



The CLEAN PAPER

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

AI ya kliniki iliyoonekana salama na kuboresha paperwork — lakini haikuboresha patient outcomes

Pragmatic, cluster-randomized trial katika primary care facilities 16 Kenya ilijaribu generative-AI decision-support tool iliyoongezwa kwenye record inayotumiwa na clinical officers. Kati ya patients 9,691, strict 14-day expert-adjudicated composite ya treatment failure ilikuwa 2.2% na tool dhidi ya 2.0% bila (adjusted odds ratio 0.77, 95% CI 0.55–1.08, $P = 0.13$) — hakuna significant difference. Tool haikuonyesha safety signal, iliboresha documentation katika domains zote, haikubadilisha prescribing au satisfaction, na ilielekea lower antibiotic costs. Huo si utafiti ulioshindwa; ni mmoja wa adimu na wa uaminifu — uliopimwa kwa patients, si benchmarks — na unaonyesha ukweli usio na glamour kwamba tool iliyoonekana salama na kusaidia process haikuonyesha kwa dhahiri patient benefit ndani ya wiki mbili, na kuona benefit yoyote kungehitaji trial kubwa zaidi.

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

Kuonekana salama na kusaidia si sawa na kuwa na ufanisi kwa mgonjwa

Vichwa vingi vya habari kuhusu AI ya tiba hujengwa juu ya aina mbaya ya utafiti. Modeli inapata alama nzuri kwenye maswali ya mtihani, au inawashinda madaktari katika ya kliniki vignettes zilizochaguliwa vizuri, na simulizi linajiandika: mashine iko tayari kwa kliniki.

jaribio hili ilifanya jambo gumu na adimu zaidi. Iliweka generative-AI msaada wa maamuzi zana ndani ya huduma ya msingi halisi, katika vituo halisi, kwa wagonjwa halisi, na kuuliza swali lenye maana kweli: je, wagonjwa walikuwa na matokeo bora?

Jibu la uaminifu ni hapana — si kwa namna iliyoweza kupimwa, si ndani ya siku 14. zana haikuonyesha ishara ya hatari ya usalama. Iliboresha ubora wa nyaraka za kliniki. Huenda hata ilipunguza baadhi ya gharama za dawa. Lakini haikupunguza kushindwa kwa matibabu kwa significance, na waandishi wako makini kusema kwamba faida yoyote, kama ipo, huenda ni ndogo.

Hilo si kushindwa kwa utafiti. Ni utafiti kufanya kazi yake. Hivi ndivyo ushahidi inayowajibika kuhusu AI katika tiba inavyoonekana inapopimwa kwa wagonjwa badala ya kipimo linganishi.

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

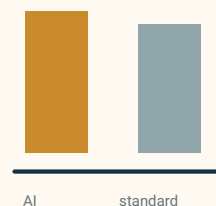
Process signal, not outcome proof.



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.

zana haikuonyesha ishara ya hatari ya usalama na iliboresha nyaraka za kliniki, lakini prespecified 14-day matokeo ya mgonjwa haikuboreka kwa significance. Msaada katika process si sawa na proven faida kwa mgonjwa.

Walichofanya waandishi

Timu iliendesha pragmatic, cluster-jaribio la nasibu katika **huduma ya msingi facilities 16** zinazoendeshwa na private health network (Penda Health) katika kaunti za Nairobi na Kiambu, Kenya. Huduma katika vituo hivi hutolewa kwa kiasi kikubwa na **maafisa wa kliniki** — wahudumu wa ngazi ya kati wenye diploma ya miaka mitatu — mara nyingi bila access rahisi kwa senior consultation.

Kitengo cha ugawaji wa nasibu kilikuwa mhudumu wa afya, si mgonjwa. **Maafisa wa kliniki 103** waligawanywa kwa nasibu: 52 katika kundi la uingiliaji na 51 katika kundi dhibiti. Makundi yote yalitumia rekodi ileile ya afya ya kielektroniki inayotegemea wingu. Kundi la uingiliaji pia lilikuwa na **AI Consult** (toleo la 2.0), zana ya msaada wa maamuzi iliyojengwa juu ya modeli kubwa ya lugha **GPT-4o ya OpenAI** na kuunganishwa ndani ya rekodi hiyo. Zana ilisoma taarifa alizoandika mhudumu wa afya na ingeweza kubainisha matatizo yanayowezekana katika utambuzi au mpango wa matibabu. Maafisa wa kliniki walibaki na mamlaka kamili ya kukubali, kurekebisha au kupuuza mapendekezo yake.

