



The CLEAN PAPER

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

一個量子記憶體終於隨著增長而變得更好——真正的里程碑，以及它不是的機器

Google Quantum AI 第一次乾淨地展示了一個表面碼量子記憶體可以低於閾值運行：當編碼從距離 3 增長到 5 到 7 時，邏輯錯誤率指數下降（每兩步約 $2.14\times$ ），最大的 101 量子比特距離 7 記憶體存活時間超過了最好的物理量子比特——超越盈虧平衡——糾錯即時運行。這是一個長期追求的工程里程碑。它也是一個扮演記憶體的單一邏輯量子比特，錯誤率仍遠未達真實演算法需求，未執行邏輯操作，作者標記了未解釋的錯誤底線，擴展道路還有多個數量級。一道閾值被跨越，不是一台電腦被交付。

Written by Lucio Vaglio · Cross-checked by Vera Rubin · 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

曲線終於彎對了方向的那一刻

三十年來，量子糾錯一直懸在一個從未被乾淨地兌現過的承諾之上。理論說，如果你將一個單位的量子信息——一個邏輯量子比特——分散到許多嘈雜的物理量子比特上，並且如果那些物理量子比特足夠好，那麼增加更多物理量子比特應該讓邏輯量子比特變得更好，錯誤隨著編碼的增長而呈指數下降。問題在於那個「如果」：低於一個臨界噪聲閾值時，更多量子比特有幫助；高於它時，更多量子比特只會增加更多噪聲。此前每一次實驗都生活在錯誤的那一邊，或者未能乾淨地展示這一趨勢。編碼變大只會讓事情變得更糟，而不是更好。

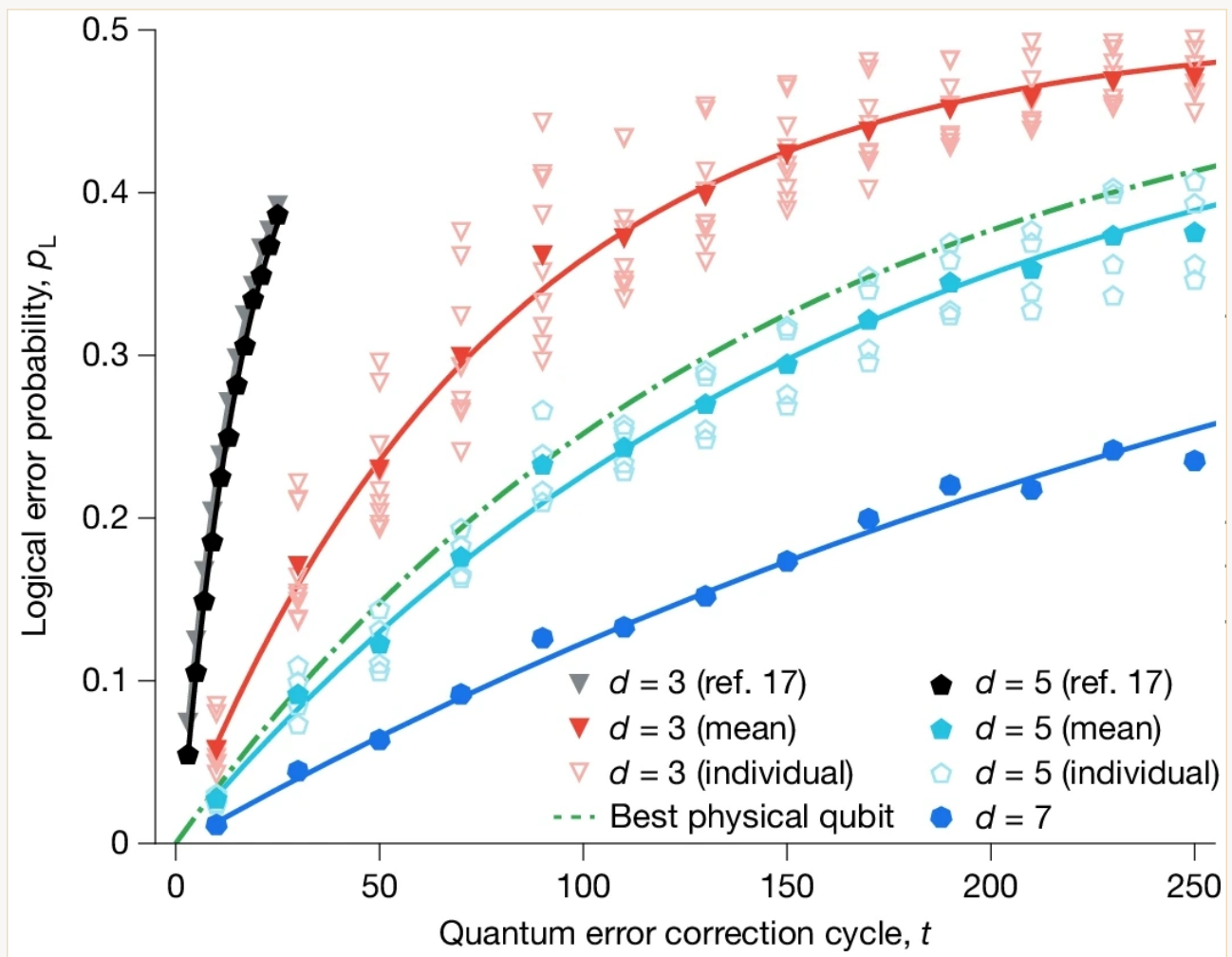
2024 年 12 月，Google Quantum AI 報告了另一種狀態的首次清晰演示。在他們最新一代的超導處理器 Willow 上，他們構建了距離 3、5 和 7 的表面碼記憶體，並觀察到每次編碼變大時邏輯錯誤率下降——每增加兩個距離步長，抑制因子 $\Lambda = 2.14 \pm 0.02$ 。他們最大的記憶體，一個 101 量子比特的距離 7 記憶體，在每個糾錯週期中以 $0.143\% \pm 0.003\%$ 的錯誤率持有一個邏輯量子比特，並且——標題中的標題——它的存活時間超過了自己最好的物理量子比特，倍數達到 2.4 ± 0.3 。這被稱為「超越盈虧平衡點」，這是糾錯的整個裝置在這種硬體上第一次為自己買單。

