



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

سوال کا جواب دینا اور پورا کلینیکل کیس چلانا ایک ہی بات نہیں

autonomous MIRA اے آئی ایجنٹ ہے جو EHR sandboxed میں تاریخ لیتا، microbiology/امیجنگ/labs ترتیب اور interpret کرتا، تشخیص تک پہنچتا اور علاج orders لکھتا ہے۔ عوامی MIMIC-IV کے 574 retrospective کیسز، آٹھ diagnoses، pre-selected پر اس نے 9% overall تشخیصی درستگی حاصل کی؛ 311-کیس براہ راست موازنہ میں senior، 8% بورڈ سے تصدیق شدہ physicians کے 1% اور mixed-seniority گروپ کے 1% کے مقابلے میں۔ guideline-concordant largely Decisions تھے اور مصنفین نے بلند medication severity غلطیاں رپورٹ نہیں کیں۔ مگر جائزہ سیمولیشن تھی، past ریکارڈز پر، متن-only؛ advantage کا بڑا حصہ واضح ٹیسٹ والی حالات میں تھا؛ MIRA نے ڈاکٹرز کے مقابلے میں far مزید analytes blood استعمال کیے (~28% vs 51%)؛ کئی علاج نتائج chart recorded کے خلاف اسکور ہوئے؛ اور عوامی MIMIC-IV تربیتی ڈیٹا contamination کا حقیقی خطرہ رکھتا ہے — جسے مصنفین ممکن بالائی bound کہتے ہیں۔ حقیقی advance کام-whole کا بہاؤ ایجنٹ ہے، نہ یہ ثبوت کہ اے آئی ڈاکٹرز سے بہتر diagnose کرتی ہے یا deployment clinic کے لیے ready ہے۔

Written by Lucio Vaglio · Cross-checked by Vera Rubin · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

Paper: Towards autonomous medical artificial intelligence agents, Dyke Ferber, Lars Hilgers, Christiane Höper and colleagues; senior author Jakob Nikolas Kather, Nature (2026) · DOI: 5-10675-026-s41586/1038.10

سوال کا جواب دینا اور پورا کلینیکل کیس چلانا ایک ہی بات نہیں

طبی اے آئی کی زیادہ تر کہانیوں ایسے ماڈل کے بارے میں ہوتی ہیں جو جواب دیتا ہے۔ آپ اسے vignette یا exam سوال دیتے ہیں، وہ تشخیص یا advice کا paragraph واپس دیتا ہے، اور اچھا اسکور کر لیتا ہے۔ مفید، مگر محدود: ایک ذہین مشیر جو chart پر خود کچھ نہیں کرتا۔

MIRA مختلف چیز کرنے کے لیے بنایا گیا ہے، اور اصل خبر یہی فرق ہے۔ ایک isolated سوال کا جواب دینے کے بجائے یہ پورا کیس چلاتا ہے: مریض ریکارڈ پڑھتا ہے، فیصلہ کرتا ہے کہ مزید کون سی تاریخ درکار ہے، لیبارٹری، امیجنگ اور microbiology ٹیسٹ ترتیب کرتا ہے، نتائج پڑھتا ہے، differential تشخیص محدود کرتا ہے، اور پھر اگلے مراتب لکھتا ہے — prescriptions، surgery admission، referral۔ actions الیکٹرانک صحت ریکارڈ کے اندر لیتا ہے، معالج کی طرح، نہ کہ صرف آزاد متن generate کر کے کسی انسانی سے نقل کرواتا ہے۔

یہ حقیقی قدم ہے: advice دینے والے نظام سے کام کا بہاؤ میں actions لینے والے نظام تک۔ اور یہی وہ قدم ہے جو coverage میں جلد "اے آئی ڈاکٹرز سے بہتر" بن جاتا ہے۔ اس لیے ضروری ہے کہ کیا ٹیسٹ ہوا، اور کہاں، اسے درست رکھا جائے۔

