



The CLEAN PAPER

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

Klīniskais MI izskatījās drošs un uzlaboja dokumentāciju — bet neuzlaboja pacientu iznākumus

Pragmatiskā, klasteros randomizētā pētījumā 16 Kenijas primārās aprūpes iestādēs pārbaudīja ģeneratīva MI lēmumu atbalsta rīku, kas pievienots klīnisko darbinieku elektroniskajam ierakstam. 9691 pacientam stingrais ekspertu vērtētais 14 dienu ārstēšanas neveiksmes kombinētais galapunkts bija 2,2% ar rīku pret 2,0% bez tā (koriģētā izredžu attiecība 0,77; 95% TI 0,55–1,08; $P = 0,13$) — statistiski nozīmīgas atšķirības nebija. Rīkam netika konstatēts drošības signāls, tas uzlaboja dokumentāciju visās vērtētajās jomās, nemainīja izrakstīšanu vai apmierinātību un uzrādīja tendenci uz zemākām antibiotiku izmaksām. Tas nav “neveiksmīgs” pētījums, bet rets, godīgs pētījums uz pacientiem, ne etaloniem.

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

Drošs un noderīgs nav tas pats, kas efektīvs

Lielākā daļa virsrakstu par medicīnisko MI balstās nepareizā pētījumu tipā. Modelis labi atbild uz eksāmena jautājumiem vai pārspēj ārstus rūpīgi atlasītās klīniskās situācijās, un stāsts uzrakstās pats: mašina ir gatava klīnikai.

Šis pētījums izdarīja ko grūtāku un daudz retāku. Tas ieviesa ģeneratīva MI lēmumu atbalsta rīku reālajā primārajā aprūpē, īstās iestādēs un ar īstiem pacientiem, un uzdeva jautājumu, kas patiešām ir svarīgs: vai pacientiem klājās labāk?

Godīgā atbilde ir — nē, vismaz ne izmērāmi 14 dienu laikā. Rīkam netika konstatēts drošības signāls. Tas uzlaboja klīniskās dokumentācijas kvalitāti. Iespējams, tas pat neredz samazināja zāļu izmaksas. Taču tas statistiski nozīmīgi nesamazināja ārstēšanas neveiksmes, un autori rūpīgi norāda, ka jebkurš ieguvums, ja tāds ir, visticamāk, ir neliels.

Tas nav pētījuma “neveiksmes” stāsts. Tas ir pētījums, kas darbojas, kā tam jādarbojas. Šādi izskatās atbildīgi pierādījumi par MI medicīnā, ja vērtē pacientu iznākumus, nevis etalontestu rezultātus.

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

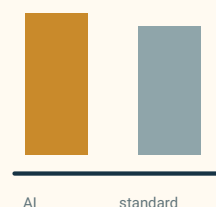
Process signal, not outcome proof.



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.

Rīkam šajā pētījumā netika konstatēts drošības signāls un tas uzlaboja klīnisko dokumentāciju, taču iepriekš definētais 14 dienu pacientu iznākums statistiski nozīmīgi neuzlabojās. Procesa uzlabošana nav tas pats, kas pierādīts ieguvums pacientam.

Ko autori darīja

Komanda veica pragmatisku, klasteros randomizētu pētījumu **16 primārās aprūpes iestādēs**, ko Nairobi un Kiambu apgabalos Kenijā pārvalda privāts veselības aprūpes tīkls Penda Health. Aprūpi šajās iestādēs lielā mērā sniedz **klīniskie darbinieki** (*clinical officers*) — vidējā līmeņa praktiķi ar trīs gadu profesionālo izglītību —, kuriem bieži nav viegli pieejamas vecāka speciālista konsultācijas.

Randomizācijas vienība bija klīnicists, nevis pacients. **103 klīniskie darbinieki** tika sadalīti grupās: 52 intervences un 51 kontroles grupā. Abas izmantoja vienu un to pašu mākoņplatformas elektronisko medicīnisko ierakstu. Intervences grupa papildus saņēma **AI Consult** (versija 2.0) — lēmumu atbalsta rīku, kas balstīts uz **OpenAI GPT-4o** lielo valodas modeli un bija integrēts ierakstā. Tas nolasīja klīnicista ievadīto informāciju un varēja norādīt uz iespējamām problēmām diagnozē vai ārstēšanas plānā. Klīnicists saglabāja pilnu kontroli: ieteikumus varēja pieņemt, mainīt vai ignorēt.

