

VEILLE INTELLIGENCE ARTIFICIELLE

Radiologie et imagerie medicale

Semaine 32 | 2 au 8 aout 2026

Dr Sergey Morozov | drsergeymorozov.com

Compile a partir de sources publiquement disponibles (publications scientifiques indexees, communiqués officiels, bases de données réglementaires).
Ne constitue pas un avis médical, réglementaire ou financier.

8

Articles evalues par les pairs

4

Annonces secteur / FDA

3

Points mediatiques

3

Leaders d'opinion

SECTION 0

RESUME EXECUTIF

Tendances cles, semaine du 2 au 8 aout 2026

[BENCHMARK LLM] LES LLM MULTIMODAUX PLAFONNENT SOUS L'UTILITE CLINIQUE EN RADIOLOGIE

Wu Q et al. (Journal of Medical Internet Research) ont construit RadM-Bench, 720 cas bilingues dans 9 surspecialites, et note 10 LLM multimodaux de 0 a 3. Tous restent sous 1,5. Les images cles 2D choisies par un radiologue aident chaque modele (+19,8% a +139,2%), mais l'apport volumique degrade tous les modeles evaluables, et la performance sur les cas d'enseignement ne se transfere pas aux cas cliniques chinois de routine (O3 avec images chute de 1,14 a 0,79). Traduire les anamneses cliniques en anglais degrade les 10 modeles.

[PREDICTION CLINIQUE] UN MODELE TRANSFUSIONNEL QUI CONNAIT SA PROPRE INCERTITUDE

Li X et al. (JMIR Medical Informatics) ont developpe MCGB, un cadre de gradient boosting sous contraintes cliniques qui predit le besoin et la dose de transfusion chez 849 patients avec hemorragie digestive haute dans 3 hopitaux. AUROC 0,97, sensibilite 0,99, specificite 0,87 sur un hopital entierement exclu du developpement, avec des doses assorties d'intervalles de prediction conformes a 95% (couverture 0,94). Contraintes monotones et incertitude quantifiee repondent aux deux objections standard au ML au lit du patient.

[BIAS] LA DIVERSITE DES DONNEES BAT LEUR VOLUME EN SEGMENTATION DU GLIOBLASTOME

Lee RS et al. (AJNR) ont audite quatre modeles nnUNet de segmentation du glioblastome sur 480 patients independants. Le modele entraine uniquement sur des hommes blancs non hispaniques obtient les scores les plus bas (Dice 0,943 FLAIR, 0,909 rehaussement) avec des biais d'age et de statut tabagique ; le modele BraTS 2024, de taille moderee mais divers, obtient les meilleurs scores (0,996, 0,999) et bat le modele FeTS pourtant bien plus grand. L'heterogeneite reduit le biais meme a taille d'entrainement identique.

[OPERATIONS] DES MODELES DE LANGAGE TROUVENT L'ARGENT QUI FUT DES NOTES D'ECHOGRAPHIE

Nguyen K et al. (Journal of Medical Internet Research) ont utilise BioClinBERT (exactitude 0,97) pour detecter les echographies au point de service dans 559 029 consultations d'obstetrique-gynecologie, puis montre qu'un modele de documentation standardise fait passer la recuperation de charges manquees de 10,0% a 2,4% (OR 0,22, P<0,001), 75,4% des codes de facturation provenant automatiquement des modeles. L'IA en auditeur du cycle de revenus, avec verification que le gain vient du flux de travail et non de plus d'actes.

[POINT DE VIGILANCE] CETTE SEMAINE SE JOUE AU POSTE DE LECTURE, PAS DANS LE MODELE

Lus ensemble, les articles de la semaine déplacent la frontière de l'exactitude vers l'implémentation. Wu Q et al. montrent des LLM multimodaux encore sous le niveau cliniquement exploitable des que les cas cessent d'être sélectionnés. Zhang R et al. constatent que la stratégie de prompt doit être adaptée à l'architecture du modèle, et que modes de raisonnement et nombre de paramètres peuvent retirer de la valeur en classification de comptes rendus. Moturu A et al. observent deux radiologues utiliser une surcouche IRM corps entier pédiatrique : l'outil les ralentit tous les deux et modifie leur marquage dans des directions opposées, d'où la conclusion que les outils d'IA pourraient nécessiter une calibration par lecteur. Même le résultat clinique le plus solide, le modèle transfusionnel MCGB de Li X et al., tire sa crédibilité des contraintes, de la calibration et du test inter-sites plutôt que de la seule discrimination. La question d'achat de la semaine n'est pas ce que score un modèle, mais ce qui arrive aux humains, au flux de travail et à la facturation quand il arrive.

