



The CLEAN PAPER

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

ويكل ذكاء اصطناعي أدار حالات مرضى كاملة بمفرده — في محاكاة وعلى سجلات قديمة

MIRA نوع جديد من الذكاء الاصطناعي الطبي: بدل الإجابة عن سؤال واحد، يدير حالة كاملة داخل سجل مستشفى محاكي — يأخذ التاريخ المرضي، ويطلب الفحوص ويقرأها، ويصل إلى تشخيص، ويكتب الأوامر. في 574 حالة استعادية من قاعدة بيانات عامة وعبر ثمانية تشخيصات مختارة مسبقاً، أفاد المؤلفون بأنه تفوق على الأطباء في دقة التشخيص واتخذ قرارات متوافقة إلى حد كبير مع الإرشادات وآمنة دوائياً. لكن لكل قيد هنا وزن: عمل داخل بيئة معزولة وعلى سجلات قديمة وبالنص فقط، وجاء جزء كبير من تفوقه في الحالات ذات نتائج الفحوص الواضحة؛ وطلب نحو ضعفي تحاليل الدم التي طلبها الأطباء؛ وقيمت عدة نتائج وفق ما كان مسجلاً في الملف الأصلي. التقدم الحقيقي هو ويكل يتصرف عبر سير العمل كله — وليس دليلاً على أن آلة أصبحت تشخص أفضل من الطبيب. والمؤلفون يوضحون أن التعميم والسلامة والحوكمة ما زالت تتطلب دراسات مستقبلية في العالم الحقيقي.

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

الإجابة عن سؤال ليست هي نفسها إدارة الحالة كاملة

معظم قصص الذكاء الاصطناعي الطبي تدور حول نموذج يجيب. تعطيه وصف حالة أو سؤال امتحان، فيعيد لك تشخيصاً أو فقرة من النصائح ويحصل على نتيجة جيدة. هذا مفيد، لكنه ضيق: مستشار ذكي لا يلمس السجل الطبي أبداً.

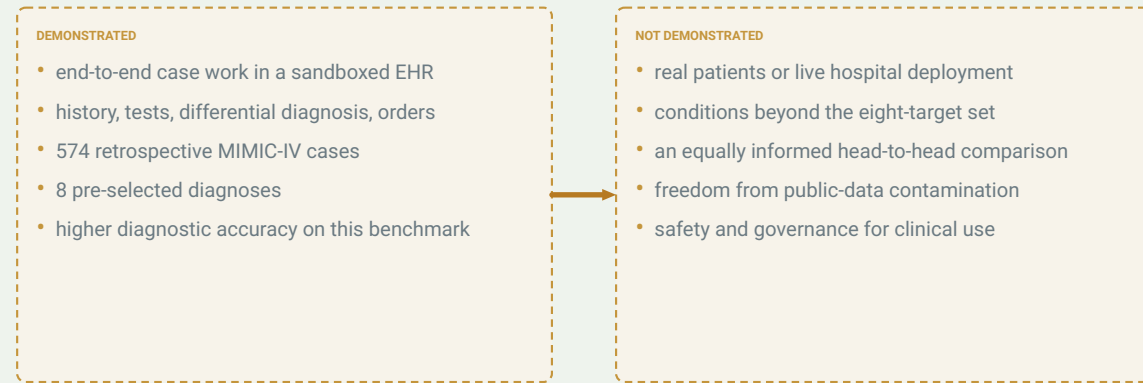
صُمم MIRA ليفعل شيئاً مختلفاً، وهذا الاختلاف هو الخبر الحقيقي. بدل الإجابة عن سؤال واحد، يدير حالة كاملة: يقرأ سجل المريض، ويقرر ما التاريخ المرضي الذي ما زال يحتاج إليه، ويطلب تحاليل مخبرية وتصويراً وفحوص أحياء دقيقة، ويقرأ النتائج، ويضيق التشخيص التفريقي، ثم يكتب الأوامر التالية – وصفات دوائية، أو إدخالاً إلى المستشفى، أو إحالة إلى الجراحة. ويفعل ذلك باتخاذ إجراءات داخل السجل الصحي الإلكتروني، كما يفعل الطبيب، بدل إنتاج نص حر ينسخه إنسان إلى النظام.

