



The CLEAN PAPER

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

Un agent IA a traite des cas patients entiers tout seul - dans un simulateur, sur des dossiers passes

MIRA est un nouveau type d'IA medicale : au lieu de repondre a une seule question, il traite un cas entier dans un dossier hospitalier simule - recueillir l'histoire, commander et lire des examens, poser un diagnostic, rediger les prescriptions. Sur 574 cas retrospectifs issus d'une base publique, couvrant huit diagnostics preselectionnes, les auteurs rapportent qu'il a depasse des medecins en precision diagnostique et produit des decisions largement conformes aux recommandations et sures cote medicaments. Chaque qualificatif compte : bac a sable, dossiers passes, texte seul; avantage surtout dans les pathologies aux resultats d'examens nets; environ deux fois plus d'analyses sanguines que les medecins; plusieurs criteres notes contre ce que le dossier original contenait. L'avance reelle est un agent qui agit dans tout le workflow - pas la preuve qu'une machine diagnostique mieux qu'un medecin.

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

Répondre à une question n'est pas la même chose que prendre en charge un cas

La plupart des récits sur l'IA médicale parlent d'un modèle qui *répond*. On lui fournit une vignette clinique ou une question d'examen ; il renvoie un diagnostic ou un paragraphe de conseil, et obtient un bon score. C'est utile, mais limité : un consultant intelligent qui ne touche jamais au dossier.

MIRA est conçu pour faire autre chose, et c'est là que réside l'information importante. Au lieu de répondre à une question isolée, il prend en charge un cas entier : il lit le dossier d'un patient, détermine quelles informations lui manquent, demande des examens de laboratoire, d'imagerie ou de microbiologie, lit les résultats, affine le diagnostic différentiel, puis rédige les ordres qui suivent — prescriptions, admission ou orientation chirurgicale. Il agit *dans* un dossier de santé électronique, à la manière d'un clinicien, plutôt que de produire un texte libre qu'un humain devrait ensuite transcrire.

C'est une véritable étape : passer d'un système qui conseille à un système qui agit dans le flux de travail clinique. C'est aussi exactement le type d'avancée que les titres simplifient en « l'IA dépasse les médecins ». Il faut donc être précis sur ce qui a réellement été testé — et dans quelles conditions.

En une phrase : MIRA a pris en charge seul des cas complets et, sur cette évaluation, a obtenu une meilleure précision diagnostique que des médecins — mais dans un environnement simulé, sur des dossiers rétrospectifs, avec seulement huit diagnostics ciblés, en texte uniquement, et avec un avantage surtout marqué pour les pathologies dont les examens donnent des réponses très nettes. Le résultat est réel. Le score n'est pas la clinique.

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 a démontré une capacité de prise en charge de bout en bout dans un dossier électronique simulé. Il n'a pas démontré un déploiement clinique réel, une couverture diagnostique large ni une comparaison parfaitement équitable, à information égale, avec des médecins.

Original Aurora diagram — The Clean Paper CC BY 4.0

Ce que les auteurs ont fait

MIRA — *Medical Intelligence for Reasoning and Action* — est un agent autonome construit à partir de modèles d'OpenAI : **GPT-4o** assure la partie conversationnelle et les actions, tandis qu'une étape distincte de planification utilise le modèle de raisonnement **o1**. L'ensemble est intégré à une architecture personnalisée conforme au standard HL7 FHIR, largement utilisé par les hôpitaux pour échanger les données de dossiers médicaux, avec une boîte à outils de plus de quatre-vingt mille actions cliniques possibles.

Dans cet environnement simulé, MIRA peut recueillir l'histoire du patient, demander et interpréter des analyses, de l'imagerie et de la microbiologie, générer un diagnostic différentiel et proposer un plan de traitement — prescription, chirurgie ou admission. Détail important : les évaluateurs automatiques utilisés pour noter MIRA sont eux-mêmes basés sur GPT-4o, donc sur la même famille de modèles que celle testée ; une revue par un médecin certifié a ensuite été ajoutée après qu'un évaluateur externe a signalé ce point.

Le jeu de test provient de **MIMIC-IV**, une grande base publique de dossiers de santé désidentifiés. Parmi environ 300 000 patients traités au Beth Israel Deaconess Medical Center entre 2008 et 2019, les auteurs ont retenu **574 cas** couvrant **huit diagnostics cibles** — notamment appendicite, pancréatite, pneumonie et infection urinaire. Pour chaque cas, MIRA partait d'une information initiale limitée et devait décider, étape après étape, de la suite : demander un examen, interroger davantage le patient, interpréter un résultat, puis avancer vers un diagnostic et un plan. La structure ressemble donc à un véritable bilan clinique, mais sur un dossier historique dont l'issue est déjà connue. La performance de MIRA a ensuite été comparée à celle de médecins sur les mêmes cas.