SECTION 1

ARTICLES EVALUES PAR LES PAIRS

Source : PubMed | Indexation vérifiée | Aucun preprint | 2 au 8 août 2026

Prehospital Injury Severity Estimate (PHISE) matches in-hospital trauma scores when embedded in AI models.

Sigle M, et al. | *NPJ digital medicine* | 2026-08-07

PHISE, un score préhospitalier de gravité lésionnelle construit en projetant les descriptions lésionnelles codées AIS de la National Trauma Data Bank américaine sur huit régions corporelles et quatre niveaux de gravité, avec une inférence assistée par LLM pour identifier les lésions nécessitant une imagerie et celles évaluables cliniquement. L'objectif est une composante anatomique de gravité utilisable avant tout scanner, conçue pour être intégrée dans des modèles d'apprentissage automatique avec les autres variables disponibles en préhospitalier.

Indicateurs clés: L'abstract rapporte les comparaisons de manière qualitative, sans valeurs numériques : PHISE montre une forte corrélation et une bonne calibration avec ISS et NISS dans la cohorte NTDB, mais une performance prédictive inférieure pour la mortalité hospitalière et le besoin transfusionnel. En validation externe sur le TraumaRegister DGU multinational, il fait moins bien que ISS/NISS pour la mortalité, aussi bien pour la prédiction transfusionnelle, et l'écart se réduit encore quand PHISE est intégré dans des modèles ML avec les variables préhospitalières.

Pertinence clinique: Les décisions de triage traumatologique se prennent avant toute imagerie, alors que les scores de référence en dépendent tous. Un score anatomique indépendant de l'imagerie, conçu d'emblée comme variable d'entrée pour des modèles d'IA et valide en externe sur un registre d'un autre système de traumatologie, ouvre une voie crédible vers une stratification plus précoce et une activation plus intelligente des ressources d'imagerie et de transfusion.

PMID 42562846 DOI 10.1038/s41746-026-03074-7

Prediction of Blood Transfusion Need and Dose in Patients With Upper Gastrointestinal Bleeding: Retrospective Multicenter Prediction Model Study.

Li X, et al. | *JMIR medical informatics* | 2026-08-04

Un cadre de gradient boosting en deux étapes sous contraintes cliniques (MCGB) qui prédit d'abord si un patient avec hémorragie digestive haute confirmée par endoscopie aura besoin d'une transfusion, puis en estime la dose avec des intervalles de prédiction conformes à 95%. Développé sur une cohorte rétrospective multicentrique de 849 adultes de 3 hôpitaux de Chongqing (2019 à 2025), avec contraintes cliniques monotones, test sur site entièrement exclu du développement et prototype d'interface au lit du patient.

Indicateurs clés: AUROC 0,97 et AUPRC 0,91 pour le besoin transfusionnel ; au seuil de 0,50, sensibilité 0,99, spécificité 0,87, F1 0,85. La prédiction de dose chez les patients transfusés atteint un R2 de 0,95 avec une erreur absolue moyenne de 0,04, et une couverture des intervalles de prédiction à 95% de 0,94. L'évaluation a utilisé un hôpital complet comme site de test indépendant, avec un test externe alterné entre hôpitaux comme contrôle de robustesse.

Pertinence clinique: Les seuils transfusionnels dans l'hémorragie digestive haute restent débattus et les seuils fixes d'hémoglobine ignorent le patient individuel. Un modèle calibre qui quantifie sa propre incertitude, respecte la monotonie clinique et a été testé entre sites répond aux deux reproches chroniques faits à l'aide à la décision par ML, l'opacité et l'excès de confiance, avec des implications directes pour la banque de sang comme pour la décision au lit du patient.

PMID 42550984 DOI 10.2196/83889

Evaluating Sociodemographic Biases in Artificial Intelligence-Based Glioblastoma Response Assessment Algorithms.

