یک آستانه بحرانی نویز، کیوبیت بیشتر کمک می‌کند؛ بالای آن، کیوبیت بیشتر فقط نویز بیشتری می‌آورد. همه آزمایش‌های قبلی در سمت نادرست این خط بودند یا نتوانسته بودند روند را تمیز نشان دهند. بزرگ کردن کد وضع را بدتر می‌کرد، نه بهتر.

در دسامبر ۲۰۲۴، AI Quantum Google گزارش کرد نخستین نمایش روشن از رژیم دیگر را به دست آورده است. Willow، روی نسل جدید پردازنده‌های ابرسانی آن‌ها، حافظه‌های code surface با فاصله کد ۳، ۵ و ۷ ساختند و دیدند هر بار کد بزرگ‌تر شد نرخ خطای منطقی کاهش یافت – به اندازه ضریب $\Lambda = 14.2 \pm 0.2$ به ازای هر دو واحد افزایش فاصله. بزرگ‌ترین نمونه، حافظه فاصله-۷ با ۱۰۱ کیوبیت، یک کیوبیت منطقی را با خطای $143\% \pm 0.0\%$ در هر چرخه تصحیح نگه داشت و – تیتز درون تیتز – از بهترین کیوبیت فیزیکی خودش هم بیشتر عمر کرد، به ضریب 4.2 ± 3.0 . این وضعیت «beyond» (این وضعیت «beyond» breakeven نامیده می‌شود و نخستین بار است که کل هزینه تصحیح خطا روی این سخت‌افزار از نظر طول عمر خودش را جبران کرده است.

این یک نقطه عطف واقعی است و باید دقیق گفت چه نوع نقطه عطفی. اثباتی است که scaling اکنون در جهت درست می‌رود. یک رایانه کوانتومی عملیاتی نیست و مقاله هم چنین ادعایی ندارد.

«surface»، «distance» code و «below» threshold یعنی چه

کدگذاری فیزیکی کیوبیت زیادی تعداد روی که است کوانتومی اطلاعات از محافظت‌شده‌ای واحد منطقی کیوبیت یک اضافی کیوبیت‌های آن در است؛ دویعدی شبکه یک روی کدگذاری این برای خاصی روش surface code می‌شود. کمترین نیست: فیزیکی فاصله d کد فاصله می‌کنند. بررسی ذخیره‌شده اطلاعات زدن برهم بدون را خطاها مدام «اندازه‌گیری» بزرگ‌تر d کنند. خراب شود متوجه کد آنکه بدون را منطقی کیوبیت می‌توانند که است درست جای گرفته‌ای خطاهای تعداد تا بیشتری، هم‌زمان خطاهای ۱ و $2d^2$ (تقریباً می‌کند مصرف بیشتری فیزیکی کیوبیت – مقاوم‌تر و بزرگ‌تر patch یعنی را هم‌زمان خطای ۳ و ۲، به ترتیب ۷، و ۳، فاصله‌های آزموده‌شده، اندازه سه بنابراین می‌کند. اصلاح را $1/2$ ، $d - 2$) کیوبیت ۱۰۱ عمل در Google فاصله-۷ حافظه – می‌کنند مصرف فیزیکی کیوبیت ۹۷ و ۴۹، تقریباً و می‌کنند اصلاح گابی. حداقل این از بیشتر کمی داشت، عبارت «below» threshold بخش حیاتی است. تصحیح خطا فقط وقتی کمک می‌کند که نرخ خطای فیزیکی زیر یک مقدار بحرانی باشد؛ آنجا هر افزایش فاصله نرخ خطای منطقی را به صورت نمایی سرکوب می‌کند. ضریب سرکوب Λ این را اندازه می‌گیرد – $1 < \Lambda$ یعنی بزرگ کردن کد کمک می‌کند و هرچه Λ بزرگ‌تر باشد بهتر است. Google مقدار $\Lambda \approx 14.2$ گزارش می‌کند، یعنی هر دو واحد افزایش فاصله تقریباً نرخ خطای منطقی را نصف کرده است. اینکه Λ با فاصله امنی بالاتر از ۱ است، تمام نکته نتیجه است.



چگونگی انباشته شدن خطای منطقی در چرخه‌های تصحیح برای حافظه‌های فاصله ۳، ۵ و ۷ (از بالا به پایین). خطی که باید نگاه کرد خط سبز نقطه چین است — بهترین کیوبیت فیزیکی منفرد روی تراشه. حافظه فاصله ۷ (آبی، پایین‌ترین) خطا را کندتر از آن خط جمع می‌کند، بنابراین کیوبیت کدگذاری شده از بهترین کیوبیت فیزیکی سازنده خودش بیشتر عمر می‌کند «beyond — breakeven» به ضریب $4.2 \times$. این نتیجه طول عمر است؛ خود سرکوب below-threshold ($14.2 = \Lambda$) هنگام رشد کد از فاصله ۳ به ۵ به ۷ در داده‌های متن دیده می‌شود.

