



The CLEAN PAPER

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

एक clinical AI सुरक्षति दखीऔर paperwork बेहतर हुआ — लेकिन मरीजोंके outcomes बेहतर नहींहुए

Kenya की 16 primary-care facilities में pragmatic cluster-randomized trial ने clinical officers के record system में जोड़े गए generative-AI decision-support tool को परखा 9,691 मरीजोंमें 14 दिन का सख्त expert-adjudicated composite treatment-failure outcome AI tool के साथ 2.2% और बना tool 2.0% था (adjusted odds ratio 0.77, 95% CI 0.55–1.08, $P = 0.13$) — महत्वपूर्ण अंतर नहीं। Safety signal नहीं मिला; documentation सभी rated domains में बेहतर हुई; prescribing और satisfaction नहीं बदले; antibiotic costs घटने की दिशा में trend था यह “failed study” नहीं बल्कि दुर्लभ ईममदार अध्ययन है — benchmarks पर नहीं मरीजों पर ममागया — और इसका समझण नषिर्ष यही है कि tool सुरक्षति दखीऔर प्रक्रियामें मदद की पर दो हफ्तोंमें मरीजोंको demonstrably बेहतर नहीं किया। इतना छोटा clinical benefit पकड़ने के लिए बहुत बड़ा trial चाहिए होगा

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

सुरक्षा और सहायक देखना प्रभावी होने के बरबर नहीं है

चिकित्सकीय AI पर अधिकांश सुरखियाँ गलत तरह के अध्ययनों से बनती हैं। कई मॉडल परीक्षा के प्रश्नों में अच्छा स्कोर करता है, या चुने हुए नैदानिक उद्देश्यों पर डॉक्टरों से बेहतर देखता है, और कहती खुद लखी हुई लगती है: मशीन क्लिनिक के लिए तैयार है।

इस परीक्षण ने कुछ अधिक कठिन और कहीं अधिक दुर्लभ किया। इसने जनरेटिव-AI आधारित निर्णय-सहायता उपकरण को वास्तविक प्रणाली देखभाल केंद्रों में, वास्तविक मरीजों के साथ इस्तेमाल किया और वही सबल पूछा जो वास्तव में मारने रखता है: क्या मरीजों के नतीजे बेहतर हुए?

ईमानदार उत्तर है—नहीं कम से कम 14 दिनों में मारने योग्य रूप से नहीं। उपकरण से कोई सुरक्षा संकेत नहीं मिला। उसने चिकित्सकीय दस्तवेजकरण की गुणवत्ता बेहतर की। संभव है कि उसने कुछ दवा मिलात भी घटई हो। लेकिन उपचार-वफाई में सांख्यिकीय रूप से महत्वपूर्ण कमी नहीं आई, और लेखक सभ्यता से कहते हैं कि यदि कोई लाभ है भी तो शायद वह छोटा है।

यह अध्ययन की वफाई तान नहीं है। यही अध्ययन का काम है। चिकित्सा में AI पर ज़िम्मेदार सक्षम ऐसे देखते हैं जब उन्हें बेचमार्क की बजाय मरीजों पर मारा जाता है।

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

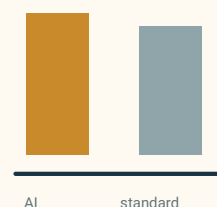
Process signal, not outcome proof.



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.

उपकरण से कोई सुरक्षा संकेत नहीं मिला और चिकित्सकीय दस्तवेजकरण बेहतर हुआ, लेकिन पहले से तय 14-दिन का मरीज-परणाम सांख्यिकीय रूप से बेहतर नहीं हुआ। प्रक्रियामें मदद और सदिध मरीज-लभ एक ही बात नहीं है।

लेखकोने क्याकिया

टीम ने केन्याके नैरेबीऔर कआम्बू कण्ट्रियोमें नजीस्वस्थ नेटवर्क Penda स्वस्थ्य द्वाराचलाए जमे वाले 16 प्रथमकि देखभल केंद्रोंमें एक व्यवहारकि, क्लस्टर-रैंडमइज्ड परीक्षण किया। इन केंद्रोंमें देखभल मुख्यतः क्लिनिकल officers देते हैं—तीन वर्षीय डप्लेमावाले मध्यम-स्तरिय चिकित्सकीय पेशेवर—जन्हें अक्सर वरिष्ठ वशिषज्ज से तुरंत सलह उपलब्ध नहीहोती।

