

# WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 31 | 26 July to 1 August 2026

Dr. Sergey Morozov | drsergeymorozov.com

Compiled from publicly available sources (indexed peer-reviewed literature, official press releases, regulatory databases). Does not constitute medical, regulatory, or financial advice.

8

Peer-reviewed papers

5

Industry / FDA items

4

Media highlights

4

Key opinion leaders

## SECTION 0

## EXECUTIVE SUMMARY

*Key trends, week of 26 July to 1 August 2026*

### [FOUNDATION MODEL] PATHOLOGY MODEL LEARNS FROM REPORT DIALOGUE, NOT LABELS

Vorontsov E et al. (Nature Medicine) present PRISM2, a slide-level pathology foundation model trained on 2.3 million whole-slide images and 14 million question-answer pairs derived from 700,000 pathology reports. Supervised by clinical dialogue rather than labels alone, it matches or exceeds clinical-grade products for cancer detection in prostate, breast and breast lymph node ( $P < 0.05$ ), and its embeddings never statistically underperform previous foundation models. The abstract publishes comparative claims without the underlying numbers, so the margin cannot be independently appraised.

### [LLM EVALUATION] TWO CHATBOTS READ THYROID REPORTS, AND STABILITY SPLITS THEM

Xie Y et al. (Journal of Medical Internet Research) tested ChatGPT-4o and DeepSeek-R1 on 1,063 thyroid ultrasound reports from 3 centres. DeepSeek-R1 was more sensitive (0.879 versus 0.692) and more accurate (0.729 versus 0.644), but the decisive gap was reproducibility: on benign versus malignant calls its stability was kappa 0.869 versus 0.609. Both stayed below senior radiologists (AUC 0.865). A model that answers the same report differently across runs cannot be governed, whatever its mean accuracy.

### [OPPORTUNISTIC SCREENING] DISCARDED ATTENUATION MAPS CARRY PROGNOSTIC SIGNAL

Zhou J et al. (European Journal of Nuclear Medicine and Molecular Imaging) applied AI to CT attenuation maps acquired alongside myocardial perfusion imaging in 4,552 patients drawn from 43,099 studies across 15 sites, all with mild calcium and normal perfusion. Proximal calcium, present in 38%, independently predicted major adverse cardiovascular events (adjusted HR 1.24) and death (adjusted HR 1.25) over 3.6 years, adding 12% net reclassification beyond the calcium score. Modest effect sizes, but in patients a normal scan would otherwise reassure.

### [EVIDENCE QUALITY] TWO SYNTHESSES SHOW WHERE THE LITERATURE OVERSTATES ITSELF

Wong LM et al. (Journal of Magnetic Resonance Imaging) pooled 38 MRI studies of AI for nasopharyngeal carcinoma across 23,398 patients: pooled sensitivity 97.1% and specificity 87.8%, yet only 7 studies were applicable to screening and cohorts ran roughly 8 cases to 1 control. Richter S et al. (Journal of Medical Internet Research) meta-analysed LLM effects on physician-patient communication, finding a large pooled empathy effect (SMD 1.02, 95% CI 0.44 to 1.60) from only 4 poolable studies. Both headline numbers rest on narrow foundations.

**[WATCHPOINT] THE WEEK'S REAL SUBJECT IS VALIDATION, NOT ACCURACY**

Four of this week's eight papers are, in different ways, about the distance between a benchmark and a clinic. Galiana-Bordera A et al. (European Radiology Experimental) describe the CHAIMELEON platform, EUCAIM-aligned infrastructure that measures clinicians before AI, with AI, and against ground truth, precisely to capture automation bias that single-shot reader studies miss. Wong LM et al. show that diagnostic-cohort performance does not transfer to screening once prevalence is realistic. Xie Y et al. show that reproducibility is a model property that has to be measured separately from accuracy. Yang H et al. (Academic Radiology) reach an external AUC of 0.875 for nodal metastasis in cervical cancer without evidence that any surgical decision would change. Taken together the signal for procurement is that the discriminating question is no longer how accurate a model is, but what evidence exists that it behaves the same way, in the intended population, in the hands of the people who will actually use it.

## SECTION 1

**PEER-REVIEWED PAPERS**

Source: PubMed | Verified indexing | No preprints | 26 July to 1 August 2026

**End-to-end multimodal pathology foundation model with clinical dialogue.**

Vorontsov E, et al. | [Nature medicine](#) | 2026-07-31

PRISM2, a multimodal slide-level pathology foundation model trained on 2.3 million whole-slide images and 14 million question-answer pairs derived from 700,000 pathology reports. Rather than encoding image patches alone, it is supervised with clinical dialogue so that histomorphology is aligned with diagnostic reasoning, producing representations that support both prompt-based inference and transferable embeddings for downstream tasks.

