



The CLEAN PAPER

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

Kliininen tekoäly näytti turvalliselta ja paransi kirjaamista — mutta ei potilastuloksia

Pragmaattisessa, klusterisatunnaistetussa tutkimuksessa 16 kenialaisessa perusterveydenhuollon yksikössä testattiin generatiiviseen tekoälyyn perustuvaa päätöksentukityökalua, joka lisättiin clinical officereiden käyttämään potilastietojärjestelmään. 9,691 potilaalla tiukka, asiantuntijoiden sokkona arvioima 14 päivän hoidon epäonnistumisen yhdistelmäpäätetapahtuma oli työkalulla 2.2% ja ilman sitä 2.0% (korjattu odds ratio 0.77, 95%:n luottamusväli 0.55-1.08, $P = 0.13$) — ei merkitsevää eroa. Työkalu ei tuottanut turvallisuussignaalia, paransi dokumentointia kaikilla arvioiduilla osa-alueilla, ei muuttanut lääkemääräyksiä tai tyytyväisyyttä ja viittasi pienempiin antibioottikustannuksiin. Tämä ei ole epäonnistunut tutkimus vaan harvinainen rehellinen tutkimus: tulos mitattiin potilailla, ei vertailutesteissä. Arkinen johtopäätös on, että turvalliselta näyttänyt ja prosessia auttanut työkalu ei osoitettavasti auttanut potilaita kahdessa viikossa, ja mahdollisen hyödyn havaitseminen vaatisi paljon suuremman tutkimuksen.

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

Turvallinen ja hyödyllinen ei ole sama asia kuin vaikuttava

Useimmat lääketieteellistä tekoälyä koskevat otsikot perustuvat vääränlaiseen tutkimukseen. Malli menestyy tenttikysymyksissä tai voittaa lääkärit huolellisesti valituissa tapauskuvauksissa, ja tarina kirjoittaa itse itsensä: kone on valmis klinikalle.

Tämä tutkimus teki jotain vaikeampaa ja harvinaisempaa. Generatiiviseen tekoälyyn perustuva päätöksentukityökalu vietiin oikeaan perusterveydenhuoltoon, oikeisiin yksiköihin ja oikeiden potilaiden hoitoon. Sitten kysyttiin kysymys, jolla on oikeasti väliä: voivatko potilaat paremmin?

Rehellinen vastaus on ei — ei mitattavasti, ainakaan 14 päivässä. Työkalusta ei löytynyt turvallisuussignaalia. Se paransi kliinisen dokumentoinnin laatua. Se saattoi jopa pienentää joitakin lääkekustannuksia. Mutta se ei vähentänyt hoidon epäonnistumisia merkittävästi, ja tutkijat sanovat varovasti, että mahdollinen hyöty on todennäköisesti vaatimaton.

Se ei ole tutkimuksen epäonnistuminen. Se on tutkimuksen onnistumista. Tältä vastuullinen näyttö tekoälystä lääketieteessä näyttää, kun tulosta mitataan potilaista eikä vertailutesteistä.

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

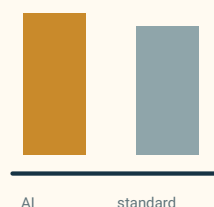
Process signal, not outcome proof.



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.

Työkalu ei tuottanut turvallisuussignaalia ja paransi kliinistä dokumentointia, mutta ennalta määritelty 14 päivän potilastulos ei parantunut merkittävästi. Prosessin helpottaminen ei ole sama asia kuin osoitettu potilashyöty.

Mitä tutkijat tekivät

Ryhmä toteutti pragmaattisen, klusterisatunnaistetun tutkimuksen **16 perusterveydenhuollon yksikössä**, joita yksityinen Penda Health -verkosto ylläpitää Nairobissa ja Kiambun piirikunnissa Keniassa. Hoitoa antavat suurelta osin **clinical officerit** — kolmen vuoden tutkinnon suorittaneet keskitason terveydenhuollon ammattilaiset — joilla ei aina ole helppoa pääsyä kokeneemman lääkärin konsultaatioon.

