



The CLEAN PAPER

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

একটি AI agent পুরনো patient case নজিচে চালিয়েছে — simulator-এ, পুরনো record-এর ওপর

Autonomous medical-AI agent historical case-based simulator-এ পুরনো clinical workflow চালিয়েছে — test order, result interpretation, iterative diagnosis ও management। 574 case overall benchmark এবং 311 case human comparison-এ strong performance দেখোলেও agent বেশি test order করছে, environment sandbox ছিল এবং MIMIC-IV contamination/retrospective limitation আছে। Result long-horizon medical reasoning capability দেখায়; real patient care-এর autonomy বা safety নয়।

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: 10.1038/s41586-026-10675-5

একটি প্রশ্নের উত্তর দেওয়া আর পুরো ক্লিনিকিয়াল ক্ষেত্র চালানো এক জিনিস নয়

চিকিৎসাবৈজ্ঞানিক AI গল্পের বেশির ভাগ মডেলে **উত্তর দিয়ে**: সংক্ষিপ্ত ক্লিনিকিয়াল বর্ণনা/পরীক্ষা প্রশ্ন দিলে রোগ নির্ণয় বা পরামর্শ অনুচ্ছেদ দেয়। উপযোগী, কনিত্ত সংকীরণ—চার্টে কোনো ক্রিয়া নেই না।

MIRA (Medical Intelligence for Reasoning and Action)—চিকিৎসাবৈজ্ঞানিক যুক্তিও সন্ধিধান্তের জন্য তৈরি একটি AI ব্যবস্থা অন্য কিছু পরীক্ষা করে। এটি একটি ক্লিনিকিয়াল কেসে শুরু থেকে শেষ পর্যন্ত পর্যালোচনা করে: রোগী নথি পড়ে, আর কী ইতিহাস দরকার ঠিক করে, পরীক্ষাগার/ইমেজিং/মাইক্রোবায়োলজি নির্দেশ দেয়, ফল পড়ে, পার্থক্যমূলক রোগ নির্ণয় সংকীরণ করে এবং প্রসেসক্রিপশন, ভর্তি বা অস্ত্রোপচার রফোরলে মতো নির্দেশ লিখে। মুক্ত পাঠ্য মানবকে অনুলপি করত দেওয়ার বদলে স্যান্ডবক্সে সীমাবদ্ধ ইলেকট্রনিক স্বাস্থ্য নথি ভেতর চিকিৎসকের মতো ধাপে ধাপে কাজ করে।

এটি সত্যিই সক্ষমতা ধাপ: পরামর্শদাতা ব্যবস্থা থেকে কর্মপ্রবাহে সক্রিয় অংশগ্রহণকারী ব্যবস্থা। কনিত্ত “AI চিকিৎসকদের ছাড়িয়ে যায়” শরিরে নামে ফল flatten হয়। এক বাক্যে বললে: MIRA পশ্চাদৃষ্টমূলক স্যান্ডবক্স মানদণ্ডে পুরো ক্ষেত্রের স্বয়ংক্রিয়ভাবে চালিয়ে রোগ নির্ণয়ের নিরীহুলতায় চিকিৎসক দলের ওপরে স্কোর করেছে—কনিত্ত **বাস্তব রোগী নয়, আর্টটি আগ থেকে নির্বাচিত রোগ নির্ণয়, শুধু পাঠ্যভিত্তিক নথি**, এবং সুবিধার বড় অংশ স্পষ্ট পরীক্ষা-ফল শর্তে; আরও রক্ত পরীক্ষা ব্যবহার করেছে। অগ্রগতি বাস্তব; স্কোরবোর্ড ক্লিনিক নয়।

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 স্যান্ডবক্সে সীমাবদ্ধ EHR-এ শুরু থেকে শেষ পর্যন্ত কর্মপ্রবাহ সক্ষমতা দেখায়। এটি বাস্তব জগতের ব্যবহারিক মতোয়ানে, বসিত্ত রোগ নির্ণয়মূলক আবরণ বা সমান তথ্যপ্রাপ্ত ন্যায্য চিকিৎসক head-to-head দেখানো করণে।

