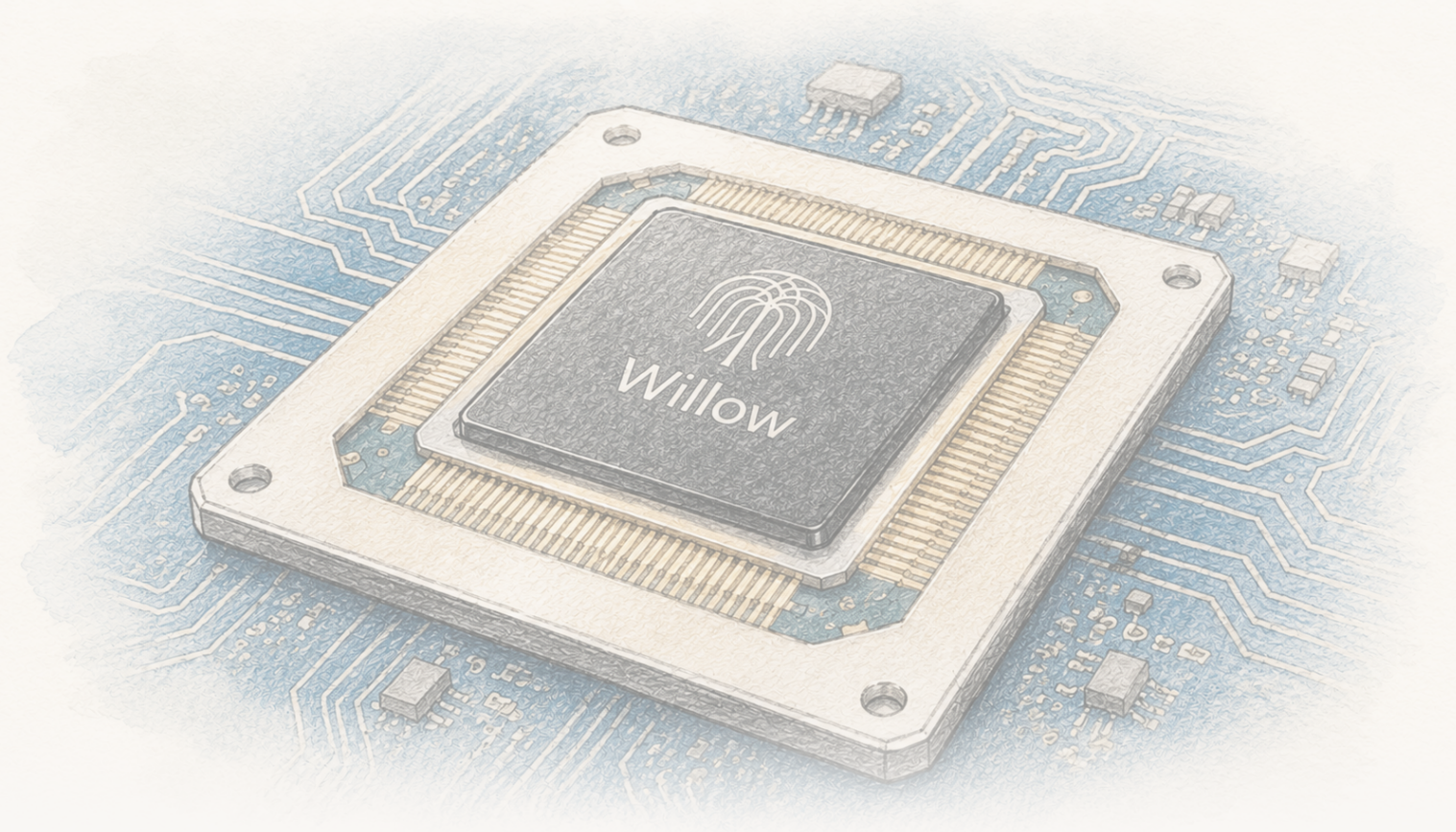




The CLEAN PAPER

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

量子記憶體終於隨著變大而變好——真正的里程碑，以及它還不是什麼

Google Quantum AI 首次清楚展示，表面碼量子記憶體可以在錯誤閾值以下運作：當程式碼距離從 3 增加到 5、再到 7，邏輯錯誤率呈指數下降（每增加兩級約改善 2.14 倍）；最大的一個距離 7 記憶體用了 101 個量子位元，壽命超過同晶片上最佳的實體量子位元——跨過 *breakeven*——而錯誤校正還能即時運作。這是長期追求的工程里程碑。但它也只是單一邏輯量子位元作為記憶體，錯誤率仍遠高於實際演算法需要的程度，沒有執行邏輯閘，作者自己還標出一個原因未明的錯誤底限，而且前面仍有好幾個數量級的擴充。跨過的是閾值，不是交付了一台量子電腦。

