



The CLEAN PAPER

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

Clinical AI paperwork উন্নত করছে এবং safe মনে হয়েছে — কিন্তু patient outcome উন্নত করেনি

Randomized primary-care trial-এ AI-supported workflow documentation ও কিছু care-process measure improve করছে এবং trial-এ clear safety signal পাওয়া যায়নি। কিন্তু prespecified primary patient outcome significantভাবে improve করেনি। তাই result হলো † “AI care improve করছে” নয়; বরং process-level benefit with no demonstrated patient-level benefit in this trial। Real-world medical AI evaluation-এ paperwork, safety ও health outcome আলাদা endpoint — এই study সটোই পরিস্কার করে।

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

নরিপদ ও কাজে লাগা মানইে কার্যকর নয়

চকিৎসাবিধিক AI নিয়ে বহু শরিে নাম ভুল ধরনরে প্রমাণ থেকে তরৈহি়। কনেো মডলে পরীক্ষার প্রশ্নে ভালোে স্কোর করে, অথবা সাজানোে কলনিকিয়াল সংক্শপিত কলনিকিয়াল বরণনায় চকিৎসকদরে ছাড়িয়ে যায়, আর সখোন থেকে খুব দ্রুত সিদ্ধান্ত আসে—যন্তর নাকি কলনিকিরে জন্য প্রস্তুত।

এই পরীক্ষা আরও কঠনি প্রশ্নন করছে। একটি জনোরটেভি AI সিদ্ধান্ত-সহায়ক সরঞ্জামকে বাস্তব প্রাথমিক চকিৎসাসবো কনেদ্ররে, বাস্তব চকিৎসিক ও বাস্তব রোগীর সঙ্গে ব্যবহার করে দেখো হয়েছে: **রোগীদের ফলাফল কিসত্যহি ভালোে হয়েছে?**

সংক্শপে উত্তর হলো—14 দিনরে মধ্যে তা পরিমাপযোগ্যভাবে ভালোে হয়নি। সরঞ্জামটির সঙ্গে কনেো নরিপত্তা-সংকতে ধরা পড়নে; চকিৎসা নথবিদ্ধকরণ উন্নত হয়েছে; অ্যান্টবিয়োটিকিরে কিছু খরচ কমতওে পারে। কনিতু প্রবনরিধারতি চকিৎসা-ব্যরথতা শেষেবনিতু উল্লেখযোগ্যভাবে কমনে, এবং লখেকরোও বলনে—উপকার থাকলওে সম্ভবত ছোট।

এটি গবেষণার ব্যরথতা নয়; বরং ভালোে পরীক্ষার কাজ। চকিৎসাবিধিক AI-কে মানদণ্ড নয়, রোগীর ফলাফল দিয়ে মাপলে প্রমাণ দেখতে এমনই হয়।

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

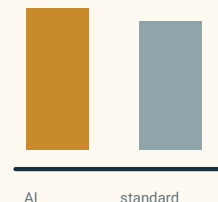
Process signal, not outcome proof.



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.

সরঞ্জামটির সঙ্গে কনেো নরিপত্তা-সংকতে ধরা পড়নে এবং নথবিদ্ধকরণ উন্নত হয়েছে, কনিতু আগ থেকে নরিধারতি 14-দিনে রোগীর ফলাফল উল্লেখযোগ্যভাবে ভালোে হয়নি। সবো প্রক্রিয়ায় সাহায্য আর রোগী উপকার প্রমাণ করা এক জনিসি নয়।

গবেষকরা কী করছেন

Kenya-র Nairobi ও Kiambu কাউন্টিতে Penda Health-এর বেসরকারি নিটেওয়ারকরে 16টি প্রাথমিক-চিকিৎসা কেন্দ্রে ব্যবহারিক গুচ্ছ-রয়ানডমাইজড পরীক্ষা চালানো হয়। এসব ক্লিনিকে চিকিৎসার বড় অংশ দেন **clinical officer**—তিনি বছরে ডিপ্লোমা প্রাপ্ত মধ্যস্তরের চিকিৎসাকর্মী—যাদের কাছে সব সময় জ্যেষ্ঠ চিকিৎসকের পরামর্শ সহজে পাওয়া যায় না।