Google Quantum AI and Collaborators / Nature CC BY-NC-ND 0.4

نویسندگان چه کردند

- حافظه‌های code surface را روی دو تراشه Willow ساختند: یک پردازنده ۱۰۵ کیوبیتی که کدهای فاصله ۳، ۵ و ۷-آزمون scaling را اجرا کرد (بزرگ‌ترین، حافظه فاصله ۷ با ۱۰۱ کیوبیت و ۴۹ کیوبیت داده)، و یک پردازنده ۷۲ کیوبیتی که حافظه فاصله ۵-با decoder بلادرنگ و همچنین repetition های code فاصله بالا را اجرا کرد.
- اندازه گرفتند خطای منطقی در هر چرخه با افزایش فاصله کد از ۳ به ۵ به ۷ چگونه تغییر می‌کند و ضریب سرکوب Λ را استخراج کردند.
- طول عمر کیوبیت منطقی را با بهترین کیوبیت فیزیکی منفرد روی همان تراشه مقایسه کردند تا breakeven را بیازمایند.
- کد فاصله ۵-را با decoder بلادرنگ اجرا کردند — سخت‌افزار کلاسیکی که چک‌های خطا را همان‌قدر سریع که تولید می‌شوند تفسیر می‌کند — تا یک میلیون چرخه، برای نشان دادن اینکه تصحیح خطا می‌تواند با دستگاه همگام بماند.
- code repetition های ساده‌تر را تا فاصله ۲۹ گسترش دادند تا منابعی نادر و عمیق خطا را که کف عملکرد را تعیین می‌کنند

پیدا کنند.

چه پیدا کردند

- کد زیر آستانه است. خطای منطقی در هر چرخه با هر دو واحد افزایش فاصله به ضریب $\Lambda = 14.2 \pm 02.0$ کاهش یافت — سرکوب نمایی تمیز، رفتاری که نظریه وعده داده بود و هیچ پردازنده‌ای پیش تر به طور قطعی نشان نداده بود.
- حافظه فاصله-7 به خطای $143\% \pm 003\% \cdot 0$ در هر چرخه رسید و 4.2 ± 3.0 برابر بیشتر از بهترین کیوبیت فیزیکی اش عمر کرد — breakeven, beyond
- **decoder** بلا درنگ از پس کار برآمد. در فاصله 5، میانگین latency آن 63 میکروثانیه در برابر زمان چرخه 11 میکروثانیه بود و برای یک میلیون چرخه پایدار ماند — تصحیح خطا زنده اجرا شد، نه فقط در تحلیل پس از آزمایش.
- یک منبع خطای نادر و عمیق باقی مانده است. در آزمون‌های code repetition، عملکرد در نهایت با های burst خطای همبسته محدود شد که تقریباً ساعتی یک بار رخ می‌دادند (حدود یک مورد در هر 3×10^9 چرخه) و کف خطایی نزدیک 10^{-10} ایجاد می‌کردند؛ منشأ آن به گفته نویسندگان هنوز فهمیده نشده است.

این مقاله چه چیزی را اثبات نمی‌کند

- این یک رایانه کوانتومی در حال محاسبه نیست. یک حافظه کوانتومی است: یک کیوبیت منطقی را ذخیره و محافظت می‌کند. عملیات منطقی (gate) میان کیوبیت‌های منطقی انجام نمی‌دهد و الگوریتمی اجرا نمی‌کند.
- فقط یک کیوبیت با ماشین مفید فاصله ندارد. یک کیوبیت منطقی فاصله-7 حدود 10^1 کیوبیت فیزیکی مصرف می‌کند؛ نرخ خطای حدود 1٪. در هر چرخه هنوز بسیار بالاتر از حدود 10^{-6} تا 10^{-10} مورد نیاز الگوریتم‌های واقعی است. بستن این شکاف یعنی رفتن به فاصله‌های بسیار بزرگ‌تر — کیوبیت فیزیکی بسیار بیشتر برای هر کیوبیت منطقی — و الگوریتم‌های مفید به هزاران کیوبیت منطقی هم‌زمان نیاز دارند. بودجه کیوبیت فیزیکی آن به میلیون‌ها می‌رسد.
- عبارت «اگر مقیاس داده شود» واقعاً بار نتیجه را حمل می‌کند. نتیجه‌گیری خود مقاله این است که عملکرد دستگاه، اگر مقیاس داده شود، می‌تواند نیازهای الگوریتم‌های بزرگ را برآورده کند. درست بودن روند برای یک کیوبیت منطقی با ساخت دستگاه مقیاس‌یافته یکی نیست و هیچ چیز تضمین نمی‌کند روند تا اندازه‌های بسیار بزرگ‌تر حفظ شود.
- کف خطای توضیح داده نشده یک مشکل زنده است. های burst همبسته‌ای که عملکرد code repetition را محدود می‌کنند، به گفته نویسندگان چند مرتبه بزرگی پیش از انتظارند و تا وقتی فهمیده نشوند کاربردهای تحمل خطای بزرگ‌تر را ناممکن می‌کنند — نقیصی باز و صریحاً بیان شده، نه جزئیاتی حل شده.
- درباره شکستن رمزنگاری یا «quantum supremacy» برای کارهای مفید چیزی نمی‌گوید. آن‌ها به ماشین تحمل خطای کامل نیاز دارند که این نتیجه فقط یکی از سنگ‌های بنیاد است.