Kurš modelis tika izmantots un kāpēc detaļas ir svarīgas

Rīks bija **AI Consult 2.0**, kas darbojās ar **OpenAI GPT-4o** (2025. gada maija versija), izmantojot OpenAI komerciālo API uzņēmuma licenci un zemu nejaušības pakāpi (temperature 0,1). Tas bija integrēts īpaši izstrādātā elektroniskajā ierakstā (EasyClinic EMR) un vadīts ar sistēmas uzvednēm, kas veidotas atbilstoši Kenijas valsts ārstēšanas vadlīnijām; autori publicēja pilno instrukciju uzvedni.

Kāpēc tik precīzi? Tāpēc, ka rezultāts attiecas uz *vienu konkrētu sistēmu* — vienu modeļa versiju, vienu uzvedni, vienu ieraksta sistēmu un vienu iestatījumu —, nevis uz “LLM medicīnā” kopumā. Paši autori uzsver to pašu, savu rezultātu saucot par *laika momenta etalonu, ne fiksētu spēju novērtējumu*. Jaunāks modelis, cita uzvedne vai mazāk digitalizēta klīnika varētu dot citu rezultātu.

Par neatkarību: vēlāk OpenAI sniedza atbalstu natūrā — mākoņskaitļošanas kredītus un tehniskas konsultācijas API lietošanai —, taču autori norāda, ka lēmums izmantot OpenAI tika pieņemts *pirms* šī piedāvājuma un ka OpenAI nepiedalījās pētījuma dizainā, datu vākšanā, analizē vai lēmumā publicēt.

No 2025. gada 22. aprīļa līdz 16. jūlijam tika iekļauts **9691 pacients**. **Primārais iznākums bija apzināti stingrs un vērsts uz pacientu**: ekspertu akli vērtēts *ārstēšanas neveiksmes* kombinētais galapunkts 14 dienu laikā pēc vizītes — klīnicistu komisija, nezinot grupu, vērtēja, vai pacientam bijis nelabvēlīgs iznākums, piemēram, nepārgājusi vai pasliktinājusies sasilšana. Pētījums bija iepriekš reģistrēts Pan-African Clinical Trials Registry (202502499779176).

Tieši šī dizaina izvēle ir būtiskā. Ir viegli parādīt, ka MI rīks maina to, ko klīnicists ieraksta. Daudz grūtāk un daudz nozīmīgāk ir parādīt, ka tas maina to, kas notiek ar pacientu.

Ko viņi atklāja

Primārais iznākums neuzlabojās. Ārstēšanas neveiksme bija 102 no 4693 pacientiem (2,2%) MI grupā un 94 no 4654 (2,0%) kontroles grupā. Neapstrādātajos procentos MI grupa bija nedaudz sliktāka, taču pēc korekcijas punkta novērtējums sliecās ieguvuma virzienā: **koriģētā izredžu attiecība 0,77 (95% ticamības intervāls 0,55–1,08; P = 0,13)** — statistiski nenoīmīgi. Atšķirība starp neapstrādāto un koriģēto rezultātu nav matemātiska kļūda; korekcija ņem vērā atšķirības starp klīnicistu klasteriem. Jebkurā gadījumā ticamības intervāls ērti ietver “nav efekta”, tāpēc ieguvumu apgalvot nevar, un absolūtais efekts bija ļoti mazs.

Vienkāršu skaidrojumu, kā kopā lasīt izredžu attiecību, ticamības intervālu un P vērtību, skatiet ceļvedī par klīniskā rezultāta lasīšanu.