Randomization রাগীকে নয়, চিকিৎসককে করা হয়। মোট 103 ক্লিনিকিয়াল officer: 52 জন হস্তক্ষেপে শাখায় এবং 51 জন ন্যূনতর শাখায়। উভয় দল একই মধে-ভিত্তিক ইলেকট্রনিক চিকিৎসাবিষয়ক নথি ব্যবহার করেছে। হস্তক্ষেপে শাখা অতিরিক্তভাবে **AI পরামর্শ নেওয়া 2.0** পয়েছেলি—**OpenAI GPT-4o**-ভিত্তিক জনোরেডি সদিধানত-সহায়ক সরঞ্জাম, যা নথি ভেতরেই চিকিৎসকের লেখা তথ্য পড়ে রাগনরিণ্য বা চিকিৎসা পরিকল্পনায় সম্ভাব্য সমস্যা চিহ্নিত করতে পারত। চূড়ান্ত সদিধানত চিকিৎসকেরই: ইঙগতি গ্রহণ, বদলানো বা উপেক্ষা—সবই তার হাতের।

কোন মডেলে ছিল, আর নরিদষ্টি বরিণর্গটকনে গুরুত্বপূর্ণ

AI পরামর্শব্যবস্থা 2.0 OpenAI-এর **GPT-4o**-র মে 2025 সংস্করণ ব্যবহার করেছে, এন্টারপ্রাইজ-স্তরে বাণিজ্যিক API-এর মাধ্যমে এবং কম এলোমেলোতা (temperature 0.1) রেখে। এটি EasyClinic-এর বিশেষভাবে তৈরি EMR-এর ভেতরে বসানো ছিল এবং Kenya-র জাতীয় চিকিৎসা নরিদশেকার সঙ্গে সামঞ্জস্য রাখতে system prompt দিয়ে নরিদশেতি করা হয়েছিল; লেখকরা পুরো নরিদশেনা-প্রম্পট প্রকাশ করছেন।

ফলটি তাই “চিকিৎসাবিজ্ঞানে LLM” সম্পর্কে চরিস্থায়ী কোনো রাগ নয়। এটি নরিদষ্টি একটি মডেলে-সংস্করণ, প্রম্পট, ইলেকট্রনিক নথি ও ক্লিনিকিয়াল পরিবিশেলে ফল। লেখকরা একে *সময়-নরিদষ্টি মানদণ্ড* বলছেন—নতুন মডেলে, ভিন্ন প্রম্পট বা কম ডিজিটাইজড ক্লিনিকে ফল বদলাতে পারে। OpenAI পরে ক্লাউড-কমপটিবিলিটি ও API-সংক্রান্ত প্রযুক্তিগত সহায়তা দিয়েছিল, তবে লেখকদের মতে মডেলেটি সেই সহায়তার প্রস্তাবের আগেই বছে নেওয়া হয়েছিল এবং OpenAI পরীক্ষা-নকশা, তথ্য সংগ্রহ, বিশ্লেষণ বা প্রকাশনার সদিধানতে ভূমিকা রাখেনি।

22 এপ্রিল থেকে 16 জুলাই 2025 পর্যন্ত 9,691 রাগী অন্তর্ভুক্ত হয়। প্রধান ফলাফল ইচ্ছাকৃতভাবে রাগী-centered ও কঠোর ছিল: সাক্ষাৎ 14 দিনের মধ্যে বিশেষজ্ঞ-adjudicated **সমন্বিত চিকিৎসা ব্যর্থতা**। গবেষণা শাখা না-জনে চিকিৎসক প্যানেলে বচার করেছে অসুস্থতা অমীমাংসতি বা খারাপ হওয়ার মতো খারাপ ফলাফল হয়েছে কিনা। পরীক্ষা আগে থেকেই Pan-African ক্লিনিকিয়াল পরীক্ষাগুলো Registry-তে 202502499779176 নম্বরে registered ছিল।

এটা ইনকশার মূল শক্তি। AI চিকিৎসক কী লিখেছে তা বদলেছে কিনা দেখানো সহজ; রাগীর কী হয়েছে তা বদলেছে কিনা দেখানো অনেক কঠিন এবং বেশি অর্থবহ।

