



The CLEAN PAPER

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

זיכרון קוונטי סוף-סוף השתפר ככל שגדל — ציון הדרך האמיתי, והמכונה שהוא עדיין לא

AI Quantum Google הראתה לראשונה באופן נקי שזיכרון קוונטי מסוג *code surface* יכול לפעול מתחת לסף. כאשר הקוד גדל מ-*distance 3* ל-*5* ל-*7*, שיעור השגיאה הלוגית ירד אקספוננציאלית (בערך פי 2.14 לכל שתי דרגות), והגדול ביותר, זיכרון *distance-7* של 101 קיוביטים, חי יותר מהקיוביט הפיזי הטוב ביותר שלו — *breakeven beyond* — כאשר תיקון השגיאות פעל בזמן אמת. זה ציון דרך הנדסי שחיכו לו זמן רב. זה גם קיוביט לוגי יחיד שמתפקד כזיכרון, עם שיעור שגיאה שעדיין רחוק מדרישות אלגוריתמים אמיתיים, בלי פעולות לוגיות ועם רצפת שגיאה לא מוסברת. סף שנחצה, לא מחשב שנמסר.

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

הרגע שבו העקומה סוף-סוף התכופפה בכיוון הנכון

במשך שלושים שנה, תיקון שגיאות קוונטי נשען על הבטחה שמעולם לא קוימה בצורה נקייה בניסוי. התאוריה אומרת שאם מפזרים יחידה אחת של מידע קוונטי — קיוביט לוגי אחד — על פני קיוביטים פיזיים רועשים רבים, ואם הקיוביטים הפיזיים טובים מספיק, אז הוספת עוד מהם אמורה להפוך את הקיוביט הלוגי ל-טוב יותר, כאשר שיעור השגיאות יורד אקספוננציאלית ככל שהקוד גדל. המלכוד הוא ה"אם": מתחת לסף רעש קריטי, עוד קיוביטים עוזרים; מעליו, עוד קיוביטים רק מוסיפים עוד רעש. כל ניסוי קודם היה בצד הלא נכון של הקו הזה, או לא הצליח להראות את המגמה בצורה נקייה. הגדלת הקוד הפכה את הדברים לגרועים יותר, לא טובים יותר.