ایک sentence میں: MIRA نے پورے کیسز شروع سے آخر تک خود چلائے اور اس بینچ مارک پر تشخیصی درستگی میں معالجین سے بہتر اسکور کیا — مگر یہ سب sandbox میں، retrospective ریکارڈز پر، پہلے سے چنی گئی آٹھ تشخیصات کے اندر، صرف متن کے ذریعے ہوا، اور اس کی بڑی برتری ان حالات میں آئی جن کے ٹیسٹ زیادہ صاف اور فیصلہ کن تھے۔ پیش رفت حقیقی ہے۔ Scoreboard کلینک نہیں۔

MIRA showed a workflow capability, not a clinic verdict

A sandboxed EHR lets an agent act through a whole retrospective case. It is still not a real patient trial.

DEMONSTRATED

- end-to-end case work in a sandboxed EHR
- history, tests, differential diagnosis, orders
- 574 retrospective MIMIC-IV cases
- 8 pre-selected diagnoses
- higher diagnostic accuracy on this benchmark

NOT DEMONSTRATED

- real patients or live hospital deployment
- conditions beyond the eight-target set
- an equally informed head-to-head comparison
- freedom from public-data contamination
- safety and governance for clinical use

A capability shown in a sandbox -- not a diagnostician proven in a clinic.

The figure separates the workflow result from the deployment claim.

MIRA نے EHR sandboxed کے اندر شروع سے آخر تک کام کا بہاؤ صلاحیت دکھائی۔ اس نے حقیقی دنیا کلینیکل عملی تعیناتی، وسیع تشخیصی، coverage یا ڈاکٹرز کے ساتھ برابر معلومات پر fair براہ راست مقابلہ ثابت نہیں کیا۔

مصنفین نے کیا کیا

MIRA — نام کا مطلب “طبی ذہانت برائے استدلال و عمل” ہے — OpenAI کے ماڈلز پر بنایا گیا خودکار ایجنٹ ہے۔ گفتگو اور عملی اقدامات کے لیے GPT-4o استعمال ہوتا ہے، جبکہ الگ منصوبہ بندی کے مرحلے میں OpenAI کا o1 استدلالی ماڈل کام کرتا ہے۔ یہ ماڈلز ایک مخصوص، طبی معیاروں کے مطابق فریم ورک کے اندر رکھے گئے ہیں جو FHIR HL7 پر مبنی ہے، یعنی اسی معیار پر جس سے حقیقی ہسپتال طبی ریکارڈ کا تبادلہ کرتے ہیں، اور اس کے پاس 80,000 سے زیادہ ممکنہ طبی اقدامات کا مجموعہ ہے۔ محفوظ تجرباتی ماحول میں MIRA مریض کی تاریخ نکال سکتا ہے، لیبارٹری جانچ، امیجنگ اور خرد حیاتیاتی ٹیسٹ طلب اور ان کی تشریح کر سکتا ہے، متبادل تشخیصات بنا سکتا ہے اور علاج کا منصوبہ تیار کر سکتا ہے — مثلاً دوائیں تجویز کرنا، جراحی طریقہ کار طے کرنا یا داخلے کی منصوبہ بندی کرنا۔ ایک تفصیل شروع ہی میں اہم ہے: MIRA کے جوابوں کو نمبر دینے والے خودکار “جج” بھی GPT-4o تھے، یعنی آزمائے گئے ماڈل ہی کے خاندان سے؛ ہم مرتبہ جائزہ لینے والے نے اس پر تشویش ظاہر کی، جس کے بعد مصنفین نے بورڈ سے تصدیق شدہ ڈاکٹروں کا جائزہ بھی شامل کیا۔

