



The CLEAN PAPER

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

Clinical AI safety signal አላሳየም እና paperworkን አሻሽሏል — ግን patient outcomesን አላሻሻለም

በ16 Kenyan primary care facilities ውስጥ pragmatic cluster-randomized trial በclinical officers መዝገብ ላይ generative-AI decision-support tool ፈተነ። በ9,691 patients strict 14-day expert-adjudicated treatment-failure composite 2.2% በtool arm እና 2.0% በcontrol ነበር (adjusted OR 0.77, 95% CI 0.55–1.08, $P = 0.13$) — significant difference የለም። Toolው safety signal አላሳየም፣ documentationን አሻሽሏል፣ prescribing እና satisfactionን አልቀየረም፣ antibiotic costs ወደ ታች አመለከቱ። ይህ failed study አይደለም፤ patients ላይ የተለካ አልፎ አልፎ የሚገኝ ትክክለኛ study ነው፣ እና process የሚያሻሽል tool በሁለት ሳምንት patient benefitን እንደማያሳይ ያሳያል።

Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

Paper: *Generative AI-enabled clinical decision support system in primary care: a pragmatic, cluster-randomized trial*, Ambrose Agweyu, Paul Mwaniki, Vaishnavi Menon, Robert Korom, Lynda Isaaka, Conrad Wanyama, Xiaoxuan Liu & Bilal A. Mateen (and colleagues), *Nature Medicine* (2026) · DOI: 10.1038/s41591-026-04503-6

ደህንነቱ ጥሩ እና ጠቃሚ መሆን በታካሚው ላይ ውጤታማ መሆን ጋር አንድ አይደለም

ስለ የሕክምና AI የሚወጡ ብዙ ርዕሶች ከተሳሳተ ዓይነት ጥናት ይጀምራሉ። ሞዴል nexam questions ላይ ከፍተኛ ነጥብ ያገኛል፤ ወይም curated vignettes ላይ doctorsን ይበልጣል፤ ከዚያ ታሪኩ “ማሽኑ ለclinic ዝግጁ ነው” ይሆናል።

ይህ ሙከራ ይበልጥ ከባድ እና እጅግ አልፎ አልፎ የሚደረግ ነገር አደረገ። Generative-AI decision-ድጋፍ መሣሪያውን በእውነተኛ የመጀመሪያ ደረጃ ሕክምና፣ በእውነተኛ የጤና ተቋማት፣ በእውነተኛ ታካሚዎች ላይ አስገብቶ በእውነት የሚጠቅመውን ጥያቄ ጠየቀ፡- ታካሚዎች የተሻለ ውጤት አገኙን?

ትክክለኛው መልስ “አይ” ነው — ቢያንስ በ14 ቀናት ውስጥ ሊለካ የሚችል ልዩነት አልታየም። መሣሪያው ደህንነት ምልክት አላሳየም። ክሊኒካዊ ሰነድ ማዘጋጀትን አሻሽሏል። ምናልባት አንዳንድ የdrug ወጪዎችንም ቀንሷል። ግን ሕክምና ውድቀቶችን በስታቲስቲካዊ ጉልህ ሁኔታ አልቀነሰም፤ ደራሲዎቹም ጥቅም ካለ ምናልባት መጠኑ ትንሽ እንደሆነ በጥንቃቄ ይጽፋሉ።

ይህ የጥናቱ ውድቀት አይደለም። ጥናቱ ትክክለኛ ሥራውን እየሠራ ነው። የሕክምና AIን በመመዘኛዎች ሳይሆን በታካሚዎች ላይ ሲለኩ ኃላፊነት ያለው ማስረጃ ይህን ይመስላል።

Process improved; patient outcome not proven

Safe-looking decision support is not the same as a demonstrated clinical benefit.

No safety signal

Looked safe in this trial

Not powered to settle rare harms.



Workflow improved

Documentation quality rose

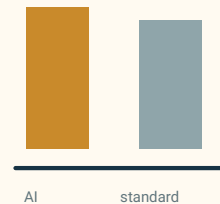
Process signal, not outcome proof.



Primary outcome null

14-day treatment failure

2.2% vs 2.0%; CI includes no effect.



Boundary: useful to the clinical process; not proven to improve patient outcomes within 14 days.