**Key metrics:** With prompt-based inference, PRISM2 achieves or exceeds the balanced accuracy of clinical-grade products calibrated for cancer detection in the prostate, breast and breast lymph node ( $P < 0.05$ ). Across diagnostic, biomarker and survival benchmarks its embeddings never statistically underperform previous foundation models under linear probing ( $P < 0.05$ ). Task-specific fine-tuning on survival prediction outperforms training from scratch on the same large survival dataset. The abstract states these as comparative claims only and publishes no accuracy or AUC values, so the size of the margin cannot be judged from it.

**Clinical relevance:** The strategically interesting move is the supervision signal, not the parameter count. By training on report-derived dialogue rather than labels alone, the model treats the pathologist's written reasoning as the training target, which is the same asset every department already generates as a by-product of routine work. Matching clinical-grade single-purpose products with a single promptable model, if it holds outside the benchmark, changes the procurement question from buying one algorithm per indication to deploying one model per department. The caveat is that comparative claims without published numbers cannot be independently appraised, and prompt-based inference has yet to be tested in a prospective clinical workflow.

PMID [42538427](#) DOI [10.1038/s41591-026-04521-4](#)

## The Performance of ChatGPT-4o and DeepSeek-R1 in Interpreting Thyroid Nodule Ultrasound Text Reports: Multicenter Study.

Xie Y, et al. | *Journal of medical Internet research* | 2026-07-28

Multicenter evaluation of two general-purpose large language models, ChatGPT-4o and DeepSeek-R1, interpreting thyroid nodule ultrasound text reports across three clinical tasks: benign versus malignant differentiation, Chinese Thyroid Imaging Reporting and Data System (C-TIRADS) classification, and management recommendation. Reports were submitted through the consumer web interfaces with task-specific prompts, five repetitions per model, final output by mode voting, with output stability measured alongside accuracy.

**Key metrics:** 1,063 ultrasound text reports from 3 medical centers, 306 with histopathological confirmation. For benign versus malignant differentiation DeepSeek-R1 reached sensitivity 0.879 versus 0.692 ( $P < .001$ ) and accuracy 0.729 versus 0.644 ( $P = .008$ ). Senior radiologists, with access to images and clinical context the models did not have, performed better (AUC 0.865, accuracy 0.804). C-TIRADS agreement with radiologists was higher for DeepSeek-R1 (kappa 0.770 versus 0.688; delta kappa 0.082, 95% CI 0.048 to 0.122). Management agreement was comparable (kappa 0.606 versus 0.608). Stability was near perfect for C-TIRADS (alpha 0.864 versus 0.866) and management (kappa 0.853 versus 0.849), but diverged sharply for benign versus malignant differentiation (kappa 0.869 versus 0.609; delta kappa 0.260, 95% CI 0.191 to 0.321).

**Clinical relevance:** The stability result matters more than the accuracy result. A model that answers the same report differently across five runs cannot be governed, audited or safely placed in front of a patient, and the 0.26 kappa gap between the two models on the hardest task shows that reproducibility is a model property that must be measured separately rather than assumed. Both models stayed below senior radiologists even though radiologists had more information, which supports a decision-support framing rather than autonomous triage. Note the design limits: text reports only, no images, consumer web interfaces rather than a validated clinical deployment, and a Chinese classification system that does not transfer directly to ACR TI-RADS.

PMID 42520216 DOI 10.2196/93890

## Multi-modal Stacking Machine Learning Model for Predicting Lymph Node Metastasis in Cervical Cancer.

Yang H, et al. | *Academic radiology* | 2026-07-29

Multi-modal stacking ensemble that combines intratumoral and peritumoral habitat radiomics from multiparametric MRI (T2-weighted, diffusion-weighted and contrast-enhanced T1-weighted) with clinicopathological variables to predict lymph node metastasis preoperatively in early-stage cervical cancer. Tumour voxels were clustered into homogeneous habitats by K-means and features extracted from the tumour and a 3 mm peritumoral expansion, with five base classifiers feeding an XGBoost meta-learner.

**Key metrics:** Retrospective study with prospective validation in 623 patients with early-stage cervical cancer, split into training ( $n = 311$ ), internal validation ( $n = 187$ ) and external validation ( $n = 125$ ) cohorts. The best base learner (SVM-based ITH integrated score) reached AUCs of 0.882, 0.865 and 0.845 across the three cohorts. The meta-learner reached 0.915, 0.895 and 0.875 and significantly outperformed every base learner on DeLong testing (all  $P < 0.05$ ), with the highest clinical net benefit on decision curve analysis and significant net reclassification improvement and integrated discrimination improvement over each base learner.