তারা কী পয়েছেন

প্রধান ফলাফল উন্নত হয়নি। AI শাখায় 4,693 জনের মধ্যে 102 জনের (2.2%) চিকিৎসা ব্যর্থতা হয়; ন্যূনতর শাখায় 4,654 জনের মধ্যে 94 জনের (2.0%)। কাঁচা শতাংশ AI শাখায় সামান্য বেশি হলেও গুচ্ছ পার্থক্য বদলানো করার পর বিন্দু অনুমান উপকারের দিকে যায়: **সমন্বিত odds অনুপাত 0.77 (95% CI 0.55–1.08, P = 0.13)**। এটি পরিসংখ্যানগতভাবে তাৎপর্যপূর্ণ নয়। কাঁচা ও সমন্বিত দিকের পার্থক্য arithmetic ত্রুটি নয়; adjustment চিকিৎসক-গুচ্ছের পার্থক্য হিসাব করে। আস্থার ব্যবধান “না প্রভাব” অন্তর্ভুক্ত করে, তাই উপকার দাবি করা যায় না; পরম পার্থক্য-ও খুব ছোট।

Odds অনুপাত, আস্থার ব্যবধান ও P মান একসঙ্গে পড়ার সাধারণ ব্যাখ্যার জন্য ক্লিনিকিয়াল ফলাফল নরিদশেকি দেখা যতে পারে।

নরিপত্তা-সংকতে পাওয়া যায়নি—কিন্তু সীমার মধ্যে। সরঞ্জামের সঙ্গে সম্পর্কিত বলে বচার করা কোনো গুরুতর বরিণ ঘটনা ছিল না, এবং

স্বাধীন পর্যালোচনা-ও নরিপত্তা-সংকতে পায়নি। তবে বরিল গুরুতর কষতধিরার মতটে পরিসংখ্যানগত কষমতা পরীক্ষায় ছিল না; পূর্বনরিধারতি noninferiority বা আনুষ্টানকি নরিপত্তা কাঠামে টে-ও ছিল না। তাই uncommon কষতির কষতেরে সরঞ্জাম “নরিপদ প্রমাণতি”—এ কথা বলা যায় না।

নথবিদ্ধকরণ ভালটে হচ্ছে। Blinded বিশেষজ্ঞে দ্বারা পর্যালোচনা করা 2,000 সাক্ষাৎ AI পরামর্শ নেওয়া ব্যবহারকারী চকিৎসকদের নথবিদ্ধকরণ সব rated কষতেরেই উন্নত ছিল—রেকর্ড করা রেগনরিণয়, চকিৎসা পরকিল্পনা এবং সামগ্রিক পূর্ণতা।

ওষুধ দেওয়া প্রায় বদলায়নি। সংশোধন করা অ্যান্টবিয়াটেটিকি ব্যবহার-সহ ওষুধ দেওয়ায় তাৎপর্যপূর্ণ পার্থক্য ছিল না (সমন্বয়কৃত odds অনুপাত 0.86, 95% CI 0.48–1.55)। অ্যান্টবিয়াটেটিকি ওষুধ দেওয়া হার সরঞ্জাম বদলায়নি।

রেগী সন্তুষ্টবিদলায়নি। সমীক্ষা শেষে করা 826 রেগীর সন্তুষ্ট দুই শাখায় প্রায় অভিনি; consultation সময়ও কাছাকাছি ছিল।

খরচ সামান্য নচিরে দকিই ইঙগতি করছে। সমন্বয়কৃত বিশ্লেষণে AI শাখার অ্যান্টবিয়াটেটিকি-সম্প্রকতি খরচ কম ছিল—সম্ভবত প্রসেক্রপিশন কমানোর বদলে কম দামের বকিল্প বছে নেওয়ার কারণে। প্রতরিগী অ্যান্টবিয়াটেটিকি শারয় সরঞ্জাম চালানোর প্রতরিগী খরচের চয়েও বেশি মনে হচ্ছে। কনিতু লখেকরো এটকিই ইঙগতিবহ বলছেন; সম্প্রণ মেটি-খরচ-of-ownership অর্থনৈতিক মূল্যায়ন পরীক্ষার পরধিতে ছিল না।

