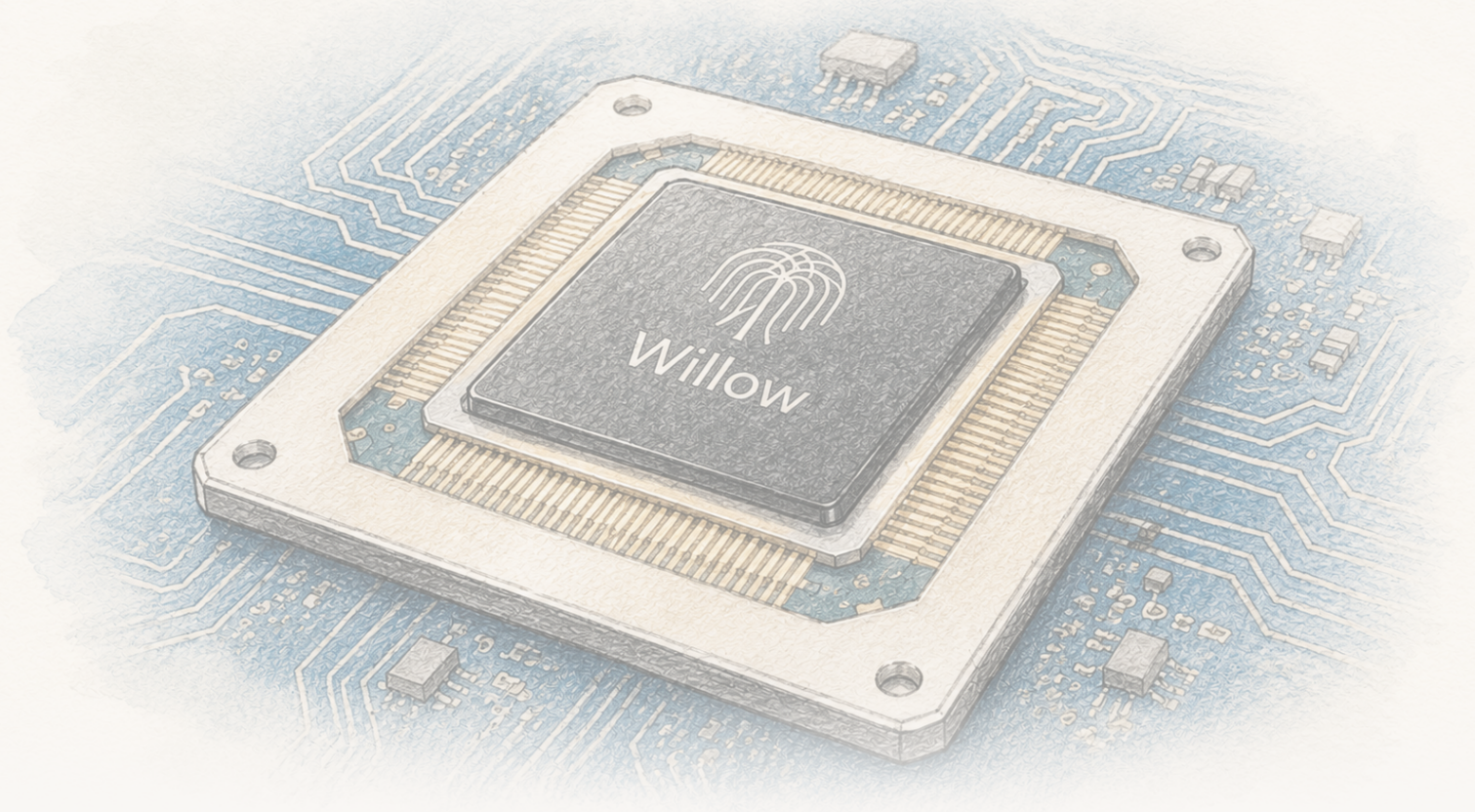




The CLEAN PAPER

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

አንድ quantum memory እያደገ ሲሄድ በመጨረሻ ተሻለ — እውነተኛው ምልክት፣ እና ያልሆነው ማሸን

Google Quantum AI ለመጀመሪያ ጊዜ በግልጽ surface-code quantum memory below threshold መስራት እንደሚችል አሳየ፤ ኮዱ ከdistance 3 ወደ 5 እና 7 ሲያድግ logical error rate በexponential መንገድ ቀነሰ (በሁለት ደረጃ ወደ 2.14x)፤ ትልቁ 101-qubit distance-7 memory ደግሞ ከምርጥ physical qubit የበለጠ ጊዜ ኖረ — beyond breakeven — error correctionም በreal time ተካሄደ። ለረጅም ጊዜ የተፈለገ የምህንድስና ምልክት ነው። ግን አንድ logical qubit እንደ memory የሚሰራ ስርዓት ብቻ ነው፤ የስህተት መጠኑ እውነተኛ algorithms ከሚፈልጉት አሁንም በጣም ይርቃል፤ logical operations አልተደረጉም፤ ያልተባራራ error floor አለ፤ እና ብዙ orders of magnitude የማስፋት መንገድ ይቀራል። Threshold ተሻግሯል፤ ኮምፒውተር ግን አልተሰጠም።