Original Aurora diagram — The Clean Paper CC BY 4.0

গবেষকরো কী করছেন

MIRA-র আলাপ ও ক্রিয়া অংশে **GPT-4o**, আর আলাদা পরিকল্পনা ধাপে OpenAI o1 reasoning model ব্যবহার করা হয়েছে। এগুলো HL7 FHIR-সম্মত বিশেষভাবে তৈরিকার্টামের মধ্যকার কাজ করে, যেখানে **80,000-এর বেশি** কলনিকিয়াল ক্রিয়া রয়েছে। স্যান্ডবক্সে এজেন্ট রোগীর ইতিহাস সংগ্রহ করতে, পরীক্ষা নির্দেশে ও ব্যাখ্যা করতে, পার্থক্যমূলক রোগনির্ণয় তৈরিকরতে, চিকিৎসা পরিকল্পনা ও প্রসেসকরপিশন দিতে, এমনকি অস্ত্রের পচারণের সময়সূচি বা ভরতি পরিকল্পনাও করতে পারে। স্বয়ংক্রিয় স্কোরিংয়েও GPT-4o পরিবারের মডলে ছিল—একই মডলে পরিবার ব্যবহারের পক্ষপাতের আশঙ্কা সহকর্মী-পর্যায়ে চিনায় ওঠার পর লেখকরো বোর্ড-সনদপ্রাপ্ত চিকিৎসকদের পর্যায়ে চিনা দিয়ে তা যাচাই করেন।

পরীক্ষা-সটে **অনুকরণ করা-IV** জনসাধারণের পরিচয়মুক্ত EHR ডটোবসে। Beth Israel Deaconess চিকিৎসাবিশেষক কেন্দ্রে 2008–2019-এর প্রায় **300,000** রোগী থেকে **574** ক্ষেত্র, **8** লক্ষ্য রোগনির্ণয়—অ্যাপেন্ডিসাইটিস, প্যানক্রিয়াটাইটিস, নডিমেট্রিয়া, UTI-সহ abdominal/অভ্যন্তরীণ-চিকিৎসাবজ্রিএগন emergency—বহুে নেওয়া হয়। MIRA সীমিত প্রারম্ভিক চিত্র থেকে ধাপ-by-ধাপ পরবর্তী ক্রিয়া সিদ্ধান্ত নেয়; ঐতিহাসিক নথিতে বাস্তব ফলাফল আগ থেকেই আছে।

চিকিৎসকদের একই ক্ষেত্র/সরঞ্জাম পরিবেশে তুলনা করা হয়, কিন্তু সব 574-তে head-to-head নয়।

“স্বায়ত্তশাসিত এজেন্ট in a স্যান্ডবক্সে সীমাবদ্ধ EHR” মানের কী

এজেন্ট: এটি একবারের prompt-এর উত্তর নয়। ব্যবস্থা বর্তমান অবস্থা দখে একটি পদক্ষেপে বহুে নেয়, ফল পায়, তারপর সেই ফলে ভিত্তিতে পরবর্তী পদক্ষেপে ঠিক করে—এভাবে রোগনির্ণয় ও পরিকল্পনা পর্যন্ত এগোয়। ভাষা মডেলটি যুক্তি ও ভাষাগত সক্ষমতা দেয়; চারপাশের সফটওয়্যার কার্টামে। সেই সিদ্ধান্তকে অনুমোদিত EHR কাজের মধ্যে সীমাবদ্ধ রাখে।

স্যান্ডবক্সে সীমাবদ্ধ EHR: এটি চলমান হাসপাতাল ব্যবস্থা নয়, বরং ঐতিহাসিক নথি দিয়ে তৈরি বন্ধ simulation। MIRA কানে। পরীক্ষা “নির্দেশে” দিলে ফল কেবল তখনই ফরে, যদি বাস্তব রোগীর ইতিহাসে সেই পরীক্ষা সত্যি করা হয়ে থাকে। না হলে উত্তর আসে “N/A—could not be performed”; কানে। মান বানিয়ে দেওয়া হয় না। নথিতে তথ্য না থাকলে simulated patient-ও সটে “জানেনা”।