সব মলিয়ি লখেকদের সংকষপিত ব্যাখ্যাই পরধিকার: এই সীমার মধ্যে LLM assistance-এর নরিপত্তা-সংকতে ছিল না, কনিতু 14 দিনের চকিৎসা ব্যর্থতা কমনে; উপকার থাকলে সম্ভবত মাঝারি।

“তাৎপর্যপূর্ণ পার্থক্য নহে” মানে “কাজ করে না” নয়

শূন্য প্রধান ফলাফল দুই দকিই অতিরিক্ত পড়া যায়। দুটকারণ সহজ রায় আটকায।

প্রথমত, **পরীক্ষা য়ে প্রভাব শনাক্ত করার জন্য নকশা হয়েছিল, প্রযবকেষতি প্রভাব তার চয়ে ছে টি।** প্রাথমিক চকিৎসাসবো-তে গুরুতর খারাপ ফলাফল বরিল—এখানে প্রায় 2%। শব্দদূষণ থেকে ছে টি বাস্তব উপকার আলাদা করতে খুব বড় নমুনা লাগে। লখেকদের পরবর্তী বিশ্লেষণ কষমতা হিসাব অনুযায়ী প্রযবকেষতি মাত্রার প্রভাব ধরতে 100,000-এরও বেশি রেগীর নরিদশে দেওয়ার পরীক্ষা দরকার হতে পারে। এই নমুনায় অ-তাৎপর্যপূর্ণ ফল ছে টি উপকার অসম্ভব প্রমাণ করে না; শুধু বলে এই গবষণে সটেসিমাধান করতে পারেনি।

দ্বিতীয়ত, **দুই শাখার পার্থক্য কঙ্কিটা ঝাপসা হয়ে গিয়েছিল।** বনিয়াসগত ত্রুটির কারণে কঙ্কি সময় নয়নতরণ-শাখার চকিৎসকরোও AI পরামর্শে প্রবশোধকির পয়েছিলেন। একই নটেওয়ারকরে চকিৎসকরো আবার পরস্পরের সঙগে কথা বলনে এবং কাজরে অভ্যাস ভাগ করে নতিে পারনে। দুটটে কারণই intervention ও control গে ষ্ঠী বাস্তবে আরও কাছাকাছি হয়ে যতে পারে; সত্যকিররে কনে টি পার্থক্য থাকলে এতে মাপা প্রভাব শূন্যেরে দকি টেনে নামবে। উপরন্তু, Penda নটেওয়ারকে আগই তুলনামূলক উচ্চ মানেরে সবো ছিল, তাই উন্নতির সুযোগও সীমতি ছিল।

এগুলো টি “বড় অগ্রগতি” শরিে নাম উদ্ধার করে না। সঠিক পাঠ আরও ক্যালকিটেডে: সবচয়ে কঠনি রেগী শেষবিন্দুতে সরঞ্জাম দুই সপ্তাহে দেখোনে টি সম্ভব উপকার দেখোনি; কনিতু নথি-keeping উন্নত করছে এবং নরিপত্তা-সংকতে দয়েনি।

এই গবষণে কী প্রমাণ করে না

- AI রেগীর ফলাফল উন্নত করছে—এটি দেখোয় না। 14-দিন প্রধান প্রবনরিধারতি ফলসূচকে তাৎপর্যপূর্ণ উপকার ছিল না।
- AI অকার্যকর—এটিও দেখোয় না। সমন্বয়কৃত বিন্দু অনুমান উপকারেরে দকি, নথবিদ্ধকরণ উন্নত, ওষুধ খরচ কমার সংকতে আছে; ছে টি উপকার নশ্চিতি করার জন্য নমুনা অপর্যাপ্ত হতে পারে।
- বরিল কষতির কষতেরে সরঞ্জাম নরিপদ প্রমাণতি—এ কথা বলা যায় না। বরিল গুরুতর ঘটনা ধরার মতটে নকশা/কষমতা ছিল না।
- “AI চকিৎসককে হারায়” বা চকিৎসক প্রতসিধাপন করে—এটি নিয়। এটি সিদ্ধান্ত-সহায়তা; চূড়ান্ত কর্তৃত্ব চকিৎসকরে।

- Kenya-র একটি শহুরে বেসরকারি স্বাস্থ্যসেবা নেটেওয়ারকে পাওয়া ফল গ্রামীণ, peri-urban বা উচ্চ-আয়রে দেশেরে পরিশেষে স্বয়ংক্রিয়ভাবে প্রযোজ্য ধরে নেওয়া যায় না।
- খরচ সাশ্রয় প্রত্যাশিত হয়নি; সংকটে আছে, সম্পূর্ণ অর্থনৈতিক মূল্যায়ন নাই।

প্রমাণ কতটা শক্ত?

