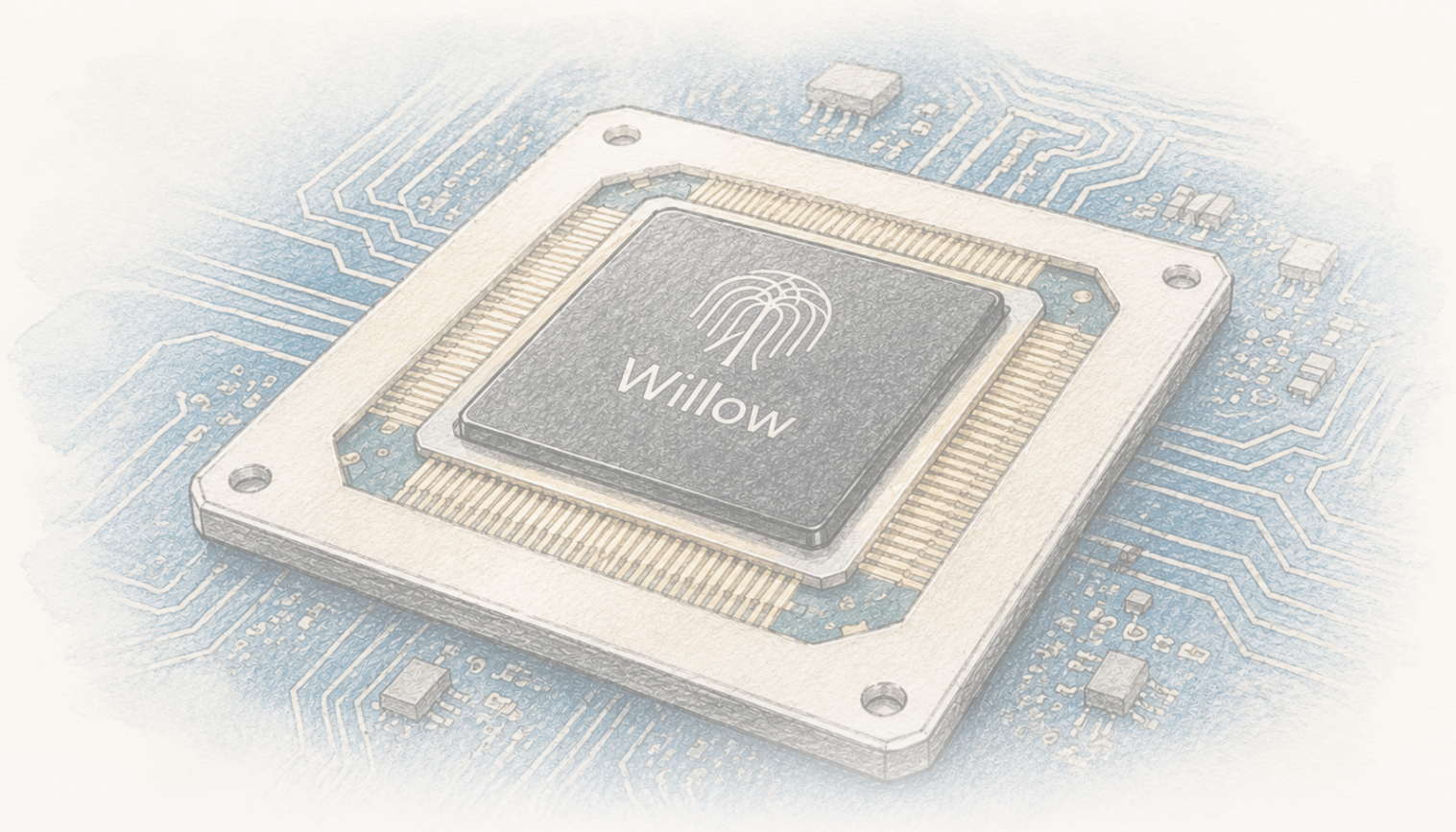




The CLEAN PAPER

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

Memori kuantum akhirnya membaik saat diperbesar — tonggak yang nyata, dan mesin yang belum ada