স্বায়ত্তশাসিত: মানব-in-লুপ ছাড়া স্যান্ডবক্স ক্ষেত্রের শুরু থেকে শেষ পর্যন্ত সম্পূর্ণ। বাস্তব রোগী unsupervised manage করা নয়। শরিরে নাম শব্দটির এই দুই অর্থ আলাদা রাখা জরুরি।

তারা কী পয়েছেন

রোগনির্ণয়ের নিরীভুলতায় MIRA মানদণ্ডে চিকিৎসকের ওপরে ছিল, কিন্তু **তিনটি সংখ্যা আলাদা**।

সব **574** ক্ষেত্র-এ MIRA ছাড়পত্রকালীন রোগনির্ণয়ের সঙ্গে **88.9%** মিলি করে। মানব head-to-head subset **311** ক্ষেত্র। ওই একই 311-এ MIRA **87.8%**।

জ্যেষ্ঠ তুলনা: 4 বোর্ড-সনদপ্রাপ্ত চিকিৎসক, **7–11 বছর অভিজ্ঞতা**, স্কোর **78.1%**। মশ্রি-seniority দল—মূলত কনিষ্ঠ resident + 2 বিশেষজ্ঞ—স্কোর **71.1%**। একই ক্ষেত্রগুলো/সরঞ্জামে নির্দেশে দেওয়া: MIRA, জ্যেষ্ঠ, কনিষ্ঠ-ভারী। গবেষণাপত্র নজিরে “ধারাবাহিকভাবে সমতুল্য to, এবং প্রায়ই exceeded” ভাষা ব্যবহার করে, সব আটটি রোগে জুড়ে।

Downstream সিদ্ধান্ত মূলত guideline-concordant; ওষুধ-নিরাপত্তা শ্রুতিতে উচ্চ-তীব্রতা ত্রুটি প্রতবিদেন হয়নি। প্রসেসকরপিশন **467/468** ক্ষেত্রে সংশোধন করা, যদিও লেখকরো স্পষ্ট বলেন **100% নিরীভরণে গ্যতা নয়**।

কিন্তু সাফল্যের আকৃতি গুরুত্বপূর্ণ। Science মাধ্যম কেন্দ্র বিশেষজ্ঞ commentary-তে Dr Wei Xing note করেন শরিরে নাম সুবধি মূলত **স্পষ্ট নির্ণায়ক পরীক্ষা ফল**-যুক্ত শ্রুতি—অ্যাপেন্ডিসাইটিস/প্যানক্রিয়াটাইটিস—এ বেশি। নডিমেট্রিয়া/UTI-তে AI ও চিকিৎসক

উভয়ই সবচেয়ে খারাপ এবং ব্যবধান ছোট।

তথ্য অসমতা: MIRA নথি-সবো উপলভ্য পরীক্ষাগার বিশ্লেষণে প্রায় 51% ব্যবহার করেছে; বোর্ড-সনদপ্রাপ্ত চিকিৎসক প্রায় 28%—মধ্যমা ক্ষেত্রে সাতটি বেশিরকত প্রারামটির। লেখকরা বলেন এটি ডটোসটে নথি-প্রারামভিক মানের নচি এবং ব্যবহুল আন্তঃ-sectional ইমজিং ধারাবাহিকভাবে বাড়ায়নি। তবু বেশি প্রমাণ নজিই রেগনরিণয়মূলক নরিভুলতা বাড়তে পারে; সমান তথ্যপ্রাপ্ত contest নয়।

“চিকিৎসককে হারিয়েছে” মানের “ভালো চিকিৎসক” নয় কেনে

1. স্কোর কীসে বন্দিধে। Julie Jacko note করনে কিছু চিকিৎসা-সামঞ্জস্য মটেরিক ঐতিহাসিক চারটে নথিবিদ্য আচরণ পুনরুত্পাদন করাকে পুরস্কার করে। ঐতিহাসিক নথি উত্তর মূলরে অংশ। রেগনরিণয়মূলক নরিভুলতা অবশ্য ছাড়পত্রকালীন রেগনরিণয়রে বন্দিধে, তাই বন্দিধ নথিবিদ্যকরণ echo-এর চয়ে শক্তিশালী; তবু চিকিৎসা মানের পরম রেফারেন্স সত্য নয়।