هذه خطوة حقيقية: من نظام يقدم المشورة إلى نظام يتصرف عبر سير العمل. وهي أيضاً بالضبط النوع من الخطوات الذي يُحتزل في التغطية الإعلامية إلى «الذكاء الاصطناعي يتفوق على الأطباء». لذلك يستحق الأمر أن نكون دقيقين بشأن ما اختبر وأين.

الخلاصة في جملة واحدة: أدار MIRA حالات كاملة بمفرده، وفي هذا الاختبار المعياري تفوق على الأطباء في دقة التشخيص – لكنه فعل ذلك داخل بيئة معزولة، وعلى سجلات استعادية، وعبر ثمانية تشخيصات مختارة مسبقاً، وبالتواصل النصي فقط، وكان جزء كبير من تفوقه في الحالات التي تملك نتائج اختبارات أوضح. التقدم حقيقي. لكن لوحة النتائج ليست العيادة.

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.



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

The figure separates the workflow result from the deployment claim.

أثبت MIRA قدرته على تنفيذ سير عمل كامل من البداية إلى النهاية داخل سجل صحي إلكتروني معزول. ولم يثبت نشرًا سريريًا في العالم الحقيقي، ولا تغطية تشخيصية واسعة، ولا مقارنة عادلة مع أطباء يملكون القدر نفسه من المعلومات.

ما الذي فعله الباحثون

MIRA — اختصار Action and Reasoning for Intelligence Medical — هو وكيل مستقل مبني على نماذج OpenAI الجزء الذي يتحاور ويتصرف يعمل على GPT-4o، بينما تستخدم خطوة تخطيط منفصلة نموذج الاستدلال o1 من OpenAI. تعمل هذه المكونات داخل إطار مخصص ومتوافق مع المعايير، مبني على HL7، FHIR، وهو معيار البيانات الذي تستخدمه المستشفيات الفعلية لنقل السجلات، مع صندوق أدوات يضم أكثر من ثمانين ألف إجراء سريري ممكن. داخل هذه البيئة المعزولة يستطيع MIRA جلب تاريخ المريض، وطلب التحاليل والتصوير والأحياء الدقيقة وتفسيرها، وإنشاء تشخيص تفريقي، وصياغة خطط علاج — بما في ذلك وصف الأدوية، وجدولة إجراءات جراحية، والتخطيط للإدخال إلى المستشفى. وهناك تفصيل يستحق التنبيه منذ البداية: «الحكام» الآليون الذين قيموا إجابات MIRA كانوا هم أنفسهم GPT-4o — من عائلة النموذج نفسها التي يجري اختبارها — وقد دعم المؤلفون ذلك بمراجعة من طبيب معتمد بعد أن أثار أحد المحكمين هذه المخاوف.

جاءت مجموعة الاختبار من MIMIC-IV، وهي قاعدة بيانات كبيرة ومتاحة للعامة لسجلات صحية إلكترونية منزوع منها ما يعرف المرضى. ومن نحو 300 ألف مريض عولجوا في Beth Israel Medical Center بين 2008 و2019، اختار المؤلفون 574 حالة مريض تغطي ثمانية تشخيصات مستهدفة — أمراضاً بطنية وطوارئ في الطب الباطني، منها التهاب الزائدة الدودية والتهاب البنكرياس والالتهاب الرئوي والتهاب المسالك البولية. في كل حالة بدأ MIRA بصورة محدودة واضطر إلى أن يقرر، خطوة بعد خطوة، ماذا يفعل بعد ذلك — شكل يشبه الاستقصاء السريري الحقيقي، لكنه يُجرى على سجل نهايته الواقعية مسجلة سلفاً.

ثم قورن أدائه بأداء أطباء على الحالات نفسها.

