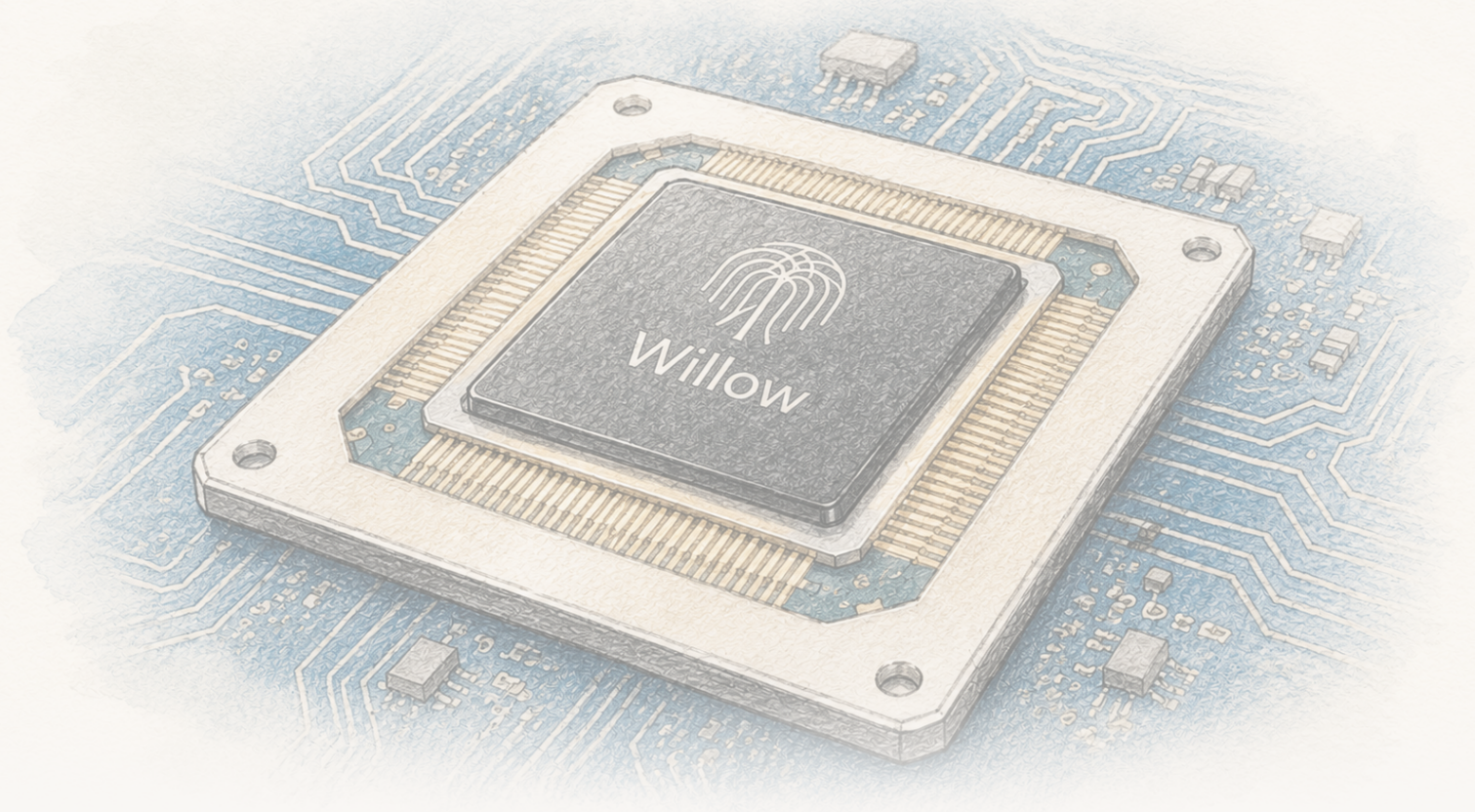




The CLEAN PAPER

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

ذاكرة كمومية أصبحت أخيراً أفضل كلما كبرت — المحطة الحقيقية، والآلة التي ليست

أظهرت *AI Quantum Google* بوضوح للمرة الأولى أن ذاكرة كمومية من نوع *code surface* تستطيع العمل تحت العتبة: كلما نما الرمز من المسافة 3 إلى 5 إلى 7 انخفض معدل الخطأ المنطقي أسياً (نحو $14.2 \times$ لكل درجتين)، وعاشت أكبر ذاكرة، *distance-7* من 101 كيوبت، أطول من أفضل كيوبت فيزيائي لديها — متجاوزة نقطة التعادل — مع تشغيل تصحيح الأخطاء في الزمن الحقيقي. إنها محطة هندسية طال انتظارها. لكنها كيوبت منطقي واحد يعمل كذاكرة، بمعدل خطأ لا يزال بعيداً عما تحتاجه الخوارزميات الحقيقية، ومن دون عمليات منطقية، ومع أرضية خطأ غير مفسرة ينه إليها الباحثون، ومراتب مقدار كثيرة من التوسع ما زالت أماناً. عتبة عُبرت، لا حاسوباً سُلّم.