2. তথ্য অসমতা। 51% বনাম 28% উপলভ্য পরীক্ষাগার বিশ্লেষণ—MIRA বেশি তথ্য কনিছে/ব্যবহার করেছে। খরচ, অস্বস্তি, বলিম্ব বিবেচনা না করে চিকিৎসক-ও বেশি কম খরচের পরীক্ষা পলে নরিভুলতা বাড়তে পারে।

3. পরবিশে। পশ্চাৎদৃষ্টমূলক স্যান্ডবক্স, আটটি নরিবাচতি রেগনরিণয়, **শুধু পাঠ্যভিত্তিক**। বাস্তব মূল্যায়নে ভেঁত পরীক্ষা, tone, আচরণ, দহে ভাষা, বরিবর্তনশীল উপসর্গ আছে। আটটি রেগনরিণয়রে বাইরে অস্পষ্ট চিকিৎসাবজ্রিএগন নিয়ে গবেষণা কথা বলে না।

4. সম্ভাব্য প্রশিক্ষণ-তথ্য দূষণ। যেকেসেগুলে অনুকরণ করা হয়েছে সেগুলে জনসাধারণের জন্য উন্মুক্ত এবং বহুল আলোচিত। LLM প্রশিক্ষণের সময় গবেষণাপত্র, কেস বা সেগুলে থেকে তৈরি উপাদান দেখে থাকতে পারে। Dr Midhun Parakkal Unni এই উদ্বেগেটি তুলছেন। লেখকরাও সতর্কভাবে বলেন, ফলটি **সম্ভাব্য উন্নয়নসীমা** হতে পারে এবং অন্য প্রকাশ্য কেসে সাধারণীকরণের ক্ষমতাকে বেশি দেখাতে পারে। স্বাধীন ব্যক্তিগত বা ভবিষ্যতমুখী পুনরাবর্তি ছাড়া মুখস্থ করা আর প্রকৃত যুক্তি আলাদা করা কঠিন।

এই সতর্কতা সক্ষমতাকে মুছে দেয় না। সংশোধন করা পাঠ: পশ্চাৎদৃষ্টমূলক মানদণ্ডে পুরো-নথি এজেন্ট চিকিৎসককে ওপরে স্কোর করেছে, বিশেষ করে পরিস্কার-পরীক্ষা রেগে; বেশি পরীক্ষাগার info ব্যবহার করেছে এবং ঐতিহাসিক চারটের সঙ্কে আংশিক সামঞ্জস্য পেয়েছে। এটি সক্ষমতা প্রদর্শন, যন্ত্র ভালো চিকিৎসক—রায় নয়।

এই গবেষণা কী প্রমাণ করে না

- **বাস্তব রেগিতে কাজ করে**—না; সব ক্ষেত্রে পশ্চাৎদৃষ্টমূলক/সমিলটেডে।
- **আটটি রেগনরিণয়রে বাইরে সাধারণীকরণ করে**—না।
- **Deploy নরিপদে করা যায়**—না; ভবিষ্যতমুখী নরিপত্তা/শাসনব্যবস্থা দরকার।
- **ন্যায্য সমান তথ্যপ্রাপ্ত চিকিৎসক head-to-head**—না; MIRA বেশি রকত বিশ্লেষণ ব্যবহার করেছে এবং চিকিৎসা মটেরিক চারটের বন্দিধে আংশিক স্কোর করেছে।
- **মুখস্থকরণ-মুক্ত**—না; জনসাধারণের অনুকরণ করা-IV প্রশিক্ষণ ওভারল্যাপ exclude হয়নি।
- **AI replaces চিকিৎসক**—না। বাস্তব ব্যবহার চিকিৎসক partnership/কর্তৃত্ব এবং পাঠ্য নথি বাইরে মানব পর্যবেক্ষণ দরকার।

প্ৰমাণ কতটা শক্ত?