Ilikuwa modeli gani, na kwa nini details zake ni muhimu

zana ilikuwa **AI Consult 2.0**, ikiendesha **GPT-4o ya OpenAI** (release ya Mei 2025), kupitia commercial API ya OpenAI chini ya enterprise licence na kwa low-randomness settings (temperature 0.1). Ilikaa ndani ya bespoke electronic record (EasyClinic's EMR) na kuongozwa na mfumo prompts zilizoandikwa ili kulinganisha na Kenyan national matibabu guidelines; waandishi walichapisha instruction prompt kamili.

Kwa nini maelezo hayo ni muhimu? Kwa sababu matokeo yanahusu *mfumo mmoja maalum* — toleo moja la modeli, maagizo yake, rekodi moja na mazingira moja — si “LLM katika tiba” kwa ujumla. Waandishi wenyewe wanaita matokeo haya kipimo cha uwezo katika kipindi fulani, si makadirio yasiyobadilika ya uwezo. Modeli mpya zaidi, maagizo tofauti au kliniki yenye mifumo michache ya kidijitali inaweza kutoa matokeo tofauti.

Kuhusu independence: OpenAI baadaye ilitoa in-kind support (cloud-compute credits na technical guidance kuhusu API yake), lakini waandishi wanasema uamuzi wa kutumia OpenAI ulifanywa *kabla* ya offer hiyo, na OpenAI haikuwa na role katika jaribio design, data collection, uchanganuzi au decision to publish.

Kati ya 22 Aprili na 16 Julai 2025, **wagonjwa 9,691** waliandikishwa. **Kipimo kikuu cha matokeo kililenga mgonjwa na kiliwekwa kuwa kigumu kwa makusudi**: wataalamu waliokagua bila kujua kundi la mgonjwa waliamua kama kulikuwa na kushindwa kwa matibabu ndani ya siku 14 baada ya ziara, kwa mfano ugonjwa kutotatuliwa au kuwa mbaya zaidi. Jaribio lilisajiliwa mapema katika Pan-African majaribio ya kliniki Registry (202502499779176).

Chaguo hilo la design ndilo hoja. Ni rahisi kuonyesha kwamba AI zana inabadilisha kile clinician anaandika. Ni vigumu zaidi, na ina maana zaidi, kuonyesha kwamba inabadilisha kile kinachotokea kwa mgonjwa.

Walichogundua

Kipimo kikuu cha matokeo hakikuboreka. Kushindwa kwa matibabu kulitokea kwa wagonjwa 102 kati ya 4,693 (2.2%) katika kundi la AI na 94 kati ya 4,654 (2.0%) katika kundi dhibiti. Asilimia ghafi zilikuwa juu kidogo katika kundi la AI, lakini baada ya kurekebisha tofauti kati ya makundi ya wahudumu, makadirio yalielekea upande wa faida: **odds ratio iliyorekebishwa ilikuwa 0.77 (95% CI 0.55–1.08, P = 0.13)**. Hiyo haikuwa muhimu kitakwimu. Kipindi cha kujiamini kinajumuisha kwa urahisi uwezekano wa kutokuwa na athari, hivyo faida haiwezi kudaiwa; na hata kwa ukubwa halisi, athari iliyokadiriwa ilikuwa ndogo.

Kwa maelezo rahisi ya kusoma odds ratios, confidence intervals na P values pamoja, tazama mwongozo wa kusoma matokeo ya kliniki.

Hakukuwa na ishara ya hatari ya usalama — ndani ya mipaka ya utafiti. Hakuna tukio kubwa lisilofaa lililohukumiwa kuwa limesababishwa na zana, na ukaguzi huru haukupata ishara ya usalama. Lakini waandishi wanaweka kikomo wazi: sampuli haikuwa kubwa vya kutosha kugundua madhara makubwa yaliyo nadra, na jaribio halikuwa na mpango uliowekwa mapema wa kuthibitisha kutokuwa duni wala mfumo rasmi wa usalama. Kwa hiyo haliwezi *kuthibitisha* usalama kwa matukio adimu.