Ce que signifie vraiment « agent autonome dans un dossier électronique simulé »

Trois mots portent ici beaucoup de sens.

Agent signifie que le système n'est pas un chatbot qui répond à une seule consigne. Il fonctionne en boucle : examiner l'état du dossier, choisir une action — demander un test, rechercher une information —, lire le résultat, puis choisir l'étape suivante jusqu'au diagnostic et au plan. Le modèle de langage fournit le raisonnement ; l'architecture de l'agent transforme ses décisions en actions autorisées dans le dossier électronique et lui renvoie les résultats.

Dossier électronique simulé signifie qu'il s'agit d'une copie isolée d'un système clinique, pas d'un hôpital en activité. MIRA peut « demander » un examen et recevoir le résultat historique correspondant parce que le cas provient d'un dossier déjà achevé. Aucune action n'affecte un patient réel.

Autonome signifie que l'agent mène le cas de bout en bout sans intervention humaine *dans cet*

environnement simulé. Cela ne signifie **pas** qu'il a été autorisé à prendre en charge des patients sans supervision. Seule la première affirmation a été testée.

Autre détail important : MIRA peut demander n'importe quel examen, mais le système ne lui renvoie un résultat que si cet examen a réellement été réalisé pour ce patient. Demander un scanner absent du dossier renvoie « N/A », pas une valeur inventée. Il en va de même pour l'interrogatoire : si l'information n'existe pas dans le dossier, le patient simulé ne peut pas la fournir.

Cette distinction est essentielle, car le mot « autonome » est précisément celui qui risque le plus d'être lu comme une affirmation beaucoup plus large que ce que l'étude montre.

Ce qu'ils ont trouvé

Sur cette évaluation, **MIRA a obtenu une meilleure précision diagnostique que les médecins** — mais le titre cache plusieurs chiffres distincts. Sur les **574 cas**, MIRA a identifié correctement le diagnostic de sortie dans **88,9 %** des cas. Pour la comparaison humaine, les médecins n'ont pas traité les 574 dossiers : ils ont travaillé sur un sous-ensemble commun de **311 cas**, avec les mêmes dossiers et les mêmes outils. Sur ces 311 cas, MIRA obtient **87,8 %** de diagnostics corrects.

Deux groupes de médecins ont été recrutés séparément. Un **groupe senior**, composé de quatre médecins certifiés ayant 7 à 11 ans d'expérience, atteint **78,1 %**. Un second groupe, de **niveau d'expérience mixte**, surtout constitué d'internes plus juniors auxquels s'ajoutaient deux spécialistes, atteint **71,1 %**. Autrement dit : mêmes cas et mêmes outils, MIRA en tête, puis les médecins expérimentés, puis l'équipe plus junior.

Les décisions en aval étaient elles aussi généralement conformes aux recommandations : prescriptions, sécurité médicamenteuse et décisions d'admission. Les auteurs ne rapportent aucune erreur médicamenteuse de gravité élevée dans les cinq catégories étudiées, et 467 prescriptions correctes sur 468 — tout en soulignant que le système n'a pas atteint une fiabilité parfaite. Pris isolément, le résultat est impressionnant : un agent peut mener tout un cas et obtenir un meilleur score diagnostique que les groupes humains étudiés.

Mais la *nature* de cet avantage compte autant que son existence. Plusieurs spécialistes indépendants ont souligné où il se concentrait. Dans un commentaire réuni par le Science Media Centre, le Dr Wei Xing a noté que l'avantage de MIRA était **surtout porté par les pathologies dont les examens fournissent des résultats très discriminants**, comme l'appendicite ou la pancréatite, où une imagerie ou une valeur biologique peut fortement orienter le diagnostic. Pour la pneumonie et l'infection urinaire, deux motifs très fréquents aux urgences, l'IA et les médecins obtenaient des résultats moins bons et l'écart entre eux était plus faible.

Il existait aussi une asymétrie d'information. MIRA utilisait davantage de données de laboratoire : environ **51 %** des analytes disponibles en pratique courante, contre environ **28 %** pour les médecins certifiés, soit une médiane de sept paramètres sanguins supplémentaires par cas. Les auteurs précisent que cela reste inférieur au niveau d'utilisation des examens observé dans l'ensemble historique

MIMIC-IV et qu'ils ne voient pas d'augmentation systématique de l'imagerie coûteuse. Mais davantage d'informations peuvent naturellement améliorer la précision : la comparaison n'est donc pas parfaitement à armes égales.

