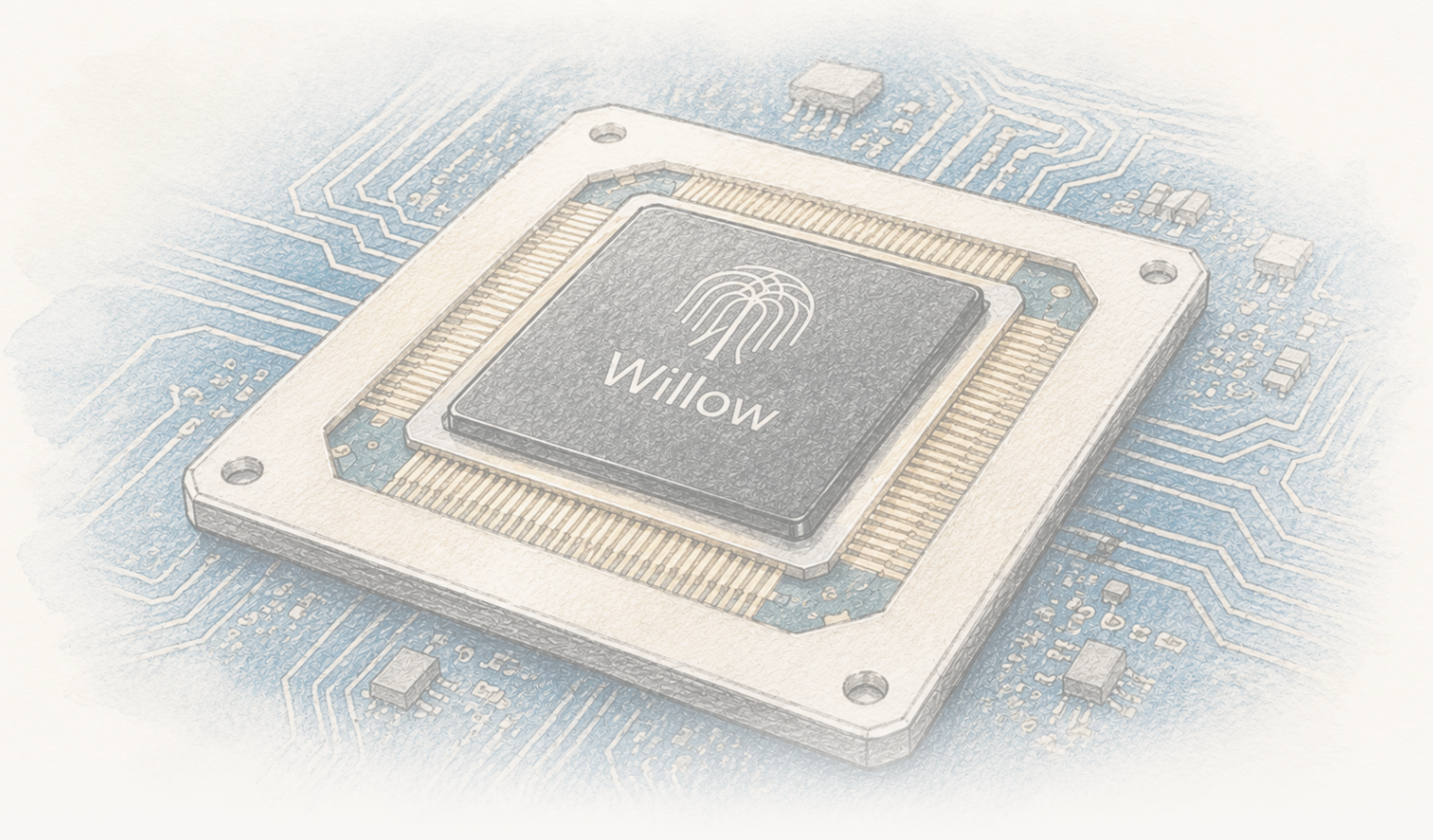




The CLEAN PAPER

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

Bir kuantum belleđi büyüdüķe sonunda daha iyi hale geldi — asıl kilometre taşı ve onun olmadığı makine