ماذا يعني فعلياً «وكيل مستقل داخل سجل صحي إلكتروني معزول»

تفكيكها. يجدر لذلك المعنى، من كثيراً هنا تحمل كلمات ثلاث وكيل (Agent) يعني أن النظام ليس روبوت دردشة يجيب عن توجيه واحد. إنه يعمل في حلقة: ينظر إلى الحالة الراهنة، ويختار فعلاً (اطلب هذا الفحص، اسأل عن هذا الجزء من التاريخ المرضي)، ويرى النتيجة، ثم يختار الفعل التالي — إلى أن يصل إلى تشخيص وخطوة. «الذكاء» هو نموذج اللغة؛ أما القدرة على الفعل فهي البنية المحيطة به التي تحول نص النموذج إلى عمليات مسموح بها داخل السجل الصحي وتعيد النتائج إليه.

سجل صحي إلكتروني معزول (Sandboxed EHR) يعني نسخة محاكاة ومغلقة من نظام سجلات طبية، لا نظام مستشفى حياً. يستطيع MIRA أن «يطلب» فحصاً وأن يتلقى النتيجة التي حصل عليها ذلك المريض فعلاً، لأن الحالة تاريخية والجواب مسجل بالفعل. لا شيء مما يفعله يمس مريضاً حقيقياً أو عيادة حقيقية.

مستقل (Autonomous) يعني أنه يكمل الحالة من البداية إلى النهاية من دون إنسان في الحلقة — داخل البيئة المعزولة. ولا يعني أنه أطلق لإدارة مرضى من دون إشراف. هذان ادعاءان مختلفان جداً، ولم يُختبر إلا الأول.

وهناك تفصيل آخر عن البيئة المعزولة، لأن من السهل تخيلها خطأ: يمكن لـ MIRA أن يطلب أي فحص يريده، لكن النظام لا يستطيع إعطاؤه نتيجة إلا إذا كان ذلك الفحص قد أُجري فعلاً لهذا المريض في الواقع. إذا طلب مجموعة تحاليل دم أجراها المريض، يحصل على القيم الحقيقية؛ وإذا طلب تصويراً لم يطلبه أحد أصلاً، يحصل على «N/A» — تعذر إجراؤه، ولا يحصل أبداً على رقم مخترع. وينطبق الشيء نفسه على التاريخ المرضي: إذا لم يُحتوِ السجل على ما يسأل عنه، MIRA يقول «المريض» المحاكى ببساطة إنه لا يعرف. لذلك لا يستطيع MIRA استحضار الفحص الحاسم الوحيد الذي لم يُجرَ أصلاً؛ بل يستطيع العمل فقط بما تضمنه الاستقصاء الحقيقي المسجل.

هذا الفرق مهم لأن كلمة العنوان — «مستقل» — هي الأكثر قابلية لأن تُفهم بالمعنى الثاني.

ماذا وجدوا

في الاختبار المعياري، تفوق MIRA على الأطباء في دقة التشخيص — لكن العنوان يخفي ثلاثة أرقام مختلفة يسهل خلطها، لذا من الأفضل فصلها.

عبر جميع الحالات الـ 574، سُمي MIRA التشخيص الصحيح في 98.8% من الحالات — بحسب تشخيص الخروج المسجل فعلياً لكل مريض في MIMIC-IV. أما في المقارنة المباشرة مع البشر، فلم يعمل الأطباء على الحالات الـ 574 كلها؛ بل على مجموعة فرعية مشتركة من 311 حالة، مستخدمين السجلات والأدوات نفسها المتاحة لـ MIRA. في هذه الحالات الـ 311 نفسها حصل MIRA على 87.8%، وقورن بمجموعتين جرى تجنيدهما بصورة منفصلة. الأولى مجموعة أطباء كبار: أربعة أطباء معتمدين بخبرة من 7 إلى 11 سنة، حققوا 78.1%، والثانية مجموعة مختلطة الخبرة: معظمها أطباء مقيمون مبتدئون ومعهم اختصاصيان، وهي أقرب إلى طريقة تشكيل قسم طوارئ فعلي، وحقت 71.1%. الحالات نفسها، والأدوات نفسها — فجاء الذكاء الاصطناعي أولاً، ثم الأطباء الأكثر خبرة، ثم الفريق الذي يغلب عليه الأطباء الأصغر خبرة. والصياغة الحذرة للورقة نفسها هي أن أداء MIRA كان «مكافئاً باستمرار للأطباء، وكثيراً ما تجاوزهم» عبر الأمراض الثمانية.