रैंडमइजेशन कीइकई मरीज नही चिकित्सक था। 103 क्लिनिकल officers कोरैंडमइज कियागया 52 हस्तक्षेप समूह में और 51 नियंत्रण समूह में। दोनोंसमूहोंने एक हीक्लइड-आधारित इलेक्ट्रॉनिक मेडिकल रिकॉर्ड इस्तेमाल किया। हस्तक्षेप समूह कोइसके अतिरिक्त “AI Consult” (संस्करण 2.0) मिला जोOpenAI के GPT-4o बड़े भाषामॉडल पर बनानरिणय-सहयताउपकरण थाऔर उसीरिकॉर्ड में जुड़ाहुआ था। यह चिकित्सक द्वारादर्ज जप्तकरीपढ़कर नदिता या उपचार योजनामें संभवति समस्याओं कीओर ध्यान दलासकताथा। अंतमि नरिणय पूरीतरह चिकित्सक के पास रहा वे सुझाव स्वीकार, बदल याअनदेखाकर सकते थे।

कौन-सामॉडल था और यह वविरण क्योमयने रखताहै

उपकरण AI Consult 2.0 था जोOpenAI के GPT-4o (मई 2025 रलीज) पर चलताथा। इसे OpenAI के वणजियकि API से एंटरप्राइज लइसेंस के तहत चलायागयाऔर randomness कम रखीगई (तममस 0.1)। यह EasyClinic के खस इलेक्ट्रॉनिक मेडिकल रिकॉर्ड (EMR) के भीतर थाऔर इसके प्रणालीप्रमैप्ट केन्याकीराष्ट्रीय उपचार-नरिदेशकिओं के अनुरूप लखे गए थे; लेखकोने पूराinstruction प्रमैप्ट प्रकशति किया।

यह सब स्पष्ट करनाक्योजरूरीहै? क्योकि परिणम एक खस प्रणालीके बारे में है—एक मॉडल संस्करण, एक प्रमैप्ट, एक रिकॉर्ड और एक सेटगि—न कि समस्त रूप से “चिकित्सामें LLMs” के बारे में। लेखक भीयहीबत कहते हैं: वे अपने नषिकर्ष कोक्षमताकास्थरि अनुमम नही बल्कि समय-वशिषिट बेंचमार्क कहते हैं। नया मॉडल, अलग प्रमैप्ट याकम डजिटइज्ड क्लिनिकि परिणम बदल सकते हैं।

स्वतंत्रताके बारे में: बढ में OpenAI ने वस्तुरूप सहयतादी—cloud-compute credits और API उपयोग पर तकनीकी मार्गदर्शन—लेकिन लेखकोके अनुसार OpenAI कोचुनने कानरिणय उस प्रस्तम से पहले लयाजाचुकाथा और OpenAI कीपरीक्षण-डजिइन, डेटासंग्रह, वशि्लेषण याप्रकशन के नरिणय में कोई भूमकानहीथी।

22 अप्रैल से 16 जुलाई 2025 के बीच 9,691 मरीज शामिल किए गए। प्रथमकि परिणम जप्तबूझकर मरीज-केंद्रति और कठोर रखागयाथा वजिटि के 14 दनिके भीतर उपचार-वफिलताकासंयुक्त परिणम, जिसे वशिषज्जोंने तय किया। चिकित्सकोके एक पैनल ने अध्ययन-समूह जमे बनायह आंकाकि मरीज कोबीमरीन सुधरने यावगिडने जैसे प्रतकिल परिणम हुए यानही। परीक्षण पहले से पंजीकृत था(Pan-African क्लिनिकल परीक्षण Registry 202502499779176)।