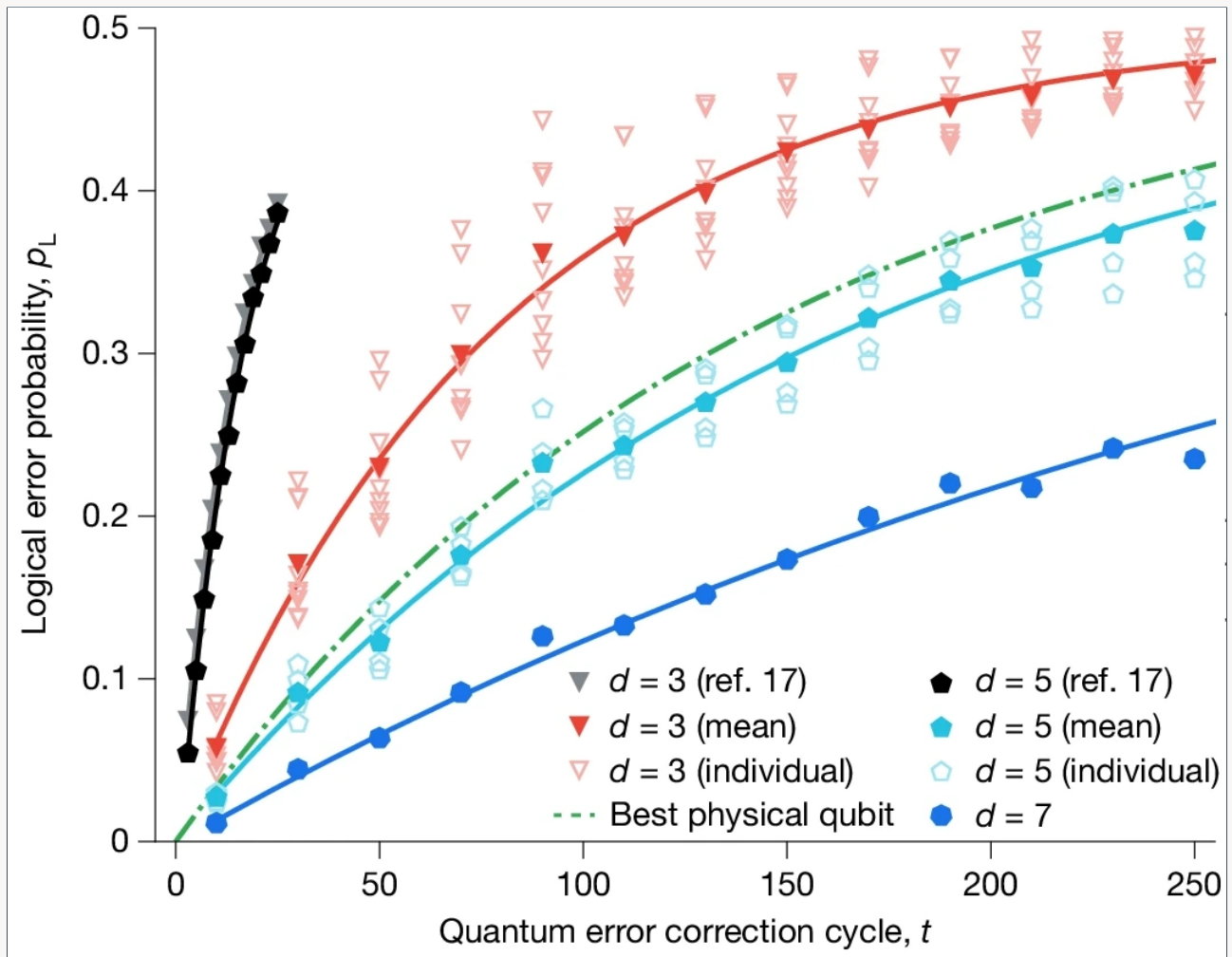
בדצמבר 2024 AI Quantum Google דיווחה על ההדגמה הברורה הראשונה של המשטר האחר. על Willow, הדור החדש של המעבדים העל-מולכים שלה, היא בנתה זיכרונות code surface במרחקי קוד 3, 5 ו-7, וצפתה בשיעור השגיאה הלוגית יורד בכל פעם שהקוד גדל — בפקטור $\Lambda = 2.14 \pm 0.02$ לכל עלייה של שתי יחידות במרחק. הגדול ביותר, זיכרון distance-7 עם 101 קיוביטים, החזיק קיוביט לוגי עם שגיאה של $0.143\% \pm 0.003\%$ לכל מחזור תיקון, ו — הכותרת שבתוך הכותרת — שרד יותר זמן מהקיוביט הפיזי הטוב ביותר שלו, בפקטור 2.4 ± 0.3 . לזה קוראים "beyond", "breakeven", וזו הפעם הראשונה שבה כל מנגנון תיקון השגיאות החזיר את המחיר שהוא גובה בחומרה הזאת.

זהו ציון דרך אמיתי, וכדאי לדייק איזה סוג של ציון דרך. זו הוכחה שה-scaling סוף-סוף נע בכיוון הנכון. זה אינו מחשב קוונטי עובד, ו-המאמר אינו טוען שכן.

מה פירוש "surface code", "distance code", ו-"threshold" below?