Lee RS, et al. | *AJNR. American journal of neuroradiology* | 2026-08-03

Un audit de biais de quatre modèles nnUNet de segmentation IRM du glioblastome dont les jeux d'entraînement diffèrent par la taille et la composition démographique : le modèle postopératoire FeTS (plus de 10 000 examens, multinational), le modèle BraTS 2024 (plus de 2 000 examens, nord-américain) et deux petits modèles monocentriques (plus de 200 examens chacun), l'un démographiquement homogène, l'autre hétérogène. Tous ont été évalués sur un jeu indépendant, corrigé manuellement, de 480 patients recrutés prospectivement.

Indicateurs clés : Le modèle entraîné exclusivement sur des hommes blancs non hispaniques a obtenu les scores de Dice les plus bas (0,943 FLAIR, 0,909 pour la prise de contraste) avec des biais liés à l'âge et au statut tabagique. Le modèle BraTS a obtenu les scores les plus élevés (0,996 FLAIR, 0,999 rehaussement) et le moins de biais, surpassant le modèle FeTS malgré un jeu d'entraînement bien plus petit. Les effets sociodémographiques ont été analysés par régression bêta.

Pertinence clinique : Deux résultats comptent pour qui constitue des jeux d'entraînement : l'hétérogénéité démographique réduit le biais même sans augmenter la taille du jeu de données, et une cohorte modérée mais diverse bat une cohorte beaucoup plus grande. Le biais des modèles de segmentation est globalement faible mais concentre là où les données d'entraînement sont homogènes, un argument pour auditer la composition, et pas seulement le volume, dans les achats et la validation.

PMID 41663205 DOI 10.3174/ajnr.A9217

Machine Learning to Identify Point-of-Care Ultrasound and Evaluate Standardized Documentation: Retrospective Operational Cohort Study.

Nguyen K, et al. | *Journal of medical Internet research* | 2026-08-07

Une étude opérationnelle dans un grand centre académique qui a utilisé le ML (LightGBM et BioClinBERT) pour détecter les échographies au point de service enfouies dans les notes libres d'obstétrique-gynécologie, puis a mesuré l'effet d'un modèle standardisé de documentation des procédures (ProcDoc) sur la capture de la facturation. L'analyse couvre 559 029 consultations de 109 776 patientes dans 11 sites de 2018 à 2024, le modèle ayant été introduit en février 2023.

Indicateurs clés : BioClinBERT atteint une exactitude de 0,97 (F1 0,55 à 0,63) pour identifier les procédures documentées et manquantes. L'adoption de ProcDoc a atteint 75,1% en 12 mois. La récupération de charges manquantes est passée de 10,0% avant intervention à 2,4% après (odds ratio 0,22, IC 95% 0,17 à 0,30, $P < .001$), avec une facturation POCUS globale en hausse de 0,6%. 1 812 des 2 404 codes CPT post-intervention (75,4%) provenaient des modèles standardisés.

Pertinence clinique : L'IA sert ici d'instrument d'audit plutôt que de diagnostic : les modèles ont quantifié la fuite de revenus liée aux échographies non documentées, puis vérifié que le gain venait du changement de flux de travail et non d'une hausse des volumes. Pour les directions d'imagerie, la leçon transférable est le couple : un modèle de langage pour mesurer le problème dans le texte libre, plus une intervention de documentation structurée pour le corriger.

PMID 42566781 DOI 10.2196/84454

A Bilingual Benchmark for Evaluating Diagnostic Performance of Multimodal Large Language Models in Radiology (RadM-Bench): Evaluation Development and Validation.

Wu Q, et al. | *Journal of medical Internet research* | 2026-08-07

RadM-Bench, un benchmark radiologique bilingue de 720 cas répartis sur 9 surspécialités : 360 cas d'enseignement publics en anglais enrichis en maladies rares et 360 cas cliniques courants en chinois. Dix LLM multimodaux (4 propriétaires dont GPT-4o, O3 et des variantes Gemini ; 6 open source) ont été testés avec l'anamnèse seule, des images clés 2D choisies par un radiologue, et un apport volumique à 2 et 10 images par seconde, notes de 0 à 3 par deux radiologues certifiées en aveugle.