ٹیسٹ سیٹ MIMIC-IV سے آیا، ایک بڑا عوامی، شناخت-de شدہ EHR۔ Deaconess Israel Beth database۔ طبی مرکز میں 2008 سے 2019 کے دوران treated تقریباً 300,000 مریض میں سے مصنفین نے 574 مریض کیسز منتخب کیے، جو آٹھ ہدف تشخیصات پر مشتمل تھے — pathologies abdominal اور داخلی-طب، emergencies، مثلاً pneumonia pancreatitis، appendicitis، urinary-tract انفیکشن۔ ہر کیس میں MIRA initial limited تصویر سے شروع کرتا اور مرحلہ وار طے کرتا کہ اگلا عمل کیا ہو — حقیقی تشخیصی عمل جیسی شکل، مگر ماضی کے ریکارڈ کے خلاف، جس کا اصل ending پہلے سے موجود ہے۔

پھر اسی کیس سیٹ پر معالجین کے ساتھ کارکردگی موازنہ کیا گیا۔

“Autonomous” sandboxed a in agent EHR کا واقعی مطلب کیا ہے؟

ہیں۔ رہے اٹھا کچھ بہت الفاظ تین یہ

ایجنٹ کا مطلب chatbot نہیں جو ایک ہدایت کا جواب دے۔ نظام لوپ چلاتا ہے: موجودہ حالت دیکھو، عمل چنو (یہ ٹیسٹ ترتیب کرو، وہ تاریخ پوچھو)، نتیجہ دیکھو، اگلا عمل چنو — تشخیص اور plan تک۔ “ذہانت” زبان ماڈل ہے، agency وہ scaffolding ہے جو ماڈل کے متن کو EHR permitted operations میں تبدیل کرتی ہے اور نتائج واپس خوراک لینا کرتی ہے۔

EHR Sandboxed کا مطلب سیمولیشن شدہ، walled-off طبی-ریکارڈ نظام ہے، hospital live نہیں۔ MIRA ٹیسٹ “ترتیب” کر سکتا ہے اور مریض کا وہی نتیجہ حاصل کر سکتا ہے جو حقیقی زندگی میں موجود تھا، کیونکہ کیس ماضی کے ہے۔ اس کا کوئی عمل حقیقی مریض یا کلینک کو نہیں چھوتا۔

Autonomous کا مطلب ہے کہ انسانی-in-the-loop کے بغیر کیس شروع سے آخر تک مکمل کرتا ہے — sandbox کے اندر۔ اس کا مطلب یہ نہیں کہ اسے حقیقی مریض کو manage unsupervised کرنے کے لیے چھوڑ دیا گیا تھا۔ صرف پہلی چیز ٹیسٹ ہوئی۔

Sandbox کے بارے میں ایک اور اہم تفصیل: MIRA کوئی بھی ٹیسٹ ترتیب کر سکتا ہے، مگر نظام نتیجہ صرف اسی ٹیسٹ کا دے سکتا ہے جو اس مریض کے حقیقی ریکارڈ میں واقعی ہوا تھا۔ وہ blood پینل مریض نے کروایا تھا تو حقیقی قدریں ملیں گی، ایسا اسکین ترتیب کیا جو کبھی ہوا ہی نہیں، تو “N/A” — be could نہیں کارکردگی دکھائی “ملے گا، fabricated عدد نہیں۔ تاریخ بھی اسی طرح: ریکارڈ میں جواب نہ ہو تو سیمولیشن شدہ مریض کہتا ہے معلوم نہیں۔ یعنی MIRA وہ فیصلہ کن ٹیسٹ جادو سے summon نہیں کر سکتا جو حقیقی تشخیصی عمل میں کبھی لیا ہی نہیں گیا۔

سرخ لفظ “autonomous” سب سے زیادہ اسی لیے غلط پڑھا جا سکتا ہے۔

کیا ملا

بینچ مارک پر MIRA تشخیصی درستگی میں معالجین سے آگے تھا - مگر سرخی کے پیچھے تین الگ اعداد ہیں جنہیں مرکب نہیں کرنا چاہیے۔