كما صمدت قراراته اللاحقة: كانت إلى حد كبير متوافقة مع الإرشادات، وآمنة دوائياً، ومناسبة في قرارات الإدخال. لا يذكر المؤلفون أخطاء دوائية شديدة الخطورة عبر خمس فئات سلامة (التداخلات الدوائية، وضبط الجرعات بحسب الكلى، والحساسيات، ومخاطر QT، ووصف المواد الأفيونية)، وكانت الصفات صحيحة في 467 من أصل 468 حالة — مع تنبيههم إلى أن النظام «لم يحقق موثوقية 100%». عند أخذ ذلك كما هو، تبدو نتيجة لافتة: وكل يدير الحالة كلها وينتهي متقدماً على الأطباء في التشخيص.

لكن شكل هذا التفوق مهم بقدر التفوق نفسه. أشار اختصاصيون مستقلون قرأوا الورقة إلى مصدر جزء مهم من الأفضلية. في لوحة خبراء Centre Media Science، أشار الدكتور Xing Wei من Sheffield of University إلى أن الرقم الرئيسي — تفوق الذكاء الاصطناعي على الأطباء في دقة التشخيص — كان مدفوعاً في الغالب بالحالات ذات نتائج الاختبارات الواضحة، مثل التهاب الزائدة والتهاب البنكرياس، حيث يحسم تصوير أو تحليل مخبري قاطع السؤال. أما في الالتهاب الرئوي والتهابات المسالك البولية — وهما من أكثر الأسباب شيوعاً لحدوث الناس فعلياً إلى الطوارئ — فقد كان أداء كل من الذكاء الاصطناعي والأطباء هو الأسوأ، وكانت الفجوة بينهما الأصغر.

وكان هناك أيضاً عدم تماثل في طريقة لعب الطرفين، وتظهر أرقامه في الورقة نفسها. اعتمد MIRA أكثر على المختبر: استخدم نحو 51% من تحليل المختبر المتاحة في الرعاية الروتينية، مقابل نحو 28% للأطباء المعتمدين — بوسيط سبعة مؤشرات دم إضافية لكل حالة. يحرص المؤلفون على وصف ذلك بأنه أقل من خط أساس الرعاية الروتينية في مجموعة البيانات نفسها، لا استراتيجية «اطلب كل شيء»، ولا يبلغون عن زيادة منهجية في التصوير المقطعي الأعلى تكلفة. ومع ذلك، قد تؤدي المعلومات الأكثر بحد ذاتها إلى دقة تشخيصية أعلى؛ لذا فهذه ليست مقارنة متكافئة تماماً بين طرفين يملكان المعلومات نفسها، وهي فجوة أشار إليها Xing مباشرة. ولو أتيح لطبيب أن يطلب كل تحليل دم رخيص من دون موازنة التكلفة أو الإزعاج أو التأخير، فسيبدو أدق على الورق أيضاً.

لماذا «تفوق على الأطباء» لا تعني «طبيب أفضل»

من السهل الإفراط في تفسير الفوز في اختبار معياري. هناك ثلاثة أمور بين «حصل MIRA على نتيجة أعلى» و«MIRA أفضل».

أولاً، ما الذي قيس الأداء بالنسبة إليه. أشارت البروفيسورة Jacko Julie من Edinburgh of University إلى أن عدة نتائج رئيسية لـ MIRA تُعرف نسبة إلى ما كان موثقاً في مجموعة البيانات الأساسية — أي إن النظام يُكافأ على إعادة إنتاج السلوك السريري المسجل، وليس بالضرورة على إظهار أفضل رعاية ممكنة. يصبح السجل التاريخي مفتاح الإجابة. وهذه طريقة معقولة لبناء اختبار معياري، لكنها تقيس الاتفاق مع ما تم فعله، لا الصحة بمعنى مطلق. وللإنصاف، يؤثر ذلك أكثر في مقاييس توافق العلاج — هل طابقت أوامر MIRA ما في

السجل؟ — بينما تُقاس دقة التشخيص الرئيسية مقابل تشخيص الخروج، وهو أقرب إلى نتيجة فعلية من مجرد صدى للتوثيق.