Drošības signāla nebija — ar ierobežojumiem. Neviena nopietns nevēlams notikums netika atzīts par saistītu ar rīku, un neatkarīgā pārskatā drošības signālu nekonstatēja. Autori atklāti norāda šā secinājuma robežu: pētījumam nebija pietiekamas jaudas retu smagu kaitējumu noteikšanai, un tam nebija iepriekš definētas ne-zemākas efektivitātes vai formālas drošības shēmas, tāpēc tas nevar pierādīt drošību reti notikumiem.

Dokumentācija kļuva labāka. 2000 vizītēs, ko akli pārskatīja eksperti, AI Consult lietojošie klīnicisti uzrādīja labāku dokumentācijas kvalitāti visās vērtētajās jomās — diagnozes ierakstā, ārstēšanas plānā un kopējā pilnīgumā.

Zāļu izrakstīšana gandrīz nemainījās. Statistiski noīmīgas atšķirības nebija arī pareizā antibiotiku lietošanā (koriģētā izredžu attiecība 0,86; 95% TI 0,48–1,55). Rīks nemainīja antibiotiku izrakstīšanas biežumu.

Pacienti atšķirību nepamanīja. No 826 pacientiem, kas aizpildīja apmierinātības aptauju, apmierinātība abās grupās bija būtībā vienāda, un konsultāciju ilgums neatšķīrās.

Izmaksas nedaudz sliecās uz leju. Koriģētā analizē ar antibiotikām saistītās izmaksas MI grupā bija zemākas — iespējams, izvēloties lētākas zāles, nevis izrakstot mazāk —, un antibiotiku ietaupījums uz pacientu šķita lielāks par rīka izmantošanas izmaksām uz pacientu. Autori to raksturo kā daudz-sološu, ne galīgu rezultātu: pilna īpašumtiesību izmaksu analīze nebija pētījuma daļa.

Pašu autoru vienas rindas kopsavilkums ir visprecīzākais: LLM atbalsts šajās robežās izskatījās drošs, bet nesamazināja ārstēšanas neveiksmi 14 dienu laikā, un jebkurš ieguvums, visticamāk, ir neliels.

Kāpēc “nav statistiski noīmīgas atšķirības” nenoīmē “tas nedarbojas”

Nulles primāro iznākumu viegli pārinterpretēt abos virzienos. Divas lietas neļauj stāstu vienkāršot.

Pirmkārt, **pētījums bija plānots lielāka efekta noteikšanai nekā tas, kas tika novērots.** Nopietni nelabvēlīgi iznākumi primārajā aprūpē ir reti — šeit ap 2% —, tāpēc neliela īsta ieguvuma nošķiršanai

no trokšņa vajag milzīgu dalībnieku skaitu. Pašu autoru post-hoc jaudas aprēķini liecina, ka tāda efekta noteikšanai, kāds novērots šajā pētījumā, būtu vajadzīgs **daudz lielāks pētījums — vairāk nekā 100 000 pacientu mērogā**. Nenožīmīgs rezultāts šāda izmēra pētījumā neizslēdz nelielu īstu ieguvumu; tas nozīmē, ka šis pētījums nespēja to droši izšķirt.

Otrkārt, **salīdzinājums daļēji izplūda**. Konfigurācijas kļūda uz īsu laiku deva dažiem kontroles grupas klīnicistiem piekļuvi AI Consult, un viena tīkla klīnicisti savā starpā sarunājas un pārņem darba ieradumus pāri grupu robežām. Abi faktori padara grupas līdzīgākas un jebkuru īstu atšķirību velk nulles virzienā. Turklāt veselības tīkls jau strādāja salīdzinoši augstā kvalitātes līmenī, atstājot rīkam mazāk vietas uzlabojumam.

Nekas no tā neglābj “izrāviena” virsrakstu. Taču pareizais lasījums nav arī cinisks: uz stingrākā un godīgākā galapunkta šis rīks 14 dienās nepierādīja pacientu ieguvumu, vienlaikus izmērāmi uzlabojot dokumentāciju un neuzrādot drošības signālu.