কেন্দ্রীয় দাবি—সংক্ষিপ্ত-পর্যায় রোগী কৃতি কমছে বলে প্রমাণ নাই—এর জন্য নকশা শক্তিশালী, উপসংহার যথাযথভাবে সংযত। ভবিষ্যতমুখী, আগে থেকে নবিন্ধতি, গুচ্ছর্যান্ডমাইজড পরীক্ষা এবং blinded রোগী-স্তর সমন্বিত ফলাফল বাস্তব ক্লিনিক্যাল AI-এর জন্য মানদণ্ড/সংক্ষিপ্ত ক্লিনিক্যাল বর্ণনা গবেষণার তুলনায় অনেকে বেশি informative।

গঠন সিদ্ধান্ত—নথিবিধকরণ উন্নতি, ওষুধ দেওয়া অপরিবর্তিত, সন্তুষ্ট একই, অ্যান্টিবায়োটিক খরচ কিছুটা কম—ভালো প্রমাণ, কিন্তু গঠন হিসেবেই পড়তে হবে। একই দৃশ্য ও একক-নেটেওয়ারক সীমাবদ্ধতা এগুলো তেও আছে।

নিরাপত্তা প্রমাণ আশ্বস্তকর কিন্তু সীমাবদ্ধ: কোনো সংকটে পাওয়া যায়নি, তবে বরিল কৃতিধরার গবেষণা নয়।

সবচেয়ে সঠিক অবস্থান “কাজ করো” বা “ব্যর্থ”—দুটোর কোনোটিই নয়। বাস্তব জগতেরে র্যান্ডমাইজড পরীক্ষায় সরঞ্জাম নিরাপত্তা-সংকটে দিয়েনি এবং সেবা প্রক্রিয়া উন্নত করেছে, কিন্তু দুই সপ্তাহে রোগী উপকার দেখানো করেনি; এমন ছোট উপকার ধরতে অনেকে বড় পরীক্ষা লাগতে পারে।

কেন এটি গুরুত্বপূর্ণ

চিকিৎসাবিষয়ক AI নিয়ে বতিরকে এই ধরনের প্রমাণ খুব কম। পরীক্ষা স্কোর, মানদণ্ড এবং tidy সংক্ষিপ্ত ক্লিনিক্যাল বর্ণনায় চিকিৎসক-মলিযুক্ত গবেষণাপত্র হাজার হাজার আছে; বাস্তব রোগী ভালো হলো কনি মাপা বড় ব্যবহারিক র্যান্ডমাইজড পরীক্ষা খুব অল্প। এই গবেষণা তাদরে একটি। তার অনাড়ম্বর শিক্ষা: **পরীক্ষা পাস করা আর রোগীকে সাহায্য করা এক জনিসি নয়।**

একটি সরঞ্জাম চিকিৎসকের কাজে বাস্তবভাবে উপযোগী হতে পারে—ভালো note, দ্বিতীয় জোড়া of চোখ, কম ওষুধ খরচ—তবু 14 দিনে কঠিন রোগীর ফলাফল নড়াতে না-ও পারে। দুই সত্য একসঙ্গে থাকতে পারে। পরিণিত স্বাস্থ্য ব্যবস্থাকে সুবিশাজনক একটি বিয়ান বছে নেওয়ার বদলে দুটোই ধরে রাখতে হবে।

এটি বোঝা of প্রমাণ-ও সঠিক জায়গায় ফরোয়। ক্লিনিক্যাল AI সেবা উন্নত করে—এ দাবি করলে প্রাসঙ্গিক প্রমাণ লভিরবোড় নয়; রোগী-প্রাসঙ্গিক ফলাফল সহ পরীক্ষা। আর মাঝারি উপকার সত্যি থাকলে তা দেখতে সম্ভবত আরও বড় পরীক্ষা দরকার।