Indicateurs clés : Les scores moyens restent sous 1,5 sur 3 pour les 10 modèles sur les deux jeux de données. L'ajout d'images clés 2D améliore chaque modèle (+19,8% à +139,2%), mais l'apport volumique à 10 ips dégrade les 8 modèles évaluables par rapport à la référence 2D (-5,3% à -31,4%). Les modèles propriétaires reculent des cas d'enseignement aux cas cliniques (O3 avec images : 1,14 à 0,79) tandis que les modèles open source sinocentres progressent. Traduire les anamneses chinoises en anglais abaisse les scores des 10 modèles, significativement pour 9.

Pertinence clinique : Le benchmark isole exactement les écarts qu'un déploiement clinique rencontrerait : les modèles ne savent pas encore exploiter la 3D, la performance sur du matériel d'enseignement sélectionné surestime la performance en routine, et l'avantage sur les maladies rares observé sur les jeux pédagogiques s'inverse en clinique. Le résultat linguistique contredit l'hypothèse qu'une entrée en anglais est toujours optimale. Une référence utile avant tout pilote de LLM multimodal en compte rendu ou en triage.

PMID 42566748 DOI 10.2196/92183

Large Language Models for Classifying Usual Interstitial Pneumonia from Radiology Reports: Native Reasoning Versus Structured Prompting.

Zhang R, et al. | *Journal of imaging informatics in medicine* | 2026-08-05

Une comparaison systematique de 10 LLM open source des familles Llama, Qwen et Gemma (8 a 405 milliards de parametres) pour classer les patterns de pneumopathie interstitielle commune (PIC) a partir de 270 comptes rendus de TDM haute resolution selon les criteres Fleischner, avec le consensus de deux radiologues thoraciques seniors comme reference. Trois strategies de prompting ont ete croisees avec l'architecture, et quatre modeles a raisonnement natif ont aussi ete testes en mode reflexion.

Indicateurs clés: La meilleure configuration atteint un kappa de Cohen de 0,70 et 82% d'exactitude en quatre classes. Le prompting a raisonnement structure ameliore tous les modeles Llama mais degrade tous les modeles a capacite de raisonnement natif (Qwen 3.5, Gemma 4). Le mode reflexion nuit aux prompts fondees sur des criteres et ne produit jamais la meilleure configuration. Llama 3.1 405B n'apporte aucun avantage sur Llama 3.3 70B, et le Qwen de 397B fait moins bien que le modele dense de 27B.

Pertinence clinique: Pour les pipelines d'exploitation des comptes rendus, les conclusions operationnelles sont que le prompt doit etre adapte a l'architecture, que les modes de raisonnement peuvent activement nuire a une classification par criteres, et que le nombre de parametres est un mauvais critere d'achat. Les equipes qui etiquettent des comptes rendus pour entrainer des modeles peuvent obtenir l'etat de l'art avec des modeles ouverts de taille moyenne, a condition de tester le couple prompt-architecture plutot que de supposer que plus recent et plus gros font mieux.

PMID 42557432 DOI 10.1007/s10278-026-02178-6

Automatic Patient Eligibility for Photon-counting CT Using Discriminative and Generative AI Models in Neuroradiology.

Martin-Noguerol T, et al. | *Academic radiology* | 2026-08-05

Une preuve de concept de routage NLP qui lit les demandes de scanner neuroradiologique en espagnol et decide si l'examen requiert un protocole complet de scanner a comptage photonique (PCCT) ou peut aller sur un scanner conventionnel a detecteur a integration d'energie. 800 demandes (2012 a 2025) ont ete etiquetees par deux neuroradiologues (kappa de Cohen 0,74) ; des transformers discriminatifs affines ont ete compares a des modeles generatifs en zero-shot (ChatGPT 5.2, Gemini 3) en validation croisee a cinq plis.

Indicateurs clés: RoBERTa biomedical arrive en tete avec une exactitude de 0,925, une precision de 0,906, un rappel de 0,967, un F1 de 0,935 et une AUC de 0,920, sans difference significative avec le BERT espagnol mais significativement meilleur que Llama-3.1-8B, Mistral-7B et les deux LLM generatifs, qui obtiennent les moins bonnes performances globales.

Pertinence clinique: La capacite PCCT est rare et aucune recommandation ne definit quand la privilegier : un routeur d'eligibilite automatique repond donc a un vrai probleme d'allocation de ressources. Le resultat porte aussi un message plus large pour l'informatique radiologique : de petits modeles discriminatifs affines battent encore les LLM generatifs generalistes sur une classification bien definie, pour une fraction du cout et de la charge de gouvernance.