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

اللحظة التي انحنى فيها المنحنى في الاتجاه الصحيح

على مدى ثلاثين عاماً، استند تصحيح الأخطاء الكمومية إلى وعد لم يتحقق بوضوح قط. تقول النظرية إنه إذا وزعت وحدة واحدة من المعلومات الكمومية — كيوبت منطقي واحد — على عدد كبير من الكيوبتات الفيزيائية الصاخبة، وكانت تلك الكيوبتات الفيزيائية جيدة بما يكفي، فإن إضافة المزيد منها ينبغي أن تجعل الكيوبت المنطقي أفضل، مع انخفاض الأخطاء أسياً كلما كبر الرمز. المشكلة في كلمة «إذا»: تحت عتبة ضوضاء حرجية تساعد الكيوبتات الإضافية؛ وفوقها تضيف فقط ضوضاء أكثر. كانت كل التجارب السابقة على الجانب الخطأ من هذا الخط، أو لم تظهر الاتجاه بوضوح. كان تكبير الرمز يجعل الأمور أسوأ لا أفضل.

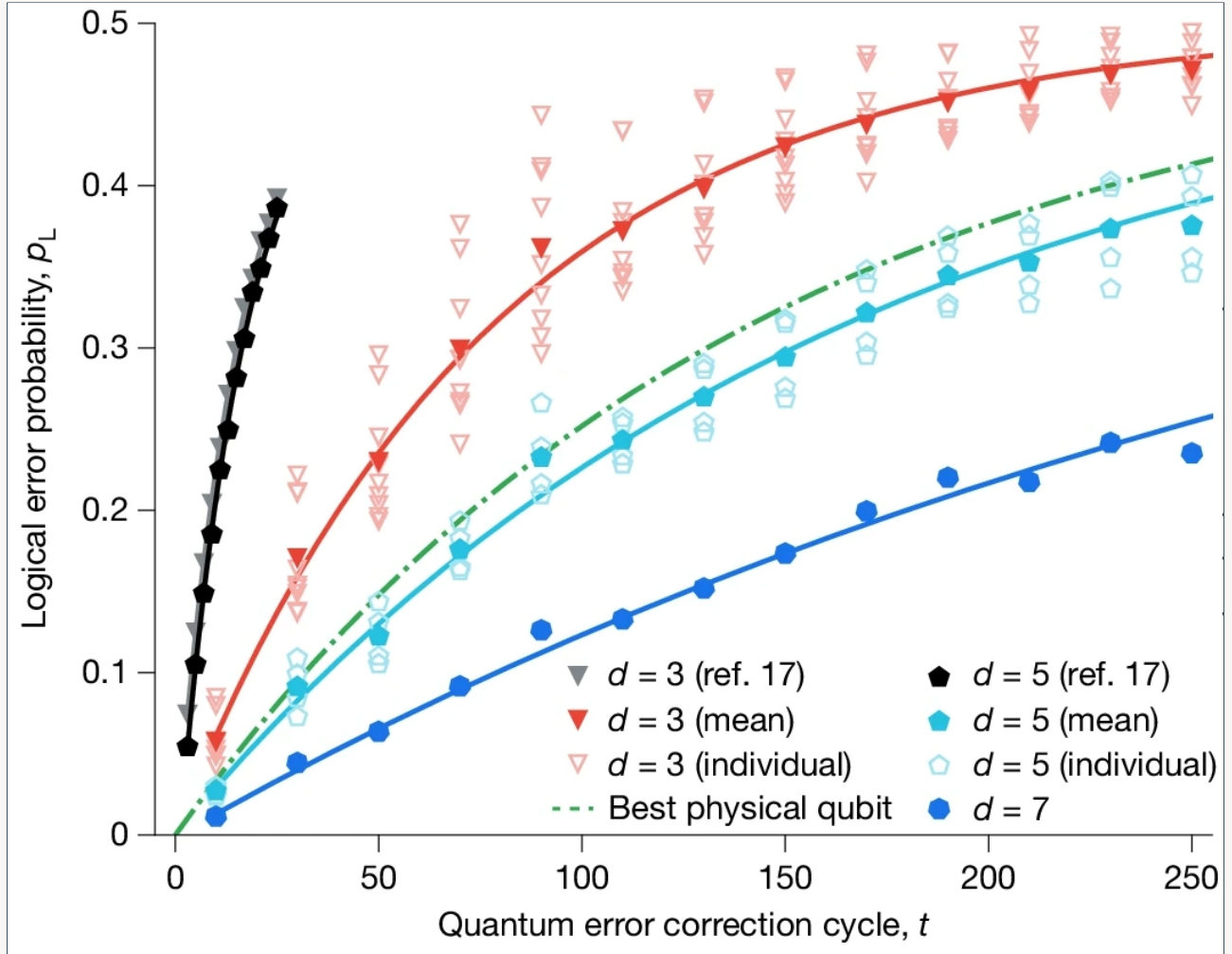
في ديسمبر 2024، أعلنت Google AI Quantum أول عرض واضح للنظام الآخر. على Willow أحدث جيل لديها من المعالجات فائقة التوصيل، بنت ذاكرات code surface بمسافات رمز 3 و 5 و 7، وشاهدت معدل الخطأ المنطقي ينخفض كلما كبر الرمز — بعامل $\Lambda = 14.2 \pm 0.2$ لكل زيادة بمقدار درجتين في المسافة. وكانت أكبر ذاكرة، distance-7 من 101 كيوبت، تحفظ كيوبتاً منطقياً بخطأ $143\% \pm 0.03\%$ لكل دورة تصحيح، والأهم أنها عاشت أطول من أفضل كيوبت فيزيائي مكوّن لها بعامل 4.2 ± 3.0 . يسمى هذا «تجاوز نقطة التعادل» beyond، breakeven وهي المرة الأولى التي يدفع فيها جهاز تصحيح الأخطاء كامل تكلفته على هذه العتاد.