দুটি দাবি আলাদা।

সক্ষমতা প্ৰদৰ্শন—ভাষা-মডলে এজেন্ট স্যান্ডবক্সে সীমাবদ্ধ EHR-এ ইতিহাস→পৰীক্ষাগুলো→ৰোগনিৰ্ণয়→নিৰ্দেশে পুরো লুপ চালাতে পাৰে—নতুন ও যথেষ্ট বিশ্বাসযোগ্য। এটিহি সত্যহি নতুন অংশ।

শ্ৰেষ্ঠত্বৰে দাবি—MIRA চিকিৎসকদৰে চয়ে ভালো ৰোগনিৰ্ণয়কাৰী—খুব সীমিতভাবে পড়তে হবে। এটি নিৰ্দিষ্ট retrospective benchmark, মাত্ৰ আটটি diagnosis, অসম তথ্যপ্ৰাপ্তি, ঐতিহাসিক উত্তৰকে gold standard ধৰা এবং সম্ভাব্য training-data contamination-এৰ মধ্যমে মাপা ফল। তাই “এই benchmark-এ impressive performance” বলা যায়; “মানুষৰে চয়ে ভালো চিকিৎসক” বলা যায় না—লেখককোও দ্বিতীয় দাবি কৰনে না।

সঠিক অবস্থানটি “AI চিকিৎসকদৰে হাৰিয়ে দিছে” কহিবা “এটি খেলনা”—কোনোটিহি নয়। নথি ভেতৰে বাস্তব ক্ৰিয়া নতি পাৰে এমন চিকিৎসাবিষয়ক এজেন্টটি কঠিন পশ্চাদৃষ্টমূলক মানদণ্ডে ভালো কৰছে; এখন দৰকাৰ বাস্তব, অস্পষ্ট ৰোগীৰ ক্ষেত্ৰে ভবিষ্যতমুখী এবং যথায় তত্ত্বাবধান-সহ পৰীক্ষা।

কনে এটি গুৰুত্বপূৰ্ণ

চিকিৎসাবিষয়ক AI-এৰ আকৰ্ষণীয় প্ৰশ্ন বদলাচ্ছে। আগে “মডলে ৰোগনিৰ্ণয় ঠিক বলতে পাৰে?”—এখন আৰও কঠিন প্ৰশ্ন “কৰ্মপ্ৰবাহ চালাতে পাৰে? ইতিহাস নওয়া, পৰীক্ষা নিৰ্দেশে দওয়া, ফল ব্যাখ্যা কৰা, সিদ্ধান্ত, ক্ৰিয়া?” MIRA এই সৰণ দেখায়। Acting-এ উপযোগিতা এবং ঝুঁকি দুটেই বাড়ে।

শৰিণাম “AI চিকিৎসকদৰে ছাড়িয়ে যায়” সৰ্বনমিন নতুন/সৰ্বনমিন সমৰ্থতি অংশ। সত্যহি গুৰুত্বপূৰ্ণ অংশ: **এজেন্ট পুরো নথিকৰ্মপ্ৰবাহে পদক্ষেপে কৰতে পাৰে**। সক্ষমতা টকি থোকা কৰলে প্ৰথমকে চিকিৎসক থাকবোতা নয়, কৰ্মপ্ৰবাহ—কৰ্ম, ফলাফল chase, ব্যাখ্যা—বদলাবে।

Leaderboard-এ ভালো ফল ভবিষ্যতমুখী clinical testing-এৰ বকিল্প নয়। ৰোগী-স্তৰে বাস্তব পৰিবেশে prospective গবেষণাই পৰবৰ্তী বড় প্ৰমাণৰে ধাপ।

