



The CLEAN PAPER

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

一个量子记忆体终于随着增长而变得更好——真正的里程碑，以及它不是的机器

Google Quantum AI 第一次干净地展示了一个表面码量子记忆体可以低于阈值运行：当编码从距离 3 增长到 5 到 7 时，逻辑错误率指数下降（每两步约 $2.14\times$ ），最大的 101 量子比特距离 7 记忆体存活时间超过了最好的物理量子比特——超越盈亏平衡——纠错实时运行。这是一个长期追求的工程里程碑。它也是一个扮演记忆体的单一逻辑量子比特，错误率仍远未达真实算法需求，未执行逻辑操作，作者标记了未解释的错误底线，扩展道路还有多个数量级。一道阈值被跨越，不是一台计算机被交付。

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

曲线终于弯对了方向的那一刻

三十年来，量子纠错一直悬在一个从未被干净地兑现过的承诺之上。理论说，如果你将一个单位的量子信息——一个逻辑量子比特——分散到许多嘈杂的物理量子比特上，并且如果那些物理量子比特足够好，那么增加更多物理量子比特应该让逻辑量子比特变得更好，错误随着编码的增长而呈指数下降。问题在于那个“如果”：低于一个临界噪声阈值时，更多量子比特有帮助；高于它时，更多量子比特只会增加更多噪声。此前每一次实验都生活在错误的那一边，或者未能干净地展示这一趋势。编码变大只会让事情变得更糟，而不是更好。

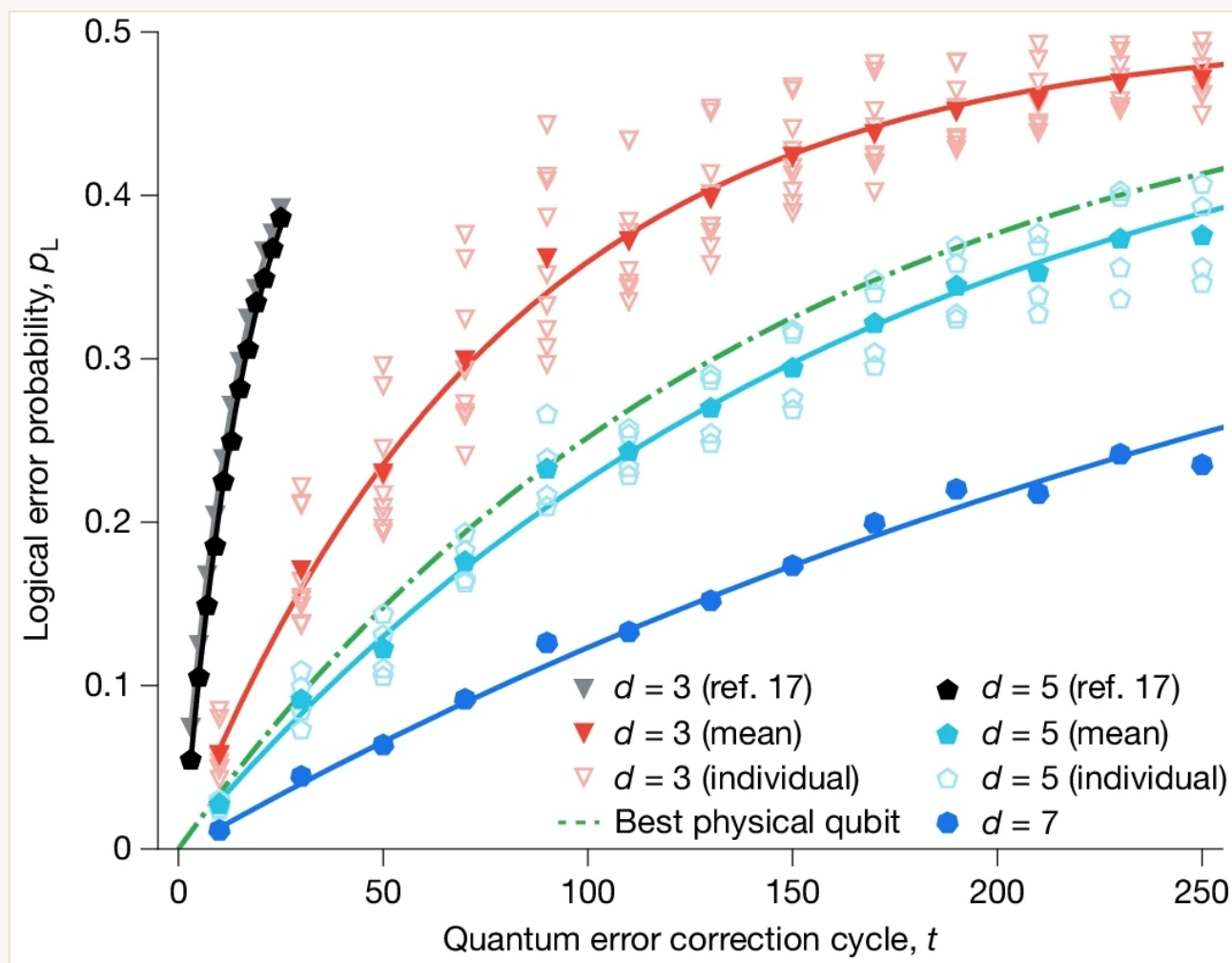
2024 年 12 月，Google Quantum AI 报告了另一种状态的首次清晰演示。在他们最新一代的超导处理器 Willow 上，他们构建了距离 3、5 和 7 的表面码记忆体，并观察到每次编码变大时逻辑错误率下降——每增加两个距离步长，抑制因子 $\Lambda = 2.14 \pm 0.02$ 。他们最大的记忆体，一个 101 量子比特的距离 7 记忆体，在每个纠错周期中以 $0.143\% \pm 0.003\%$ 的错误率持有一个逻辑量子比特，并且——标题中的标题——它的存活时间超过了自己最好的物理量子比特，倍数达到 2.4 ± 0.3 。这被称为“超越盈亏平衡点”，这是纠错的整个装置在这种硬件上第一次为自己买单。