Ko tas nepierāda

- Tas **neparāda**, ka MI uzlaboja pacientu iznākumus. Primārajā 14 dienu galapunktā statistiski nozīmīga ieguvuma nebija.
- Tas **neparāda**, ka MI ir bezjēdzīgs. Punkta novērtējums sliecās ieguvuma virzienā, dokumentācija uzlabojās visās jomās un zāļu izmaksas nedaudz kritās; nulles rezultāts ir saderīgs ar nelielu īstu ieguvumu, ko pētījums bija par mazu, lai apstiprinātu.
- Tas **nepierāda**, ka rīks ir drošs retiem kaitējumiem. Drošības signāls netika atrasts, bet pētījums nebija veidots retu smagu notikumu sertificēšanai.
- Tas **neparāda**, ka “MI pārspēj ārstus” vai aizstāj klīnicistus. Tas ir lēmumu *atbalsts*; pilna lēmumu vara palika cilvēkam.
- Rezultāti **nepārnesas automātiski**. Pētījums notika vienā privātā pilsētu tīklā Kenijā; lauku, piepilsētu vai augstāku ienākumu sistēmās efekts var būt citāds jebkurā virzienā.
- Tas **nepierāda izmaksu ietaupījumu**. Izmaksu signāls ir daudzsološs, bet nav pabeigta ekonomiskā analīze.

Cik pārliecinoši ir pierādījumi?

Centrālajam secinājumam — nav pierādīta īstermiņa pacienta kaitējuma samazināšanās — pierādījumi ir spēcīgi *pētījuma dizaina ziņā* un pareizi piesardzīgi *secinājuma ziņā*. Prospektīvs, iepriekš reģistrēts, klasteros randomizēts pētījums ar akli vērtētu pacienta līmeņa kombinēto galapunktu ir tuvu labākajam reālās pasaules pierādījumu tipam, ko šādam rīkam var iegūt. Tas ir nesalīdzināmi informatīvāks par etalona punktiem vai klīnisku vinješu testu.

Sekundārajiem rezultātiem — labākai dokumentācijai, nemainītai izrakstīšanai, līdzīgai apmierinātībai, zemākām antibiotiku izmaksām — pierādījumi ir labi, taču tie jālasa kā sekundāri signāli, ne virsraksts, un tos ietekmē tie paši grupu piesārņojuma un viena tīkla ierobežojumi.

Drošības dati ir nomierinoši, bet ierobežoti: signāls netika atrasts, taču pētījums nebija būvēts retu kaitējumu atklāšanai.

Noderīgākā attieksme nav ne “tas darbojas”, ne “tas izgāzās”. Tā ir: rūpīgs reālās pasaules pētījums atklāja, ka šis MI rīks neuzrādīja drošības signālu un uzlaboja aprūpes procesu, bet 14 dienu laikā nepierādīja pacientu iznākuma ieguvumu — un jebkura tāda ieguvuma noteikšanai vajadzētu daudz lielāku pētījumu.

Kāpēc tas ir svarīgi

Diskusijā par medicīnisko MI ļoti trūkst tieši šādu pierādījumu. Ir tūkstošiem darbu, kur modeļi labi kārtoti eksāmenus un līdzinās klīnicistiem tīri sagatavotos gadījumos. Ir ļoti maz lielu, pragmatisku, randomizētu pētījumu, kas mēra, vai īstiem pacientiem klājas labāk. Šis ir viens no tiem, un tas nonāk pie neglamūrīgas patiesības: nokārtot testu nav tas pats, kas palīdzēt pacientam.

Tieši šī plaista ir visa stāsta būtība. Rīks var būt klīnicistam patiesi noderīgs — skaidrākas piezīmes, mazāki zāļu izdevumi, “otrs acu pāris” — un tomēr divu nedēļu laikā nemainīt stingru pacienta iznākumu. Abi fakti var būt patiesi vienlaikus, un nobriedušai veselības sistēmai tie jātur kopā, ne jāizvēlas ērtākais.

Pētījums arī pareizi pārbīda pierādīšanas pienākumu. Ja uzņēmums grib apgalvot, ka tā klīniskais MI uzlabo aprūpi, atbilstošais pierādījums nav līderu tabula. Tas ir šāda veida pētījums ar pacientiem svarīgiem iznākumiem — un ideālā gadījumā vēl lielāks, jo godīgā mācība šeit ir tāda, ka mērenu ieguvumu noteikšanai vajag lielus paraugus.