Surface code רבים. פיזיים קיוביטים פני על המקודדת אחת מוגנת קוונטי מידע יחידת הוא לוגי קיוביט הרף ללא בודקים נוספים "מדידה" קיוביט שבה דו-ממדית, רשת גבי על הקידוד את לבצע מסוימת דרך הוא שגיאות של ביותר הקטן המספר זהו פיזי: מרחק אינו d הקוד מרחק המאוחסן. למידע להפריע בלי שגיאות patch פירושו יותר גדול d בכך. יבחין שהקוד בלי הלוגי הקיוביט את להשחית שיכולות מתאימים במיקומים d (עד בו-זמנית, שגיאות יותר ומתקן 1) $2d^2 -$ (בקירוב פיזיים קיוביטים יותר צורך הוא — יותר ועמיד גדול וצורכים בו-זמנית שגיאות ו-3 2, 1, מתקנים ו-7, 5, 3, מרחקים כאן, שנבדקו הגדלים שלושת לכן מהן. $1/2 -$ המינימום מעל מעט ב-101, השתמש Google של distance-7 זיכרון — פיזיים קיוביטים ו-17 49 97, בקירוב הלימוד. מספר

"מתחת לסף" הוא הביטוי הקריטי. תיקון שגיאות עוזר רק אם שיעור השגיאות הפיזי נמצא מתחת לערך קריטי; שם כל עלייה במרחק מדכאת את שיעור השגיאה הלוגית באופן אקספוננציאלי. פקטור הדיכוי Λ מודד זאת — $\Lambda < 1$ פירושו שהגדלת הקוד עוזרת, וככל ש- Λ גבוה יותר, כך טוב יותר. Google מדווחת על $\Lambda \approx 2.14$, כלומר כל עלייה של שתי יחידות במרחק חתכה בערך לחצי את שיעור השגיאה הלוגית. העובדה ש- Λ נמצא בנוחות מעל 1 היא לב התוצאה.



כך השגיאה הלוגית מצטברת לאורך מחזורי התיקון בזיכרונות, distance-3, 5-1-7 (מלמעלה למטה). הקו שכדאי לעקוב אחריו הוא הירוק המקווקו — הקיוביט הפיזי היחיד הטוב ביותר על השבב. זיכרון distance-7 (כחול, הנמוך ביותר) צובר שגיאה לאט יותר מהקו הזה, ולכן הקיוביט המקודד חי יותר מהקיוביט הפיזי הטוב ביותר שממנו הוא בנוי — “beyond breakeven”, “beyond breakeven” מופיע $\times 2.4$. זו תוצאת אורך החיים; דיכוי השגיאות below-threshold עצמו ($\Delta = 2.14$ כאשר הקוד גדל מ-3 ל-5 ול-7) מספיק במספרים שבטקסט.

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

מה עשו המחברים

- בנו זיכרונות code surface על שני שבבי Willow: מעבד 105 קיוביטים שהריץ קודים, distance-3, 5-1-7 למבחן scaling-ה (הגדול שבהם זיכרון distance-7 של 101 קיוביטים, מהם 49 קיוביטי נתונים), ומעבד 72 קיוביטים שהריץ זיכרון distance-5 עם decoder בזמן אמת וגם codes repetition במרחקים גבוהים.
- מדדו כיצד השגיאה הלוגית למחזור משתנה כאשר מרחק הקוד עולה מ-3 ל-5 ל-7, והפיקו את פקטור הדיכוי Δ .
- השוו את אורך חיי הקיוביט הלוגי לזה של הקיוביט הפיזי היחיד הטוב ביותר על אותו שבב, כדי לבדוק “breakeven”.
- הריצו קוד distance-5 עם decoder בזמן אמת — חומרה קלאסית שמפרשת את בדיקות השגיאה בקצב שבו הן נוצרות — עד מיליון מחזורים, כדי להראות שתיקון השגיאות מסוגל לעמוד בקצב המכונה.
- דחפו codes repetition פשוטים יותר עד distance 29 כדי לחפש מקורות שגיאה נדירים ועמוקים שקובעים רצפת ביצועים.