यहीडजिइन-चयन अहम है। यह दखिमाआसत है कि AI उपकरण चिकित्सक के लखिने के तरीके कोबदलताहै। यह दखिमाकहीकठनि—और कहीअधिक सार्थक—है कि वह मरीज के सग हमें वालीचीज कोबदलताहै।

उन्हें क्यामला

प्रथमकि परिणम बेहतर नहीहुआ। AI समूह के 4,693 मरीजोंमें 102 (2.2%) और नियंत्रण समूह के 4,654 मरीजोंमें 94 (2.0%) में उपचार-वफिलताहुई। कच्चे प्रतशित AI के सग थोड़ाअधिक थे, लेकिन समयोजन के बढ बिदु-अनुमम लम कीदशामें गया समयोजति odds अनुपत 0.77 था(95% वशि्वस-अंतरल 0.55 से 1.08, $P = 0.13$)—जोसंख्यकीय रूप से महत्वपूर्ण नहीथा। कच्चे और समयोजति परिणमोकादशिबदलनागणनाकीगलतीनहीहै; समयोजन चिकित्सक-क्लस्टरों

के बीच अंतर को ध्यान में रखता है। किसी भी तरह विश्वास-अंतराल आरम्भ से “कई प्रभाव नहीं” को शामिल करता है, इसलिए लक्ष्य का दखानही किया जा सकता और पूर्ण अंतर बहुत छोटा था।

Odds अनुपात, विश्वास-अंतराल और P मान को सप्र पढ़ने का आसन्न तरीका क्लिनिकल परीक्षण पढ़ने की माहगदर्शकियों में है।

सीमाओं के भीतर कोई सुरक्षा संकेत नहीं मिला। किसी गंभीर प्रतिकूल घटना को उपकरण से संबंधित नहीं माना गया और स्वतंत्र समीक्षा में भी सुरक्षा संकेत नहीं मिला। लेखक इस बारे में की सीमा साफ़ बताते हैं: अध्ययन दुर्लभ गंभीर हानियों को पकड़ने के लिए पर्याप्त शक्ति वाला नहीं था और इसमें पहले से तय noninferiority या औपचारिक सुरक्षा दखानही था। इसलिए यह असमर्थ घटनाओं के लिए सुरक्षा को सदिध नहीं कर सकता।

दस्तवेजीकरण बेहतर हुआ। 2,000 मुलक्तों की blinded विशेषज्ञ समीक्षा में AI Consult इस्तेमाल करने वाले चिकित्सकों के रिकॉर्ड हर आँके गए क्षेत्र में बेहतर थे—दर्ज नदिम, उपचार योजना और कुल पूर्णता।

दवा लिखने के तरीके में बहुत कम बदलाव आया। Prescribing में कोई महत्वपूर्ण अंतर नहीं था जिसमें सही antibiotic उपयोग भी शामिल था (समयोक्ति odds अनुपात 0.86, 95% CI 0.48 से 1.55)। उपकरण ने antibiotic prescribing दर नहीं बदली।

मरीजों ने कोई अंतर महसूस नहीं किया। संतुष्टि सर्वे पूरा करने वाले 826 मरीजों में दोनों समूहों की संतुष्टि लगभग समान थी और consultation समय भी लगभग समान रहा।

लगात में हल्की गिरावट का संकेत मिला। समयोक्ति विश्लेषण में AI समूह में antibiotics से जुड़ी लगातार कम थी—संभवतः कम prescriptions की वजह से सस्ते विकल्प चुने जाने से—और प्रति मरीज antibiotic बचत उपकरण चलने की प्रति मरीज लगातार से अधिक दिखी। लेखक इसे संकेत मानते हैं, स्थिति नष्ट नहीं कुल ownership लगातार का पूरा आर्थिक लेखा परीक्षण के दायरे में नहीं था।

लेखकों का अपना एक-पंक्ति सर सबसे सफ़ है: इन सीमाओं के भीतर LLM सहायता सुरक्षित दिखी लेकिन उसने 14 दिनों में उपचार-वफाता कम नहीं की यदिलम है, तो शायद वह छोटा है।