መሣሪያው በዚህ ሙከራ ደህንነት ምልክት አላሳየም እና ክሊኒካዊ ሰነድ ማዘጋጀትን አሻሽሏል፤ ግን ቀድሞ የተወሰነው 14-day ታካሚ ውጤት በጉልህ ሁኔታ አልተሻሻለም። Process ማሻሻል ከተረጋገጠ ታካሚ ጥቅም ጋር አንድ አይደለም።

ደራሲዎቹ ምን አደረጉ?

ቡድኑ በKenya፣ Nairobi እና Kiambu counties በPenda ጤና የሚንቀሳቀሱ 16 የመጀመሪያ ደረጃ ሕክምና የጤና ተቋማት ውስጥ pragmatic፣ cluster-በዘፈቀደ የተመደበ ሙከራ አካሄደ። በእነዚህ የጤና ተቋማት አብዛኛውን እንክብካቤ ክሊኒካዊ officers — የሦስት-ዓመት diploma ያላቸው mid-level practitioners — ይሰጣሉ፣ ብዙ ጊዜ senior ምክክር በቀላሉ አይገኝም።

Randomization በክሊኒኬያን ደረጃ ነበር፣ በታካሚ አይደለም። 103 ክሊኒካዊ officers በዘፈቀደ የተመደበ ሆኑ፡- 52 ወደ ጣልቃ ገብነት arm፣ 51 ወደ የቁጥጥር ቡድን። ሁለቱም arms ተመሳሳይ cloud-based ኤሌክትሮኒክ የሕክምና መዝገብ ተጠቀሙ። Intervention arm በተጨማሪ AI Consult 2.0 የተባለ፣ በOpenAI GPT-4o ላይ የተገነባ decision-ድጋፍ መሣሪያ ተጠቀመ። መሣሪያው ክሊኒኬያን የጻፈውን መረጃ አንብቦ ምርመራ ወይም ሕክምና plan ላይ ሊኖሩ የሚችሉ ችግኞችን ያሳይ ነበር። Clinicians ሙሉ autonomy ነበራቸው፡- suggestionን ሊቀበሉ፣ ሊቀይሩ ወይም ሊተዉ ይችላሉ።

የትኛው ሞዴል ነበር? እና ዝርዝሩ ለምን ያስፈልጋል?

መሣሪያው AI Consult 2.0 ነበር፣ OpenAI GPT-4o (May 2025 release) ይጠቀም ነበር። OpenAI commercial API በenterprise licence ተጠቅሞ፣ low-የዘፈቀደነት አውድ (temperature 0.1) ላይ ነበር። በEasyClinic bespoke EMR ውስጥ ተካትቶ ነበር፣ ሥርዓት prompts ከKenyan national ሕክምና guidelines ጋር እንዲስማማ ተጻፉ፤ ደራሲዎቹ full instruction promptን አሳትመዋል።

ዝርዝሩ አስፈላጊ ነው ምክንያቱም ውጤቱ ስለ አንድ specific ሥርዓት — አንድ ሞዴል version፣ prompt፣ መዝገብ ሥርዓት እና አውድ — ነው፣ “LLMs in medicine” በአጠቃላይ አይደለም። ደራሲዎቹም ውጤታቸውን temporal መመዘኛ rather than a fixed estimate of capability ይሉታል። አዲስ ሞዴል፣ ሌላ prompt ወይም ያነሰ digitized clinic ውጤቱን ሊቀይር ይችላል።

OpenAI በኋላ in-kind ድጋፍ — cloud-compute credits እና API technical guidance — ሰጥቷል። ደራሲዎቹ ግን OpenAIን የመጠቀም ውሳኔ ከዚያ እርዳታ በፊት እንደተደረገ እና OpenAI በሙከራ ንድፍ፣ መረጃ collection፣ ትንተና ወይም publish ለማድረግ ውሳኔ ላይ ሚና እንዳልነበረው ይገልጻሉ።

ከ22 April እስከ 16 July 2025 9,691 ታካሚዎች ተመዝግበዋል። Primary ውጤት ሆነ ተብሎ ታካሚ-centered እና ጥብቅ ተደርጎ ተመርጧል፡- ከvisit በኋላ በ14 ቀናት ውስጥ expert-adjudicated composite of ሕክምና ውድቀት። የክሊኒኬያኖች panel የጥናት armን ሳያውቅ እያንዳንዱ ታካሚ unresolved ወይም worsening illness ያለ መጥፎ ውጤት እንዳጋጠመው ገምግሟል። ሙከራው ቀድሞ ተመዝግቧል (Pan-African ክሊኒካዊ ሙከራዎች Registry 202502499779176)።