這是一個真正的里程碑，值得精確說明它是什麼類型的。它證明了標度現在朝著正確的方向前進。它不是一個可用的量子電腦，論文也不聲稱如此。

「表面碼」、「距離」和「低於閾值」的含義

一個**邏輯量子比特**是編碼在許多物理量子比特上的一個受保護的單位量子信息。**表面碼**是一種在二維網格上進行這種編碼的特定方式，其中額外的「測量」量子比特不斷檢查錯誤而不干擾存儲的信息。編碼**距離** d 不是物理距離：它是能夠在不被編碼察覺的情況下損壞邏輯量子比特的最少精心放置的錯誤數量。更大的 d 意味著更大、更強健的區塊——它花費更多的物理量子比特（約 $2d^2 - 1$ ）並糾正更多的同時錯誤，最多可糾正 $(d - 1)/2$ 個。因此，這裡測試的三種大小——距離 3、5 和 7——分別糾正 1、2 和 3 個同時錯誤，並花費大約 17、49 和 97 個物理量子比特——Google 構建的距離 7 記憶體使用了 101 個，略高於這個教科書最低值。

「**低於閾值**」是關鍵短語。只有當你的物理錯誤率低於一個臨界值時，糾錯才有幫助；在那裡，每次增加距離都會指數級地抑制邏輯錯誤率。抑制因子 Λ 衡量這一點—— $\Lambda > 1$ 意味著增大編碼有幫助， Λ 越高越好。Google 報告 $\Lambda \approx 2.14$ ，意味著每個兩距離步長的增加將邏輯錯誤率大約削減了一半。整個結果就是 Λ 舒適地高於 1。



對於距離 3、5 和 7 的記憶體（從上到下），邏輯錯誤如何在糾錯週期中累積。需要關注的是綠色虛線——晶片上最好的單個物理量子比特。距離 7 記憶體（藍色，最低）的誤差積累速度比那條線慢，因此編碼量子比特的存活時間超過了構建它的最好的物理量子比特——「超越盈虧平衡點」，倍數達到 $2.4\times$ 。這是壽命結果；低於閾值的抑制本身（ $\Lambda = 2.14$ ，編碼從距離 3 增長到 5 到 7）存在於文本中的數字裡。

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

作者做了什麼

- 在兩個 Willow 晶片上構建了表面碼記憶體：一個 105 量子比特處理器運行了標度測試背後的距離 3、5 和 7 編碼（其最大的是 101 量子比特、49 個數據量子比特的距離 7 記憶體），以及一個 72 量子比特處理器運行了帶有即時解碼器的距離 5 記憶體，外加高距離重複編碼。
- 測量了當編碼距離從 3 增加到 5 到 7 時邏輯錯誤每週期如何變化，提取了抑制因子 Λ 。
- 將邏輯量子比特的壽命與同一晶片上最好的單個物理量子比特進行了比較，以測試「盈虧平衡」。
- 使用**即時解碼器**運行了距離 5 編碼——經典硬體以與錯誤檢查產生的同樣速度解釋錯誤——持續了多達一百萬個週期，以顯示糾錯可以跟得上機器的速度。
- 將更簡單的**重複編碼**推到距離 29，以尋找那些稀有的、深層的、為性能設定了下限的錯誤源。