تمام 574 کیسز پر MIRA نے درست discharge تشخیص 98.88% وقت شناخت کیا۔ براہِ راست انسانی براہِ راست مقابلہ میں ڈاکٹرز نے تمام 574 کیسز نہیں کیے؛ انہوں نے مشترک 311-کیس subset پر کام کیا، وہی ریکارڈز/ٹولز استعمال کرتے ہوئے جو MIRA کے پاس تھے۔ ان 311 پر MIRA 88.7% تھا۔ انسانی موازنہ کے دو گروہس تھے۔ پہلا senior گروپ: چار بورڈ سے تصدیق شدہ معالجین، 7-11 سال experience، اسکور 1.78%۔ دوسرا ملے جلے seniority گروپ: زیادہ تر residents، junior ساتھ دو specialists، اسکور 1.71%۔ ایک ہی کیسز، ایک ہی ٹولز: اے آئی پہلا، experienced ڈاکٹرز دوسرا، بھاری junior گروپ تیسرا۔ مقالے کی محتاط phrasing ہے کہ MIRA "consistently" often and to, equivalent "exceeded" معالجین میں تمام اٹھ diseases۔

بعد کے فیصلے بھی زیادہ تر رہنما اصولوں کے مطابق، دواؤں کے لحاظ سے محفوظ اور داخلے کے فیصلوں میں مناسب تھے۔ مصنفین نے پانچ حفاظتی زمروں – دواؤں کے باہمی اثرات، گردوں کے مطابق خوراک، الرجی، QT خطرہ اور opioid تجویز – میں شدید نوعیت کی دوائی کی غلطی رپورٹ نہیں کی، اور 468 کیسز میں 467 نسخوں کو درست قرار دیا، ساتھ ہی یہ بھی لکھا کہ نظام نے “100% قابلِ اعتماد کارکردگی” حاصل نہیں کی۔ پہلی نظر میں یہ متاثر کن ہے: ایجنٹ پورا کیس چلاتا ہے اور تشخیصی پینچ مارک میں ڈاکٹروں سے آگے نکلتا ہے۔

مگر win کی شکل خود win جتنی اہم ہے۔ آزاد specialists نے نوٹ کیا کہ برتری کہاں سے آرہی تھی۔ سائنس میڈیا مرکز ماہر پینل پر (University Xing Wei Dr) of Sheffield نے کہا کہ سرخی فائدہ زیادہ تر ایسی حالات سے ڈرائیو ہوئی جن میں واضح ٹیسٹ نتائج تھے، جیسے appendicitis اور pancreatitis جہاں فیصلہ کن اسکین یا لیب قدر سوال settle کر دیتی ہے۔ Pneumonia اور urinary-tract infections — departments emergency میں بہت عام — میں آئی اور ڈاکٹرز دونوں نے worst کارکردگی دی اور خلا سب سے چھوٹا تھا۔

دونوں sides نے برابر معلومات استعمال بھی نہیں کی۔ مقالے کے اپنے اعداد کے مطابق MIRA معمول نگہداشت میں دستیاب لیبارٹری analytes کے تقریباً 51% کو استعمال کرتا تھا، جبکہ بورڈ سے تصدیق شدہ معالجین تقریباً 28% – median طور پر کیس کے لیے سات اضافی blood پیرامیٹرز۔ مصنفین اسے ڈیٹا سیٹ کے معمول-نگہداشت بنیادی موازنہ سے بھی کم کہتے ہیں، نہ کہ ”ترتیب everything حکمت عملی، اور زیادہ-لاگت cross-sectional امیجنگ میں منظم اضافہ نہیں ملا۔“ پھر بھی زیادہ معلومات خود تشخیصی درستگی بڑھا سکتی ہے۔ اس لیے یہ مکمل contest like-for-like نہیں۔ Cheap blood ٹیسٹ کی لاگت/delay/discomfort کم سمجھنے والا معالج بھی مقالہ پر زیادہ درست دکھ سکتا ہے۔