Īss kopsavilkums

Pragmatiskā, klasteros randomizētā pētījumā 16 Kenijas primārās aprūpes iestādēs pārbaudīja ģeneratīva MI lēmumu atbalsta rīku AI Consult, kas pievienots klīnisko darbinieku elektroniskajam ierakstam. 9691 pacientam ekspertu akli vērtētais ārstēšanas neveiksmes kombinētais galapunkts 14 dienu laikā bija 2,2% ar rīku un 2,0% bez tā (korigētā izredžu attiecība 0,77; 95% TI 0,55–1,08; $P = 0,13$) — statistiski nozīmīgas atšķirības nebija. Rīkam netika konstatēts drošības signāls, tas uzlaboja klīnisko dokumentāciju visās vērtētajās jomās, nemainīja zāļu izrakstīšanu vai pacientu apmierinātību un bija saistīts ar nelielu zemākām antibiotiku izmaksām. Autori secina, ka rīks šajās robežās izskatījās drošs, bet nesamazināja ārstēšanas neveiksmi; jebkurš ieguvums, visticamāk, ir neliels, un novērotā efekta noteikšanai būtu vajadzīgs daudz lielāks pētījums — vairāk nekā 100 000 pacientu mērogā. Rezultāts neparāda ne to, ka klīniskais MI uzlabo pacientu iznākumus, ne to, ka tas ir bezjēdzīgs; tas rāda, ka rīks, kas palīdzēja procesam un neuzrādīja drošības signālu, 14 dienās nepierādīja izmērāmu pacienta ieguvumu vienā pilsētas privātā tīklā.

Bez pārspīlējumiem

Ko pētījums parāda: reālās pasaules randomizētā pētījumā LLM lēmumu atbalsta rīka pievienošana primārās aprūpes ierakstam neuzrādīja drošības signālu, uzlaboja klīniskās dokumentācijas kvalitāti, nemainīja izrakstīšanu un statistiski nozīmīgi nesamazināja stingru 14 dienu ārstēšanas neveiksmes kombinēto pacienta galapunktu.

Kas ir ticams, bet nav pierādīts: ka rīks rada nelielu īstu ārstēšanas neveiksmes samazinājumu, ko šis pētījums bija par mazu noteikt; ka pilnā izmaksu analizē tas ietaupa naudu; ka labāka dokumentācija ilgākā laikā pārvēršas labākā aprūpē.

Ko tas neparāda: ka klīniskais MI uzlabo pacientu iznākumus; ka tas ir nedrošs reti notikumiem; ka tas aizstāj vai pārspēj klīnicistus; ka rezultāti pārnēs uz lauku vai turīgāku valstu sistēmām; ka izmaksu ietaupījums ir pierādīts.

Galvenie ierobežojumi: pētījums bija plānots lielākam efektam nekā novērots (retu iznākumu gadījumā vajag milzīgus paraugus); konfigurācijas kļūda īslaicīgi “piesārņoja” kontroles grupu un kopējā tīkla paradumi izpludināja salīdzinājumu, abi faktori velkot atšķirību uz nulli; viens privāts pilsētas tīkls ar jau augstiem standartiem; nebija iepriekš definētas non-inferiority vai drošības shēmas; tikai 14 dienu novērošanas horizonts.

Cik lielai pārliecībai vajadzētu būt vispārīgam lasītājam? Augstai par to, ka šajā pētījumā rīks neuzrādīja drošības signālu un uzlaboja dokumentāciju. Augstai arī par to, ka šeit tas nepierādīja 14 dienu pacienta iznākuma uzlabojumu. Zemai par jebkuru kategorisku secinājumu, ka tas “darbojas” vai “nedarbojas” pacientu ieguvumam — šis jautājums patiesi paliek neatrisināts un prasa daudz lielāku pētījumu. Pareizā attieksme: rūpīgs reālās pasaules rezultāts par rīku, kas palīdzēja aprūpes procesam, ne spriedums, ka MI pārveido vai sagrauj primāro aprūpi.

Avoti

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>