ይህ የንድፍ ምርጫ ዋናው ነጥብ ነው። AI መሣሪያ ክሊኒኬያን የሚጽፈውን መቀየሩን ማሳየት ቀላል ነው። በታካሚ ላይ የሚከሰተውን መቀየሩን ማሳየት ይበልጥ ከባድ እና ይበልጥ ጠቃሚ ነው።

ምን አገኙ?

Primary ውጤት አልተሻሻለም። ሕክምና ውድቀት በAI arm 102/4,693 ታካሚዎች (2.2%)፣ በየቁጥጥር ቡድን 94/4,654 (2.0%) ነበር። Raw percentages በAI ትንሽ ከፍ ነበሩ፣ hadjustment በኋላ ግን point estimate ወደ ጥቅም አቀና፡- የተስተካከለ odds ratio 0.77 (95% CI 0.55–1.08, P = 0.13)። ይህ ስታቲስቲካዊ ጉልህ አይደለም። Raw እና የተስተካከለ ቁጥሮች የተለያዩ አቅጣጫ ማሳየታቸው ስህተት አይደለም፤ adjustment በክሊኒኬያን

clusters መካከል ያሉ ልዩነቶችን ይቆጥራል። Confidence interval “no ተፅዕኖ”ን በግልጽ ያካትታል፤ ስለዚህ ጥቅም ማለት አይቻልም፤ absolute ተፅዕኖም ትንሽ ነበር። Confidence interval፤ odds ratio እና P valueን አንድ ላይ እንዴት ማንበብ እንደሚቻል የክሊኒካዊ ውጤቶች መመሪያን ይመልከቱ።

ደህንነት ምልክት አልታየም — በዚህ ገደብ ውስጥ። Serious adverse ክስተቶች ከመሣሪያው ጋር ተያይዘው አልተመዘኑም፤ independent ግምገማም ደህንነት ምልክት አላገኘም። ግን ሙከራው አልፎ አልፎ የሚከሰት severe harmsን ለመያዝ powered አልነበረም፤ prespecified noninferiority ወይም መደበኛ ደህንነት ማዕቀፍም አልነበረውም። ስለዚህ uncommon ክስተቶች ላይ “ደህንነት ተረጋግጧል” ማለት አይቻልም።

ሰነድ ማዘጋጀት ተሻሻለ። በblinded ባለሙያዎች ከተገመገሙ 2,000 encounters ውስጥ AI Consult የተጠቀሙ ክሊኒሺያኖች በሁሉም የተገመገሙ ዘርፎች — recorded ምርመራ፤ ሕክምና plan እና overall completeness — የተሻለ ክሊኒካዊ ሰነድ ማዘጋጀት አሳዩ።

Prescribing ብዙ አልተቀየረም። Prescribing ላይ ጉልህ ልዩነት አልነበረም፤ correct antibiotic አጠቃቀምን ጨምሮ (የተስተካከለ OR 0.86, 95% CI 0.48–1.55)። Antibiotic prescribing መጠኖችም አልተቀየሩም።

ታካሚዎች ልዩነት አላስተዋሉም። 826 እርካታ survey የሞሉ ታካሚዎች መካከል እርካታ በሁለቱ arms በመሠረቱ ተመሳሳይ ነበር፤ ምክክር timesም ተመሳሳይ ነበሩ።

Costs ትንሽ ወደ ታች አመለከቱ። Adjusted ትንተና ውስጥ antibiotic-related ወጪዎች በAI arm ዝቅ ነበሩ — ምናልባት የprescription ብዛትን በመቀነስ ሳይሆን ርካሽ choices በመምረጥ። Per-ታካሚ antibiotic saving መሣሪያውን ለማስኬድ ከሚወጣው per-ታካሚ ወጪ ይበልጥ መስሎ ታይቷል። ደራሲዎቹ ይህን suggestive ብለው ይጠሩታል፤ full total-ወጪ-of-ownership ግምገማ ከሙከራው ውጭ ነበር።