Documentation iliboreka. Kati ya encounters 2,000 zilizopitiwa na blinded experts, clinicians walio-tumia AI Consult waliandika nyaraka za kliniki bora katika domains zote zilizopimwa — recorded utambuzi, mpango wa matibabu na overall completeness.

Prescribing karibu haikubadilika. Hakukuwa na tofauti muhimu kitakwimu katika prescribing, ikiwemo correct antibiotic use (adjusted odds ratio 0.86, 95% CI 0.48 hadi 1.55). zana haikubadilisha antibiotic prescribing rates.

Wagonjwa hawakuona tofauti. Kati ya wagonjwa 826 waliokamilisha satisfaction survey, satisfaction ilikuwa karibu sawa katika arms zote, na consultation times zilifanana.

Gharama zilielekea chini kidogo. Katika adjusted uchanganuzi, antibiotic-related gharama zilikuwa chini katika AI arm — pengine kwa kuchagua dawa za bei nafuu badala ya prescriptions chache — na per-mgonjwa antibiotic saving ilionekana kuzidi per-mgonjwa gharama ya kuendesha zana. Waandishi wanaitaja kama suggestive, si settled: full total-gharama-of-ownership accounting ilikuwa nje ya jaribio.

One-line summary ya waandishi wenyewe ndiyo safi zaidi: LLM assistance ilikuwa salama ndani ya mipaka hiyo lakini haikupunguza kushindwa kwa matibabu ndani ya siku 14, na faida yoyote huenda ni ndogo.

Kwa nini “no tofauti muhimu kitakwimu” si “haifanyi kazi”

Null matokeo makuu ni rahisi kuisoma kupita kiasi pande zote. Mambo mawili yanazuia simulizi rahisi.

Kwanza, **jaribio liliandaliwa kugundua athari kubwa kuliko ile iliyopatikana.** Matokeo mabaya makubwa katika huduma ya msingi ni nadra — karibu 2% hapa — hivyo kutenganisha faida ndogo ya kweli na kelele ya kitakwimu kunahitaji idadi kubwa sana ya wagonjwa. Mahesabu ya baada ya utafiti yanaonyesha kwamba kugundua athari ya ukubwa ulioonekana kungehitaji **jaribio kubwa zaidi sana, la zaidi ya wagonjwa 100,000.** Matokeo yasiyo muhimu kitakwimu katika utafiti huu hayaondoi uwezekano wa faida ndogo ya kweli; yanaonyesha tu kwamba utafiti huu haukuwa mkubwa vya kutosha kuitenganisha na kelele.

Pili, **ulinganisho ilichanganywa kiasi.** Configuration error kwa muda mfupi iliwapa baadhi ya control-arm clinicians access kwa AI Consult, na clinicians katika shared network huzungumza na kubeba habits kuvuka boundary. athari zote mbili hufanya arms zionekane sawa zaidi, zikisukuma real difference yoyote kuelekea zero. Zaidi ya hapo, host network tayari ilikuwa na relatively high standards, hivyo nafasi ya zana kuonyesha improvement ilikuwa ndogo.

Hakuna kati ya hayo kinachookoa “breakthrough” kichwa cha habari. Lakini inamaanisha reading sahihi ni calibrated, si dismissive: kwenye kigezo cha mwisho ngumu na ya uaminifu zaidi, zana hii haikuonyesha kwa dhahiri kuwa inawasaidia wagonjwa ndani ya wiki mbili — huku ikisaidia record-keeping kwa measurable way na kuonekana salama.