Google Quantum AI, ilk kez net biçimde, yüzey kodu bir kuantum belleğinin eşik altında çalışabildiğini gösterdi: kodu mesafe 3'ten 5'e, 5'ten 7'ye büyüttükçe mantıksal hata oranı üstel olarak düştü (her iki adımda yaklaşık 2.14x) ve en büyükleri olan 101 kübitlik mesafe-7 belleđi, hata düzeltmesi gerçek zamanlı çalışırken en iyi fiziksel kübitinden daha uzun yaşadı — başabaşın ötesinde. Bu, uzun süredir aranan bir mühendislik kilometre taşı. Aynı zamanda bellek işlevi gören tek bir mantıksal kübit; hata oranı gerçek algoritmaların gereksinim duyduğunun hâlâ çok uzağında, hiçbir mantıksal işlem gerçekleştirilmemiş, yazarların işaret ettiği açıklanamayan bir hata tabanı var ve önünde büyüklük mertebeleriyle ölçülen bir ölçekleme yolu uzanıyor. Aşılan bir eşik, teslim edilen bir bilgisayar değil.

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

Eğrinin doğru yöne büküldüğü an

Otuz yıldır kuantum hata düzeltme, hiçbir zaman tam anlamıyla yerine getirilememiş bir söze dayanıyordu. Teoriye göre, bir birim kuantum bilgisini — bir *mantıksal* kübiti — çok sayıda gürültülü fiziksel kübite yayarsanız ve bu fiziksel kübitler yeterince iyiye, *daha fazla* fiziksel kübit eklemek mantıksal kübiti *daha iyi* hale getirmelidir; kod büyüdükçe hatalar üstel olarak azalır. İşin püf noktası bu “eğer”de: kritik bir gürültü eşiğinin altında daha fazla kübit yardımcı olur; üstünde ise daha fazla kübit yalnızca daha fazla gürültü ekler. Önceki her deney bu çizginin yanlış tarafında kalmış ya da bu eğilimi net biçimde gösterememişti. Kodu büyütme işleri iyileştirmek yerine kötüleştiriyordu.