“कई महत्वपूर्ण अंतर नहीं” का अर्थ “यह कम नहीं करता” नहीं है

नकारात्मक प्रथमिक परीक्षण को किसी भी दिशा में ज़्यादा पढ़ना आसन्न है। दोनों सरल कहानी को रोकती हैं।

पहली परीक्षण उस प्रभाव से बड़ा प्रभाव पकड़ने के लिए बनाया गया था जो वास्तव में दिखा। प्रथमिक देखभाल में गंभीर खराब परीक्षण दुर्लभ है—यहाँ लगभग 2%—इसलिए छोटे वास्तविक लक्ष्य को शोर से अलग करने के लिए बहुत बड़े नमूने चाहिए। लेखकों की वक्तव्य की शक्ति गणना बताती है कि उनके देखे गए आकार का प्रभाव पकड़ने के लिए कहीं बड़ा परीक्षण, लगभग 100,000 से अधिक मरीजों के स्तर का चाहिए होगा। इस आकार के अध्ययन में non-significant परीक्षण छोटे, वास्तविक लक्ष्य को खोज नहीं करता इसका अर्थ है कि यह अध्ययन उसे स्पष्ट नहीं कर सका।

दूसरी तुलना कुछ हद तक धुंधली होगई थी। configuration त्रुटि के कारण कुछ समय के लिए नियंत्रण समूह के कुछ चिकित्सकों को AI Consult मिला गया और एक ही नेटवर्क के चिकित्सक आपस में बात करते हैं तथा आदतें समूह-सीमा पर ले जा सकते हैं। दोनों समूहों को अधिक समान दिखती हैं और वास्तविक अंतर को शून्य की ओर धकेल सकती हैं। इसके अलावा मेजबान नेटवर्क पहले से अपेक्षित उच्च मान को पर चलता था इसलिए सुधार के लिए जगह कम थी।

इनमें से कोई बात “breakthrough” वाली सुरक्षा को बचती नहीं है। लेकिन सही पढ़ाई नरिषाद्वी भी नहीं है: सबसे कठिन और सबसे ईमानदार एंडपॉइंट पर यह उपकरण दोहफ्तों में मरीजों को मानने योग्य लक्ष्य नहीं दे सका—जबकि उसने रिकॉर्ड-खराब को

स्पष्ट रूप से बेहतर किया और सुरक्षा दिखा।

यह क्या सिद्ध नहीं करता

- यह नहीं दिखाता कि AI ने मरीजों के परिणाम बेहतर किए। प्रथम 14-दिन एंडपॉइंट पर महत्वपूर्ण लाभ नहीं था।
- यह नहीं दिखाता कि AI बेकार है। बटु-अनुमत उसके पक्ष में था दस्तवेजीकरण हर क्षेत्र में सुधरा और दवा लागत नीचे जमे का संकेत मिला null परिणाम छोटे वास्तविक लाभ के साथ संगत है जिसे यह परीक्षण पुष्टि करने के लिए बहुत छोटा था।
- यह दुर्लभ हमियों के लिए उपकरण की सुरक्षा सिद्ध नहीं करता। कई सुरक्षा संकेत नहीं मिला, लेकिन परीक्षण दुर्लभ गंभीर घटनाओं की सुरक्षा प्रमाणित करने के लिए डिज़ाइन या powered नहीं था।
- यह नहीं दिखाता कि “AI डॉक्टरों को हरा देता है” या चिकित्सक को बदल देता है। यह decision समर्थन था सुझाव स्वीकार या अस्वीकार करने का पूरा अधिकार चिकित्सक के पास था।
- यह अपने-आप समर्थित नहीं होता। परीक्षण के न्याय के एक नजिहाहरी नेटवर्क में हुआ; ग्रामीण, periurban या अधिक आय वाले परिवार अलग हो सकते हैं।
- यह लागत-बचत स्थापित नहीं करता। लागत का संकेत आशाजनक है, पूर्ण आर्थिक मूल्यंकन नहीं।

संक्षेप कतिनामजबूत है?