قدرت شواهد چقدر است؟

- ادعای مرکزی محکم و مهم است. کارکرد زیر آستانه با سرکوب نمایی تمیز در سه فاصله کد، همراه با طول عمر beyond-breakeven و decoder بلا درنگ عملی، دقیقاً ترکیبی است که حوزه مدت‌ها در پی آن بود و مستقیماً نشان داده شده نه استنباط. این مصنوع hype نیست؛ یک نتیجه مهندسی واقعی از گروهی پیشرو است.

- نویسندگان درباره دامنۀ نتیجه محتاط‌اند. آن را حافظۀ *below-threshold* معرفی می‌کنند، خودشان کف خطای همبسته ناشناخته را برجسته می‌کنند و آینده را با همان «اگر مقیاس داده شود» آشکار مشروط می‌گذارند. اغراق، وقتی رخ می‌دهد، بیشتر در پوششی است که «حافظۀ تصحیح خطا با بزرگ شدن بهتر شد» را به «رایانش کوانتومی رسید» گرد می‌کند.
- وضعیت صادقانه یک گام بنیادی خوب برداشته شده است. یک کیوبیت منطقی، آن قدر خوب محافظت شده که افزودن افزودنی بالاخره کمک می‌کند — با راهی بلند، دشوار و تضمین نشده برای *scaling* های *gate* منطقی و خطاهای توضیح داده نشده در پیش.

چرا مهم است

رایانش کوانتومی تحمل خطا همیشه حس مرغ و تخم مرغ داشته است: ماشین مفید به نرخ خطایی نیاز دارد که هیچ کیوبیت فیزیکی نمی‌تواند به آن برسد، و راه حل — تصحیح خطا — فقط وقتی کار می‌کند که سخت افزار از قبل آن قدر خوب باشد که زیر آستانه قرار گیرد. عبور از این خط، حتی یک بار و حتی روی فقط یک کیوبیت منطقی، پرسش را از «اصلاً ممکن است؟» به «تا جفا و با چه سرعتی می‌شود مقیاسش داد؟» تغییر می‌دهد. این تغییر واقعی و معنادار است و دلیل توجه به نتیجه همین است.

اما همان دقتی که نتیجه را قابل اعتماد می‌کند باید داستان پیرامونش را هم مهار کند. این نخستین آجریک پی است، خوب گذاشته شده. ساختمان نیست و کسانی که آجر را گذاشتند اولین کسانی‌اند که همین را می‌گویند. راه درست دنبال کردن رایانش کوانتومی در چند سال آینده همین منحنی کم‌زرق و برق است: آیا ضریب سرکوب با بزرگ شدن کدها باقی می‌ماند، آیا های *gate* منطقی به تمیزی حافظۀ منطقی انجام می‌شوند، و آیا آن خطای مرموز ساعتی بالاخره توضیح داده می‌شود.

خلاصۀ تمیز

AI Quantum Google برای نخستین بار به طور تمیز نشان داد یک حافظۀ کوانتومی *code surface* می‌تواند زیر آستانه کار کند: با بزرگ شدن کد از فاصله ۳ به ۵ به ۷، نرخ خطای منطقی به طور نمایی کاهش یافت (حدود $14.2 \times$ به ازای هر دو واحد افزایش)، و بزرگ ترین نمونه، حافظۀ فاصله ۷- با ۱۰۱ کیوبیت، از بهترین کیوبیت فیزیکی خودش بیشتر عمر کرد — *breakeven beyond* — در حالی که تصحیح خطا به صورت بلادرنگ اجرا می‌شد. این یک نقطۀ عطف واقعی و دیرانتظار در مهندسی رایانه های کوانتومی است. اما همچنان فقط یک کیوبیت منطقی در نقش حافظه است، با نرخ خطایی بسیار بالاتر از نیاز الگوریتم های واقعی، بدون عملیات منطقی، با کف خطای ناشناخته ای که خود نویسندگان برجسته کرده اند و چندین مرتبۀ بزرگی *scaling* در پیش. یک آستانه واقعاً رد شد — نه اینکه رایانه تحویل شده باشد.

منابع

920 ,638 Nature threshold, code surface the below correction error Quantum Collaborators, and AI Quantum Google (2025)

<https://doi.org/10.1038/s41586-024-08449-y>

a corrects — (2026) E5 ,653 Nature threshold, code surface the below correction error Quantum Correction: Author results (1 (Fig. below-threshold the affect not does keys; legend and label axis 3a Fig.

<https://www.nature.com/articles/s41586-026-10559-8>