”ڈاکٹرز کو سے بہتر رہا کیا“ کا مطلب ”بہتر ڈاکٹر“ نہیں

بینچ مارک میں برتری کو ضرورت سے زیادہ معنی دینا آسان ہے۔ MIRA نے زیادہ اسکور کیا اور MIRA بہتر ڈاکٹر ہے کے درمیان کم از کم تین الگ فاصلے ہیں۔

پہلا سوال یہ ہے کہ نظام کو کس چیز کے مقابل اسکور کیا گیا؟ Edinburgh of University کی Jacko Julie Professor نے نشاندہی کی کہ MIRA کے کئی اہم نتائج بنیادی ڈیٹا سیٹ میں درج طبی رویے کے نسبت متعین کیے گئے ہیں - یعنی نظام کو لازماً بہترین نگہداشت ثابت کرنے پر نہیں، بلکہ تاریخی ریکارڈ میں درج طبی فیصلوں کو دوبارہ پیدا کرنے پر انعام مل سکتا ہے۔ یوں تاریخی چارٹ ہی جواب نامہ بن جاتا ہے۔ بینچ مارک بنانے کا یہ معقول طریقہ ہے، مگر یہ اس بات سے مطابقت ناپتا ہے کہ

حقیقت میں کیا کیا گیا تھا؟ یہ مطلق معنوں میں بہترین نگہداشت کی درستگی ثابت نہیں کرتا۔ انصافاً، یہ اعتراض سب سے زیادہ علاج سے مطابقت والے پیمانوں پر لاگو ہوتا ہے۔ کیا MIRA کے احکامات چارٹ سے ملتے تھے؟ — جبکہ نمایاں تشخیصی درستگی کو ڈیپجارج تشخیص کے مقابل ناپا گیا، جو محض دستاویز کی بازگشت کے مقابل ایک زیادہ حقیقی نتیجہ ہے۔

دوسرا، معلومات عدم توازن: دستیاب blood ٹیسٹ کا 51% بمقابلہ 28% شواہد کی مختلف مقدار ہے۔ درستگی خلا کا کچھ حصہ بہتر استدلال کے بجائے زیادہ شواہد خریدنے سے آسکا ہے۔

تیسرا، کہاں چلایا گیا۔ یہ retrospective sandbox، کیسز، pre-selected آٹھ تشخیصات، اور متن-صرف ترتیب تھی۔ Dominic Dr (University Oliver) of (Oxford) نے یاد دلایا کہ حقیقی کلینیکل جائزہ صرف مریض کیا کہتا ہے پر نہیں؛ کیسے کہتا ہے، طبیعی exam، مشاہدہ شدہ رویہ اور جسم زبان بھی اہم ہیں — جو finished متن ریکارڈ پڑھنے والے ایجنٹ کے پاس نہیں۔ اور اگر مسئلہ ان آٹھ تشخیصات میں مطابقت نہ ہو تو مطالعہ اس کے بارے میں کچھ نہیں کہتی۔

ایک اور quiet تشویش ڈیٹا آلودگی ہے۔ MIMIC-IV عوامی ہے اور discussed heavily ہے، اس لیے تربیت-internet یافتہ زبان ماڈل شاید مقالات، کیس discussions، یا ڈیٹا itsefl دیکھ چکا ہو۔ (University Unni Parakkal Midhun Dr) of (Sheffield) نے یہ خطرہ براہ راست نشان کیا۔ اگر جوابات تربیت ڈیٹا میں تھے تو کارکردگی کا کچھ حصہ یادداشت ہو سکتا ہے، کلینیکل استدلال نہیں۔ آزاد replication ہی فرق واضح کرے گی۔ اہم بات یہ ہے کہ مصنفین خود تشویش dismiss نہیں کرتے: وہ نتائج کو “ممکن بالائی حد” کہتے ہیں اور لکھتے ہیں کہ عوامی کیسز پر عمومی نتیجہ کو overestimate کر سکتے ہیں۔