केंद्रीय दवा—कम समय में मरीज-हम घटने का सिद्ध प्रमाण नहीं मिला—के लिए अध्ययन-डिज़ाइन मजबूत है और नष्ट उचित रूप से संयमति है। भर्ती pre-registered, समूह हिमाम्यहृच्छिकीकृत परीक्षण, जिसमें blinded और मरीज-स्तरिय संयुक्त परिणाम है, ऐसे उपकरण के लिए वास्तविक दुनिया के सर्वोत्तम प्रकार के संक्षेपों में आता है। यह बेंचमार्क स्कोर या केस-विवरण अध्ययन से कहीं अधिक जमा कर देता है।

द्वितीयक नष्ट—बेहतर दस्तवेजीकरण, अपरिवर्तित prescribing, समान संतुष्टि, कम antibiotic लागत—अच्छे संक्षेप हैं, लेकिन उन्हें द्वितीयक ही पढ़ना चाहिए: सहायक संकेत, मुख्य सुखी नहीं और वही दूषण तथा single-network सिमाँ इन पर भी लागू होती हैं।

सुरक्षा के लिए संक्षेप आश्वस्त करने वाला लेकिन सीमा है: कई संकेत नहीं मिला पर अध्ययन दुर्लभ हमियाँ खोजने के लिए बना नहीं था।

सबसे उपयोगी रुख न “यह कम करता है” है, न “यह विफल हुआ”। सीख है: समझी पूर्वक वास्तविक दुनिया के परीक्षण में इस AI उपकरण से कई सुरक्षा संकेत नहीं मिला और देखभाल प्रक्रिया बेहतर हुआ, लेकिन दोहफ्तो में मरीज-परिणाम कालम सिद्ध नहीं हुआ—और ऐसा लाभ पकड़ने के लिए कहीं बड़ा अध्ययन चाहिए।

यह क्यों मानने रखता है

चिकित्सकीय AI की बहस को ठीक इसी तरह के संक्षेप की कमी है। हजरो शोध पत्र हैं जिनमें मॉडल परीक्षा में अव्वल आते हैं या सुथरे क्लिनिकल मामले पर चिकित्सक से मेल खाते हैं। वास्तविक मरीजों को फायदा हुआ या नहीं यह मानने वाले बड़े pragmatic रैंडमइज्ड परीक्षण बहुत कम हैं। यह उनमें से एक है, और इसका संचरण-सासत्य महत्वपूर्ण है: परीक्षा पास

करना मरीज की मदद करने के बराबर नहीं है।

यही अंतर पूरी कहानी है। कई उपकरण चिकित्सक के लिए सच में उपयोगी हो सकता है—सफ़ा notes, कम दवा खर्च, एक अतिरिक्त जाँच—और फिर भी पंद्रह दिनों से भी कम समय में किसी कठोर मरीज परणाम कोन बदल पाए। दोहोबतें एक सच हो सकती हैं, और परपिक्व स्वस्थ-व्यवस्था को सुविधाजनक कहानी चुनने की बजाय दोहो को सच रखना होगा।

यह प्रमाण का भार भीसही दिशा में रीसेट करता है। यदि कई कंपनी दवा करती है कि उसकी क्लिनिकल AI देखभाल बेहतर करती है, तो प्रसंगिक सक्षम लीडरबोर्ड नहीं है। वह ऐसा परीक्षण है, उन परणाम पर जो मरीजों के लिए मगने रखते हैं—और आदर्श रूप से इससे बड़ा क्योफ़ी यहाँई मगदा सबक यही है कि छोटे लम देखने के लिए बड़े अध्ययन चाहिए।