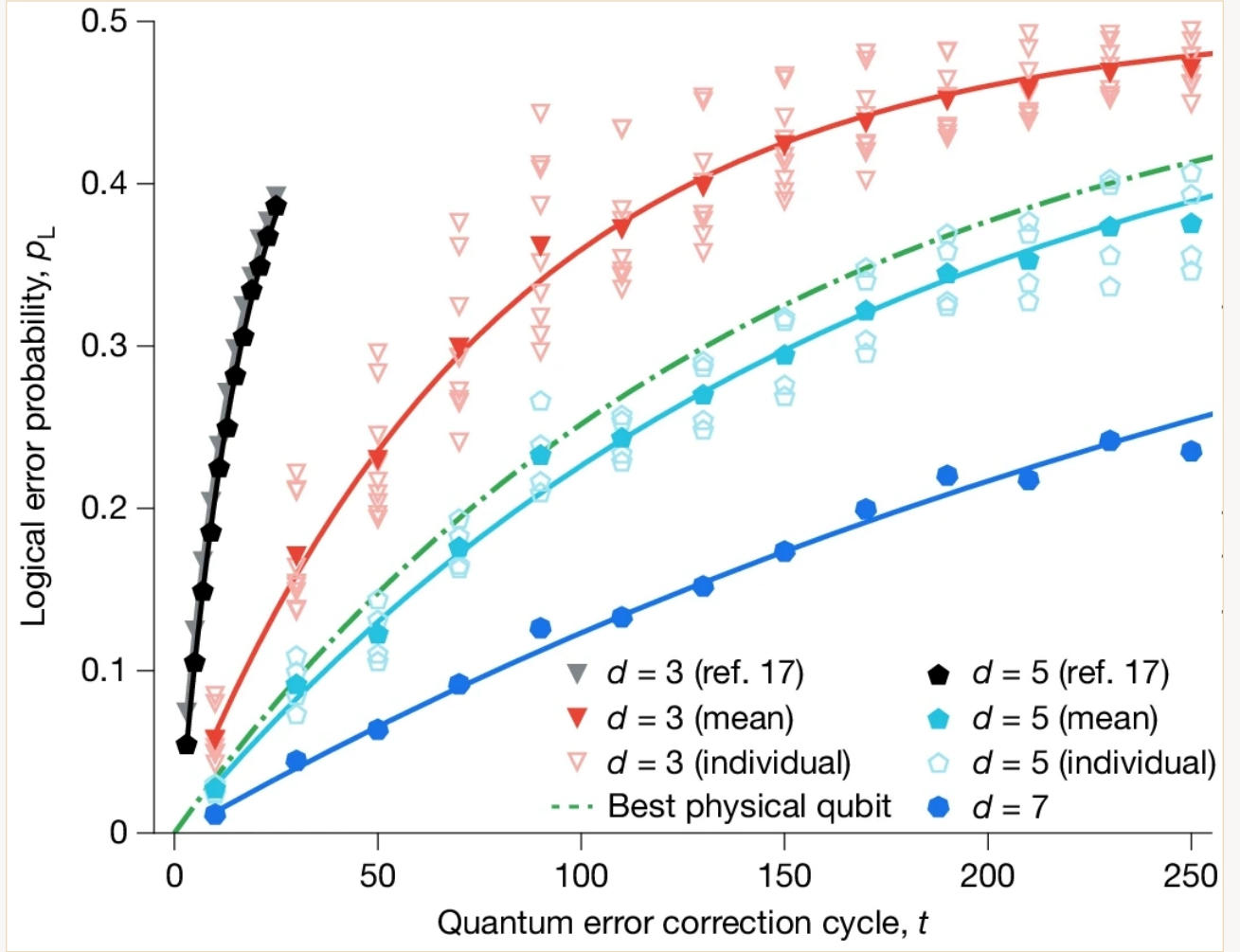
Aralık 2024’te Google Quantum AI, diğer rejimin ilk net gösterimini bildirdi. Süperiletken işlemcilerinin en yeni nesli olan Willow’da, 3, 5 ve 7 kod mesafelerinde yüzey kodu bellekleri kurdular ve kod her büyüdüğünde mantıksal hata oranının *düşüğünü* gözlemlediler — mesafede yukarı doğru her iki adımda $\Lambda = 2.14 \pm 0.02$ çarpanı kadar. En büyükleri olan 101 kübitlik mesafe-7 belleği, düzeltme döngüsü başına $0.143\% \pm 0.003\%$ hatayla bir mantıksal kübiti tuttu ve — asıl haberin içindeki haber — *kendi en iyi fiziksel kübitinden daha uzun yaşadı*; hem de 2.4 ± 0.3 çarpanı kadar. Buna “başabaşın ötesine geçmek” denir ve hata düzeltmenin tüm düzeneğinin bu donanımda kendi masrafını çıkardığı ilk sefer bu.

Bu gerçek bir kilometre taşı ve ne tür bir kilometre taşı olduğu konusunda kesin olmakta yarar var. Ölçeklemenin artık doğru yönde ilerlediğinin bir kanıtı. Çalışan bir kuantum bilgisayarı değil ve makale böyle bir iddiada bulunmuyor.