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

ግራፉ በመጨረሻ ወደ ትክክለኛው አቅጣጫ የተጣመመበት ጊዜ

ለሠላሳ ዓመታት የኳንተም ስህተት ማስተካከያ ሐሳብ በንድፈ ሐሳብ የተሰጠ ነገር ግን በሙከራ በግልጽ ያልተፈጸመ አንድ ተስፋ ላይ ተመሥርቶ ነበር። ንድፈ ሐሳቡ እንዲህ ይላል፡- አንድ የኳንተም መረጃ ክፍል — አንድ ሎጂካዊ ኩቢት — ብዙ ጫጫታ ያላቸው አካላዊ ኩቢቶች ላይ ቢከፋፈል፤ እና እነዚያ አካላዊ ኩቢቶች በቂ ጥራት ካላቸው፤ ተጨማሪ ኩቢቶች ማከል ሎጂካዊ ኩቢትን ማሻሻል አለበት። ኮዱ እየተስፋፋ ሲሄድ ስህተቶች በexponential መንገድ መቀነስ አለባቸው። ችግሩ “ካላቸው” በሚለው ውስጥ ነው። የአካላዊ ኩቢቶች የስህተት መጠን ከአንድ ወሳኝ የnoise ገደብ በታች ከሆነ፤ ተጨማሪ ኩቢቶች ይረዳሉ፤ ከዚያ በላይ ከሆነ ግን ተጨማሪ ኩቢቶች ተጨማሪ ጫጫታ ብቻ ያመጣሉ። ቀደም ያሉ ሙከራዎች ሁሉ በዚህ መስመር የተሳሳተ ወገን ላይ ነበሩ፤ ወይም አቅጣጫውን በግልጽ ማሳየት አልቻሉም። ኮዱን ማሳደግ ሁኔታውን የበለጠ ያባብሰው ነበር፤ አያሻሽለውም።