সংক্ষিপ্ত সারাংশ

MIRA স্যান্ডবক্সে সীমাবদ্ধ EHR-এ স্বাযতশাসতি এজেন্ট: ইতিহাস, পৰীক্ষাগাৰ/ইমজেথি/মাইক্ৰোবায়োজি নিৰ্দেশে দওয়া/ব্যাখ্যা কৰা, ৰোগনিৰ্ণয় এবং চিকিৎসা নিৰ্দেশে দওয়া শুরু থেকে শেষে পৰ্যন্ত। জনসাধাৰণৰে অনুকৰণ কৰা-IV-এৰ 574 পশ্চাদৃষ্টমূলক ক্ষেত্ৰ, 8 ৰোগনিৰ্ণয়-এ সামগ্ৰিক ৰোগনিৰ্ণয়মূলক নিৰ্ভুলতা 88.9%। যোথ 311-ক্ষত্ৰ চিকিৎসক তুলনায় MIRA 87.8%, জ্যেষ্ঠ বোৰ্ড-সনদপ্ৰাপ্ত দল 78.1%, কনষ্টি-ভাৰী দল 71.1%। ওষুধ সিদ্ধান্ত মূলত নিৰাপদ/guideline-concordant, 467/468 প্ৰসেক্ৰিপশন সংশোধন কৰা, উচ্চ-তীব্ৰতা ওষুধ তৰুটি প্ৰতিবিদেন নহে। কিন্তু স্যান্ডবক্স/শুধু পাঠ্যভিত্তিক, নিৰ্বাচতি ৰোগনিৰ্ণয়, স্পষ্ট-পৰীক্ষা শৰত প্ৰান্ত বড়, MIRA ~51% বনাম চিকিৎসক ~28% উপলভ্য পৰীক্ষাগাৰ বিশ্লেষণ ব্যবহার কৰছে, কিছু ফলাফল ঐতিহাসিক চাৰ্টৰে বৰিদ্ধে স্কোৰ কৰছে, এবং জনসাধাৰণৰে অনুকৰণ কৰা-IV দুষণ ঝুঁকি আছে। বাস্তব অগ্ৰগতি হলো পুরো-কৰ্মপ্ৰবাহ actor; বাস্তব-ক্লিনিক শ্ৰেষ্ঠত্ব/ব্যবহারিক মতোয়নে প্ৰমাণ নয়।

সরাসরি যাচাই

পপোর যা দখোয়: ভাষা-মডলে এজেন্ট স্যান্ডবকসে সীমাবদ্ধ EHR-এ পুরো ক্ষেত্র চালাতে পারে এবং নির্দিষ্ট পশ্চাদৃষ্টমূলক মানদণ্ডে উচ্চ রোগনির্ণয়মূলক স্কোর করেছে।

যা সম্ভব, কনিত্ত প্রমাণতি নয়: বাস্তব রোগী উপকার; সুবধি বশিদ্ধ ববিচেনামূলক চন্তির কারণে; ব্যক্তিগিত অদখো তথ্যেও একই কার্যসম্পাদন।

যা দখোয় না: বসিত্ত চকিৎসাবজ্রিএগন সাধারণীকরণ, ব্যবহারকি মতোতায়নে নরিপত্তা, ন্যায্য সমান তথ্যপ্রাপ্ত চকিৎসক শ্রেষ্টত্ব, চকিৎসক প্রতস্থাপন, দূষণ-মুক্ত ফল।

প্রধান সীমাবদ্ধতা: পশ্চাদৃষ্টমূলক সম্মিলশেন, আটটি রোগনির্ণয়, শুধু পাঠ্যভিত্তিকি, আরও পরীক্ষাগার পরীক্ষাগুলো, চার্ট-ভিত্তিকি স্কোরিং elements, জনসাধারণের ডটোসটে দূষণ সম্ভাবনা।

সাধারণ পাঠকরে কতটা আস্থা রাখা উচতি? স্যান্ডবকসে শুরু থেকে শেষে পর্যন্ত কাজ সম্পাদনরে সক্ষমতায়—উচ্চ। “AI চকিৎসকদরে চয়ে ভালো রোগনির্ণয় করে” এই বৃহত্তর দাবতি—কম; ফলটি benchmark-এর সীমানার মধ্যে।

উৎস

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

PubMed 42310457 — peer-reviewed abstract
<https://pubmed.ncbi.nlm.nih.gov/42310457/>

Science Media Centre — expert reaction to MIRA and AMIE (2026)
<https://www.sciencemediacentre.org/expert-reaction-to-presentation-of-two-new-medical-ai-models-for->

News-Medical (2026) — illustrates the ‘outperforms doctors’ framing
<https://www.news-medical.net/news/20260621/Autonomous-medical-AI-outperforms-doctors-in-simulated-EH.aspx>