“Yüzey kodu”, “mesafe” ve “eşik altı” ne anlama geliyor

Mantıksal kübit, çok sayıda fiziksel kübite kodlanmış, korunan bir birim kuantum bilgisidir. **Yüzey kodu**, bu kodlamayı 2 boyutlu bir ızgara üzerinde yapmanın belirli bir yoludur; burada fazladan “ölçüm” kübitleri, depolanan bilgiyi bozmadan sürekli olarak hataları denetler. Kod **mesafesi** d , fiziksel bir mesafe değildir: kodun fark etmeden mantıksal kübiti bozabilecek, iyi yerleştirilmiş en küçük hata sayısıdır. Daha büyük bir d , daha büyük ve daha sağlam bir yamadır — daha fazla fiziksel kübit harcar (kabaca $2d^2 - 1$) ve daha fazla eşzamanlı hatayı, en fazla $(d - 1)/2$ tanesini düzeltir. Dolayısıyla burada test edilen üç boyut — 3, 5 ve 7 mesafeleri — 1, 2 ve 3 eşzamanlı hatayı düzeltir ve kabaca 17, 49 ve 97 fiziksel kübit harcar — Google’ın kurduğu mesafe-7 belleği 101 kübit kullandı; bu ders kitabı asgarisinin biraz üzerinde.

“Eşik altı”, buradaki kilit ifadedir. Hata düzeltme yalnızca fiziksel hata oranınız kritik bir değerin altındaysa işe yarar; bu durumda mesafedeki her artış mantıksal hata oranını üstel olarak bastırır. Bastırma faktörü Λ bunu ölçer — $\Lambda > 1$, kodu büyütmenin işe yaradığı anlamına gelir ve Λ ne kadar yüksekse o kadar iyidir. Google $\Lambda \approx 2.14$ bildiriyor; bu, mesafedeki her iki adımlık artışın mantıksal hata oranını kabaca yarıya indirdiği anlamına geliyor. Λ ’nın 1’in rahatça üzerinde olması, sonucun ta kendisidir.



Mesafe-3, -5 ve -7 bellekleri için (yukarıdan aşağıya) mantıksal hatanın düzeltme döngüleri boyunca nasıl biriktiği. İzlenmesi gereken çizgi yeşil kesikli olan — çipteki *en iyi tek fiziksel kübit*. Mesafe-7 belleği (mavi, en alttaki) hatayı bu çizgiden daha yavaş biriktirir; böylece kodlanmış kübit, kendisini oluşturan en iyi fiziksel kübitten daha uzun yaşar — $2.4\times$ çarpanıyla “başabaşın ötesinde”. Bu, yaşam süresi sonucudur; eşik altı bastırmanın kendisi (kod 3’ten 5’e, 5’ten 7’ye büyürken $\Lambda = 2.14$) metindeki sayılarda yer alır.

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

Yazarlar ne yaptı

- İki Willow çipi üzerinde yüzey kodu bellekleri kurdular: ölçekleme testinin arkasındaki mesafe-3, -5 ve -7 kodlarını çalıştıran 105 kübitlik bir işlemci (en büyüğü 101 kübitlik, 49 veri kübitli mesafe-7 belleği) ve gerçek zamanlı bir kod çözücüyle bir mesafe-5 belleğini ve yüksek mesafeli yineleme kodlarını çalıştıran 72 kübitlik bir işlemci.
- Kod mesafesini 3’ten 5’e, 5’ten 7’ye artırdıkça *döngü başına* mantıksal hatanın nasıl değiştiğini ölçtüler ve bastırma faktörü Λ ’yı çıkardılar.
- “Başabaşı” sınamak için mantıksal kübitin yaşam süresini aynı çipteki en iyi tek fiziksel kübitle karşılaştırdılar.
- Hata düzeltmenin makineye ayak uydurabildiğini göstermek için mesafe-5 kodunu **gerçek za-**

manlı bir kod çözücüyle — hata denetimlerini üretildikleri hızda yorumlayan klasik donanım — bir milyon döngüye kadar çalıştırdılar.

- Performansa bir taban koyan nadir, derin hata kaynaklarını avlamak için daha basit **yineleme kodlarını** mesafe 29'a kadar zorladılar.