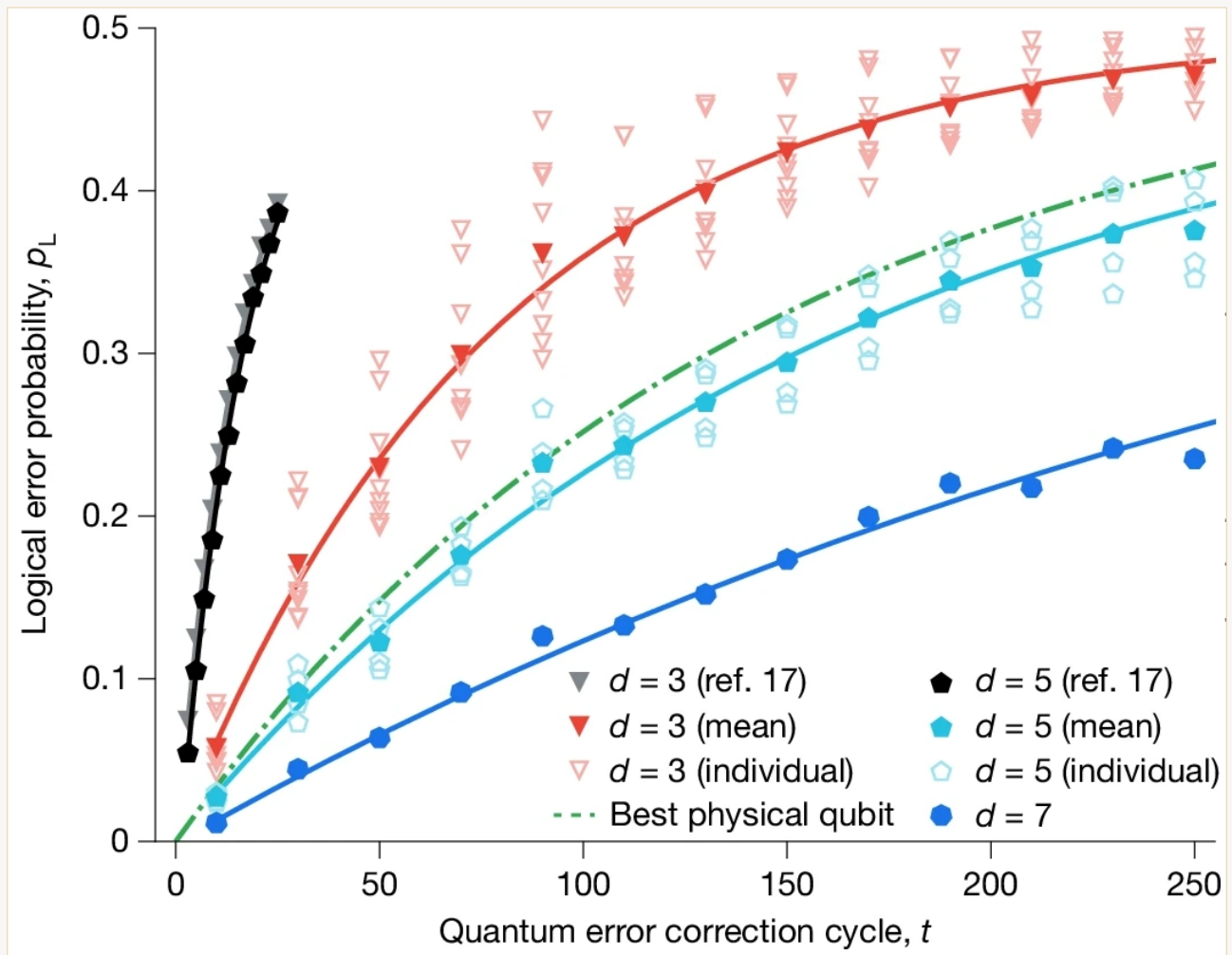
በታላላቅ 2024፤ Google ኳንተም AI የሌላኛውን ሁኔታ የመጀመሪያውን ግልጽ ማሳያ ዘገበ። በWillow፤ በአዲሱ ትውልድ ሱፐርኮንዳክቲቪቲንግ ፕሮሰሰሮቻቸው፤ በኮድ ርቀት 3፣ 5 እና 7 ገጽ-ኮድ ማህደሮችን ገነቡ። ኮዱ በየጊዜው ሲጨምር ሎጂካዊ ስህተት መጠን ሲቀንስ አዩ — ርቀት በሁለት ደረጃ በጨመረ ቁጥር በ $\Lambda = 2.14 \pm 0.02$ እጥፍ። ትልቁ፤ 101-ኩቢት ርቀት-7 ትውስታ፤ በእያንዳንዱ ስህተት-ማስተካከያ cycle $0.143\% \pm 0.003\%$ ስህተት ያለውን ሎጂካዊ ኩቢት ጠበቀ። ከዚያም — በርዕሱ ውስጥ ያለው ትልቁ ርዕስ — ከራሱ ምርጥ አካላዊ ኩቢት የበለጠ ጊዜ ኖረ፤ 02.4 ± 0.3 እጥፍ። ይህ “beyond breakeven” ተብሎ ይጠራል፤ እና በዚህ hardware ላይ የስህተት ማስተካከያ ሙሉ ስርዓት የሚጠቀምበትን ተጨማሪ ሀብት ከራሱ ጥቅም በማሳየት የከፈለበት የመጀመሪያው ጊዜ ነው።

ይህ እውነተኛ የእድገት ምልክት ነው፤ ግን ምን ዓይነት ምልክት እንደሆነ በትክክል መግለጽ አስፈላጊ ነው። የማስፋት አቅጣጫው አሁን ወደ ትክክለኛው አቅጣጫ እንደሄደ ያሳያል። የሚሰራ ኳንተም ኮምፒውተር አይደለም፤ ጽሑፉም እንደዚያ አይናገርም።

“ገጽ ኮድ”፤ “ርቀት” እና “below ገደብ” ምን ማለት ናቸው?