Satunnaistamisen yksikkö oli klinikko, ei potilas. **103 clinical officeria** satunnaistettiin: 52 interventiohaaraan ja 51 kontrolliin. Molemmat ryhmät käyttivät samaa pilvipohjaista sähköistä potilastietojärjestelmää. Interventoryhmällä oli lisäksi **AI Consult** (versio 2.0), päätöksentukityökalu joka perustui **OpenAI:n GPT-4o**-suureen kielimalliin ja oli upotettu potilastietojärjestelmään. Se luki klinikon kirjaamia tietoja ja saattoi huomauttaa mahdollisista diagnoosin tai hoitosuunnitelmaan liittyvistä ongelmista. Kliinikko säilytti täyden päätösvallan: ehdotuksen sai hyväksyä, muuttaa tai sivuuttaa.

Mikä malli oli käytössä ja miksi yksityiskohdilla on väliä

Työkalu oli **AI Consult 2.0**, joka käytti **OpenAI:n GPT-4o**-mallia (toukokuun 2025 julkaisu) OpenAI:n kaupallisen API:n kautta enterprise-lisenssillä ja pienellä satunnaisuudella (temperature 0.1). Se toimi räätälöidyn EasyClinic-potilastietojärjestelmän sisällä ja sitä ohjasivat järjestelmäkehotteet, jotka oli kirjoitettu Kenian kansallisten hoitosuosituksen mukaisiksi. Tutkijat julkaisivat koko ohjekehotteen.

Miksi tämä kerrotaan näin tarkasti? Koska tulos koskee *yhtä tiettyä järjestelmää* — yhtä malliversiota, yhtä kehotetta, yhtä potilastietojärjestelmää ja yhtä käyttöympäristöä — ei ”LLM:iä lääketieteessä” yleisesti. Tutkijat korostavat samaa: he kutsuvat havaintoa *ajalliseksi vertailukohdaksi, ei pysyväksi arvioksi kyvykkyydestä*. Uudempi malli, erilainen kehote tai vähemmän digitalisoitu klinikka voisi muuttaa tulosta.

Riippumattomuudesta: OpenAI antoi myöhemmin luontoissuorituksena tukea, kuten pilvilaskentakrediittejä ja teknistä API-ohjausta. Tutkijoiden mukaan päätös OpenAI:n käyttämisestä oli kuitenkin tehty *ennen* tätä tarjousta, eikä OpenAI osallistunut tutkimusasetelmaan, tiedonkeruuseen, analyysiin tai julkaisupäätökseen.

22. huhtikuuta – 16. heinäkuuta 2025 tutkimukseen osallistui **9,691 potilasta**. **Ensisijainen päätetapahtuma oli tarkoituksella potilaskeskeinen ja tiukka**: asiantuntijoiden arvioima *hoidon epäonnistumisen yhdistelmäpätetapahtuma* 14 päivän kuluessa käynnistä. Kliinikkopaneeli, joka ei tiennyt tutkimushaaraa, arvioi oliko potilaalla huono lopputulos, kuten jatkuva tai paheneva sairaus. Tutkimus rekisteröitiin etukäteen (Pan-African Clinical Trials Registry 202502499779176).

Juuri tämä suunnitteluratkaisu on olennaista. On helppo osoittaa, että tekoäly muuttaa sitä, mitä klinikko kirjoittaa. On paljon vaikeampaa — ja merkityksellisempää — osoittaa, että se muuttaa sitä, mitä potilaalle tapahtuu.

Mitä he löysivät

Ensisijainen päätetapahtuma ei parantunut. Hoito epäonnistui 102:lla 4,693 potilaasta (2.2%) tekoälyhaarassa ja 94:llä 4,654 potilaasta (2.0%) kontrollissa. Raakaprosentti oli tekoälyn kanssa hieman suurempi, mutta korjauksen jälkeen piste-estimaatti kallistui hyödyn puolelle: **korjattu odds ratio oli 0.77 (95%:n luottamusväli 0.55–1.08, P = 0.13)** — ei tilastollisesti merkitsevää. Raaka- ja korjattujen lukujen suunnan vaihtuminen ei ole laskuvirhe: korjaus huomioi klinikkoklustereiden välisiä eroja. Joka tapauksessa luottamusväli sisältää selvästi ”ei vaikutusta” -vaihtoehdon, joten hyötyä ei voida väittää, ja absoluuttinen vaikutus oli hyvin pieni.

Odds ratioiden, luottamusvälien ja P-arvojen yhteislukemisesta tavallisella kielellä voi lukea kliinisten tulosten oppaasta.