Hili halithibitishi nini

- **Halionyeshi kwamba AI iliboresha matokeo ya wagonjwa.** Katika primary 14-day kigezo cha mwisho, hakukuwa na faida muhimu kitakwimu.
- **Halionyeshi kwamba AI haina maana.** makadirio ya nukta iliielekea, documentation iliboreka kote, na drug gharama ziliielekea chini; null matokeo inaendana na small real faida ambayo jaribio lilikuwa ndogo sana kuthibitisha.
- **Halithibitishi zana ni safe kwa madhara adimu.** Haikuonyesha ishara ya hatari ya usalama, lakini haikuwa powered au designed kuthibitisha usalama kwa uncommon severe events.
- **Halionyeshi kwamba “AI beats madaktari” au inachukua nafasi ya clinicians.** Hii ni decision support; clinician alibaki na full authority kukubali au kukataa.
- **Halisambai automatically kila mahali.** jaribio ilifanyika katika single private urban network Kenya; rural, periurban na higher-income settings zinaweza kutofautiana pande zote.
- **Haliestablish gharama savings.** gharama ishara ni suggestive, si completed economic evaluation.

Ushahidi una nguvu kiasi gani?

Kwa dai kuu — kwamba hakuna upungufu uliothibitishwa wa madhara ya muda mfupi kwa mgonjwa — ushahidi ni thabiti kutokana na muundo wa utafiti, na hitimisho limewekwa kwa tahadhari inayofaa. Jaribio la mbele, lililosajiliwa mapema na kugawa wahudumu kwa nasibu, lenye kipimo cha matokeo kwa kila mgonjwa kilichohukumiwa na wataalamu wasiojua kundi, ni aina ya ushahidi wa mazingira halisi wenye uzito mkubwa kwa zana kama hii. Lina taarifa nyingi zaidi kuliko alama ya kipimo linganishi au utafiti wa vignettes.

Kwa secondary matokeo — better documentation, unchanged prescribing, similar satisfaction, lower antibiotic gharama — ushahidi ni nzuri lakini inapaswa kusomwa kama secondary: supportive ishara, si kichwa cha habari, na vulnerable kwa contamination na single-network limits zilezile.

Kwa usalama, ushahidi ni reassuring lakini bounded: hakuna ishara iliyopatikana, lakini si utafiti iliyojengwa kuona madhara adimu.

Msimamo wenye manufaa zaidi si “inafanya kazi” wala “imeshindwa.” Ni: katika ulimwengu halisi jaribio iliyofanywa kwa umakini iliona AI zana hii haikutoa ishara ya hatari ya usalama na iliboresha mchakato wa huduma, bila kuonyesha mgonjwa-outcome faida ndani ya wiki mbili — na kwamba kuona faida kama ipo kungehitaji utafiti mkubwa zaidi.

Kwa nini ni muhimu

Mjadala kuhusu AI ya kitabibu una njaa ya ushahidi wa aina hii. Kuna maelfu ya makala yanayoonyesha modeli zikifanya vizuri kwenye exams na kulingana na clinicians katika cases safi. Kuna pragmatic randomized majaribio chache sana zinazopima kama wagonjwa halisi wanakuwa bora. Hii ni mojawapo, na inafika kwenye ukweli usio na glamour: kupita test si sawa na kusaidia mgonjwa.

Hilo ndilo pengo muhimu. Zana inaweza kuwa na manufaa halisi kwa wahudumu — kumbukumbu bora, gharama ndogo za dawa, au kuwa jozi ya pili ya macho — na bado isibadilishe matokeo magumu ya mgonjwa ndani ya wiki mbili. Mambo yote mawili yanaweza kuwa kweli kwa wakati mmoja, na mfumo wa afya uliokomaa unapaswa kuyashikilia pamoja badala ya kuchagua linalofaa kichwa cha habari.

Pia inaweka mzigo wa uthibitisho mahali panapofaa. Kampuni ikitaka kudai AI yake ya kliniki inaboresha huduma, ushahidi unaofaa si alama ya jedwali la majaribio bali jaribio kama hili, lenye matokeo yenye maana kwa wagonjwa — na ikiwezekana jaribio kubwa zaidi. Somo hapa ni kwamba faida ndogo zinahitaji tafiti kubwa ili kuonekana kwa uhakika.