这是一个真正的里程碑，值得精确说明它是什么类型的。它证明了标度现在朝着正确的方向前进。它不是一个可用的量子计算机，论文也不声称如此。

“表面码”、“距离”和“低于阈值”的含义

一个**逻辑量子比特**是编码在许多物理量子比特上的一个受保护的单位量子信息。**表面码**是一种在二维网格上进行这种编码的特定方式，其中额外的“测量”量子比特不断检查错误而不干扰存储的信息。编码**距离** d 不是物理距离：它是能够在不被编码察觉的情况下损坏逻辑量子比特的最少精心放置的错误数量。更大的 d 意味着更大、更强健的区块——它花费更多的物理量子比特（约 $2d^2 - 1$ ）并纠正更多的同时错误，最多可纠正 $(d - 1)/2$ 个。因此，这里测试的三种大小——距离 3、5 和 7——分别纠正 1、2 和 3 个同时错误，并花费大约 17、49 和 97 个物理量子比特——Google 构建的距离 7 记忆体使用了 101 个，略高于这个教科书最低值。

“低于阈值”是关键短语。只有当你的物理错误率低于一个临界值时，纠错才有帮助；在那里，每次增加距离都会指数级地抑制逻辑错误率。抑制因子 Λ 衡量这一点—— $\Lambda > 1$ 意味着增大编码有帮助， Λ 越高越好。Google 报告 $\Lambda \approx 2.14$ ，意味着每个两距离步长的增加将逻辑错误率大约削减了一半。整个结果就是 Λ 舒适地高于 1。



对于距离 3、5 和 7 的记忆体（从上到下），逻辑错误如何在纠错周期中累积。需要关注的是绿色虚线——芯片上最好的单个物理量子比特。距离 7 记忆体（蓝色，最低）的误差积累速度比那条线慢，因此编码量子比特的存活时间超过了构建它的最好的物理量子比特——“超越盈亏平衡点”，倍数达到 2.4×。这是寿命结果；低于阈值的抑制本身（ $\Lambda = 2.14$ ，编码从距离 3 增长到 5 到 7）存在于文本中的数字里。

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

作者做了什么

- 在两个 Willow 芯片上构建了表面码记忆体：一个 105 量子比特处理器运行了标度测试背后的距离 3、5 和 7 编码（其最大的是 101 量子比特、49 个数据量子比特的距离 7 记忆体），以及一个 72 量子比特处理器运行了带有实时解码器的距离 5 记忆体，外加高距离重复编码。
- 测量了当编码距离从 3 增加到 5 到 7 时逻辑错误每周期如何变化，提取了抑制因子 Λ 。
- 将逻辑量子比特的寿命与同一芯片上最好的单个物理量子比特进行了比较，以测试“盈亏平衡”。
- 使用**实时解码器**运行了距离 5 编码——经典硬件以与错误检查产生的同样速度解释错误——持续了多达一百万个周期，以显示纠错可以跟得上机器的速度。
- 将更简单的**重复编码**推到距离 29，以寻找那些稀有的、深层的、为性能设定了下限的错误源。