Ei turvallisuussignaalia — rajatusti. Mitään vakavaa haittatapahtumaa ei arvioitu työkalun aiheuttamaksi, eikä riippumaton tarkastus löytänyt turvallisuussignaalia. Tutkijat ovat rehellisiä siitä, kuinka pitkälle tämä rauhoittaa: tutkimuksessa ei ollut riittävästi voimaa harvinaisten vakavien haittojen havaitsemiseen eikä ennalta määriteltä noninferiority- tai muodollista turvallisuuskehystä. Siksi se ei voi *todistaa* turvallisuutta harvinaisissa tapahtumissa.

Dokumentointi parani. Sökkoutettujen asiantuntijoiden arvioimissa 2,000 potilaskäynnissä AI Consultia käyttäneiden klinikkojen kirjaukset olivat parempia kaikilla arvioituilla osa-alueilla: kirjattu diagnoosi, hoitosuunnitelma ja kokonaisuuden täydellisyys.

Lääkemääräykset muuttuivat tuskin lainkaan. Lääkemääräyksissä ei ollut merkitsevää eroa, mukaan lukien antibioottien oikea käyttö (korjattu odds ratio 0.86, 95%:n luottamusväli 0.48–1.55). Työkalu ei muuttanut antibioottien määräämisastetta.

Potilaat eivät huomanneet eroa. Tyytyväisyyskyselyn täyttäneiden 826 potilaan tyytyväisyys oli käytännössä sama molemmissa haaroissa, samoin vastaanottojen kesto.

Kustannukset osoittivat hieman alaspäin. Korjatussa analyysissä antibiootteihin liittyvät kustannukset olivat tekoälyhaarassa pienemmät — mahdollisesti halvempien valintojen, eivät harvempien reseptien vuoksi — ja potilaskohtainen antibioottisäästö näytti ylittävän työkalun potilaskohtaisen käyttökustannuksen. Tutkijat merkitsevät tämän vihjaavaksi, eivät ratkaistuksi: täydellinen kokonaisomistuskustannusten arviointi ei kuulunut tutkimukseen.

Tutkijoiden oma yhden lauseen yhteenveto on puhtain: LLM-avusteisuus oli näissä rajoissa turvallista, mutta ei vähentänyt hoidon epäonnistumista 14 päivässä, ja mahdollinen hyöty on todennäköisesti vaatimaton.

Miksi ”ei merkitsevää eroa” ei tarkoita ”se ei toimi”

Nollatuloksena ensisijaisessa päätetapahtumassa on helppo ylitulkita kumpaankin suuntaan. Kaksi asiaa estää liian yksinkertaisen tarinan.

Ensinnäkin **tutkimus oli mitoitettu löytämään suurempi vaikutus kuin mitä havaittiin**. Vakavat huonot tulokset ovat perusterveydenhuollossa harvinaisia — täällä noin 2% — joten pienen todellisen hyödyn erottaminen kohinasta vaatii valtavia osallistujamääriä. Tutkijoiden jälkikäteinen voimalaskelma arvioi, että havaitun kokoisen vaikutuksen osoittamiseen tarvittaisiin **paljon suurempi tutkimus, yli 100,000 potilaan luokkaa**. Ei-merkittävä tulos tämän kokoisessa tutkimuksessa ei sulje pois pientä todellista hyötyä; se tarkoittaa, ettei tutkimus pystynyt ratkaisemaan sitä.

Toiseksi **vertailu osittain sekoittui**. Konfiguraatiovirhe antoi lyhyeksi ajaksi joillekin kontrolliryhmän klinikoille pääsyn AI Consultiin, ja saman verkoston klinikot keskustelevat keskenään ja siirtävät opittuja toimintatapoja ryhmien rajojen yli. Molemmat ilmiöt tekevät ryhmistä keskenään samankaltaisempia ja painavat mahdollista todellista eroa kohti nollaa. Lisäksi verkosto toimi jo valmiiksi melko korkealla laatutasolla, joten parantamisen varaa oli vähemmän.

Mikään tästä ei pelasta ”läpimurto”-otsikkoa. Oikea tulkinta on kuitenkin kalibroitu eikä lannistava: vaikeimmalla ja rehellisimmällä päätetapahtumalla työkalu ei osoitettavasti auttanut potilaita kahdessa viikossa — samalla kun se paransi kirjaamista mitattavasti ja näytti turvalliselta.