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

曲線終於往對的方向彎了

三十年來，量子錯誤校正一直建立在一個一直未被明確實現的承諾上。理論說，如果把一個量子資訊單位——一個邏輯量子位元——分散編碼到許多有雜訊的實體量子位元上，而且這些實體量子位元已經夠好，那麼加入更多量子位元後，邏輯量子位元反而應該變得更好；隨著錯誤校正碼變大，錯誤應呈指數下降。

問題就在那個「如果」。低於某個臨界雜訊閾值，增加量子位元會有幫助；高於它，更多量子位元只會帶來更多雜訊。過去所有實驗不是還待在這條線的錯誤一側，就是無法清楚展示所需趨勢。把編碼做大，結果通常反而更糟。

2024 年 12 月，Google Quantum AI 報告了第一個清楚進入另一個區域的示範。在最新一代超導處理器 Willow 上，他們建立了距離 3、5、7 的表面碼記憶體，並觀察到每次把程式碼做大，邏輯錯誤率都會下降——距離每增加兩級，改善因子為 $\Lambda = 2.14 \pm 0.02$ 。最大的一個距離 7 記憶體使用 101 個量子位元，每個錯誤校正循環的邏輯錯誤率為 $0.143\% \pm 0.003\%$ ；而更關鍵的是，它比同一晶片上最佳的實體量子位元活得更久，壽命高出 2.4 ± 0.3 倍。這叫做超過 breakeven，意思是：在這套硬體上，整套錯誤校正的成本第一次真的換回了淨收益。

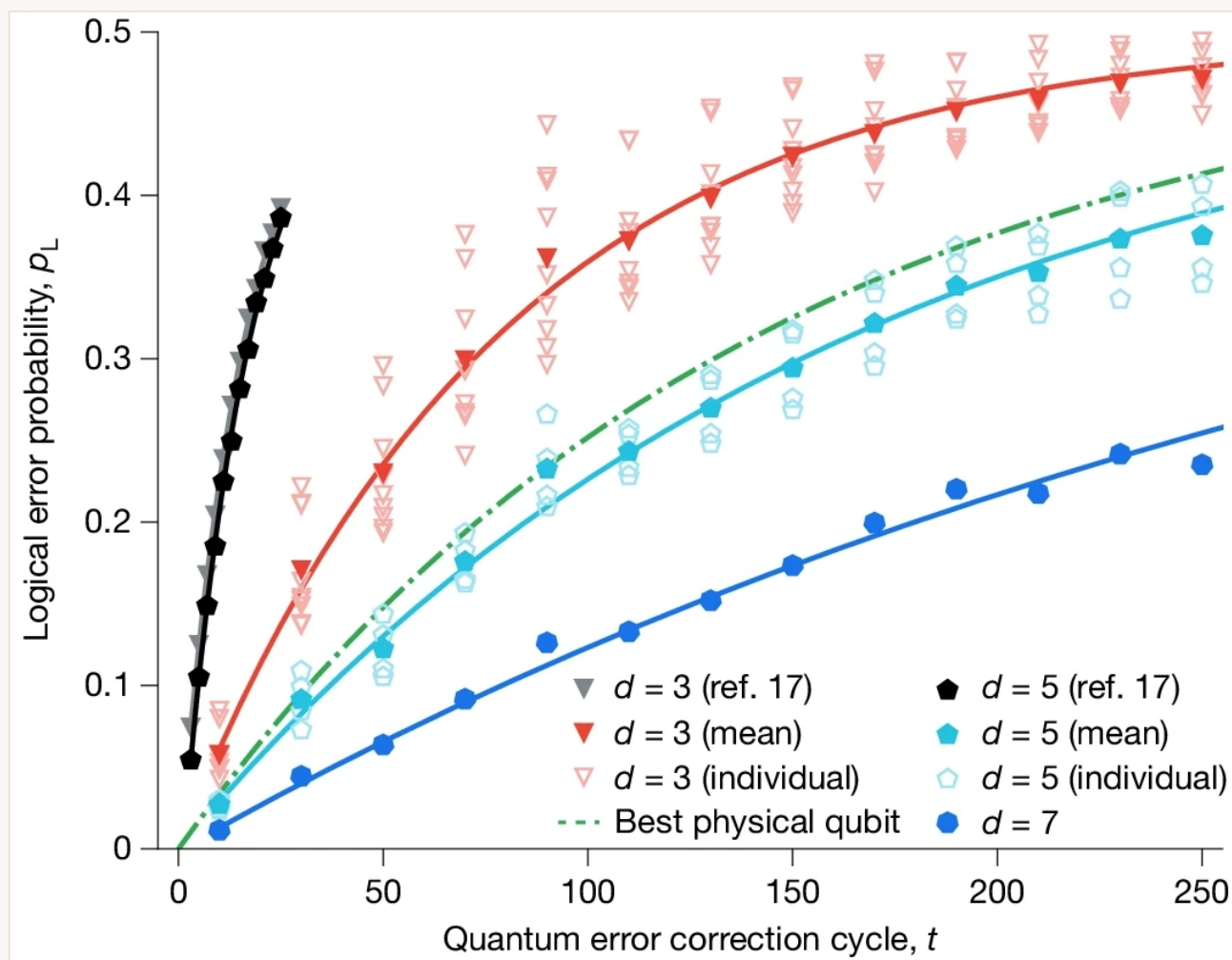
這是一個真正的里程碑，而重要的是把它屬於哪一類里程碑說清楚。它證明擴充趨勢現在終於朝正確方向前進。它不是一台可工作的量子電腦，論文也沒有聲稱它是。