PMID 42557204 DOI 10.1016/j.acra.2026.07.026

Human-AI Interaction With AI-Assisted Tumor Overlays in Pediatric Whole-Body Magnetic Resonance Imaging: Exploratory Reader Study.

Moturu A, et al. | *JMIR human factors* | 2026-08-07

Une etude exploratoire d'interaction humain-IA sur une surcouche de segmentation par patches qui met en evidence les regions d'allure tumorale sur l'IRM corps entier de surveillance pediatrique dans le syndrome de Li-Fraumeni, entraine sur des coupes 2D augmentees issues de 675 volumes d'IRM corps entier. Deux radiologues ont lu les cas avec et sans surcouche pendant que le temps d'evaluation, le marquage des lesions candidates, les recommandations de suivi et la charge percue etaient enregistres.

Indicateurs cles: L'abstract rapporte des resultats directionnels sans valeurs numeriques : le temps moyen d'evaluation par cas augmente avec l'outil pour les deux lecteurs ; un lecteur marque plus de localisations candidates avec l'outil et l'autre moins ; un lecteur rapporte une demande mentale plus faible et les deux un stress reduit. La variabilite inter-lecteurs est manifeste tout au long de l'etude.

Pertinence clinique: Un temoignage lucide sur ce que fait reellement l'assistance IA au poste de lecture : elle a modifie les comportements dans des directions propres a chaque lecteur et a ralenti les deux, tout en reduisant le stress. La conclusion des auteurs, des outils d'IA a calibrer individuellement par radiologue, remet en cause le deploiement uniforme, et la surveillance pediatrique des maladies rares est precisement le contexte ou une seconde paire d'yeux a les enjeux les plus eleves.

PMID 42567688 DOI 10.2196/81066

SECTION 2

SECTEUR ET REGLEMENTATION

Sources : Registres officiels | Communiques de presse | Bases FDA

[MARKET] GE HealthCare lance Invenia ABUS Prime et le viewer ABUS StreamVue autorise par la FDA

GE HealthCare | 3 aout 2026

GE HealthCare a annonce le lancement d'Invenia ABUS Prime, son nouveau systeme d'echographie mammaire automatisee pour le depistage complementaire des seins denses, avec ABUS StreamVue, un viewer d'entreprise en navigateur, sans empreinte locale, qui vient de recevoir l'autorisation FDA 510(k) et le marquage CE. ABUS Prime embarque l'evaluation de qualite d'acquisition par IA, la detection automatique du mamelon et l'aide a la lecture QVCAD ; selon les documents de l'entreprise, la revue assistee par IA atteint 93% de sensibilite et peut reduire les temps d'interpretation jusqu'a 33%, Fast Scan reduisant les temps d'acquisition jusqu'a 40%. StreamVue, developpe avec Candelis, vise la lecture a distance et la standardisation multi-sites ; l'une des premieres installations americaines d'Invenia ABUS Prime est chez BeSound a Los Angeles.

Source: GE HealthCare

[REGULATION] Cortechs.ai etend son marquage CE MDR a l'ensemble de sa plateforme d'imagerie quantitative

Cortechs.ai | 4 aout 2026

Cortechs.ai a annonce l'extension de son marquage CE sous MDR europeen pour ses derniers produits : la version 5 de NeuroQuant (dont NeuroQuant MS et NeuroQuant ARIA), NeuroQuant Brain Tumor et OnQ Prostate. La certification permet la commercialisation de la generation actuelle de ses outils de volumetrie cerebrale automatisee, de suivi des tumeurs cerebrales et d'imagerie prostatique fondee sur la RSI dans tout l'Espace economique europeen. L'entreprise presente cette etape comme le socle de son deploiement europeen, apres plus d'un an de travail reglementaire sous MDR.

Source: PR Newswire

[PARTNERSHIP] 4DMedical investit 3,4 M\$ dans RevealDx et prend la distribution mondiale de RevealAI-Lung

4DMedical | 3 aout 2026

4DMedical a annonce un investissement strategique de 3,4 M\$ dans RevealDx (Seattle) et un accord de distribution couvrant les Etats-Unis, l'Europe, l'Australie et la Nouvelle-Zelande pour RevealAI-Lung, un outil CADx qui caracterise les nodules pulmonaires avec un Malignancy Similarity Index (mSI). RevealAI-Lung est autorise par la FDA, certifie MDR europeen et approuve TGA, valide chez plus de 1 500 patients, et rembourse aux Etats-Unis sous les codes CPT 0721T et 0722T. 4DMedical l'integrera a la plateforme de scanner thoracique contextflow recemment acquise et l'associera a son imagerie fonctionnelle CT:VQ, assemblant un ecosysteme cardiopulmonaire unique couvrant detection, stratification du risque et fonction.