Pourquoi « battre les médecins » ne signifie pas « être un meilleur médecin »

Un meilleur score sur une évaluation peut être surinterprété très facilement. Trois éléments séparent « MIRA obtient un score plus élevé » de « MIRA est un meilleur médecin ».

D'abord, **la référence utilisée pour noter le système**. Julie Jacko a relevé que plusieurs critères évaluent MIRA par rapport à ce qui avait été documenté dans le dossier historique. Le système est donc, en partie, récompensé pour **reproduire le comportement clinique enregistré**, pas nécessairement pour démontrer le meilleur soin possible dans l'absolu. Le dossier devient la clé de correction. Cette remarque concerne surtout les mesures d'alignement du traitement ; la précision diagnostique, elle, est comparée au diagnostic de sortie et constitue une référence plus proche d'un résultat clinique réel.

Ensuite, **l'asymétrie d'information** : utiliser davantage d'analyses sanguines signifie disposer de davantage d'indices. Une partie de l'écart peut donc provenir de l'accès à plus d'information, pas uniquement d'un meilleur raisonnement.

Enfin, **le contexte d'exécution**. Les cas étaient rétrospectifs, limités à huit diagnostics, traités dans un environnement simulé et uniquement sous forme textuelle. Une consultation réelle comprend aussi la manière dont le patient parle, l'examen physique, le comportement observé et le langage corporel — autant d'éléments absents d'un agent qui lit un dossier déjà constitué. Et un patient dont le problème ne correspond pas aux huit diagnostics sélectionnés est simplement hors du périmètre de cette étude.

Il existe également une inquiétude plus discrète : **la contamination des données d'entraînement**. MIMIC-IV est public et largement discuté. Un modèle entraîné sur de grandes quantités de données du Web pourrait avoir rencontré des articles, des discussions de cas ou des fragments liés à cette base. Si certaines réponses ont été vues pendant l'entraînement, une partie de la performance pourrait refléter de la mémorisation plutôt que du raisonnement clinique. Les auteurs n'ignorent pas ce risque : ils présentent leurs chiffres avec prudence comme une possible borne supérieure de performance et reconnaissent que la généralisation à d'autres données publiques pourrait être surestimée.

Aucune de ces réserves ne rend MIRA inintéressant. Elles définissent simplement la lecture correcte : sur un ensemble rétrospectif, un agent capable d'agir dans l'ensemble du dossier obtient de très bons résultats, en particulier pour des diagnostics soutenus par des examens nets, en partie en demandant davantage de tests et en partie en reproduisant les décisions du dossier historique. C'est une démonstration réelle de capacité, pas un verdict selon lequel les machines diagnostiquent mieux que les médecins.

Ce que cela ne prouve pas

- L'étude ne montre **pas** que MIRA fonctionne sur de vrais patients. Tous les cas sont rétrospectifs et simulés.
- Elle ne montre **pas** qu'il fonctionne au-delà des huit diagnostics retenus.
- Elle ne montre **pas** qu'il est sûr à déployer. Généralisation, sécurité et gouvernance nécessitent encore des études prospectives sur des patients réels.
- Elle n'établit **pas** une comparaison parfaitement équitable avec les médecins : MIRA utilise davantage de données biologiques et plusieurs critères de traitement sont notés par rapport au dossier historique.
- Elle ne démontre **pas** l'absence de mémorisation liée à MIMIC-IV.
- Elle ne signifie **pas** que « l'IA remplace les médecins ». Même un agent capable d'agir dans le dossier reste un outil d'aide dont le déploiement réel impliquerait des cliniciens responsables et tout ce qu'un dossier textuel ne contient pas.

Quelle est la solidité de la preuve ?

Il faut distinguer deux affirmations.

Comme démonstration de capacité — un agent fondé sur un modèle de langage qui mène un cas clinique entier dans un dossier électronique simulé, en enchaînant interrogatoire, examens, diagnostic et décisions — le travail est réellement nouveau et assez convaincant. C'est la partie la plus importante de l'étude.

Comme affirmation de supériorité — MIRA serait *meilleur que les médecins* — la preuve est beaucoup plus limitée. Elle vaut pour un ensemble rétrospectif précis, huit diagnostics, une asymétrie d'information, une clé de correction issue du dossier historique et un risque de contamination des données d'entraînement. C'est suffisant pour dire que l'agent a très bien réussi cette évaluation. Ce n'est pas suffisant pour conclure qu'il est un meilleur diagnosticien qu'un médecin en pratique clinique.