「表面碼」、「距離」與「低於閾值」是什麼意思

一個邏輯量子位元，是把一份需要保護的量子資訊編碼到許多實體量子位元上。表面碼 (surface code) 是在二維網格上實現這種編碼的一種方式；其中額外的「測量」量子位元會持續檢查錯誤，同時盡量不擾動被儲存的資訊。

程式碼的距離 d 不是物理距離，而是：至少需要多少個位置恰當的錯誤，才可能在編碼沒察覺的情況下破壞邏輯量子位元。較大的 d 代表更大、更耐錯的編碼區塊，但也會花掉更多實體量子位元（約為 $2d^2 - 1$ ），並能修正更多同時發生的錯誤，最多為 $(d - 1) / 2$ 個。因此，這裡測試的距離 3、5、7，分別可修正 1、2、3 個同時錯誤，理論上大約需要 17、49、97 個實體量子位元；Google 實際建立的距離 7 記憶體用了 101 個，略高於這個教科書最低值。

「低於閾值 (below threshold)」才是整件事的核心。錯誤校正只有在實體錯誤率低於某個臨界值時才有用；在那個區域裡，每增加程式碼距離，都會以指數方式壓低邏輯錯誤率。抑制因子 Λ 就是在量化這件事： $\Lambda > 1$ 表示把程式碼做大有幫助，而且 Λ 越大越好。Google 報告 $\Lambda \approx 2.14$ ，也就是距離每增加兩級，邏輯錯誤率大約減半。真正的成果就是： Λ 明確而穩定地大於 1。



這張圖顯示距離 3、5、7 記憶體（由上到下）在一輪輪錯誤校正中如何累積邏輯錯誤。最值得看的，是綠色虛線——晶片上最佳的單一實體量子位元。距離 7 記憶體（藍色、最低）累積錯誤比這條線更慢，因此編碼後的量子位元比組成它的最佳實體量子位元活得更久——也就是「超過 breakeven」，壽命約長 2.4 倍。這張圖顯示的是壽命結果；真正「低於閾值」的抑制，也就是程式碼從距離 3 增加到 5 再到 7 時 $\Lambda = 2.14$ ，則反映在正文數字裡。

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

作者做了什麼

- 在兩顆 Willow 晶片上建立表面碼記憶體：一顆 105 量子位元處理器負責距離 3、5、7 的擴充測試（其中最大的是使用 101 個量子位元、49 個資料量子位元的距離 7 記憶體）；另一顆 72 量子位元處理器則執行距離 5 記憶體、即時解碼器，以及高距離的重複碼測試。
- 測量當程式碼距離從 3 增加到 5、再到 7 時，每一循環的邏輯錯誤率如何變化，並由此求出抑制因子 Λ 。
- 把邏輯量子位元的壽命與同晶片上最佳的單一實體量子位元相比，測試是否跨過「breakeven」。
- 讓距離 5 程式碼配合**即時解碼器**運作——也就是能跟上錯誤檢查產生速度的古典硬體——最多連續執行一百萬個循環，以證明錯誤校正不是只能事後離線分析，而是真能跟機器同步運作。
- 把較簡單的重複碼一路做到距離 29，用來尋找那些罕見、但會在極低錯誤率下限制表現的深層錯

誤來源。

他們發現了什麼

- **程式碼確實低於閾值**。距離每增加兩級，每循環邏輯錯誤率就以 $\Lambda = 2.14 \pm 0.02$ 的因子下降——呈現清楚的指數抑制，也就是理論長期預言、但先前處理器沒有明確展示過的行為。
- **距離 7 記憶體達到每循環 0.143% $\pm$ 0.003% 的錯誤率**，而且壽命比同晶片上最佳實體量子位元長 2.4 ± 0.3 倍——跨過 breakeven。
- **即時解碼跟得上**。在距離 5 上，解碼器平均延遲為 63 微秒，而量子錯誤校正循環本身是 1.1 微秒；系統持續運作達一百萬循環。換言之，錯誤校正是在線即時執行，不只是事後分析。
- **仍有罕見而深層的錯誤來源**。在重複碼測試裡，最終性能受到相關錯誤爆發限制，約每小時一次（大約每 3×10^9 個循環一次），造成約 10^{-10} 的錯誤底限；作者明確表示，這個來源目前仍不清楚。