هذه محطة حقيقية، ويستحق الأمر الدقة بشأن نوعها. إنها إثبات أن التوسع أصبح يسير في الاتجاه الصحيح. وليست حاسوباً كمومياً عاملاً، والدراسة لا تدعي ذلك.

ماذا تعني «شفرة السطح» و«المسافة» و«دون العتبة»؟

لهذا محددة طريقة surface code و فيزيائية. كيوبتات عدة عبر مشفرة الكمومية المعلومات من محمية وحدة المنطقي الكيوبت مسافة أما الخزنة. المعلومة إفساد دون من باستمرار الأخطاء إضافية قياس كيوبتات تفحص حيث الأبعاد، ثنائية شبكة على الترميز من المنطقي الكيوبت إفساد يستطيع الذي المناسبة المواقع في الموضوع الأخطاء من عدد أصغر بل فيزيائية؛ مسافة فليست d الرمز وتصحيح 1، $2d^2$ تقريباً أكثر، فيزيائية كيوبتات تستهلك — متانة وأكثر أكبر الرقعة أصبح d كبرت كلما الرمز. يكتشفه أن دون أخطاء 3 و 2 وتصحيح 7 و 5 والمسافات هنا، المختبرة الثلاثة الأجام فإن لذلك منها $1/2 - d$ (حتى أكثر، متزامنة أخطاء قليلاً أعلى أي 101، عدد Google بنتها التي distance-7 ذاكرة استخدمت فيزيائياً؛ كيوبتاً 97 و 49 و 17 تقريباً وتستهلك متزامنة الأدنى. المدرسي الحد هذا من

عبارة «تحت العتبة» هي المفتاح. لا يساعد تصحيح الأخطاء إلا إذا كان معدل الخطأ الفيزيائي تحت قيمة حرجية؛ عندها تؤدي كل زيادة في المسافة إلى كبح معدل الخطأ المنطقي أسياً. يقيس عامل الكبح Λ ذلك — إذا كان $\Lambda < 1$ فإن تكبير الرمز يفيد، وكلما ارتفع كان أفضل. تعلن Google عن $\Lambda \approx 14.2$ ، أي إن كل زيادة بمقدار درجتين في المسافة خفضت الخطأ المنطقي إلى نحو النصف. كون Λ أعلى بوضوح من 1 هو النتيجة كلها.



كيف يتراكم الخطأ المنطقي عبر دورات التصحيح لذاكرات distance-3 و-5 و-7 من الأعلى إلى الأسفل. الخط الذي ينبغي مراقبته هو الأخضر المتقطع — أفضل كيوبت فيزيائي منفرد على الشريحة. ذاكرة distance-7 باللون الأزرق وفي الأسفل، تراكم الخطأ ببطء أكبر من ذلك الخط، ولذلك يعيش الكيوبت المشفر أطول من أفضل كيوبت فيزيائي بُني منه — «متجاوزاً نقطة التعادل» بعامل $4.2 \times$. هذه نتيجة العمر، أما كبح الأخطاء تحت العتبة نفسه ($\Lambda = 14.2$ عند الانتقال من المسافة 3 إلى 5 إلى 7) فتظهره الأرقام في النص.