Source: 4DMedical

[MARKET] GE HealthCare presente la famille LOGIQ e d'echographes compacts dotes d'IA

GE HealthCare | 4 aout 2026

GE HealthCare a presente les LOGIQ e Xi et LOGIQ e Si, des echographes au format portable combinant des performances de console et des fonctions d'IA pour l'imagerie au point de service et multi-specialites : etiquetage anatomique Nerveblox pour les blocs nerveux, identification des structures MSK Go, Automated Function Imaging pour l'evaluation cardiaque et biometrie foetale SonoBiometry. Les outils Scan Assistant et Voice Control reduisent les interactions manuelles, et le Xi prend en charge la gestion de flotte cloud Verisound et les sondes Vscan Air. Le lancement etend le segment de l'imagerie portable dotee d'IA aux urgences, aux soins critiques, au MSK et a l'anesthesie.

Source: MPO Magazine

SECTION 3

POINTS MEDIATIQUES

AuntMinnie | Radiology Business | The Imaging Wire | Diagnostic Imaging | ITN

Les LLM peuvent-ils combler le deficit de communication avec les prescripteurs ?

AuntMinnie.com | 4 aout 2026

Couverture d'une etude de l'UCSF publiee le 4 aout dans Radiology montrant que des LLM peuvent generer des indications d'imagerie plus completes et cliniquement plus utiles que celles fournies par les prescripteurs. A partir de 28 313 patients, 740 867 notes cliniques et 77 626 examens, l'equipe a compare des modeles proprietaires et open source ; dans une etude de lecture avec 20 radiologues, Claude 3.5 Sonnet obtient la meilleure note d'exhaustivite dans 37,14% des cas et de factualite dans 68,05%, contre 6,64% et 50% pour les prescripteurs, le modele open source Qwen 2.5-7B suivant de pres. Un editorial associe insiste sur la transparence, la surveillance et le controle humain avant tout deploiement.

Lien: <https://www.auntminnie.com/imaging-informatics/artificial-intelligence/article/1...>

Konica Minolta etend sa plateforme Exa Enterprise avec teleradiologie cloud et partenaires IA

Radiology Business | 3 aout 2026

Un reportage sur l'extension par Konica Minolta Medical Imaging USA de son PACS Exa Enterprise en plateforme de teleradiologie cloud assemblee a partir de partenaires plutot que de developpements internes : orchestration des flux NewVue, comptes rendus par IA generative RADPAIR, infrastructure cloud AWS Health Imaging et VNA Apollo. La plateforme en navigateur est facturee par examen et s'adapte du cabinet individuel aux systemes hospitaliers de 500 lits, illustrant la tendance a la plateformeisation ou les editeurs de PACS integrent l'IA de tiers plutot que de la developper.

Lien: <https://radiologybusiness.com/topics/health-it/enterprise-imaging/konica-minolta...>

Revue de la semaine : les delais de compte rendu continuent de s'allonger

[AuntMinnie.com](#) | 7 aout 2026

La note hebdomadaire de la redaction d'AuntMinnie met en avant les resultats du Harvey L. Neiman Health Policy Institute : les delais de restitution des comptes rendus radiologiques pour les patients Medicare ont continue d'augmenter en 2024, quoique plus lentement qu'en 2023, sur fond de penurie de radiologues et de volumes croissants. Elle signale aussi un pilote du UMass Memorial Health orientant les patients urgents de medecine de premier recours directement vers la radiologie plutot que vers les urgences, reduisant de pres de 50% le delai entre l'arrivee et les resultats d'imagerie, un resultat organisationnel pertinent pour l'argument de capacite qui porte l'adoption de l'IA en radiologie.