**Clinical relevance:** Nodal status drives the choice between primary surgery and chemoradiation in early cervical cancer, and current preoperative imaging assessment is weak, so a non-invasive marker holding an AUC of 0.875 in an external cohort is clinically consequential rather than merely technically interesting. The habitat approach, which subdivides the tumour into physiologically distinct regions instead of averaging features over the whole lesion, plus the peritumoral margin, is the part most likely to generalise to other tumours. Set against that, this remains retrospectively trained radiomics with a single external cohort of 125 patients, no reader comparison, and no evidence yet that surgical decisions would actually change.

PMID 42527228 DOI 10.1016/j.acra.2026.07.022

## Impact of Large Language Model-Based AI Tools on Physician-Patient Communication: Systematic Review and Meta-Analysis.

Richter S, et al. | *Journal of medical Internet research* | 2026-07-31

PRISMA-guided systematic review and random-effects meta-analysis of peer-reviewed studies published between 2020 and 2025 evaluating how large language model interventions affect physician-patient communication, covering empathy, clarity, trust and patient understanding. PubMed/MEDLINE, Embase, Scopus and Web of Science were searched, with two reviewers independently screening and extracting data across randomized, observational, cross-sectional and qualitative designs.

**Key metrics:** 10 studies were included from 312 records, all quantitative and predominantly cross-sectional. In 6 direct comparisons of AI-generated versus physician-generated responses, LLMs were rated significantly higher in empathy in 5. One large study rated chatbot replies empathetic in 45.1% of cases versus 4.6% for physician replies (odds ratio approximately 9.8,  $P < .001$ ); ChatGPT-4 answers scored 4.18 versus 2.70 on a 5-point empathy scale ( $P < .001$ ); a neurology study found a Consultation and Relational Empathy gain of 1.38 points ( $P < .01$ ). GPT-4 simplification of pathology reports raised patient comprehension from 5.23 to 7.98 out of 10 ( $P < .001$ ) and cut consultation time by 70%. Pooled across 4 studies and 2,604 evaluations, the standardized mean difference for empathy was 1.02 (95% CI 0.44 to 1.60).

**Clinical relevance:** For radiology this lands squarely on the patient-facing report, which is where imaging departments carry communication risk they rarely measure. A pooled effect above 1.0 on empathy and a measured comprehension gain on pathology reports specifically suggest that LLM-mediated report translation is one of the few LLM use cases with real evidence behind it, and the 70% consultation-time reduction is an operational argument, not just a satisfaction one. The evidence base is nonetheless thin and fragile: 10 predominantly cross-sectional studies, only 4 poolable, a wide confidence interval, no study assessing long-term trust, and a documented tendency toward lengthy or occasionally inaccurate output that keeps physician oversight mandatory.

PMID 42536971 DOI 10.2196/77307

## Hidden risk in normal myocardial perfusion scans: AI-detected proximal coronary calcium on CT attenuation maps improves prognosis.

Zhou J, et al. | *European journal of nuclear medicine and molecular imaging* | 2026-07-28

AI applied to the CT attenuation correction maps already acquired alongside SPECT and PET myocardial perfusion imaging, to automatically identify coronary artery calcium in proximal coronary segments and test whether its spatial distribution carries prognostic information beyond the calcium score itself. Outcomes were assessed in two large international multicenter registries using multivariable Cox regression, likelihood ratio testing and continuous net reclassification.

**Key metrics:** From 43,099 hybrid myocardial perfusion imaging and CT studies across 15 sites, 4,552 patients were included who had no prior coronary artery disease, AI-derived mild calcium scores of 1 to 99, and normal perfusion (stress total perfusion deficit below 5%). Mean age 65 plus or minus 12 years, 51% male; 1,730 (38%) had proximal calcium. Over a median 3.6 years of follow-up (interquartile interval 2.1 to 5.2), 599 (13%) had a major adverse cardiovascular event and 444 (10%) died. Proximal calcium was independently associated with major adverse cardiovascular events (adjusted hazard ratio 1.24, 95% CI 1.03 to 1.48,  $P = 0.02$ ) and all-cause mortality (adjusted hazard ratio 1.25, 95% CI 1.01 to 1.53,  $P = 0.04$ ) after adjustment for calcium score and density, clinical risk factors and perfusion deficit, improving risk stratification for both endpoints (likelihood ratio test  $P = 0.02$  and  $P = 0.04$ ; net reclassification 12%).