የደራሲዎቹ አንድ-ሐረግ summary ንጹሁ ነው፡- LLM assistance በእነዚህ ገደቦች ውስጥ ደህንነት ምልክት አላሳየም፤ ግን በ14 ቀናት ውስጥ ሕክምና ውድቀትን አልቀነሰም፤ ጥቅም ካለም ምናልባት መጠኑ ትንሽ ነው።

“ጉልህ ልዩነት አልነበረም” ማለት “አይሠራም” ማለት አይደለም

Null primary ውጤት በሁለቱም አቅጣጫ ከመጠን በላይ ሊተረጎም ይችላል። ሁለት ነገሮች ቀላሉን ታሪክ ያቆማሉ።

መጀመሪያ፡ ሙከራው ከታየው ይበልጥ ትልቅ ተፅዕኖ ለመያዝ ተነድፎ ነበር። በመጀመሪያ ደረጃ ሕክምና ውስጥ serious bad ውጤቶች ዝቅ ናቸው — እዚህ በግምት 2%። ትንሽ እውነተኛ ጥቅምን ከnoise ለመለየት እጅግ ብዙ ታካሚዎች ያስፈልጋሉ። የደራሲዎቹ post-hoc power calculations የታየውን ተፅዕኖ size ለማግኘት ከ100,000 በላይ ታካሚዎች የሚደርስ ትልቅ ሙከራ ሊያስፈልግ እንደሚችል ያመለክታል። Non-ጉልህ ውጤት ትንሽ እውነተኛ ጥቅም እንደሌለ አያስወግድም፤ ይህ ሙከራ መለየት አልቻለም ማለት ነው።

ሁለተኛ፡ comparisonው በከፊል ተደባልቋል። Configuration ስህተት አንዳንድ ቁጥጥር-arm ክሊኒሺያኖች ለጥቂት ጊዜ AI Consult እንዲጠቀሙ አደረገ፤ እና በአንድ shared network ውስጥ ክሊኒሺያኖች እርስ በርሳቸው ይነጋገራሉ እና habits ከአንድ arm ወደ ሌላ ሊያልፉ ይችላሉ። ሁለቱም ተፅዕኖዎች ቡድኖቹን ይበልጥ ተመሳሳይ ያደርጋሉ፤ እውነተኛ difference ካለ ወደ zero ይገፉታል። በተጨማሪም host network ቀድሞ ወደ ከፍተኛ standards ቅርብ ነበር፤ ስለዚህ improvement ለማሳየት ቦታ ያነሰ ነው።

ይህ ሁሉ “breakthrough” ርዕስን አያድንም። ግን ትክክለኛው ንባብ ዝቅ የሚያደርግ ሳይሆን calibrated መሆን አለበት፡- በከባዱ እና ትክክለኛው ታካሚ የመጨረሻ መለኪያ ላይ ይህ መሣሪያ በሁለት ሳምንት ውስጥ demonstrable

ጥቅም አላሳየም፣ ነገር ግን ሰነድ ማዘጋጀትን አሻሽሏል እና ደህንነት ምልክት አላሳየም።

ይህ ምን አያረጋግጥም?

- AI ሥርዓቱ ታካሚ ውጤቶችን እንዳሻሻለ አያሳይም። በprimary 14-day የመጨረሻ መለኪያ ላይ ጉልህ ጥቅም አልነበረም።
- AI ሥርዓቱ የማይጠቅም እንደሆነ አያሳይም። Point estimate ወደ ጥቅም አቀና፣ ሰነድ ማዘጋጀት በሁሉም ዘርፎች ተሻሻለ፣ drug ወጪዎችም ወደ ታች አመለከቱ። Null ውጤት ትንሽ ጥቅም እንዳለ ግን ሙከራው ለማረጋገጥ በቂ እንዳልነበር ጋር ይስማማል።
- Rare harms ላይ ደህንነት ተረጋግጧል ብሎ አያሳይም።
- “AI beats doctors” ወይም ክሊኒሺያኖችን ይተካል ብሎ አያሳይም። ይህ decision ድጋፍ ነው፣ ክሊኒሺያን ሙሉ authority ነበረው።
- Automatically generalize አይደረግም። ሙከራው በአንድ private urban Kenyan network ውስጥ ነበር።
- Cost savings አያረጋግጥም። የወጪ ምልክት suggestive ነው፣ full economic ግምገማ አይደለም።

ማስረጃው ምን ያህል ጠንካራ ነው?