ثانياً، عدم تماثل المعلومات المذكور: مزيد كبير من تحاليل الدم (نحو 51% من المؤشرات المتاحة مقابل 28%) يعني أدلة أكثر. ومن المرجح أن جزءاً من فجوة الدقة قد تم شراؤه بالمعلومات بدل أن يكون نتاج استدلال أفضل وحده.

ثالثاً، أين جرى الاختبار. كان ذلك في بيئة معزولة، وعلى حالات استعادية، وعبر ثمانية تشخيصات مختارة مسبقاً، وبالنص فقط. وكما قال الدكتور Oliver Dominic من Oxford، of University يعتمد التقييم السريري الحقيقي ليس فقط على ما يقوله المرضى بل على كيفية قولهم له — إلى جانب الفحص الجسدي والسلوك الملحوظ ولغة الجسد، وهي أشياء لا يراها وكيل نصي يقرأ سجلاً مكتماً. والمريض الذي لا تدرج مشكلته ضمن واحد من التشخيصات الثمانية المختارة يقع ببساطة خارج ما تستطيع هذه الدراسة الحديث عنه.

وهناك قلق أهدأ أثاره المراجعون: تلوث بيانات التدريب. قاعدة MIMIC-IV عامة وُكِّت عنها كثيراً، لذلك ربما يكون نموذج لغوي دُرِّب على الإنترنت المفتوح قد رأى من قبل أوراقاً أو مناقشات حالات أو البيانات نفسها. أشار الدكتور Unni Parakkal Midhun من Sheffield of University إلى ذلك مباشرة: إذا كانت بعض الإجابات موجودة في بيانات التدريب، فجزء من الأداء قد يكون ذاكرة لا استدلالاً سريرياً، ولا يستطيع الفصل بينهما إلا تكرار مستقل. ومن اللافت أن المؤلفين لا يتجاهلون هذه النقطة؛ بل يكتبون إن نتائجهم «يمكن تفسيرها بحذر بوصفها حدّاً أعلى محتملاً» وإنها «قد تتألف في تقدير القدرة على التعميم إلى حالات عامة أخرى» — أي إن الورقة نفسها تضع سقفاً لعنوانها.

لا يجعل أي من هذا MIRA أقل إثارة للاهتمام. بل يجعل القراءة الصحيحة أكثر معايير: في اختبار استعادي، تفوق وكيل يستطيع العمل عبر السجل كله على الأطباء في التشخيصات ذات الاختبارات الواضحة — جزئياً لأنه طلب اختبارات أكثر، وجزئياً لأنه وافق السجل المسجل. هذا عرض حقيقي لقدرة جديدة، لا حكماً بأن الآلات أصبحت تشخص أفضل من الأطباء.

ما الذي لا تثبته هذه الدراسة

- لا تُظهر أن MIRA يعمل على مرضى حقيقيين. كل حالة كانت استعادية ومحاكاة؛ ولم يدر أي مريض فعلياً بواسطة النظام.
- لا تُظهر أنه يعمل خارج التشخيصات الثمانية. الحالات الواقعة خارج المجموعة المختارة — وهي أغلبية الطب الفوضوية وغير المتميزة — لم تُختبر.
- لا تُظهر أنه آمن للنشر السريري. يكتب المؤلفون أنفسهم أن التعميم والسلامة والحكمة ما زالت تتطلب دراسات مستقبلية في العالم الحقيقي.
- لا تقدم مقارنة عادلة تماماً مع الأطباء. استخدم MIRA عدداً أكبر بكثير من تحاليل الدم (نحو 51% من المؤشرات المتاحة مقابل 28%)، كما قُيِّمت عدة نتائج علاجية جزئياً وفق ما يُجل تاريخياً بدل «أفضل رعاية» باعتبارها حقيقة أرضية مستقلة.
- لا تُظهر أن النتيجة خالية من الحفظ. قد تتداخل بيانات MIMIC-IV العامة مع تدريب النموذج، ولا يستبعد ذلك إلا تكرار مستقل.
- ولا تعني أن «الذكاء الاصطناعي يستبدل الأطباء». إنها أداة دعم قرار تستطيع التصرف عبر السجل؛ وإجماع المراجعين هو أن الاستخدام الواقعي سيكون بالشراكة مع الأطباء، الذين يحتفظون بالسلطة ويقدمون كل ما لا يلتقطه السجل النصي.