這並沒有證明什麼

- **這不是一台正在進行計算的量子電腦**。這是一個量子記憶體：它儲存並保護一個邏輯量子位元。它沒有在多個邏輯量子位元間執行邏輯運算（邏輯閘），也沒有跑任何演算法。
- **它不是只差一個量子位元就能變成實用機器**。距離 7 的單一邏輯量子位元就花掉約 101 個實體量子位元；每循環約 0.1% 的錯誤率，仍遠高於實際演算法大約需要的 10^{-6} 到 10^{-10} 。要填平這個差距，就得把距離推到更大——也就是每個邏輯量子位元花更多實體量子位元——而有用的演算法還需要成千上萬個邏輯量子位元同時存在。實體量子位元的總預算很快就會進入數百萬等級。
- **「如果能擴充」這個如果非常重要**。論文自己的結論是：目前的裝置性能，如果能順利擴充，可能達到大型演算法需求。在單一邏輯量子位元上證明趨勢正確，不等於已經打造出擴充後的機器，也沒有任何事情保證這個趨勢在大很多的規模仍然成立。
- **原因未明的錯誤底限是真正的現存問題**。這些限制重複碼性能的相關爆發錯誤，按照作者自己的說法，比預期大好幾個數量級；在找到原因前，它們會阻止更大型的容錯應用。這是一個明白寫在論文裡的未解缺陷，不是已經處理完的小細節。
- 它也沒有證明任何關於**破解加密或在實用任務上達成「量子霸權」**的事情。那些需要的是完整的容錯量子機器；這項工作只是為那台機器鋪下的一塊基石。

證據有多強

- **核心主張扎實而重要**。跨三種程式碼距離看到清楚的指數抑制、加上超過 breakeven 的壽命、再加上能運作的即時解碼器——這正是整個領域多年一直想達成的組合，而且是直接示範，不是間接推論。這不是炒作產物，而是頂尖團隊做出的真實工程結果。
- **作者對範圍相當謹慎**。他們把結果明確稱為低於閾值的記憶體，主動指出原因未明的相關錯誤底限，對未來也保留那個非常醒目的「如果能擴充」。若有誇大，多半出現在周邊報導把「錯誤校正後的記憶體量子位元，隨擴大而改善」直接四捨五入成「量子運算已經到來」。

- **最誠實的說法是：這是一個扎實完成的基礎步驟。**一個邏輯量子位元，終於被保護到增加冗餘真的有幫助；但前面仍有漫長、艱難，而且沒有保證能走完的路，包括更大規模、邏輯閘，以及那個尚未解釋的錯誤來源。

為什麼重要

容錯量子運算一直有一種先有雞還是先有蛋的困境：真正有用的機器需要任何單一實體量子位元都達不到的低錯誤率；而解法——錯誤校正——又只有在硬體本身已經好到低於閾值時才有效。哪怕只在一個邏輯量子位元上、只跨過一次這條線，也會把問題從「這到底做不做得得到？」改成「能擴到多遠？能擴得多快？」這是真正有意義的改變，也是這個結果值得被注意的原因。

但也正因為研究做得仔細，周邊敘事更應該克制。這是地基的第一塊磚，而且鋪得很好。它不是整棟建築；而且最先這樣說的人，就是鋪磚的研究者自己。接下來幾年追蹤量子運算，最值得看的其實就是這條不華麗的曲線：程式碼繼續變大時，抑制因子還能不能維持；邏輯閘能不能像邏輯記憶體一樣明確；以及那個大約每小時出現一次的神祕錯誤，究竟能不能被解釋。

重點摘要

Google Quantum AI 首次清楚展示，表面碼量子記憶體可以在低於閾值的區域運作：當程式碼距離從 3 增加到 5、再到 7，邏輯錯誤率呈指數下降（每增加兩級約改善 2.14 倍）；最大的一個距離 7 記憶體使用 101 個量子位元，壽命超過同晶片上最佳的實體量子位元——跨過 **breakeven**——而且錯誤校正可以即時執行。這是量子電腦工程上一個真正而且期待已久的里程碑。但它也只是單一邏輯量子位元作為記憶體，錯誤率仍遠高於實際演算法需求，沒有執行任何邏輯運算，作者還明確指出一個原因未明的錯誤底限，而且前方仍有好幾個數量級的擴充路程。真正跨過的是閾值，不是交付了一台量子電腦。

來源

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>