ሎጂካዊ ኩቢት ብዙ አካላዊ ኩቢቶች ላይ የተቀመጠ እና ከስህተት የተጠበቀ አንድ የኳንተም መረጃ ክፍል ነው። ገጽ ኮድ ይህን መረጃ በ2D ግራፍ ላይ ለመክተት የሚጠቅም አንድ ዘዴ ነው፤ ተጨማሪ “measure” ኩቢቶች የተከማቸውን መረጃ ሳይረብሹ ስህተቶችን በቀጣይነት ይፈትሻሉ። የኮዱ ርቀት d አካላዊ ርቀት አይደለም። ኮዱ ሳያውቅ አንድ ሎጂካዊ ኩቢትን ሊያበላሹ የሚችሉ በትክክል የተቀመጡ የስህተቶች አነስተኛ ቁጥር ነው። d ትልቅ ሲሆን ኮዱ የበለጠ ትልቅና ጠንካራ ይሆናል፤ ተጨማሪ አካላዊ ኩቢቶች ይጠቀማል (በግምት $2d^2 - 1$) እና በአንድ ጊዜ የሚከሰቱ ተጨማሪ ስህተቶችን — እስከ $(d - 1)/2$ — ማስተካከል ይችላል። ስለዚህ እዚህ የተፈተኑት ርቀት 3፣ 5 እና 7 በአንድ ጊዜ 1፣ 2 እና 3 ስህተቶችን ማስተካከል ይችላሉ፤ በግምት 17፣ 49 እና 97 አካላዊ ኩቢቶች ይጠቀማሉ። Google የገነባው ርቀት-7 ትውስታ 101 ኩቢቶች ተጠቅሟል፤ ከዚህ የመማሪያ መጽሐፍ አነስተኛ መጠን ትንሽ በላይ።

“Below ገደብ” ዋናው አገላለጽ ነው። ስህተት ማስተካከያ የሚረዳው የአካላዊ ስህተት መጠን ከአንድ ወሳኝ እሴት በታች ከሆነ ብቻ ነው። በዚያ ክልል ውስጥ ርቀት በጨመረ ቁጥር ሎጂካዊ ስህተት መጠን በexponential መንገድ ይቀንሳል። የመቀነሻ ፋክተሩ Λ ይህን ይለካል — $\Lambda > 1$ ከሆነ ኮዱን ማሳደግ ይረዳል፤ Λ ትልቅ በሆነ ቁጥርም የተሻለ ነው። Google $\Lambda \approx 2.14$ ይዘግባል፤ ይህም ርቀት በሁለት ደረጃ በጨመረ ቁጥር ሎጂካዊ ስህተት መጠን በግምት በግማሽ ተቆርጧል ማለት ነው። Λ ከ1 በግልጽ በላይ መሆኑ የዚህ ውጤት ዋና ነገር ነው።



ለርቀት-3፣ -5 እና -7 memories (ከላይ ወደ ታች) ሎጂካዊ ስህተት በስህተት-ማስተካከያ cycles ሂደት እንዴት እንደሚከማች ያሳያል። መመልከት ያለብን አረንጓዴ ተቆራረጥ መስመር ነው — በchip ላይ ያለው ምርጥ ነጠላ አካላዊ ኩቢት። ርቀት-7 ትውስታ (ሰማያዊ፣ ከሁሉ በታች) ከዚያ መስመር የበለጠ ቀስ ብሎ ስህተት ይከማቻል። ስለዚህ encoded ኩቢት ከተገነባበት ምርጥ አካላዊ ኩቢት የበለጠ ይኖራል — “beyond breakeven”፣ በ $2.4\times$ እጥፍ። ይህ የlifetime ውጤት ነው፤ below-ገደብ suppression ራሱ ($\Lambda = 2.14$ ፣ ኮዱ ከርቀት 3 ወደ 5 እና 7 ሲያድግ) በጽሑፉ ውስጥ ባሉት ቁጥሮች ይታያል።

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

ደራሲዎቹ ምን አደረጉ?