Google Quantum AI menunjukkan, untuk pertama kalinya dengan bersih, bahwa memori kuantum surface code dapat berjalan below threshold: ketika kode diperbesar dari distance 3 ke 5 ke 7, tingkat error logis turun secara eksponensial (sekitar $2,14\times$ setiap dua langkah), dan yang terbesar, memori distance-7 101-qubit, bertahan lebih lama daripada qubit fisik terbaiknya — beyond breakeven — dengan koreksi error berjalan real-time. Ini tonggak rekayasa yang lama dicari. Namun ia masih satu qubit logis sebagai memori, dengan tingkat error jauh dari kebutuhan algoritma nyata, tanpa operasi logis, dengan error floor yang belum dijelaskan dan ditandai penulis sendiri, serta scaling beberapa orde besaran di depan. Threshold dilintasi, bukan komputer diserahkan.

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

Saat kurva akhirnya membengkok ke arah yang benar

Selama tiga puluh tahun, quantum error correction bertumpu pada janji yang belum pernah dipenuhi dengan bersih. Teorinya mengatakan bahwa jika satu unit informasi kuantum — satu qubit *logis* — disebarkan ke banyak qubit fisik yang bising, dan jika qubit fisiknya cukup baik, maka menambah *lebih banyak* qubit harus membuat qubit logis *lebih baik*, dengan error turun secara eksponensial ketika kode membesar. Masalahnya ada pada kata “jika”: di bawah ambang noise kritis, lebih banyak qubit membantu; di atasnya, lebih banyak qubit hanya menambah noise. Semua eksperimen sebelumnya berada di sisi yang salah dari garis itu, atau gagal menunjukkan tren dengan bersih. Membesarkan kode justru memperburuk keadaan.

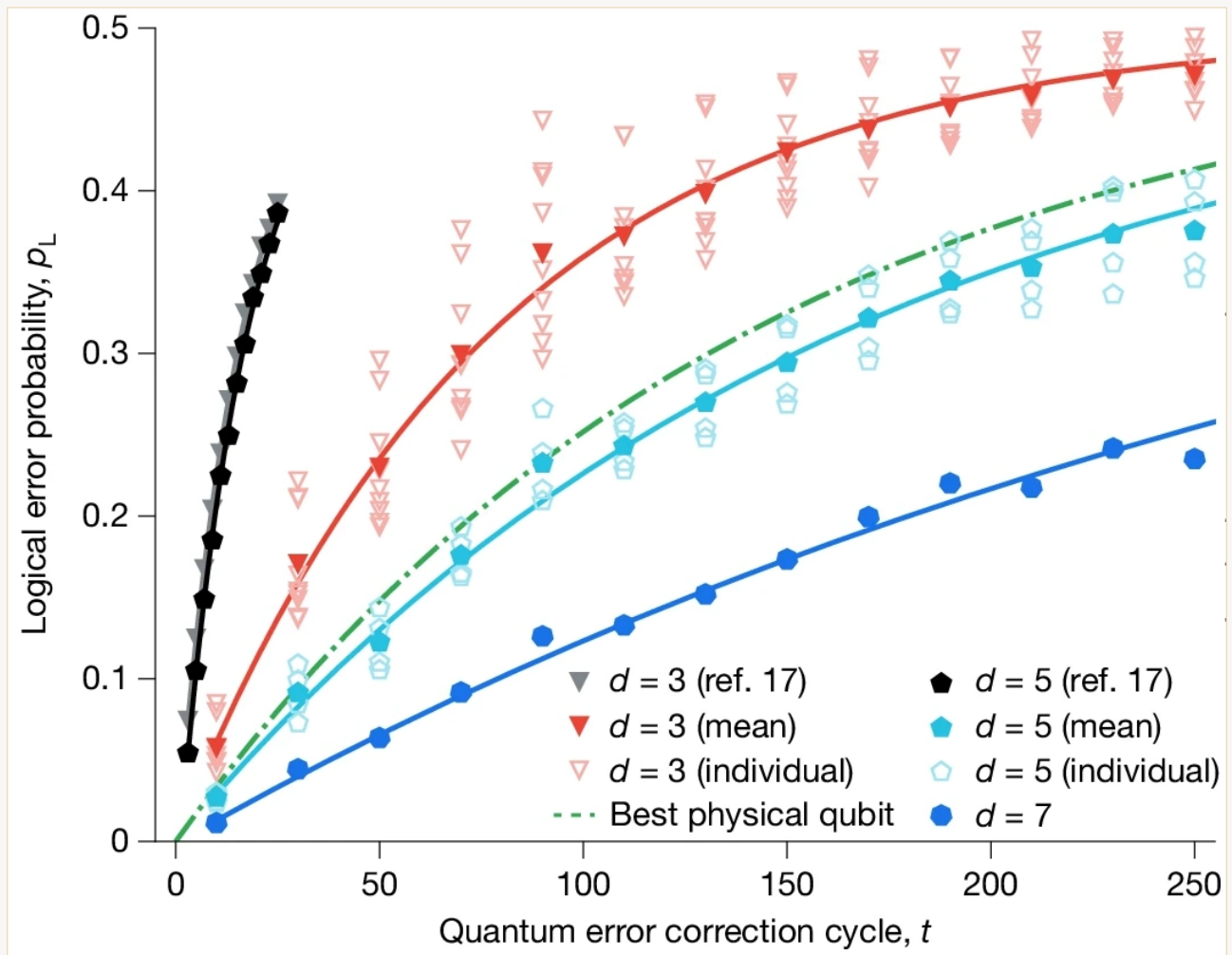
Pada Desember 2024, Google Quantum AI melaporkan demonstrasi jelas pertama dari rezim sebaliknya. Pada Willow, generasi prosesor superkonduktor terbaru mereka, tim membangun memori surface code dengan jarak kode 3, 5, dan 7, lalu melihat tingkat error logis *turun* setiap kali kode diperbesar — dengan faktor $\Lambda = 2,14 \pm 0,02$ untuk setiap kenaikan dua langkah jarak. Yang terbesar, memori distance-7 dengan 101 qubit, menyimpan satu qubit logis dengan error $0,143\% \pm 0,003\%$ per siklus koreksi, dan — judul di dalam judul — bertahan *lebih lama daripada qubit fisik terbaiknya sendiri*, dengan faktor $2,4 \pm 0,3$. Ini disebut “beyond breakeven”, dan merupakan pertama kalinya seluruh aparat koreksi error benar-benar membayar overhead-nya sendiri pada perangkat keras ini.