ما مدى قوة الدليل؟

من الأفضل تقسيم الادعاء إلى جزأين، لأن قوة الدليل مختلفة جداً بينهما.

بوصفه عرضاً لقدرة تقنية — أي أن وكيلاً قائماً على نموذج لغوي يستطيع إدارة حالة سريرية كاملة داخل سجل صحي إلكتروني معزول، رابطاً التاريخ المرضي والفحوص والتشخيص والأوامر من البداية إلى النهاية — فالعمل جديد فعلاً ومقنع إلى حد معقول. هذا هو الجزء الجديد، وهو الجزء الذي يستحق الانتباه.

وبوصفه ادعاء تفوق — أن MIRA أفضل من الأطباء — فالدليل محدود النطاق ويجب قراءته بحذر. ينطبق على اختبار استعادي محدد، وثمانية تشخيصات، مع عدم تماثل في المعلومات (اختبارات أكثر) ومفتاح إجابة مأخوذ من السجل التاريخي، ومع احتمال قائم لتلوث بيانات التدريب. هذا يكفي للقول «كان أداء الوكيل مثيراً للإعجاب في هذا الاختبار». ولا يكفي للقول «الوكيل أفضل في التشخيص من الطبيب»، والمؤلفون لا يدعون الشيء الثاني.

الموقف الأكثر فائدة ليس «الذكاء الاصطناعي يهزم الأطباء» ولا «مجرد لعبة». بل: نوع جديد من الذكاء الاصطناعي الطبي — نوع يتصرف عبر السجل بدل أن يجيب عن أسئلة فقط — أدى جيداً في اختبار استعادي صعب، والآن عليه أن يثبت نفسه حيث لم يُختبر بعد: مع مرضى حقيقيين غير مصنفين مسبقاً، بصورة مستقبلية وتحت حوكمة.

لماذا يهم ذلك

خلال السنوات القليلة الماضية تغير السؤال المثير في الذكاء الاصطناعي الطبي بهدوء. كان السؤال «هل يستطيع النموذج الحصول على التشخيص الصحيح؟» — واستمرت النماذج في الإجابة بنعم على مجموعات اختبار أكثر نظافة. يمثل MIRA انتقالاً إلى سؤال أصعب: «هل يستطيع النموذج أداء العمل — جمع المعلومات، وطلب الفحوص، وتفسيرها، واتخاذ القرار، والتصرف — عبر سير عمل كامل؟» هذا سؤال أكثر فائدة وصدقاً، لأن الفعل هو حيث تكمن الصعوبة والمخاطر معاً.

ولهذا تهم طريقة التأطير. رد الفعل السريع في العنوان — الذكاء الاصطناعي يتفوق على الأطباء — يشير إلى الجزء الأقل جدة والأضعف دعماً من النتيجة. الجديد حقاً أصغر وأكثر أهمية: وكيل يستطيع الانتقال عبر السجل كله. وإذا صمدت هذه القدرة، فإنها ستغير سير العمل قبل وقت طويل من أن تغير من يتحكم فيه. الطبيب لا يُستبدل؛ بل إن طلب الفحوص، وتفسيرها، وملاحقة النتائج — أي سير العمل — هو المكان الذي تهبط فيه أداة كهذه أولاً.

كما أنها تعيد عبء الإثبات إلى الاتجاه الصحيح. الفوز في لوحة صدارة لحالات استعادية سبب لإجراء تجربة مستقبلية، وليس بديلاً عنها. والمؤلفون يقولون ذلك. القراءة الصادقة لـ MIRA هي دعوة إلى الدراسة التالية — بالمعيار الذي يعرفه المجال بالفعل: القياس على المرضى، لا على الاختبارات المعيارية.