Lien: <https://www.auntminnie.com/editors-note/article/15831989/week-in-review-imaging-...>

SECTION 4

SOCIETES SAVANTES

SIIM | ACR | RCR | Publications et couverture officielles

La RSNA lance le Knee Abnormality Detection AI Challenge 2026

5 aout 2026

La RSNA a publie le communique officiel de son AI Challenge 2026, le premier centre sur l'IRM musculosquelettique et le premier a combiner images et texte des comptes rendus radiologiques. Le jeu d'entrainement couvre plus de 5 000 IRM du genou avec des comptes rendus en 12 langues provenant de 16 sites sur cinq continents, le jeu d'evaluation etant annote par des radiologues experts. La competition, hebergee sur Kaggle, court jusqu'au 22 octobre, avec 77 000 \$ de prix dont une premiere recompense pour les modeles les plus efficaces ; les laureats seront distingues au congres RSNA 2026 a Chicago (29 novembre au 3 decembre).

Source: <https://www.rsna.org/media/press/2026/2669>

SECTION 5

LEADERS D'OPINION

LinkedIn | Declarations publiques | Semaine du 2 au 8 aout 2026

Curt Langlotz

Professeur de radiologie, de medecine et de science des donnees biomedicales, Stanford University School of Medicine

A relaye (7 aout) la publication de MedVAL dans npj Digital Medicine, un article de Stanford dont il est coauteur : Toward expert-level medical text validation with language models, publie le 4 aout. Le travail s'attaque a un goulot pratique de l'IA generative en radiologie, verifier qu'un texte medical genere par IA est effectivement conforme a sa source, et propose les modeles de langage eux-memes comme validateurs a grande echelle. Langlotz a aussi publie cette semaine sur sa rencontre avec l'equipe Raidium a Palo Alto, en reprenant sa formule de 2017 : l'IA ne remplacera pas les radiologues, mais les radiologues qui l'utilisent remplaceront ceux qui ne l'utilisent pas.

Source: [LinkedIn](#)

Woojin Kim

Chief Strategy Officer et CMIO de HOPPR, CMO du ACR Data Science Institute, radiologue MSK

A publie le 6 aout qu'il participera a EuSoMII 2026 a Heraklion, en Crete (9 au 10 octobre) pour représenter le ACR Data Science Institute, en relayant la campagne de l'institut sur sa contribution a l'avenir de l'IA en radiologie. Plus tot dans la semaine (2 aout), il a raconte son intervention sur l'entrepreneuriat et l'IA en radiologie dans le cours Healthcare Innovation and Transformation de Geraldine McGinty a Weill Cornell, signe que le leadership de l'IA radiologique entre dans la formation des cadres de sante.

Source: [LinkedIn](#)

Amine Korchi

Radiologue et entrepreneur healthtech, Geneve ; coauteur de la revue VLM de neuroradiologie des HUG

A publie le 7 aout au sujet d'une revue systematique fraichement parue, menee avec l'equipe de neuroradiologie des Hopitaux universitaires de Geneve (HUG), sur les modeles de fondation vision-langage pour l'interpretation neuroradiologique 3D. Sa synthese des preuves : les VLM specialises en neuro-imagerie surpassent systematiquement les modeles multimodaux generalistes et rapprochent la redaction assistee de comptes rendus de la realite clinique, mais ils ne sont pas encore prêts pour une adoption en routine. Parmi les lacunes citees : trop de corrections de comptes rendus pour degager un vrai gain de flux, des besoins massifs en calcul, l'absence de metriques d'evaluation cliniquement pertinentes, hallucinations et lesions manquees, l'integration au flux de travail et la maturite reglementaire. Pour les createurs d'IA, il presente ces limites comme la feuille de route sur laquelle se construira la prochaine generation d'entreprises d'IA en radiologie.

Source: [LinkedIn](#)

NOTE D'ASSURANCE QUALITE

Compile par le Dr Sergey Morozov a partir de sources publiquement disponibles : articles evalues par les pairs indexes dans PubMed (PMID et DOI verifiees, aucun preprint ni arXiv), communiques officiels, bases reglementaires et medias specialises. Ne constitue pas un avis medical, reglementaire ou financier. Tous les articles ont des dates verifiees dans la fenetre du 2 au 8 aout 2026. La selection est deterministe (requete PubMed fixe, liste de revues Q1/phares, exclusion des revues pures [meta-analyses conservees], de la radiomique seule et de la radiologie interventionnelle, puis un score de classement transparent). Les annonces secteur proviennent de communiques officiels ; les items KOL de l'activite publique LinkedIn dans la fenetre.