Mitä tämä ei todista

- Se **ei** osoita, että tekoäly paransi potilastuloksia. Ensisijaisessa 14 päivän päätetapahtumassa ei ollut merkitsevää hyötyä.
- Se **ei** osoita, että tekoäly olisi hyödytön. Piste-estimaatti suosi sitä, dokumentointi parani kautta linjan ja lääkekustannukset osoittivat alaspäin; nollatulokset sopii myös pieneen todelliseen hyötyyn, jonka vahvistamiseen tutkimus oli liian pieni.
- Se **ei** todista työkalun turvallisuutta harvinaisten haittojen suhteen. Turvallisuussignaalia ei löytynyt, mutta tutkimusta ei mitoitettu tai suunniteltu harvinaisten vakavien tapahtumien turvallisuuden varmistamiseen.
- Se **ei** osoita, että ”tekoäly voittaa lääkärit” tai korvaa klinikot. Tämä on päätöksentuki; klinikolla säilyi täysi valta hyväksyä tai hylätä ehdotukset.
- Se **ei** yleisty automaattisesti. Tutkimus tehtiin yhdessä yksityisessä kaupunkiverkostossa Keniassa; maaseutu-, esikaupunki- ja korkeamman tulotason ympäristöt voivat erota kumpaan suuntaan tahansa.
- Se **ei** vahvista kustannussäästöjä. Kustannussignaali on vihjaava, ei valmis taloudellinen arvio.

Kuinka vahvaa näyttö on?

Keskeisen väitteen — ei osoitettua vähennystä lyhyen aikavälin potilashaitassa — osalta näyttö on tutkimusasetelmana vahvaa ja johtopäätöksenä asianmukaisen nöyrää. Prospektiivinen, ennakkoon rekisteröity, klusterisatunnaistettu tutkimus, jossa on sokkoutettu potilastason yhdistelmäpätetapahtuma, on lähellä parasta reaali maailman näyttöä, jonka tällaisesta työkalusta voi saada. Se kertoo paljon enemmän kuin vertailutestien pisteet tai tapauskuvaustutkimus.

Toissijaisten havaintojen — parempi dokumentointi, muuttumaton lääkkeenmääräys, samanlainen tyytyväisyys, pienemmät antibioottikustannukset — näyttö on hyvää mutta ne on luettava toissijaisina: tukevina signaaleina, ei pääotsikkona, ja samoille kontaminaatio- ja yhden verkoston rajoille alttiina.

Turvallisuuden näyttö rauhoittaa mutta on rajallinen: signaalia ei löytynyt, mutta tutkimusta ei rakennettu harvinaisten haittojen löytämiseen.

Hyödyllisin asenne ei ole ”se toimii” eikä ”se epäonnistui”. Se on: huolellinen reaali maailman tutkimus havaitsi, ettei tämä tekoälytyökalu tuottanut turvallisuussignaalia ja että se paransi hoitoprosessia, mutta potilashyötyä ei osoitettu kahdessa viikossa — ja mahdollisen hyödyn havaitseminen vaatisi paljon suuremman tutkimuksen.

Miksi tällä on merkitystä

Lääketieteellisen tekoälyn keskustelu kärsii juuri tällaisen näytön puutteesta. On tuhansia tutkimuksia, joissa mallit loistavat kokeissa ja saavuttavat kliinikoiden tason siisteissä tapauksissa. Suuria pragmaattisia satunnaistettuja tutkimuksia siitä, voivatko oikeat potilaat paremmin, on hyvin vähän. Tämä on yksi niistä, ja tulos on epäglamourinen totuus: kokeen läpäiseminen ei ole sama asia kuin potilaan auttaminen.

Juuri tämä kuilu on koko tarina. Työkalu voi aidosti auttaa kliinikkoa — selkeämmät kirjaukset, pienemmät lääkelaskut, toinen silmäpari — eikä silti liikuttaa kovaa potilaspäätetapahtumaa kahdessa viikossa. Molemmat voivat olla yhtä aikaa totta, ja kypsän terveydenhuoltojärjestelmän pitää pystyä pitämään ne yhdessä eikä valita vain mukavampaa.

Tutkimus myös siirtää todistustaakkaa hyödylliseen suuntaan. Jos yritys väittää kliinisen tekoälynsä parantavan hoitoa, olennainen näyttö ei ole leaderboard. Se on tällainen tutkimus, potilaalle tärkeillä päätetapahtumilla — ja mieluiten suurempi, koska tämän työn rehellinen opetus on, että vaatimattomien hyötyjen näkeminen vaatii suuria tutkimuksia.