सफ़ा सारांश

केन्या के 16 प्रथमिक देखभाल केंद्रों में एक pragmatic, समूह-हित हेतु चिकित्सक-परिष्करण ने क्लिनिकल officers द्वारा इस्तेमाल किए जाने वाले इलेक्ट्रॉनिक रिकॉर्ड में जोड़े गए generative-AI निर्णय-सहायता उपकरण ("AI Consult") को परखा। 9,691 मरीजों में 14 दिनों के भीतर expert-adjudicated treatment-failure composite उपकरण के साथ 2.2% और बनाव उपकरण 2.0% था (समय-जोड़ odds अनुपात 0.77, 95% CI 0.55–1.08, P = 0.13)—कई महत्वपूर्ण अंतर नहीं। उपकरण से कई सुरक्षा संकेत नहीं मिला हर आंके गए क्षेत्र में क्लिनिकल दस्तवेजीकरण बेहतर हुआ, prescribing नहीं बदली मरीज-संतुष्टि समान रही और antibiotic लगत कुछ कम हमें से जुड़ी थी। लेख को कानिष्कर्ष है कि इन सीमाओं के भीतर उपकरण सुरक्षा दिखाले कि उपचार failure कम नहीं हुआ; यदि लम है तो शायद वह छोटा है, और देखे गए आकार का प्रभाव पकड़ने के लिए लगभग 100,000 मरीजों के स्तर का कहीं बड़ा परीक्षण चाहिए। परणाम न यह दिखाता है कि क्लिनिकल AI मरीजों के नतीजे बेहतर करती है, न यह कि वह बेकार है—यह दिखाता है कि सुरक्षा दिखने और प्रक्रिया में मदद करने वाला उपकरण एक शहरी नेटवर्क में दोहफ्तों में मरीजों को मगने योग्य लम नहीं दे सका।

बनाव-चढ़ाकर जाँच

पेपर क्या दिखाता है: वास्तविक दुनिया के रैंडमइज्ड परीक्षण में प्रथमिक देखभाल रिकॉर्ड में LLM आधारित निर्णय-सहायता उपकरण जोड़ने से कई सुरक्षा संकेत नहीं मिला, क्लिनिकल दस्तवेजीकरण बेहतर हुआ, prescribing नहीं बदली और 14-दिन के कठोर treatment-failure composite में महत्वपूर्ण कमी नहीं आई।

क्या संभव है लेकिन सिद्ध नहीं: उपकरण उपचार failure में इतना छोटा वास्तविक लम दे कि यह परीक्षण उसे पकड़ न सका, कुल लगत जोड़ने पर यह पैसा बचाए; बेहतर दस्तवेजीकरण आगे चलकर बेहतर देखभाल में बदले।

यह क्या नहीं दिखाता: कि क्लिनिकल AI मरीजों के परणाम बेहतर करती है; कि दुर्लभ घटनाओं के लिए यह असुरक्षा है; कि यह चिकित्सक को बदलती या उनसे बेहतर है; कि यही परणाम ग्रामीण या अधिक आय वाले परिवेश में भी होगा; कि लगत-बचत स्थिति है।

मुख्य सीमाएँ: अध्ययन उस प्रभाव से बड़ा प्रभाव पकड़ने के लिए powered था जो देखा गया (दुर्लभ परणाम के लिए बहुत बड़े नमूने चाहिए); configuration त्रुटि से नियंत्रण arm में दूषण हुआ और सन्नाने टर्क की आदतें तुलना को धुंधला करती हैं—दोहो अंतर कोशून्य की ओर झुका सकते हैं; पहले से ऊँचे मग को वाला एक ही नेटवर्क; पहले से तय noninferiority या सुरक्षा फ्रेमवर्क नहीं केवल 14-दिन का horizon।

समय पठक को कतिना भरें रखना चाहिए? इस परीक्षण में उपकरण से कई सुरक्षा संकेत नहीं मिला और दस्तवेजीकरण बेहतर हुआ—इस पर उच्च भरोसा यहाँ 14 दिनों में मरीज-परणाम मगने योग्य रूप से बेहतर नहीं हुआ—इस पर भी उच्च

भरसे। लेकिन “यह कम करता है” या “यह वफिल है” जैसे व्यक्तक patient-benefit दमोपर कम भरसे—वह प्रश्न वस्तव में खुला है और कहीबड़े अध्ययन की जरूरत है। उचित रुख: ऐसा सन्नधनीपूर्वक वस्तवकि-दुनियापरणिम जसिमें उपकरण सुरक्षति दिखा और देखभल प्रक्रिया बेहतर हुआ, न कि यह फैसला कि AI प्रणमकि देखभल को बदल देती है या बर्बद कर देती है।

स्रोत

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>