他们发现了什么

- **编码处于阈值以下。**每周期逻辑错误率以 $\Lambda = 2.14 \pm 0.02$ 的幅度在每个两距离增量中下降——干净的指数抑制，理论承诺而此前没有一个处理器明确展示的行为。
- **距离 7 记忆体达到了每周期 $0.143\% \pm 0.003\%$ 的错误率**，并且存活时间比它最好的物理量子比特长 2.4 ± 0.3 倍——超越盈亏平衡点。
- **实时解码跟得上。**在距离 5 的情况下，解码器平均延迟 63 微秒，对比 1.1 微秒的周期时间，持续超过了一百万个周期——纠错是实时运行的，而不仅仅是在事后分析中。
- **一个稀有的深层错误源仍然存在。**在重复编码测试中，性能最终被大约**每小时一次**（约每 3×10^9 个周期一次）的相关错误爆发所限制，设定了接近 10^{-10} 的错误底线，作者称其起源**尚未被理解**。

这篇论文没有证明什么

- 它**不**是一台正在计算的量子计算机。这是一个量子记忆体：它存储并保护一个逻辑量子比特。它不在逻辑量子比特之间执行逻辑操作（门），也不运行任何算法。
- 它**不是**离有用的机器只差一个量子比特。一个距离 7 的逻辑量子比特花费大约 101 个物理量子比特；0.1% 每周期的错误率仍然远高于真正算法所需的大约 10^{-6} 到 10^{-10} 。填补这一差距意味着推向更大的距离——每个逻辑量子比特需要更多的物理量子比特——而有用的算法需要同时有数千个逻辑量子比特。物理量子比特预算达到数百万。
- **那个“如果扩展”在承担真正的工作。**论文自己的结论是，设备性能如果扩展，可以满足大型算法的需求。在一个逻辑量子比特上展示趋势是正确的，并不等同于造出了扩展的机器，而且这里没有任何东西保证趋势能存活到更大的尺度。
- **那未解释的错误底线是一个活跃的问题。**那些限制了重复编码性能的相关爆发，用作者的话说，比预期大了几个数量级，在未理解之前将妨碍更大的容错应用——一个公开的缺陷，坦率陈述，而非一个已解决的细节。
- 它对**破解加密或“有用任务的量子优越性”**什么也没有说。那些需要完整的容错机器，而这是其基础的一块基石，而非一个演示。

证据有多强

- **核心主张是坚实且重要的。**跨越三个编码距离的干净指数抑制的低于阈值运行，加上超越盈亏平衡点的寿命和一个正常工作的实时解码器，正是该领域一直在努力达成的组合，它被直接演示而非推断出来。这不是炒作的产物；这是一项来自领先团队的真实工程成果。
- **作者对其范围谨慎。**他们将其框架定为低于阈值的记忆体，自己标记了未解释的相关错误底线，并将未来寄托在那显眼的“如果扩展”上。过度渲染出现在周围的报道中，它们将“一个纠错记忆体量子比特随着增长而改进”四舍五入为“量子计算到来了”。
- **诚实的现状是干净地迈出的一步基础性步骤。**一个逻辑量子比特，被保护得足够好，最终增加冗余确实有帮助——前面还有漫长、艰难且尚未有保障的扩展道路、逻辑门和未解释的错误。

为什么重要

容错量子计算一直有一种鸡生蛋蛋生鸡的感觉：有用的机器需要的错误率没有物理量子比特能达到，而解决办法——纠错——只有在硬件已经足够好、低于阈值时才能起作用。跨越那道线，哪怕只有一次，哪怕只在一个逻辑量子比特上，将问题从“这到底可能吗？”“变成了”它可以扩展到多远，以多快的速度？“那是一次真正且有意义的转变，这也是该结果值得关注的原因。

但使结果可信的同样的谨慎也应缓和围绕它的叙事。这是一块基础的第一块砖，砌得很好。它不是整栋建筑，建造它是第一个这样说的人。正确的方式是在未来几年中以这种方式追随量子计算：正是这条不光彩的曲线——抑制因子能否随着编码增长而保持，逻辑门能否像逻辑记忆体一样干净地完成，以及那个神秘的每小时一次的错误是否能被解释。

简洁摘要

Google Quantum AI 第一次干净地展示了一个表面码量子记忆体可以低于阈值运行：当他们将编码从距离 3 增长到 5 到 7 时，逻辑错误率呈指数下降（每两步约 $2.14\times$ ），最大的一个——101 量子比特距离 7 记忆体——在存活时间上超过了它最好的物理量子比特——超越盈亏平衡点——同时它的纠错实时运行。这是量子计算机工程中一个真正的、长期追求的里程碑。这也是一个扮演记忆体角色的单一逻辑量子比特，错误率仍然远未达到真正算法的需求，没有执行逻辑操作，有一个作者自己标记的未解释错误底线，以及前方多个数量级的扩展道路。一道真正的阈值被跨越——而不是一台量子计算机被交付。

来源

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>