Muhtasari safi

Jaribio la kimatendo lililogawa kwa nasibu wahudumu katika vituo 16 vya huduma ya msingi nchini Kenya lilipima **AI Consult**, zana ya msaada wa maamuzi ya AI iliyoongezwa kwenye rekodi ya kielektroniki inayotumiwa na maafisa wa kliniki. Kati ya wagonjwa 9,691, kipimo cha pamoja cha kushindwa kwa matibabu ndani ya siku 14 kilikuwa 2.2% kwa waliotibiwa kwa msaada wa zana dhidi ya 2.0% bila zana (adjusted odds ratio 0.77, 95% CI 0.55–1.08, $P = 0.13$), tofauti isiyo muhimu kitakwimu. Zana haikuonyesha ishara ya hatari ya usalama, iliboresha nyaraka za kliniki katika maeneo yote yaliyopimwa, haikubadilisha sana uagizaji wa dawa, kuridhika kwa wagonjwa kulibaki karibu sawa, na gharama za antibiotiki zilielekea kuwa chini kidogo. Waandishi wanahitimisha kuwa zana ilionekana salama ndani ya mipaka ya jaribio lakini haikupunguza kushindwa kwa matibabu; faida yoyote huenda ni ndogo, na kugundua athari ya ukubwa ulioonekana kungehitaji jaribio lenye

zaidi ya wagonjwa 100,000. Matokeo haya hayaonyeshi kwamba AI ya kliniki inaboresha matokeo ya wagonjwa, wala kwamba haina maana. Yanaonyesha zana iliyosaidia mchakato wa huduma bila kuonyesha wazi faida kwa mgonjwa ndani ya wiki mbili katika mtandao mmoja binafsi wa mijini.

Ukaguzi bila kupamba

Makala inaonyesha nini: Katika jaribio la nasibu katika mazingira halisi, kuongeza zana ya msaada wa maamuzi inayotumia LLM kwenye rekodi za huduma ya msingi hakukuonyesha ishara ya hatari ya usalama, kuliboresha ubora wa nyaraka za kliniki na hakukubadilisha sana uagizaji wa dawa, lakini hakukupunguza kwa kiwango muhimu kitakwimu kipimo cha pamoja cha kushindwa kwa matibabu ndani ya siku 14.

Kinachowezelekana lakini hakijathibitishwa: Kwamba zana hupunguza kushindwa kwa matibabu kwa kiwango kidogo ambacho jaribio hili halikuwa kubwa vya kutosha kugundua; kwamba inaokoa pesa baada ya gharama zote kuhesabiwa; na kwamba nyaraka bora hatimaye huboresha huduma.

Isichoonyeshwa: Kwamba AI ya kliniki inaboresha matokeo ya wagonjwa; kwamba ni unsafe kwa rare events; kwamba inachukua nafasi au inawashinda clinicians; kwamba matokeo hizi zinahamia rural au higher-income settings; kwamba gharama saving imeestablished.

Vikwazo vikuu: Jaribio liliundwa kugundua athari kubwa kuliko iliyopatikana, kwa sababu matokeo mabaya yaliyo nadra yanahitaji sampuli kubwa sana; hitilafu ya usanidi iliruhusu baadhi ya wahudumu wa kundi dhibiti kutumia AI Consult kwa muda, na mazoea katika mtandao mmoja yanaweza kuvuka kati ya makundi, hivyo kupunguza tofauti; utafiti ulifanyika katika mtandao mmoja binafsi wa mijini wenye viwango tayari vya juu; hakukuwa na mpango uliowekwa mapema wa noninferiority au mfumo rasmi wa usalama; na ufuatiliaji ulikuwa wa siku 14 tu.

Msomaji wa kawaida anapaswa kuwa na uhakika kiasi gani? Uhakika mkubwa kwamba zana haikuonyesha ishara ya hatari ya usalama katika jaribio hili na iliboresha uwekaji kumbukumbu. Uhakika mkubwa pia kwamba haikuonyesha uboreshaji wazi wa matokeo ya wagonjwa ndani ya siku 14. Uhakika mdogo kwa kauli pana kwamba “inafanya kazi” au “imeshindwa” kama zana ya kuboresha matokeo ya wagonjwa — swali hilo bado halijatatuliwa na linahitaji utafiti mkubwa zaidi. Msimamo unaofaa ni kuona hili kama matokeo makini ya mazingira halisi kuhusu zana iliyosaidia mchakato wa huduma bila kuonyesha wazi faida kwa mgonjwa, si hukumu ya mwisho kuhusu AI katika huduma ya msingi.

Vyanzo

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>