یہ سب MIRA کو غیر دلچسپ نہیں بناتا۔ یہ صحیح تشریح calibrate کرتا ہے: retrospective بینچ مارک پر مکمل ریکارڈ ایجنٹ نے صاف ٹیسٹ والی تشخیصات میں ڈاکٹرز سے بہتر اسکور کیا — جزوی طور پر زیادہ ٹیسٹ ترتیب کر کے، جزوی طور پر ماضی کے chart کے ساتھ مطابقت میں۔ یہ حقیقی صلاحیت مظاہرہ ہے، فیصلہ نہیں کہ machines ڈاکٹرز سے بہتر diagnose کرتی ہیں۔

یہ مطالعہ کیا ثابت نہیں کرتی

- یہ نہیں دکھاتی کہ MIRA حقیقی مریض پر کام کرتا ہے۔ ہر کیس retrospective اور سیمولیشن شدہ تھا؛ کسی مریض کو MIRA نے واقعی manage نہیں کیا۔
- یہ آٹھ تشخیصات سے باہر کارکردگی نہیں دکھاتی۔ طب کی غیر واضح، majority undifferentiated ٹیسٹ نہیں ہوئی۔
- یہ عملی تعیناتی حفاظت ثابت نہیں کرتی۔ مصنفین خود لکھتے ہیں کہ عمومی نتیجہ، حفاظت اور حکمرانی کے لیے prospective حقیقی دنیا مطالعات درکار ہیں۔
- یہ ڈاکٹرز کے ساتھ fair perfectly براہ راست مقابلہ نہیں۔ MIRA نے far مزید blood ٹیسٹ استعمال کیے (~51% دستیاب analytes بمقابلہ ~28%)، اور کئی علاج نتائج ماضی کے chart کے خلاف اسکور ہوئے۔
- یہ نتیجہ حفظ سے پاک ہے، یہ نہیں دکھاتی۔ عوامی MIMIC-IV ماڈل تربیت کے ساتھ overlap کر سکتا ہے؛ آزاد replication چاہیے۔
- یہ “اے آئی ڈاکٹرز کو جگہ لینا کرتی ہے” نہیں دکھاتی۔ یہ ریکارڈ میں actions لینے والا فیصلہ-تائید ایجنٹ ہے؛ حقیقی استعمال میں معالجین اختیار اور وہ سیاق دیتے رہیں گے جو متن ریکارڈ میں نہیں۔

شواہد کتنے مضبوط ہیں؟

دعویٰ کو دو حصوں میں تقسیم کرنا بہتر ہے۔

صلاحیت مظاہرہ کے طور پر — زبان-ماڈل ایجنٹ EHR sandboxed میں پورا کلینیکل کیس تاریخ، ٹیسٹ، تشخیص اور مراتب سمیت شروع سے آخر تک چلا سکا ہے — کام novel genuinely اور معقول حد تک convincing ہے۔ یہی نئی چیز ہے اور یہی توجہ کے قابل ہے۔

برتری کے دعوے کے طور پر — کہ MIRA معالجین سے بہتر ہے — شواہد محدود ہیں۔ یہ نتیجہ ایک مخصوص retrospective بینچ مارک، آٹھ تشخیصات، معلومات کے عدم توازن، ماضی کے چارٹ پر مبنی اسکورنگ اور ممکنہ تربیتی ڈیٹا آلودگی کے اندر قائم رہتا ہے۔ اتنا شواہد یہ کہنے کے لیے کافی ہے کہ ”اس بینچ مارک پر ایجنٹ کی کارکردگی متاثر کن تھی“، مگر یہ ثابت کرنے کے لیے نہیں کہ وہ حقیقی طبی عمل میں بہتر ڈاکٹر ہے۔