מה הם מצאו

- הקוד נמצא מתחת לסף. השגיאה הלוגית למחזור ירדה בפקטור $\Lambda = 2.14 \pm 0.02$ לכל עלייה של שתיים במרחק — דיכוי אקספוננציאלי נקי, בדיוק ההתנהגות שהתאוריה הבטיחה ושום מעבד לא הראה קודם באופן חד-משמעי.
- זיכרון distance-7 הגיע ל- $0.143\% \pm 0.003\%$ שגיאה למחזור, וחי פי 2.4 ± 0.3 יותר מהקיוביט הפיזי הטוב ביותר שלו — breakeven. beyond
- ה-decoding בזמן אמת עמד בקצב. ב-distance 5 ה-decoder הציג latency ממוצע של 63 מיקרו-שניות מול זמן מחזור של 1.1 מיקרו-שנייה, לאורך מיליון מחזורים — תיקון השגיאות פעל בזמן אמת, לא רק בניתוח בדיעבד.
- נשאר מקור שגיאות עמוק ונדיר. במבחני code, repetition הביצועים הוגבלו בסופו של דבר על ידי התפרצויות שגיאה מתואמות שהתרחשו בערך פעם בשעה (כ-1 בכל 10^9 מחזורים), וקבעו רצפת שגיאה סביב 10^{-10} – 10^{-01} שמקורה, לפי המחברים, עדיין אינו מובן.

מה זה לא מוכיח

- זה לא מחשב קוונטי שמבצע חישוב. זהו זיכרון קוונטי: הוא מאחסן ומגן על קיוביט לוגי אחד. הוא אינו מבצע פעולות לוגיות (gates) בין קיוביטים לוגיים ואינו מריץ אלגוריתם.
- זה לא אומר שאנחנו "קיוביט אחד" ממכונות שימושיות. קיוביט לוגי distance-7 צורך כ-101 קיוביטים פיזיים; שיעור שגיאה של כ-0.1% למחזור עדיין גבוה בהרבה מהטווח של בערך 10^{-6} עד 10^{-10} שאלגוריתמים אמיתיים דורשים. סגירת הפער דורשת מרחקים גדולים בהרבה — הרבה יותר קיוביטים פיזיים לכל קיוביט לוגי — ואלגוריתמים שימושיים דורשים אלפי קיוביטים לוגיים בו-זמנית. תקציב הקיוביטים הפיזיים מגיע למיליונים.
- ה-"אם נרחיב" עושה כאן עבודה אמיתית. מסקנת המאמר עצמה היא שביצועי המכשיר, אם יורחבו, עשויים לעמוד בדרישות של אלגוריתמים גדולים. להראות שהמגמה נכונה בקיוביט לוגי אחד אינו זהה לבניית המכונה המורחבת, ושום דבר כאן אינו מבטיח שהמגמה תשרוד בקנה מידה גדול בהרבה.
- רצפת השגיאה הלא-מוסברת היא בעיה חיה. התפרצויות השגיאה המתואמות שמגבילות את code repetition הן, לדברי המחברים, גדולות בסדרי גודל מהצפוי וימנעו יישומים fault-tolerant גדולים יותר עד שיובנו — פגם פתוח שמוצאה בבירור, לא פרט שכבר נפתר.
- התוצאה אינה אומרת דבר על שבירת הצפנה או "עליונות קוונטית" במשימות שימושיות. אלה דורשים את המכונה ה-fault-tolerant המלאה שעבורה זו אכן יסוד, לא הדגמה שלה.