Lyhyesti

Pragmaattisessa, klusterisatunnaistetussa tutkimuksessa 16 kenialaisessa perusterveydenhuollon yksikössä testattiin generatiivisen tekoälyn päätöksentekityökalua (”AI Consult”), joka lisättiin clinical officereiden käyttämään sähköiseen potilastietojärjestelmään. 9,691 potilaalla asiantuntijoiden arvioima hoidon epäonnistumisen yhdistelmäpäätetapahtuma 14 päivän kuluessa oli työkalulla 2.2% ja ilman sitä 2.0% (korjattu odds ratio 0.77, 95%:n luottamusväli 0.55–1.08, $P = 0.13$) — ei merkitsevää eroa. Työkalu ei tuottanut turvallisuussignaalia, paransi kliinistä dokumentointia kaikilla arvioituilla osa-alueilla, ei muuttanut lääkemääräyksiä, ei muuttanut potilastyytyväisyyttä ja liittyi hieman pienempiin antibioottikustannuksiin. Tutkijat päättelevät, että se oli näissä rajoissa turvallinen mutta ei vähentänyt hoidon epäonnistumista ja että mahdollinen hyöty on

todennäköisesti pieni; havaitun kokoisen vaikutuksen osoittamiseen tarvittaisiin paljon suurempi, noin 100,000 potilaan tutkimus. Tulos ei osoita kliinisen tekoälyn parantavan potilastuloksia eikä sen olevan hyödytön — se osoittaa, että turvalliselta näyttänyt ja prosessia auttanut työkalu ei yhdessä yksityisessä kaupunkiverkostossa osoitettavasti auttanut potilaita kahdessa viikossa.

Ei-hölynpölyä-tarkistus

Mitä tutkimus osoittaa: Reaalimaailman satunnaistetussa tutkimuksessa LLM-pohjaisen päätöksen-tukityökalun lisääminen perusterveydenhuollon potilastietoihin ei tuottanut turvallisuussignaalia, paransi kliinisen dokumentoinnin laatua, ei muuttanut lääkkeenmääräyksiä eikä vähentänyt merkittävästi tiukkaa 14 päivän hoidon epäonnistumisen yhdistelmäpäätetapahtumaa.

Mikä on uskottavaa mutta todistamatta: Että työkalu vähentää hoidon epäonnistumista hieman, mutta vaikutus on liian pieni tämän tutkimuksen havaittavaksi; että se säästää rahaa, kun kaikki kustannukset lasketaan; että parempi dokumentointi muuttuu ajan mittaan paremmaksi hoidoksi.

Mitä tutkimus ei osoita: Että kliininen tekoäly parantaa potilastuloksia; että se on turvallinen harvinaisten tapahtumien suhteen; että se korvaa tai voittaa klinikot; että tulokset siirtyvät maaseudulle tai korkeamman tulotason ympäristöihin; että kustannussäästö on osoitettu.

Tärkeimmät rajoitukset: Tutkimus oli mitoitettu suuremmalle vaikutukselle kuin havaittiin, ja harvinaiset päätetapahtumat tarvitsevat valtavan otoksen; konfiguraatiovirhe kontaminoi kontrollihaaraa ja saman verkoston yhteiset käytännöt sumentavat vertailua, molemmat kohti pienempää eroa; yksi yksityinen kaupunkiverkosto jo valmiiksi korkeilla standardeilla; ei ennalta määritettyä noninferiority- tai turvallisuuskehystä; lyhyt 14 päivän seuranta.

Kuinka paljon yleislukijan kannattaa luottaa tulokseen? Korkealla varmuudella työkalusta ei tässä tutkimuksessa löytynyt turvallisuussignaalia ja se paransi dokumentointia. Korkealla varmuudella se ei osoitettavasti parantanut 14 päivän potilastuloksia tässä asetelmassa. Vähäisellä varmuudella väitteeseen, että se ”toimii” tai ”epäonnistuu” potilashyödyn interventiona — kysymys on aidosti avoin ja vaatii paljon suuremman tutkimuksen. Oikea asenne: huolellinen reaalimaailman tulos työkalusta, joka ei tuottanut turvallisuussignaalia ja paransi hoitoprosessia, ei tuomio siitä että tekoäly mullistaa tai tuhoaa perusterveydenhuollon.

Lähteet

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>