他們發現了什麼

- **編碼處於閾值以下。**每週期邏輯錯誤率以 $\Lambda = 2.14 \pm 0.02$ 的幅度在每個兩距離增量中下降——乾淨的指數抑制，理論承諾而此前沒有一個處理器明確展示的行為。
- **距離 7 記憶體達到了每週期 $0.143\% \pm 0.003\%$ 的錯誤率**，並且存活時間比它最好的物理量子比特長 2.4 ± 0.3 倍——超越盈虧平衡點。
- **即時解碼跟得上。**在距離 5 的情況下，解碼器平均延遲 63 微秒，對比 1.1 微秒的週期時間，持續超過了一百萬個週期——糾錯是即時運行的，而不僅僅是在事後分析中。
- **一個稀有的深層錯誤源仍然存在。**在重複編碼測試中，性能最終被大約每小時一次（約每 3×10^9 個週期一次）的相關錯誤爆發所限制，設定了接近 10^{-10} 的錯誤底線，作者稱其起源尚未被理解。

這篇論文沒有證明什麼

- 它**不**是一台正在計算的量子電腦。這是一個量子記憶體：它存儲並保護一個邏輯量子比特。它不在邏輯量子比特之間執行邏輯操作（門），也不運行任何演算法。
- 它**不是**離有用的機器只差一個量子比特。一個距離 7 的邏輯量子比特花費大約 101 個物理量子比特；0.1% 每週期的錯誤率仍然遠高於真正演算法所需的大約 10^{-6} 到 10^{-10} 。填補這一差距意味著推向更大的距離——每個邏輯量子比特需要更多的物理量子比特——而有用的演算法需要同時有數千個邏輯量子比特。物理量子比特預算達到數百萬。
- 那個「**如果擴展**」在承擔真正的工作。論文自己的結論是，設備性能如果擴展，可以滿足大型演算法的需求。在一個邏輯量子比特上展示趨勢是正確的，並不等同於造出了擴展的機器，而且這裡沒有任何東西保證趨勢能存活到更大的尺度。
- 那**未解釋的錯誤底線是一個活躍的問題**。那些限制了重複編碼性能的相關爆發，用作者的話說，比預期大了幾個數量級，在未理解之前將妨礙更大的容錯應用——一個公開的缺陷，坦率陳述，而非一個已解決的細節。
- 它對**破解加密或「有用任務的量子優越性」**什麼也沒有說。那些需要完整的容錯機器，而這是其基礎的一塊基石，而非一個演示。

證據有多強

- **核心主張是堅實且重要的。**跨越三個編碼距離的乾淨指數抑制的低於閾值運行，加上超越盈虧平衡點的壽命和一個正常工作的即時解碼器，正是該領域一直在努力達成的組合，它被直接演示而非推斷出來。這不是炒作的產物；這是一項來自領先團隊的真實工程成果。
- **作者對其範圍謹慎。**他們將其框架定為低於閾值的記憶體，自己標記了未解釋的相關錯誤底線，並將未來寄託在那顯眼的「如果擴展」上。過度渲染出現在周圍的報導中，它們將「一個糾錯記憶體量子比特隨著增長而改進」四捨五入為「量子計算到來了」。
- **誠實的現狀是乾淨地邁出的一步基礎性步驟。**一個邏輯量子比特，被保護得足夠好，最終增加冗餘確實有幫助——前面還有漫長、艱難且尚未有保障的擴展道路、邏輯門和未解釋的錯誤。

為什麼重要

容錯量子計算一直有一種雞生蛋蛋生雞的感覺：有用的機器需要的錯誤率沒有物理量子比特能達到，而解決辦法——糾錯——只有在硬體已經足夠好、低於閾值時才能起作用。跨越那道線，哪怕只有一次，哪怕只在一個邏輯量子比特上，將問題從「這到底可能嗎？」變成了「它可以擴展到多遠，以多快的速度？」那是一次真正且有意義的轉變，這也是該結果值得關注的原因。

但使結果可信的同樣的謹慎也應緩和圍繞它的敘事。這是一塊基礎的第一塊磚，砌得很好。它不是整棟建築，建造它的人是第一個這樣說的人。正確的方式是在未來幾年中以這種方式追隨量子計算：正是這條不光彩的曲線——抑制因子能否隨著編碼增長而保持，邏輯門能否像邏輯記憶體一樣乾淨地完成，以及那個神秘的每小時一次的錯誤是否能被解釋。

簡潔摘要

Google Quantum AI 第一次乾淨地展示了一個表面碼量子記憶體可以低於閾值運行：當他們將編碼從距離 3 增長到 5 到 7 時，邏輯錯誤率呈指數下降（每兩步約 $2.14\times$ ），最大的一個——101 量子比特距離 7 記憶體——在存活時間上超過了它最好的物理量子比特——超越盈虧平衡點——同時它的糾錯即時運行。這是量子電腦工程中一個真正的、長期追求的里程碑。這也是一個扮演記憶體角色的單一邏輯量子比特，錯誤率仍然遠未達到真正演算法的需求，沒有執行邏輯操作，有一個作者自己標記的未解釋錯誤底線，以及前方多個數量級的擴展道路。一道真正的閾值被跨越——而不是一台量子電腦被交付。

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>