Ini tonggak nyata, dan penting untuk tepat mengenai jenis tonggakunya. Ini bukti bahwa peningkatan skala akhirnya bergerak ke arah yang benar. Ini bukan komputer kuantum yang berfungsi, dan makalah juga tidak mengklaim demikian.

Apa arti “surface code”, “distance”, dan “below threshold”

Sebuah **qubit logis** adalah satu unit informasi kuantum yang dilindungi dengan mengodekannya di banyak qubit fisik. **Surface code** adalah salah satu cara melakukan pengodean itu pada kisi 2D, tempat qubit “ukur” tambahan terus memeriksa error tanpa mengganggu informasi yang disimpan. **Jarak kode** d bukan jarak fisik: ia adalah jumlah minimum error yang ditempatkan dengan tepat yang dapat merusak qubit logis tanpa terdeteksi kode. d lebih besar berarti patch lebih besar dan lebih kuat — memakai lebih banyak qubit fisik (kira-kira $2d^2 - 1$) dan memperbaiki lebih banyak error simultan, hingga $(d - 1)/2$. Jadi tiga ukuran yang diuji di sini, distance 3, 5, dan 7, memperbaiki 1, 2, dan 3 error simultan dan secara kasar memakai 17, 49, dan 97 qubit fisik — memori distance-7 Google memakai 101, sedikit di atas minimum textbook ini.

“**Below threshold**” adalah frasa kunci. Koreksi error hanya membantu jika tingkat error fisik berada di bawah nilai kritis; di sana, setiap kenaikan distance menekan error logis secara eksponensial. Faktor penekanan Λ mengukurnya — $\Lambda > 1$ berarti membesarkan kode membantu, dan semakin tinggi Λ , semakin baik. Google melaporkan $\Lambda \approx 2,14$, artinya setiap kenaikan dua langkah distance memotong error logis kira-kira menjadi separuh. Bahwa Λ nyaman berada di atas 1 adalah inti hasilnya.