সংক্ষিপ্ত সারাংশ

Kenya-র 16 প্রধান-সেবা কেন্দ্রে গুচ্ছর্যান্ডমাইজড ব্যবহারিক পরীক্ষায় ক্লিনিক্যাল officer-দের ইলেকট্রনিক নথি সঙ্গে জেনারেটিভ AI সিদ্ধান্ত-সহায়ক **AI পরামর্শ দেওয়া** যোগ করা হয়। মোট 9,691 রোগীর মধ্যে বিশেষজ্ঞ-adjudicated 14-দিন চিকিৎসা-ব্যর্থতা সমন্বিত ছিল AI শাখায় 2.2% এবং ন্যূনতর 2.0%; **সমন্বিত OR 0.77, 95% CI 0.55–1.08, P=0.13**—তৎপর্যপূর্ণ পার্থক্য নয়। সরঞ্জামের সঙ্গে নিরাপত্তা-সংকটে পাওয়া যায়নি; নথিবিধকরণ সব rated কৃতিত্বেরে উন্নত হয়েছে; ওষুধ দেওয়া বা সন্তুষ্ট বদলায়নি; অ্যান্টিবায়োটিক খরচ কিছুটা কমার সংকটে ছিল। লেখকদের উপসংহার: এই সীমার মধ্যে সরঞ্জাম আশ্বস্তকর, কিন্তু চিকিৎসা ব্যর্থতা কমানোর প্রমাণ নাই; পর্যবেক্ষিত প্রভাবের মতো উপকার ধরতে 100,000-এর নরিদশে দেওয়ার রোগী লাগতে পারে। এটি ক্লিনিক্যাল AI

রোগীর ফলাফল উন্নত করে প্রমাণ করে না, আবার অকার্যকর-ও প্রমাণ করে না।

সরাসরি যাচাই

পপোর যা দেখায়: বাস্তব র‍্যান্ডমাইজড প্রধান-সবো পরীক্ষায় LLM-ভিত্তিক সন্ধান-সহায়তার নরিপত্তা-সংকতে ধরা পড়েনি, চিকিৎসা নথিবিদ্যকরণ উন্নত হয়েছে, ওষুধ দোষে উল্লেখযোগ্যভাবে বদলায়নি এবং কঠোর 14-দিন চিকিৎসা-ব্যবস্থা শেষবিন্দু উল্লেখযোগ্যভাবে কমেনি।

যা সম্ভব, কিন্তু প্রমাণিত নয়: প্রযবেক্ষেতি ছোট প্রভাব সত্যকারের মাঝারি উপকার হতে পারে; সম্পূর্ণ খরচ গণনা টাকা বাঁচতে পারে; ভালো নথিবিদ্যকরণ ভবিষ্যতে ভালো ব্যবস্থা রূপ নতি পারে।

যা দেখায় না: কলনিকিয়াল AI রোগীর ফলাফল উন্নত করে; বরিল ঘটনায় নরিপত্তা প্রতিষ্ঠিত; AI চিকিৎসক প্রতিস্থাপন করা/outperform করে; ফল সব বিনিয়োগে সাধারণীকরণ করে; খরচ সশ্রয় নিশ্চিত।

প্রধান সীমাবদ্ধতা: প্রযবেক্ষেতি প্রভাবে তুলনায় পরীক্ষা ছোট; ন্যূনতর-শাখা দৃষণ ও যথেষ্ট-নটেওয়ারক আচরণ পার্থক্যকে শূন্যর দিকে পক্ষপাত করতে পারে; এক ব্যক্তিগত urban নটেওয়ারক; পূর্বনির্ধারিত noninferiority/নরিপত্তা কাঠামো নেই; অনুসরণকালীন 14 দিন।

সাধারণ পাঠকের আস্থা কতটা হওয়া উচিত? নথিবিদ্যকরণ improvement এবং এই পরীক্ষায় নরিপত্তা-সংকতে না-পাওয়া—উচ্চ। 14-দিন রোগীর ফলাফল demonstrably improve না-হওয়া—উচ্চ। সরঞ্জাম রোগী-উপকার হস্তক্ষেপে হিসেবে “কাজ করে” বা “ব্যর্থ”—এই দুই রায়ে জন্ম কম; প্রশ্নটি এখনও খোলা।

উৎস

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>