مفید موقف نہ ”اے آئی سے بہتر ہے ڈاکٹرز“ ہے، نہ ”toy“ ماڈل۔ ”بہتر مطالعہ: طبی اے آئی کی نئی قسم — جو isolated سوالات کے جواب کے بجائے ریکارڈ کے اندر actions لیتی ہے — مشکل retrospective بینچ مارک پر اچھی رہی، اور اب اسے وہاں ثابت کرنا ہے جہاں ابھی ٹیسٹ نہیں ہوئی: حقیقی، undifferentiated مریض، prospectively حکمرانی کے تحت۔

یہ کیوں اہم ہے

کچھ سالوں میں طبی اے آئی کا دلچسپ سوال quietly بدل گیا ہے۔ پہلے سوال تھا ”کیا ماڈل تشخیص ٹھیک کر سکا ہے؟“۔ اب زیادہ مشکل سوال ہے: ”کیا ماڈل پورا کام کر سکا ہے — gather، ترتیب، interpret، act decide — پورے کام کا بہاؤ میں؟“ یہی زیادہ مفید اور دیانت دار سوال ہے، کیونکہ difficulty اور خطرہ دونوں عمل میں رہتے ہیں۔

اسی لیے framing اہم ہے۔ سرخی reflex — اے آئی outperforms ڈاکٹرز — نتیجہ کے سب سے کم novel اور سب سے کم supported حصے کی طرف اشارہ کرتا ہے۔ نئی اور consequential چیز یہ ہے کہ ایجنٹ پورے ریکارڈ میں حرکت کر سکا ہے۔ اگر صلاحیت برقرار رہتی ہے تو یہ پہلے کام کا بہاؤ بدلتی ہے، اختیار نہیں۔ ڈاکٹر فوراً جگہ لینا نہیں ہوتا، ordering نتائج chase کرنا، تشریح اور دستاویزات جیسی کام کا بہاؤ کام پہلے بدلتی ہیں۔

اور burden of ثبوت صحیح سمت میں reset ہوتا ہے۔ Retrospective win leaderboard prospective ٹرائل چلانے کی وجہ ہے، اس کا substitute نہیں۔ مصنفین بھی یہی کہتے ہیں۔ MIRA کی دیانت دار مطالعہ اگلے مطالعے کی دعوت ہے — مریض پر ناپا گیا، معیارات پر نہیں۔

صاف خلاصہ

MIRA ایک خودکار اے آئی ایجنٹ ہے جو محفوظ EHR ماحول میں مریض کی تاریخ لیتا، لیبارٹری جانچ، امیجنگ اور خرد حیاتیاتی ٹیسٹ طلب اور ان کی تشریح کرتا، تشخیص تک پہنچتا اور علاج کے احکامات لکھتا ہے۔ عوامی MIMIC-IV کے 574 سابقہ کیسز اور آٹھ پہلے سے منتخب تشخیصات پر اس کی مجموعی تشخیصی درستگی 9% 88 تھی؛ 311 کیسز کے براہ راست موازنے میں 8% 87، جبکہ سینئر بورڈ سے تصدیق شدہ ڈاکٹروں کی 1% 78 اور مختلف سیناریو والے گروہ کی 1% 71۔ فیصلے زیادہ تر رہنما اصولوں کے مطابق تھے اور مصنفین نے شدید نوعیت کی دوائی کی غلطیاں رپورٹ نہیں کیں۔ مگر یہ حقیقی کلینک نہیں بلکہ سابقہ ریکارڈز پر، صرف متن کے ساتھ چلائی گئی سیمولیشن تھی؛ برتری کا بڑا حصہ ان حالتوں میں تھا جہاں واضح ٹیسٹ دستیاب تھے؛ MIRA نے ڈاکٹروں سے کہیں زیادہ خون کے تجزیے استعمال کیے (تقریباً 51% بمقابلہ 28%)؛ بعض علاجی فیصلوں کو تاریخی چارٹ ہی کے مطابق نمبر ملے؛ اور چونکہ MIMIC-IV عوامی ڈیٹا ہے، تربیتی آلودگی کا حقیقی خطرہ موجود ہے — جسے مصنفین ممکنہ بالائی حد کہتے ہیں۔ اصل پیش رفت مکمل کام کے بہاؤ والا ایجنٹ ہے، نہ یہ ثبوت کہ اے آئی ڈاکٹر سے بہتر تشخیص کرتی ہے یا کلینک