ዋናው አባባል — በshort-term ታካሚ harm ላይ proven reduction አልታየም — በጥናት ንድፍ አንፃር ጠንካራ ነው፣ በconclusion አንፃር ግን ተገቢ ጥንቃቄ አለው። Prospective፣ preregistered፣ cluster-በዘፈቀደ የተመደበ ሙከራ፣ blinded ታካሚ-level composite ውጤት ጋር ለእንዲህ ዓይነት መሣሪያ ሊሰበሰብ ከሚችለው ምርጥ በእውነተኛ ዓለም ማስረጃ ጋር ቅርብ ነው። ከመመዘኛ ነጥብ ወይም vignette ጥናት በጣም ይበልጣል።

Secondary findings — የተሻለ ሰነድ ማዘጋጀት፣ unchanged prescribing፣ similar እርካታ፣ lower antibiotic ወጪዎች — ማስረጃው ጥሩ ነው፣ ግን secondary ተብሎ መነበብ አለበት፣ ዋና ርዕስ አይደለም፣ contamination እና single-network limitsም ይነካቸዋል። ደህንነት ላይ ማስረጃ አበረታች ነው ግን bounded፡- ምልክት አልታየም፣ ግን አልፎ አልፎ የሚከሰት harms ለመያዝ የተነደፈ ጥናት አይደለም።

በጣም ጠቃሚው stance “ይሠራል” ወይም “ወደቀ” አይደለም። ትክክለኛው፡- በጥንቃቄ የተደረገ በእውነተኛ ዓለም ሙከራ ይህ AI መሣሪያ ደህንነት ምልክት እንዳላሳየ እና process of እንክብካቤን እንዳሻሻለ አገኘ፣ ግን በሁለት ሳምንት ታካሚ-ውጤት ጥቅም አላረጋገጠም፤ የሚኖር ትንሽ ጥቅም ለማየት ብዙ ትልቅ ጥናት ያስፈልጋል።

ለምን ይጠቅማል?

የሕክምና AI ውይይት እንዲህ ዓይነት ማስረጃ በጣም ያነሰ አለው። ሞዴሎች examsን እንደሚያልፉ እና tidy ጉዳዮች ላይ ክሊኒሺያኖችን እንደሚወዳደሩ የሚያሳዩ ሺዎች ጥናቶች አሉ። እውነተኛ ታካሚዎች ይበልጥ ጥሩ ውጤት እንዳገኙ የሚለኩ pragmatic በዘፈቀደ የተመደበ ሙከራዎች ግን ጥቂት ናቸው። ይህ ከእነሱ አንዱ ነው፣ እና የማያምር እውነት ያሳያል፡- ሙከራን ማለፍ ታካሚን ከመርዳት ጋር አንድ አይደለም።

መሣሪያ ለክሊኒሺያኖች በእውነት ጠቃሚ ሊሆን ይችላል — ግልጽ notes፣ ዝቅተኛ drug bills፣ second ጥንድ of eyes — እና በተመሳሳይ ጊዜ በ14 ቀናት hard ታካሚ ውጤትን ላይንቀሳቀስ ይችላል። ሁለቱም እውነት ሊሆኑ

ይችላሉ።

ይህ የማረጋገጫ burdenንም በትክክል ያስቀምጣል። Company የክሊኒካል AI ሥርዓቱ እንክብካቤን ያሻሽላል ማለት ከፈለገ፣ ተገቢው ማስረጃ leaderboard አይደለም። ታካሚዎች ላይ የሚጠቅሙ ውጤቶችን የሚለካ እንደዚህ ዓይነት ሙከራ ነው — እና ምናልባት ትንሽ benefitsን ለመያዝ ይበልጥ ትልቅ ሙከራ።