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

ما الذي فعله الباحثون؟

- بنوا ذاكرات code surface على شريحتي Willow: معالج بـ 105 كيوبتات شغل رموز المسافة 3 و 5 و 7 وراء اختبار التوسع — وكان أكبرها ذاكرة distance-7 من 101 كيوبت، منها 49 كيوبت بيانات — ومعالج بـ 72 كيوبتاً شغل ذاكرة distance-5 مع مفكك تشفير آني، إضافة إلى codes repetition ذات المسافات العالية.
- قاسوا كيف يتغير الخطأ المنطقي لكل دورة عند زيادة مسافة الرمز من 3 إلى 5 إلى 7، واستخرجوا عامل الكبح Λ .
- قارنوا عمر الكيوبت المنطقي بأفضل كيوبت فيزيائي منفرد على الشريحة نفسها لاختبار «نقطة التعادل».
- شغلوا رمز distance-5 مع مفكك تشفير في الزمن الحقيقي — عتاد كلاسيكي يفسر فحوص الأخطاء بالسرعة التي تُنتج بها — حتى مليون دورة، لإظهار أن تصحيح الأخطاء يستطيع مواكبة الآلة.
- دفعوا codes repetition الأبسط حتى distance 29 للبحث عن مصادر الأخطاء النادرة والعميقة التي تضع أرضية للأداء.

ماذا وجدوا؟

- الرمز تحت العتبة. انخفض الخطأ المنطقي لكل دورة بعامل $\Lambda = 14.2 \pm 0.2$ لكل زيادة مقدارها اثنان في المسافة — كبح أسي واضح، وهو السلوك الذي وعدت به النظرية ولم يظهره أي معالج بصورة حاسمة من قبل.
- وصلت ذاكرة distance-7 إلى خطأ $143\% \pm 0.0\%$ لكل دورة، وعاشت 4.2 ± 3.0 مرة أطول من أفضل كيوبت فيزيائي لديها — متجاوزة نقطة التعادل.
- واكب فك التشفير في الزمن الحقيقي. بلغ متوسط كمون مفكك التشفير 63 ميكروثانية عند distance 5 مقابل زمن دورة 1.1 ميكروثانية، واستمر على مدى مليون دورة — أي إن تصحيح الأخطاء عمل أثناء التشغيل، لا فقط في تحليل لاحق للبيانات.
- يبقى مصدر خطأ نادر وعميق. في اختبارات code، repetition حدثت من الأداء في النهاية دفعات أخطاء مترابطة تحدث تقريباً مرة في الساعة — نحو مرة في كل 3×10^9 دورات — واضحة أرضية خطأ قرب 10^{-10} يقول الباحثون إن أصلها غير مفهوم بعد.

ما الذي لا تثبته الدراسة؟

- هذه ليست حاسوباً كمومياً يجري حساباً. إنها ذاكرة كمومية: تخزين وتحمي كيوبتاً منطقياً واحداً. لا تنفذ عمليات منطقية — بوابات — بين كيوبتات منطقية، ولا تشغل خوارزمية.
- وهي ليست «كيوبتاً واحداً فقط» بعيداً عن آلات مفيدة. يستهلك الكيوبت المنطقي distance-7 نحو 101 كيوبت فيزيائي، ومعدل خطأ $1\% \sim 0.0\%$ لكل دورة ما يزال أعلى بكثير من نحو 10^{-6} إلى 10^{-10} الذي تحتاجه الخوارزميات الحقيقية. سد هذه الفجوة يعني الانتقال إلى مسافات أكبر بكثير — كيوبتات فيزيائية أكثر لكل كيوبت منطقي — والخوارزميات المفيدة تحتاج إلى آلاف الكيوبتات المنطقية في وقت واحد. ميزانية الكيوبتات الفيزيائية لذلك تصل إلى الملايين.
- عبارة «إذا أمكن التوسع» مهمة فعلاً. استنتاج الدراسة نفسها أن أداء الجهاز، إذا تم توسيعه، قد يفي بمتطلبات الخوارزميات الكبيرة. إظهار الاتجاه الصحيح على كيوبت منطقي واحد ليس هو نفسه بناء الآلة الموسعة، ولا شيء هنا يضمن بقاء الاتجاه عند أحجام أكبر بكثير.
- أرضية الخطأ غير المفهومة مشكلة حية. دفعات الأخطاء المترابطة التي تحد أداء code repetition أكبر بمراتب مقدار مما كان متوقعاً، وبحسب الباحثين ستمنع تطبيقات أكبر مقاومة للأخطاء حتى تفهم — عيب مفتوح مصرح به، لا تفصيل محلول.
- ولا تقول الدراسة شيئاً عن كسر التشفير أو «التفوق الكمومي» في مهام مفيدة. هذه تحتاج إلى آلة كاملة مقاومة للأخطاء؛ وهذه النتيجة حجر أساس لها، لا عرضاً لها.