עד כמה הראיות חזקות

- הטענה המרכזית מוצקה וחשובה. פעולה below-threshold עם דיכוי אקספוננציאלי נקי בשלושה מרחקי קוד, יחד עם אורך חיים beyond-breakeven ו-decoder עובד בזמן אמת, היא בדיוק השילוב שהתחום ניסה להגיע אליו, והוא מודגם ישירות ולא מוסק בעקיפין. זו אינה תוצאה שנוצרה מהיפך; זה הישג הנדסי אמיתי של קבוצה מובילה.
- המחברים זהירים לגבי ההיקף. הם ממסגרים זאת כ-זיכרון מתחת לסף, מסמנים בעצמם את רצפת השגיאות המתואמות הלא-מוסברת, ומתנים את העתיד ב"אם יורחב" הבולט. ההגזמה, כאשר היא מופיעה, נמצאת בסיקור שמעלה את "זיכרון מתוקן-שגיאות השתפר כאשר גדל" ל"מחשוב הקוונטי הגיע".
- הסטטוס הכנה הוא צעד יסודי שבוצע היטב. קיוביט לוגי אחד, מוגן מספיק כך שהוספת יתירות סוף-סוף עוזרת — כאשר הדרך של scaling, שערים לוגיים ושגיאות לא מוסברות עדיין ארוכה, קשה ולא מובטחת.

למה זה חשוב

מחשוב קוונטי חסין-תקלות תמיד סבל מבעיית ביצה-ותרנגולת: המכונות השימושיות דורשות שיעורי שגיאה שאף קיוביט פיזי אינו יכול להשיג, והפתרון — תיקון שגיאות — עובד רק אם החומרה כבר טובה מספיק להיות מתחת לסף. חציית הקו הזה, אפילו פעם אחת ואפילו בקיוביט לוגי יחיד, משנה את השאלה מ"האם זה בכלל אפשרי?" ל"עד כמה אפשר להרחיב את זה, ובאיזו מהירות?" זה שינוי אמיתי ומשמעותי, ולכן התוצאה ראויה לתשומת לב.

אבל אותה זהירות שהופכת את התוצאה לאמינה צריכה גם לרסן את הסיפור שסביבה. זו הלבנה הראשונה של יסוד, שהונחה היטב. זה לא הבניין, והאנשים שהניחו אותה הם הראשונים לומר זאת. הדרך הנכונה לעקוב אחר מחשוב קוונטי בשנים הקרובות היא בדיוק העקומה הלא-זוהרת הזאת: האם פקטור הדיכוי נשמר כשהקודים גדלים, האם אפשר לבצע שערים לוגיים באותה נקיות כמו זיכרון לוגי, והאם השגיאה המסתורית של פעם בשעה תוסבר אי פעם.

בקצרה

AI Quantum Google הראתה לראשונה באופן נקי שזיכרון קוונטי מסוג code surface יכול לפעול מתחת לסף: כאשר הקוד גדל מ-3 ל-5 ל-7, שיעור השגיאה הלוגית ירד אקספוננציאלית (בערך פי 2.14 לכל שתי דרגות), והגדול ביותר — זיכרון distance-7 של 101 קיוביטים — חי יותר מהקיוביט הפיזי הטוב ביותר שלו, כלומר עבר breakeven, בזמן שתיקון השגיאות פעל בזמן אמת. זה ציון דרך אמיתי ומבוקש זה זמן רב בהנדסת מחשבים קוונטיים. וזה גם קיוביט לוגי יחיד שמתפקד כזיכרון, עם שיעור שגיאה שעדיין רחוק ממה שאלגוריתמים אמיתיים דורשים, ללא פעולות לוגיות, עם רצפת שגיאה לא מוסברת שהמחברים עצמם מדגישים, ועם scaling של סדרי גודל רבים שעדיין לפנינו. סף אמיתי שנחצה — לא מחשב קוונטי שנמסר.

מקורות

920 ,638 Nature threshold, code surface the below correction error Quantum Collaborators, and AI Quantum Google (2025)

<https://doi.org/10.1038/s41586-024-08449-y>

a corrects — (2026) E5 ,653 Nature threshold, code surface the below correction error Quantum Correction: Author results (1 (Fig. below-threshold the affect not does keys; legend and label axis 3a Fig.

<https://www.nature.com/articles/s41586-026-10559-8>