Bagaimana error logis terakumulasi sepanjang siklus koreksi untuk memori distance-3, -5, dan -7 (dari atas ke bawah). Garis yang perlu diperhatikan adalah garis putus-putus hijau — *qubit fisik tunggal terbaik* pada chip. Memori distance-7 (biru, paling bawah) mengakumulasi error lebih lambat daripada garis itu, sehingga qubit terencode bertahan lebih lama daripada qubit fisik terbaik yang menyusunnya — “beyond breakeven”, dengan faktor $2,4\times$. Ini adalah hasil lifetime; penekanan below-threshold sendiri ($\Lambda = 2,14$ ketika kode tumbuh dari distance 3 ke 5 ke 7) berada pada angka dalam teks.

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

Apa yang dilakukan para penulis

- Membangun memori surface code pada dua chip Willow: prosesor 105-qubit yang menjalankan kode distance-3, -5, dan -7 untuk uji peningkatan skala (yang terbesar adalah memori distance-7 101-qubit dengan 49 data qubit), serta prosesor 72-qubit yang menjalankan memori distance-5 dengan decoder real-time dan repetition code ber-distance tinggi.
- Mengukur bagaimana error logis *per siklus* berubah saat distance kode dinaikkan dari 3 ke 5 ke 7, lalu mengekstrak faktor penekanan Λ .
- Membandingkan lifetime qubit logis dengan qubit fisik individu terbaik pada chip yang sama, untuk menguji “breakeven”.

- Menjalankan kode distance-5 dengan **decoder real-time** — perangkat keras klasik yang menafirkan pemeriksaan error secepat data dihasilkan — hingga satu juta siklus, untuk menunjukkan koreksi error dapat mengikuti mesin.
- Mendorong **repetition code** yang lebih sederhana hingga distance 29 untuk mencari sumber error langka dan dalam yang menetapkan rantai performa.

Apa yang mereka temukan

- **Kode berada di bawah threshold.** Error logis per siklus turun dengan $\Lambda = 2,14 \pm 0,02$ untuk setiap kenaikan distance sebanyak dua — penekanan eksponensial yang bersih, perilaku yang dijanjikan teori dan belum pernah ditunjukkan definitif oleh prosesor sebelumnya.
- **Memori distance-7 mencapai error $0,143\% \pm 0,003\%$ per siklus, dan hidup $2,4 \pm 0,3$ kali lebih lama** daripada qubit fisik terbaiknya — beyond breakeven.
- **Decoding real-time mampu mengikuti.** Decoder memiliki latensi rata-rata 63 mikrodetik pada distance 5 dibanding waktu siklus 1,1 mikrodetik, dipertahankan selama sejuta siklus — koreksi error berjalan secara live, bukan hanya dianalisis setelah eksperimen selesai.
- **Masih ada sumber error langka dan dalam.** Pada uji repetition code, performa akhirnya dibatasi burst error terkorelasi yang terjadi kira-kira **sekali per jam** (sekitar satu dalam setiap 3×10^9 siklus), menetapkan error floor dekat 10^{-10} yang asalnya, menurut penulis, **belum dipahami**.

Apa yang *tidak* dibuktikan

- Ini **bukan** komputer kuantum yang sedang menghitung. Ini *memori* kuantum: menyimpan dan melindungi satu qubit logis. Ia tidak melakukan operasi logis (gate) antarqubit logis dan tidak menjalankan algoritma.
- Ini **bukan** tinggal satu qubit lagi menuju mesin berguna. Satu qubit logis distance-7 memakai sekitar 101 qubit fisik; tingkat error $\sim 0,1\%$ per siklus masih jauh di atas sekitar 10^{-6} hingga 10^{-10} yang dibutuhkan algoritma nyata. Menutup kesenjangan itu berarti distance jauh lebih besar — lebih banyak qubit fisik per qubit logis — sementara algoritma berguna membutuhkan *ribuan* qubit logis sekaligus. Anggaran qubit fisiknya dapat mencapai jutaan.
- Frasa “**jika diskalakan**” benar-benar **bekerja keras**. Kesimpulan makalah adalah performa perangkat, *jika diskalakan*, dapat memenuhi kebutuhan algoritma besar. Menunjukkan tren yang benar pada satu qubit logis tidak sama dengan membangun mesin skala tersebut, dan tidak ada di sini yang menjamin tren bertahan ke ukuran jauh lebih besar.
- **Error floor yang belum dijelaskan adalah masalah aktif.** Burst terkorelasi yang membatasi repetition code, menurut penulis, beberapa orde besaran lebih besar daripada perkiraan dan akan menghalangi aplikasi fault-tolerant yang lebih besar sampai dipahami — kelemahan terbuka, dinyatakan langsung, bukan detail yang sudah selesai.
- Hasil ini tidak mengatakan apa pun tentang **memecahkan enkripsi atau “quantum supremacy” untuk tugas berguna**. Itu membutuhkan mesin fault-tolerant penuh yang sedang dibangun fondasinya oleh hasil ini, bukan yang didemonstrasikan di sini.