ما مدى قوة الأدلة؟

- الادعاء الأساسي قوي ومهم. التشغيل تحت العتبة مع كبح أسي واضح عبر ثلاث مسافات رمز، إلى جانب عمر يتجاوز نقطة التعادل ومفكك تشفير آلي عامل، هو بالضبط المزيج الذي حاول المجال الوصول إليه، وقد عُرض مباشرةً لا بالاستنتاج. هذه ليست قطعة من الضجيج الإعلامي؛ بل نتيجة هندسية حقيقية من مجموعة رائدة.
- الباحثون دقيقون بشأن نطاقها. يؤطرونها كذاكرة تحت العتبة، وينبهون بأنفسهم إلى أرضية الأخطاء المترابطة غير المفسرة، ويربطون المستقبل بعبارة «إذا تم التوسع» الواضحة. المبالغة، عندما تظهر، تأتي من التغطية التي تقرب «تحسنت ذاكرة كيوبت مصححة الأخطاء عندما كبرت» إلى «وصل الحوسبة الكمومية».
- الحالة الصداقة هي خطوة تأسيسية أخذت بوضوح. كيوبت منطقي واحد، محمي بما يكفي كي تصبح إضافة التكرار مفيدة أخيراً — مع

طريق طويل وصعب وغير مضمون بعد من التوسع والبوابات المنطقية والأخطاء غير المفسرة.

لماذا يهم هذا؟

كان للحوسبة الكمومية المقاومة للأخطاء دائماً طابع مشكلة الدجاجة والبيضة: الآلات المفيدة تحتاج إلى معدلات خطأ لا يستطيع أي كيوبت فيزيائي الوصول إليها، والحل — تصحيح الأخطاء — لا يعمل إلا إذا كان العتاد جيداً أصلاً بما يكفي ليكون تحت العتبة. عبور ذلك الخط، ولو مرة واحدة، ولو على كيوبت منطقي واحد، يغير السؤال من «هل هذا ممكن أصلاً؟» إلى «إلى أي مدى يمكن توسيعه، وبأي سرعة؟». وهذا تحول حقيقي ومهم، ولهذا تستحق النتيجة الاهتمام.

لكن العناية نفسها التي تجعل النتيجة موثوقة ينبغي أن تكبح القصة حولها. هذه أول طوبة في أساس، موضوعة جيداً. ليست المبني، ومن وضعوها هم أول من يقول ذلك. الطريقة الصحيحة لمتابعة الحوسبة الكمومية في السنوات المقبلة هي هذا المنحنى غير البراق تحديداً: هل يبقى عامل الكبح مع نمو الرموز؟ وهل يمكن تنفيذ البوابات المنطقية بنظافة الذاكرة المنطقية نفسها؟ وهل سيفهم ذلك الخطأ الغامض الذي يظهر مرة كل ساعة؟

الخلاصة الواضحة

أظهرت AI، Quantum Google، أن ذاكرة code surface كمومية تستطيع العمل تحت العتبة: عند تكبير الرمز من المسافة 3 إلى 5 إلى 7، انخفض معدل الخطأ المنطقي أسياً — بنحو $14.2 \times$ لكل درجتين — وعاشت أكبر ذاكرة، 7-distance من 101 كيوبت، أطول من أفضل كيوبت فيزيائي مكوّن لها، أي تجاوزت نقطة التعادل، بينما كان تصحيح الأخطاء يعمل في الزمن الحقيقي. هذه محطة حقيقية طال انتظارها في هندسة الحواسيب الكمومية. لكنها أيضاً كيوبت منطقي واحد يعمل كذاكرة، بمعدل خطأ ما يزال بعيداً عن احتياجات الخوارزميات الحقيقية، ومن دون أي عمليات منطقية، ومع أرضية خطأ غير مفسرة ينبه إليها الباحثون أنفسهم، وطريق توسع أمامه عدة مراتب مقدار. عتبة حقيقية عُبرت — لا حاسوب كمومي سَلَمَ جاهزاً.

المصادر

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>