- በሁለት Willow chips ላይ ገጽ-ኮድ memories ገነቡ። አንደኛው 105-ኩቢት processor ነበር፤ ርቀት-3፣ -5 እና -7 codesን ለscaling ሙከራ አስኬደ፤ ትልቁ 101-ኩቢት፣ 49-መረጃ-ኩቢት ርቀት-7 ትውስታ ነበር። ሁለተኛው 72-ኩቢት processor ደግሞ ርቀት-5 ትውስታን ከእውነተኛ-ጊዜ decoder ጋር እንዲሁም ከፍተኛ-ርቀት repetition codesን አስኬደ።
- የኮድ ርቀትን ከ3 ወደ 5 እና 7 ሲያሳድጉ በእያንዳንዱ cycle ያለው ሎጂካዊ ስህተት እንዴት እንደተቀየረ ለኩ፣ የsuppression factor Λ ንም አወጡ።
- “Breakeven” ለመፈተን የሎጂካዊ ኩቢትን የሕይወት ጊዜ በተመሳሳይ chip ላይ ካለው ምርጥ ነጠላ አካላዊ ኩቢት ጋር አነጻጸሩ።
- ርቀት-5 ኮድን እውነተኛ-ጊዜ decoder ጋር እስከ አንድ ሚሊዮን cycles ድረስ አስኬዱ። Decoder ማለት የስህተት

checks በሚመጡበት ፍጥነት የሚተረጉም የተለመደ ክላሲካል hardware ነው። ይህ የስህተት ማስተካከያ ስርዓቱ ከማሽኑ ፍጥነት ጋር መከታተል እንደሚችል ለማሳየት ነበር።

- አፈጻጸሙ የማይወርድበትን ወለል የሚያስከትሉ አልፎ አልፎ የሚከሰቱ ጥልቅ የስህተት ምንጮችን ለመፈለግ ቀለል ያሉ repetition codesን እስከ ርቀት 29 ድረስ አሳደጉ።

ምን አገኙ?

- ኮዱ below ገደብ ነው። ርቀት በሁለት ደረጃ በጨመረ ቁጥር ሎጂካዊ ስህተት per cycle በ $\Lambda = 2.14 \pm 0.02$ እጥፍ ቀነሰ — ንጹህ exponential suppression፤ ንድፈ ሐሳቡ የተነበየው እና ከዚህ በፊት ምንም processor በግልጽ ያሳየው ባህሪ።
- ርቀት-7 ትውስታ በእያንዳንዱ cycle $0.143\% \pm 0.003\%$ ስህተት ደረሰ፤ እና ከምርጥ አካላዊ ኩቢት 2.4 ± 0.3 እጥፍ የበለጠ ጊዜ ኖረ — beyond breakeven።
- እውነተኛ-ጊዜ decoding ከማሽኑ ጋር ተከታተለ። በርቀት 5 decoder በአማካይ 63 microseconds latency ነበረው፤ ከ1.1-microsecond cycle ጊዜ ጋር፤ እና ይህን እስከ አንድ ሚሊዮን cycles ቀጠለ። ስህተት ማስተካከያ ከሙከራው በኋላ በትንተና ብቻ ሳይሆን በቀጥታ ተካሄደ።
- አልፎ አልፎ የሚከሰት ጥልቅ የስህተት ምንጭ አሁንም አለ። በrepetition-ኮድ ሙከራዎች፤ አፈጻጸሙ በመጨረሻ በግምት በሰዓት አንድ ጊዜ በሚከሰቱ እርስ በርስ የተያያዙ ስህተት bursts (በ 3×10^9 cycles አንድ ገደማ) ተገደበ። ይህ ወደ 10^{-10} የሚጠጋ የስህተት floor ያስቀምጣል፤ ምንጩም ደራሲዎቹ እንደሚሉት ገና አልተረዳም።

ይህ ምን አያረጋግጥም?