ንጹህ ማጠቃለያ

በ16 Kenyan የመጀመሪያ ደረጃ ሕክምና የጤና ተቋማት ውስጥ pragmatic cluster-በዘፈቀደ የተመደበ ሙከራ በክሊኒካል officers የሚጠቀሙበት ኤሌክትሮኒክ መዝገብ ላይ “AI Consult” generative-AI decision-ድጋፍ መሣሪያውን ፈተነ። በ9,691 ታካሚዎች ውስጥ expert-adjudicated composite ሕክምና ውድቀት በ14 ቀናት 2.2% በመሣሪያ arm እና 2.0% በቁጥጥር ነበር (የተስተካከለ OR 0.77, 95% CI 0.55–1.08, P = 0.13) — ጉልህ ልዩነት የለም። መሣሪያው ደህንነት ምልክት አላሳየም፣ ክሊኒካል ሰነድ ማዘጋጀትን በሁሉም rated ዘርፎች አሻሽሏል፣ prescribing አልቀየረም፣ ታካሚ እርካታ ተመሳሳይ ነበር፣ እና antibiotic ወጪዎች ትንሽ ዝቅ ከማለት ጋር ተያይዟል። ደራሲዎቹ በእነዚህ ገደቦች ውስጥ ደህንነት ምልክት እንዳልታየ ግን ሕክምና ውድቀት እንዳልቀነሰ፣ ጥቅም ካለም ምናልባት ትንሽ እንደሆነ ይደርሳሉ። የታየውን ተፅዕኖ size ለመያዝ ከ100,000 ታካሚዎች የሚበልጥ ትልቅ ሙከራ ሊያስፈልግ ይችላል።

ያለ ማጋነን ማረጋገጫ

ጥናቱ የሚያሳየው: እውነተኛ-world በዘፈቀደ የተመደበ ሙከራ ውስጥ LLM-based decision-ድጋፍ መሣሪያውን ወደ የመጀመሪያ ደረጃ ሕክምና መዝገብ መጨመር ደህንነት ምልክት አላሳየም፣ ክሊኒካል ሰነድ ማዘጋጀት qualityን አሻሽሏል፣ prescribingን አልቀየረም፣ እና strict 14-day ታካሚ composite ሕክምና ውድቀትን ጉልህ አልቀነሰም።

ሊሆን የሚችል ግን ያልተረጋገጠ: ሙከራው ለማየት ትንሽ የሆነ እውነተኛ reduction in ሕክምና ውድቀት እንዳለ፣ full ወጪዎች ሲቆጠሩ መሣሪያው money እንደሚያድን፣ የተሻለ ሰነድ ማዘጋጀት በኋላ የተሻለ እንክብካቤ እንደሚያመጣ።

የማያሳየው: ክሊኒካል AI ታካሚ ውጤቶችን ያሻሽላል፣ አልፎ አልፎ የሚከሰት ክስተቶች ላይ ደህንነት ተረጋግጧል፣ ክሊኒኪያኖችን ይተካል ወይም ይበልጣል፣ ውጤቶች ወደ rural/high-income settings በቀጥታ ይተላለፋሉ፣ ወይም ወጪ saving ተረጋግጧል ብሎ አያሳይም።

ዋና ገደቦች: ከታየው ይበልጥ ትልቅ ተፅዕኖ ለመያዝ powered ነበር፣ አልፎ አልፎ የሚከሰት ውጤቶች ትልቅ sample ይፈልጋሉ። Configuration ስህተት የቁጥጥር ቡድንን በክፍል contaminated አደረገ፣ shared-network habitsም armsን ሊያቀራርቡ ይችላሉ፣ ሁለቱም bias toward no difference ያመጣሉ። አንድ private urban network፣ ቀድሞ ከፍተኛ standards፣ no prespecified noninferiority/ደህንነት ማዕቀፍ፣ እና 14-day horizon።

አጠቃላይ አንባቢ ምን ያህል ሊተማመን ይገባል? መሣሪያው በዚህ ሙከራ ደህንነት ምልክት እንዳላሳየ እና ሰነድ ማዘጋጀትን እንዳሻሻለ ከፍተኛ እምነት። እዚህ 14-day ታካሚ ውጤቶችን demonstrably እንዳላሻሻለ ከፍተኛ እምነት። “ይሠራል” ወይም “ወድቋል” በሚለው ታካሚ-ጥቅም verdict ላይ ዝቅተኛ እምነት — ጥያቄው ገና ክፍት ነው እና ብዙ ትልቅ ጥናት ይፈልጋል።

ምንጮች

Agweyu et al., Nature Medicine (2026)

<https://doi.org/10.1038/s41591-026-04503-6>

Pan-African Clinical Trials Registry, registration 202502499779176

<https://pactr.samrc.ac.za/>

EurekAlert / PATH press release on the trial (2026)

<https://www.eurekalert.org/news-releases/1133583>