Seberapa kuat buktinya?

- **Klaim inti kuat dan penting.** Operasi below-threshold dengan penekanan eksponensial bersih pada tiga distance, ditambah lifetime beyond-breakeven dan decoder real-time yang bekerja, adalah kombinasi tepat yang selama ini dicari bidang ini, dan semuanya didemonstrasikan langsung, bukan diinferensikan. Ini bukan artefak hype; ini hasil rekayasa nyata dari kelompok terdepan.
- **Para penulis hati-hati tentang cakupannya.** Mereka membingkainya sebagai *memori* below-threshold, menandai sendiri error floor terkorelasi yang tidak dipahami, dan menggantung masa depan pada frasa mencolok “jika diskalakan”. Overreach biasanya muncul pada pemberitaan yang membulatkan “qubit memori yang dikoreksi error membaik saat diperbesar” menjadi “komputasi kuantum sudah tiba”.
- **Status jujurnya adalah langkah fondasional yang diambil dengan bersih.** Satu qubit logis, terlindungi cukup baik sehingga redundansi tambahan akhirnya membantu — dengan jalan panjang, sulit, dan belum terjamin berupa scaling, logical gate, dan error yang belum menjelaskan masih di depan.

Mengapa ini penting

Komputasi kuantum fault-tolerant selalu terasa seperti masalah ayam dan telur: mesin yang berguna membutuhkan tingkat error yang tidak dapat dicapai satu qubit fisik, sementara solusinya — koreksi error — hanya bekerja jika perangkat keras sudah cukup baik untuk berada di bawah threshold. Melintasi garis itu, walau sekali dan hanya pada satu qubit logis, mengubah pertanyaan dari “apakah ini mungkin sama sekali?” menjadi “seberapa jauh dapat diskalakan, dan seberapa cepat?” Itu perubahan nyata dan bermakna, dan alasan hasil ini layak diperhatikan.

Namun kehati-hatian yang membuat hasilnya dapat dipercaya juga harus menahan cerita di sekelilingnya. Ini bata pertama fondasi, dipasang dengan baik. Ini bukan gedungnya, dan orang yang memasangnya adalah yang pertama mengatakannya. Cara tepat mengikuti komputasi kuantum beberapa tahun ke depan justru lewat kurva yang tidak glamor ini: apakah faktor penekanan bertahan ketika kode membesar, apakah logical gate dapat dilakukan sebersih logical memory, dan apakah error misterius sekali per jam itu akhirnya dijelaskan.

Singkatnya

Google Quantum AI menunjukkan, untuk pertama kalinya dengan bersih, bahwa memori kuantum surface code dapat beroperasi *below threshold*: ketika kode diperbesar dari distance 3 ke 5 ke 7, tingkat error logis turun secara eksponensial (sekitar $2,14 \times$ setiap dua langkah), dan yang terbesar, memori distance-7 101-qubit, bertahan lebih lama daripada qubit fisik terbaiknya — beyond breakeven — sementara koreksi error berjalan real-time. Ini tonggak nyata yang lama dicari dalam rekayasa komputer kuantum. Namun ia juga hanya satu qubit logis yang bertindak sebagai memori, dengan

tingkat error masih jauh dari kebutuhan algoritma nyata, tanpa operasi logis, dengan error floor yang belum dijelaskan dan ditandai penulis sendiri, serta jalan scaling beberapa orde besaran di depan. Sebuah threshold benar-benar dilintasi — bukan komputer kuantum yang sudah diserahkan.

Sumber

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>