- ይህ ስሌት የሚሰራ ኳንተም ኮምፒውተር አይደለም። ይህ ኳንተም ትውስታ ነው፤ አንድ ሎጂካዊ ኩቢት ያከማቻል እና ይጠብቃል። በሎጂካዊ ኩቢቶች መካከል ሎጂካዊ operations (gates) አያደርግም፤ ምንም አልጎሪዝምም አያስኬድም።
- ጠቃሚ ከሆነ ማሽን አንድ ኩቢት ብቻ አይርቅም። ርቀት-7 ሎጂካዊ ኩቢት ወደ 101 አካላዊ ኩቢቶች ይጠቀማል። በእያንዳንዱ cycle 0.1% የስህተት መጠን እውነተኛ አልጎሪዝሞች ከሚፈልጉት ወደ 10^{-6} እስከ 10^{-10} በጣም ከፍ ያለ ነው። ይህን ክፍተት ለመዝጋት ከፍተኛ የኮድ ርቀት ያስፈልጋል — በእያንዳንዱ ሎጂካዊ ኩቢት ብዙ ተጨማሪ አካላዊ ኩቢቶች። ጠቃሚ አልጎሪዝሞች ደግሞ በሺዎች የሚቆጠሩ ሎጂካዊ ኩቢቶችን በአንድ ጊዜ ይፈልጋሉ። ለዚህ የሚያስፈልጉ አካላዊ ኩቢቶች በሚሊዮኖች ሊቆጠሩ ይችላሉ።
- “If scaled” የሚለው አገላለጽ እውነተኛ ገደብ ነው። ጽሑፉ ራሱ የመሣሪያው አፈጻጸም በስፋት ከተገነባ ትላልቅ አልጎሪዝሞች የሚፈልጉትን ሊያሟላ እንደሚችል ይደመድማል። በአንድ ሎጂካዊ ኩቢት ላይ አቅጣጫው ትክክል እንደሆነ ማሳየት በስፋት የተገነባውን ማሽን መገንባት አይደለም፤ እና አቅጣጫው በጣም ትልቅ መጠን ላይ እንደሚቀጥል ምንም ዋስትና የለም።
- ያልተብራራው ስህተት floor አሁንም የሚፈታ ችግኝ ነው። የrepetition-ኮድ አፈጻጸምን የሚገድቡት የተያያዙ ስህተት bursts ደራሲዎቹ እንደሚሉት ከሚጠበቀው በብዙ orders of magnitude ትልቅ ናቸው፤ እና ምንጫቸው እስካልተረዳ ድረስ ትላልቅ fault-tolerant applicationsን ሊከለክሉ ይችላሉ። ይህ በግልጽ የተገለጸ ክፍት ችግኝ ነው፤ የተፈታ ዝርዝር አይደለም።
- Encryption መስበር ወይም ለጠቃሚ ሥራዎች “ኳንተም supremacy” ማሳየት ስለሚቻል ምንም አይናገርም። እነዚህ ይህ ሙከራ የመሠረት ድንጋይ የሆነለትን ሙሉ fault-tolerant ማሽን ይፈልጋሉ፤ ይህን ማሳያ ብቻ አይደለም።

ማስረጃው ምን ያህል ጠንካራ ነው?

- **ዋናው አባባል ጠንካራና አስፈላጊ ነው።** በሦስት ኮድ distances ላይ ግልጽ exponential suppression ያለው below-ገደብ operation፣ beyond-breakeven lifetime እና የሚሰራ እውነተኛ-ጊዜ decoder በአንድ ላይ መኖራቸው መስኩ ለማሳካት ሲሞክር የነበረው ጥምረት በትክክል ነው፤ እና ይህ በቀጥታ ታይቷል እንጂ ከሌላ ነገር አልተገመተም። ይህ የhype ውጤት አይደለም፤ ከመሪ ቡድን የመጣ እውነተኛ የምህንድስና ውጤት ነው።
- **ደራሲዎቹ በውጤቱ ወሰን ላይ ጥንቃቄ ያደርጋሉ።** ይህን below-ገደብ ትውስታ እንደሆነ ይገልጻሉ፤ ያልተብራራውን correlated-ስህተት floor ራሳቸው ያሳያሉ፤ ወደፊቱም “if scaled” በሚለው ግልጽ ገደብ ያስቀምጣሉ። “ስህተት-corrected ትውስታ ኩቢት እያደገ ሲሄድ ተሻለ” የሚለውን “ኳንታም computing ደርሷል” ብሎ የሚያጋንነው በዙሪያው ያለው ሽፋን ነው።
- **ትክክለኛው ሁኔታ መሠረታዊ እርምጃ በጥሩ ሁኔታ መወሰዱ ነው።** አንድ ሎጂካዊ ኩቢት፣ redundancy ማከል በመጨረሻ እንዲረዳው በቂ ጥበቃ ያገኘ — ግን ማስፋት፣ ሎጂካዊ gates እና ያልተብራሩ ስህተቶች ያሉበት ረጅም፣ ከባድ እና ገና ያልተረጋገጠ መንገድ ከፊት አለ።

ለምን ይጠቅማል?