خلاصة واضحة

MIRA وكيل ذكاء اصطناعي مستقل يعمل داخل سجل صحي إلكتروني معزول: يستطيع أخذ التاريخ المرضي، وطلب تحاليل وتصوير وفحوص أحياء دقيقة وتفسيرها، والوصول إلى تشخيص، وكتابة خطط العلاج. عند اختباره على 574 حالة استعادية من قاعدة MIMIC-IV العامة عبر ثمانية تشخيصات مختارة مسبقاً، تفوق على الأطباء في دقة التشخيص (9%، 88% إجمالاً؛ و8%، 87% مقابل 1%، 78% في المقارنة المباشرة) واتخذ قرارات متوافقة إلى حد كبير مع الإرشادات وآمنة دوائياً. لكن التقييم كان محاكاة على سجلات قديمة وبالنص فقط؛ وجاء جزء كبير من الأفضلية في الحالات ذات نتائج الفحوص الواضحة؛ واستخدم MIRA تحاليل دم أكثر بكثير من الأطباء (نحو 51% من المؤشرات المتاحة مقابل 28%)، وقيمت عدة نتائج علاجية وفق ما سجله الملف الأصلي؛ كما تثير مجموعة البيانات العامة خطراً حقيقياً لتلوث التدريب —

وهو ما يصفه المؤلفون أنفسهم بأنه قد يجعل الأرقام حدًا أعلى محتملاً. التقدم الحقيقي هو وكيل يتصرف عبر سير العمل كله بدل الإجابة عن أسئلة منفصلة. والمؤلفون واضحون في أن التعميم والسلامة والحكمة ما زالت تتطلب دراسات مستقبلية في العالم الحقيقي. القراءة الصحيحة إذن: عرض مثير للإعجاب لقدرة جديدة، لا دليل على أن الذكاء الاصطناعي يشخص أفضل من الأطباء، ولا نظام جاهز لعيادة حقيقية.

مراجعة بلا تهويل

ما الذي تُظهره الورقة: يستطيع وكيل قائم على نموذج لغوي (MIRA) إدارة حالة سريرية كاملة من البداية إلى النهاية داخل سجل صحي إلكتروني معزول — التاريخ المرضي، والفحوص، والتشخيص، والأوامر — وفي اختبار استعادي من 574 حالة عبر ثمانية تشخيصات، تفوق على الأطباء في دقة التشخيص (9%.88 إجمالاً؛ 8%.87 مقابل 1%.78 للأطباء المعتمدين)، من دون أخطاء دوائية عالية الشدة.

ما هو معقول لكنه غير مثبت: أن ترجم هذه القدرة إلى فائدة لمرضى حقيقيين غير مصنفين مسبقاً، وأن تعكس أفضلية الدقة استدلالاً أفضل بدل اختبارات أكثر، أو توافقاً مع السجل، أو حفظاً لبيانات عامة.

ما الذي لا تُظهره: أن MIRA يعمل خارج التشخيصات الثمانية؛ أو أنه آمن للنشر؛ أو أنه يهزم الأطباء في مقارنة عادلة متكافئة المعلومات؛ أو أنه يستبدل الأطباء؛ أو أن النتائج خالية من تلوث بيانات التدريب.

أهم القيود: تقييم استعادي ومحاكاة ونصي فقط؛ ثمانية تشخيصات مختارة مسبقاً؛ عدم تماثل في المعلومات (تحاليل دم أكثر بكثير — نحو 51% من المؤشرات المتاحة مقابل 28%، في تحاليل دم منخفضة التكلفة لا التصوير)؛ نتائج علاجية قيّمت جزئياً مقابل السجل التاريخي؛ واختبار عام (MIMIC-IV) ربما رآه نموذج لغوي أثناء التدريب — وهو ما يشير إليه المؤلفون بوصفه حدًا أعلى محتملاً للأداء.

ما مقدار الثقة المناسب للقارئ العام؟ ثقة عالية بأن MIRA عرض حقيقي وجديد لقدرة تقنية — وكيل يتصرف عبر السجل، لا مجرد روبوت دردشة. وثقة منخفضة بأنه أفضل تشخيصياً من الطبيب: هذا الادعاء محصور في اختبار استعادي واحد ومتداخل مع طلب لخصوص أكثر، والتقييم مقابل السجل، واحتمال تلوث البيانات. خلاصة المؤلفين أنفسهم هي الأنسب: يحتاج هذا إلى دراسة مستقبلية في العالم الحقيقي قبل أن يعني شيئاً لرعاية المرضى.

المصادر

(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