Ne buldular

- **Kod eşik altında.** Döngü başına mantıksal hata, mesafedeki her iki birimlik artış için $\Lambda = 2.14 \pm 0.02$ çarpanı kadar düştü — teorinin vaat ettiği ve hiçbir işlemcinin kesin biçimde gösteremediği davranış olan net üstel bastırma.
- **Mesafe-7 belleği döngü başına $0.143\% \pm 0.003\%$ hataya ulaştı** ve en iyi fiziksel kübitinden **2.4 ± 0.3 kat daha uzun** yaşadı — başabaşın ötesinde.
- **Gerçek zamanlı kod çözme ayak uydurdu.** Kod çözücü, mesafe 5'te 1.1 mikrosaniyelik döngü süresine karşılık ortalama 63 mikrosaniye gecikme gösterdi ve bunu bir milyon döngü boyunca sürdürdü — hata düzeltme, yalnızca olay sonrası analizde değil, canlı olarak çalıştı.
- **Nadir, derin bir hata kaynağı varlığını sürdürüyor.** Yineleme kodu testlerinde performans, nihayetinde kabaca **saatte bir** meydana gelen (her 3×10^9 döngüde bir civarında) birbiriyle ilişkili hata patlamalarıyla sınırlandı; bu patlamalar 10^{-10} dolayında bir hata tabanı belirledi ve yazarlar bu tabanın kökeninin **henüz anlaşılmadığını** söylüyor.

Bunun kanıtlamadıkları

- Bu, hesaplama yapan bir kuantum bilgisayarı **değil**. Bu bir kuantum *belleği*: tek bir mantıksal kübiti depolar ve korur. Mantıksal kübitler arasında mantıksal işlemler (kapılar) gerçekleştirmez ve hiçbir algoritma çalıştırmaz.
- Yararlı makinelere yalnızca bir kübit uzaklıkta **değil**. Mesafe-7'lik bir mantıksal kübit yaklaşık 101 fiziksel kübit harcar; döngü başına 0.1%'lik bir hata oranı, gerçek algoritmaların gereksinim duyduğu kabaca 10^{-6} ile 10^{-10} arasındaki değerin hâlâ çok üzerinde. Bu açığı kapatmak, çok daha büyük mesafelere geçmek — mantıksal kübit başına çok daha fazla fiziksel kübit — demek ve yararlı algoritmalar aynı anda *binlerce* mantıksal kübite gereksinim duyar. Bunun için gereken fiziksel kübit bütçesi milyonlara varır.
- **“Ölçeklenirse” ifadesi gerçekten iş görüyor.** Makalenin kendi sonucu, cihaz performansının *ölçeklenirse* büyük algoritmaların gereksinimlerini karşılayabileceğidir. Eğilimin tek bir mantıksal kübit üzerinde doğru olduğunu göstermek, ölçeklenmiş makineyi kurmuş olmakla aynı şey değildir ve buradaki hiçbir şey eğilimin çok daha büyük boyutlarda geçerli kalacağını garanti etmez.
- **Açıklanamayan hata tabanı, güncel bir sorun.** Yineleme kodu performansına tavan koyan ilişkili patlamalar, yazarların ifadesiyle beklenenden büyüklük mertebeleri kadar büyük ve anlaşılmadıkça daha büyük hata toleranslı uygulamaları olanaksız kılardı — açıkça dile getirilen, çözülmüş bir ayrıntı değil, açık bir kusur.
- **Şifrelemeyi kırma ya da yararlı görevler için “kuantum üstünlüğü”** hakkında hiçbir şey

söylemiyor. Bunlar, bu sonucun bir gösterimi değil temel taşı olduğu, tam anlamıyla hata toleranslı makineyi gerektirir.