Fault-tolerant ኳንታም computing ሁልጊዜ የ“ዶሮ እና እንቁላል” ችግኝ ያለው ይመስላል። ጠቃሚ የሚሆኑ ማሽኖች ምንም አካላዊ ኩቢት በራሱ ሊደርስበት የማይችለውን ዝቅተኛ ስህተት መጠን ይፈልጋሉ፤ መፍትሔው — ስህተት ማስተካከያ — ግን hardwareው ቀድሞ ከገደብ በታች ለመሆን በቂ ጥራት ካለው ብቻ ይሰራል። ይህን መስመር አንድ ጊዜ እንኳን፣ በአንድ ሎጂካዊ ኩቢት ላይ እንኳን ማለፍ፣ “ይህ ፈጽሞ ሊሰራ ይችላል?” የሚለውን ጥያቄ “ምን ያህል ሊሰፋ ይችላል? ምን ያህልስ ፈጣን?” ወደሚል ጥያቄ ይቀይረዋል። ይህ እውነተኛና ትርጉም ያለው ለውጥ ነው፤ ውጤቱም ትኩረት የሚገባው ለዚህ ነው።

ነገር ግን ውጤቱን እምነት የሚያስገኝለት ይኸው ጥንቃቄ ነው በዙሪያው ያለውን ታሪክ ማስተካከል ያለበት። ይህ የመሠረት የመጀመሪያው ጡብ ነው፤ በጥሩ ሁኔታ የተቀመጠ። ሕንፃው ግን አይደለም፤ እና ጡቡን ያስቀመጡት ሰዎች ራሳቸው በመጀመሪያ ይህን ይናገራሉ። በሚቀጥሉት ጥቂት ዓመታት ኳንታም ኮምፒውቲንግን ለመከታተል ትክክለኛው ነገር ይህ ብዙም የማያበራ ግራፍ ነው፤ ኮዶቹ ሲያድጉ suppression factor ይቆያልን? ሎጂካዊ gates እንደ ሎጂካዊ ትውስታ በንጽህና ሊደረጉ ይችላሉን? እና ያ በሰዓት አንድ ጊዜ የሚመጣው ሚስጥራዊ ስህተት መቼም ይብራራልን?

ንጹህ ማጠቃለያ

Google ኳንታም AI ለመጀመሪያ ጊዜ በግልጽ ገጽ-ኮድ ኳንታም ትውስታ below ገደብ መስራት እንደሚችል አሳየ። ኮዱን ከርቀት 3 ወደ 5 እና 7 ሲያሳድጉ ሎጂካዊ ስህተት መጠን በexponential መንገድ ቀነሰ (በሁለት ርቀት ደረጃ ወደ $2.14 \times$)። ትልቁ፣ 101-ኩቢት ርቀት-7 ትውስታ፣ ከምርጥ አካላዊ ኩቢት የበለጠ ጊዜ ኖረ — beyond breakeven — እና ስህተት ማስተካከያ በእውነተኛ ጊዜ ተካሄደ። ይህ በኳንታም ኮምፒውተር ምህንድስና ውስጥ ለረጅም ጊዜ የተፈለገ እውነተኛ የእድገት ምልክት ነው። ነገር ግን ይህ አንድ ሎጂካዊ ኩቢት እንደ ትውስታ የሚሰራ ስርዓት ነው፤ የስህተት መጠኑ እውነተኛ አልጎሪዝሞች ከሚፈልጉት አሁንም በጣም ይርቃል፤ ሎጂካዊ operations አልተደረጉም፤ ደራሲዎቹ ራሳቸው የሚጠቅሱት ያልተብራሩ ስህተት floor አለ፤ እና ከፊት በብዙ orders of magnitude የሚለካ የማስፋት መንገድ አለ። እውነተኛ ገደብ ተሻግሯል — ኳንታም ኮምፒውተር ግን አልተሰጠም።

ምንጮች

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>