La position utile n'est donc ni « l'IA bat les médecins » ni « ce n'est qu'un jouet ». Elle est plus précise : une nouvelle génération d'IA médicale, capable d'*agir* dans le dossier plutôt que de simplement répondre à des questions, obtient de très bons résultats sur une évaluation rétrospective exigeante et doit maintenant être testée là où elle ne l'a pas encore été — prospectivement, sur des patients réels et non sélectionnés, sous une gouvernance appropriée.

Pourquoi c'est important

Depuis quelques années, la question la plus intéressante en IA médicale a changé. Il ne s'agit plus seulement de demander si un modèle peut trouver le bon diagnostic à partir d'une vignette. MIRA

pose une question plus difficile : un modèle peut-il *effectuer le travail* — recueillir les informations, demander des examens, interpréter les résultats, décider puis agir — sur l'ensemble d'un cas ? C'est une question plus utile et plus exigeante, car l'action concentre à la fois la difficulté et le risque.

C'est pourquoi le cadrage compte. Le réflexe du titre — *l'IA dépasse les médecins* — met l'accent sur la partie la moins nouvelle et la plus facile à surinterpréter. La nouveauté réelle est plus sobre mais plus importante : un agent capable de parcourir tout le dossier. Si cette capacité tient dans des essais prospectifs, elle pourrait transformer le **flux de travail clinique** bien avant de modifier la question de la responsabilité. Ce sont d'abord les demandes d'examens, l'interprétation des résultats et la coordination du dossier qui changeraient.

Cela remet aussi la charge de la preuve au bon endroit. Un classement sur des cas rétrospectifs justifie un essai prospectif ; il ne le remplace pas. La lecture rigoureuse de MIRA est donc une invitation à cette prochaine étude, menée sur de vrais patients plutôt que sur un ensemble de référence.

Résumé clair

MIRA est un agent d'IA autonome capable d'agir dans un dossier de santé électronique simulé : recueillir des informations, demander et interpréter des examens, poser un diagnostic et rédiger des ordres de traitement. Testé sur 574 cas rétrospectifs de MIMIC-IV couvrant huit diagnostics ciblés, il obtient une meilleure précision diagnostique que les groupes de médecins comparés et produit des décisions généralement conformes aux recommandations. Mais l'évaluation reste une simulation sur dossiers historiques, en texte seulement ; l'avantage est particulièrement marqué pour les maladies dont les examens sont très discriminants ; MIRA utilise davantage d'analyses sanguines que les médecins ; plusieurs critères sont comparés à ce que le dossier original contenait ; et l'utilisation d'une base publique crée un risque de contamination des données d'entraînement. L'avancée réelle est celle d'un agent capable d'agir dans tout le dossier. Ce n'est pas la preuve que l'IA diagnostique mieux que les médecins en clinique, ni qu'un tel système est prêt à être déployé.

Mise au point

Ce que l'article montre : un agent fondé sur un modèle de langage peut mener un cas clinique complet dans un dossier électronique simulé — recueil d'informations, examens, diagnostic et ordres — et, sur cette évaluation rétrospective de 574 cas couvrant huit diagnostics, obtenir une précision diagnostique supérieure à celle des groupes de médecins comparés, sans erreur médicamenteuse de gravité élevée rapportée dans les catégories étudiées.

Ce qui est plausible mais non démontré : que cette capacité bénéficierait à de vrais patients non sélectionnés ; que l'avantage diagnostique reflète surtout un meilleur raisonnement plutôt que davantage d'examens, un meilleur alignement sur le dossier historique ou une éventuelle mémorisation.

Ce que cela ne montre pas : fonctionnement hors des huit diagnostics ; sécurité d'un déploiement

réel ; comparaison parfaitement équitable à information égale ; remplacement des cliniciens ; absence de contamination des données d'entraînement.

Principales limites : évaluation rétrospective et simulée, uniquement textuelle ; huit diagnostics ; asymétrie d'information ; certains traitements notés par rapport au dossier historique ; base publique susceptible d'avoir été rencontrée pendant l'entraînement.

Niveau de confiance pour un lecteur général : élevé sur le fait que MIRA constitue une véritable démonstration nouvelle de capacité — un agent qui *agit* dans le dossier plutôt qu'un simple chatbot. Faible sur l'idée que l'étude prouve qu'il est *un meilleur diagnosticien qu'un médecin*. Des essais prospectifs sur des patients réels sont nécessaires avant de tirer une conclusion clinique.

Sources

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>