Kanıtlar ne kadar güçlü

- **Temel iddia sağlam ve önemli.** Üç kod mesafesi boyunca net üstel bastırmayla eşik altı çalışma, buna ek olarak başabaşın ötesinde bir yaşam süresi ve çalışan bir gerçek zamanlı kod çözücü, alanın ulaşmaya çalıştığı tam da o bileşimdir ve çıkarımla değil, doğrudan gösterilmiştir. Bu bir abartı ürünü değil; önde gelen bir grubun gerçek bir mühendislik sonucu.
- **Yazarlar kapsamı konusunda dikkatli.** Bunu eşik altı bir *bellek* olarak çerçevesiyor, açıklanamayan ilişkili hata tabanını bizzat işaret ediyor ve geleceğe dair o göze çarpan “ölçeklenirse” kaydını koyuyorlar. Aşırıya kaçış, görüldüğü yerde, “hata düzeltmeli bir bellek kubitini büyüdükçe iyileşti” ifadesini “kuantum bilgisayarları geldi”ye yuvarlayan çevredeki haberlerdedir.
- **Dürüst durum, net biçimde atılmış temel bir adım.** Tek bir mantıksal kubit, fazladan kubit eklemenin sonunda işe yarayacağı kadar iyi korunmuş — önünde hâlâ uzun, zorlu ve henüz garanti edilmemiş bir ölçekleme, mantıksal kapılar ve açıklanamayan hatalar yolu var.

Neden önemli

Hata toleranslı kuantum hesaplama, hep bir “yumurta-tavuk” ikilemi havası taşımıştır: yararlı olacak makineler, hiçbir fiziksel kubitin ulaşamayacağı hata oranlarına gereksinim duyar ve çözüm — hata düzeltme — yalnızca donanım eşik altında olacak kadar iyiye işe yarar. Bu çizgiyi, tek bir mantıksal kubit üzerinde bile, bir kez bile aşmak, soruyu “bu hiç mümkün mü?”den “ne kadar ölçeklenebilir ve ne kadar hızlı?”ya dönüştürür. Bu gerçek ve anlamlı bir değişim ve sonucun dikkati hak etmesinin nedeni de bu.

Ama sonucu güvenilir kılan aynı özen, onun etrafındaki anlatıyı da dizginlemeli. Bu, iyi yerleştirilmiş bir temelin ilk tuğlası. Binanın kendisi değil ve onu koyanlar bunu ilk söyleyenler. Önümüzdeki birkaç yıl boyunca kuantum hesaplamayı izlemenin doğru yolu tam da bu gösterişsiz eğridir: kodlar büyüdükçe bastırma faktörünün geçerliliğini koruyup korumadığı, mantıksal kapıların mantıksal bellek kadar net biçimde yapılıp yapılamayacağı ve o gizemli, saatte bir görülen hatanın hiç açıklanıp açıklanmayacağı.

Kısa özet

Google Quantum AI, ilk kez net biçimde, yüzey kodu bir kuantum belleğinin *eşik altında* çalışabildiğini gösterdi: kodu mesafe 3’ten 5’e, 5’ten 7’ye büyüttükçe mantıksal hata oranı üstel olarak düştü (her iki adımda yaklaşık $2.14\times$ kadar) ve en büyükleri olan 101 kubitlik mesafe-7 belleği, hata düzeltmesi gerçek zamanlı çalışırken en iyi fiziksel kubitinden daha uzun yaşadı — başabaşın

ötesinde. Bu, kuantum bilgisayarlarının mühendisliğinde gerçek ve uzun süredir aranan bir kilometre taşı. Aynı zamanda bellek işlevi gören tek bir mantıksal kübit; hata oranı gerçek algoritmaların talep ettiğinin hâlâ çok uzağında, hiçbir mantıksal işlem gerçekleştirilmemiş, yazarların bizzat işaret ettiği açıklanamayan bir hata tabanı var ve önünde birçok büyüklük mertebesi kadar bir ölçekleme yolu uzanıyor. Aşılan gerçek bir eşik — teslim edilen bir kuantum bilgisayarı değil.

Kaynaklar

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>