میں تعیناتی کے لیے تیار ہے۔

بلا مبالغہ جانچ

مقالے کا دکھانا ہے: MIRA زبان-ماڈل ایجنٹ EHR sandboxed میں entire کیس شروع سے آخر تک چلا سکا ہے — تاریخ، ٹیسٹ، تشخیص، مراتب — اور 574-کیس retrospective بینچ مارک میں بلند تشخیصی درستگی دکھاتا ہے (98.8% مجموعی؛ 311-کیس subset میں 8.87% بمقابلہ بورڈ سے تصدیق شدہ معالجین 1.78%)، جبکہ بلند-medication severity غلطیاں رپورٹ نہیں ہوئے۔

کیا قرین قیاس ہے مگر ثابت نہیں: یہ صلاحیت حقیقی undifferentiated مریض کے لیے فائدہ میں translate ہوگی؛ درستگی برتری بہتر استدلال سے ہے، نہ کہ زیادہ ٹیسٹ، ماضی کے chart مطابقت یا memorized عوامی ڈیٹا سے۔

یہ کیا نہیں دکھاتا: آٹھ تشخیصات سے باہر کارکردگی، محفوظ عملی تعیناتی، informed equally ڈاکٹرز کے خلاف fair برتری؛ معالجین کی متبادل؛ تربیتی ڈیٹا آلودگی سے پاک نتائج۔

اہم حدود: retrospective سمولیشن شدہ متن-صرف جائزہ؛ آٹھ pre-selected تشخیصات؛ معلومات عدم توازن (دستیاب analytes کا 51% بمقابلہ 28%)؛ علاج نتائج جزوی طور پر ماضی کے chart کے خلاف؛ عوامی MIMIC-IV جسے ماڈل تربیت میں دیکھ چکا ہو سکا ہے۔

عام قاری کو کتنا اعتماد ہونا چاہیے؟ زیادہ اعتماد کہ MIRA novel صلاحیت مظاہرہ ہے — chatbot سے آگے، ریکارڈ کے اندر actions لینے والا ایجنٹ۔ کم اعتماد کہ وہ معالج سے بہتر تشخیص کار ثابت ہو چکا ہے: وہ دعویٰ ایک retrospective بینچ مارک، مزید-ٹیسٹ عدم توازن، پر-chart مبنی اسکورنگ اور ممکن آلودگی سے حد میں ہے۔ مصنفین کی اپنی bottom لائن درست ہے: مریض نگہداشت کے لیے معنی نکالنے سے پہلے prospective حقیقی دنیا مطالعہ ضروری ہے۔

ماخذ

(2026) Nature agents, intelligence artificial medical autonomous Towards al., et Ferber
<https://doi.org/10.1038/s41586-026-10675-5>

abstract peer-reviewed — 42310457 PubMed

<https://pubmed.ncbi.nlm.nih.gov/42310457/>

(2026) AMIE and MIRA to reaction expert — Centre Media Science

<https://www.sciencemediacentre.org/expert-reaction-to-presentation-of-two-new-medical-ai-r>

framing doctors' 'outperforms the illustrates — (2026) News-Medical

<https://www.news-medical.net/news/20260621/Autonomous-medical-AI-outperforms-doctors-in-si.aspx>