**Clinical relevance:** This is opportunistic screening in its purest form: the attenuation map is acquired for a technical reason, is normally discarded for diagnostic purposes, and turns out to carry prognostic signal in exactly the patients a normal perfusion study would otherwise reassure. The population matters, since 38% of patients with mild calcium and a normal scan carry a risk marker that current reporting does not mention, and the incremental value survives adjustment for the calcium score itself, meaning location adds information beyond burden. Effect sizes are nonetheless modest, with hazard ratios near 1.25 and confidence intervals approaching 1.0, so this refines risk stratification rather than reclassifying patients into treatment; whether acting on it improves outcomes is untested.

PMID 42517890 DOI 10.1007/s00259-026-08084-x

## Systematic Review of AI-Driven Applications for Screening Nasopharyngeal Carcinoma Using MRI.

Wong LM, et al. | *Journal of magnetic resonance imaging* | 2026-07-26

PRISMA-guided systematic review with meta-analysis of MRI-based AI for detecting, localizing and diagnosing nasopharyngeal carcinoma, with the specific aim of judging whether the published literature actually supports a screening use case rather than routine clinical care. PubMed, Scopus and Embase were searched from January 2009 to August 2025 by two independent reviewers, with subgroup analysis of the effect of intravenous contrast and formal assessment of applicability and risk of bias.

**Key metrics:** 38 studies covering 23,398 patients (20,693 with nasopharyngeal carcinoma and 2,705 without), at field strengths of 1.5 T to 3.0 T. Pooling 30 localization and 6 discrimination studies, the Dice similarity score was 0.80 (CI 0.78 to 0.82) and sensitivity and specificity were 97.1% (CI 88.0% to 99.3%) and 87.8% (CI 78.9% to 93.2%). Intravenous contrast did not significantly affect AI performance for either task ( $p > 0.05$ ). Only 7 of the 38 studies were fully applicable to screening.

**Clinical relevance:** The headline sensitivity of 97.1% is the least important number here. The finding that matters is that only 7 of 38 studies were applicable to the screening scenario, and that the pooled cohorts run roughly 8 cases of disease to 1 control, an inversion of screening prevalence that inflates apparent performance and makes the specificity of 87.8% the real constraint: at true screening prevalence that false-positive rate governs the programme's cost and harm. The contrast finding is directly actionable, since non-contrast MRI performing equivalently removes gadolinium, cost and time from any screening protocol. Read this as a warning that diagnostic-cohort performance does not transfer to screening, which applies well beyond nasopharyngeal carcinoma.

PMID 42503234 DOI 10.1002/jmri.70445

## Social Media Discourse on Breast Cancer Screening Barriers in Japanese and English: Cross-Sectional Infodemiology Study.

Terada M, et al. | *Journal of medical Internet research* | 2026-07-31

Cross-sectional infodemiology study using an automated natural language processing pipeline to compare barriers to breast cancer screening expressed in Japanese-language and English-language social media discourse. The pipeline combined large language model-assisted sentiment polarity scoring with strict negation handling, co-occurrence network topology analysis and distributional visualization, with subgroup comparisons before versus after screening and ultrasound with versus without mammography.

**Key metrics:** 46,823 screening-related posts from X collected in 2025 (30,027 Japanese and 16,796 English), drawn from an initial 76,955 after noise exclusion. Overall sentiment showed no meaningful cross-cultural divergence (Cohen  $d = 0.049$ ), but barrier prevalence differed sharply. In English the dominant barrier was cost ( $n = 1,239$ , 7.4%) with pain lower ( $n = 842$ , 5.0%); in Japanese pain dominated ( $n = 5,788$ , 19.3%) within a tightly interconnected pain, fear and appointment network. Japanese subgroup analysis showed anticipatory fear (739/3,805, 19.4%) and appointment logistics (1,556/3,805, 40.9%) before screening giving way to persistent pain memory afterwards (1,160/7,419, 15.6%). Sentiment toward mammography was significantly more negative than toward ultrasound ( $P < .001$ , Cohen  $d = 0.264$ ), and pain descriptors for ultrasound rose from 5.0% (106/2,110) alone to 18.2% (733/4,017) when performed with mammography. Screening adherence context: below 50% in Japan versus above 70% in the United States.

**Clinical relevance:** This is an AI method applied to the demand side of imaging rather than to the images, and it produces an operational finding a survey would likely have missed: mammographic discomfort contaminates the perception of adjacent painless modalities, with pain complaints about ultrasound tripling when it is performed in the same visit as mammography. That points at concrete levers, namely individualized compression protocols and decoupling supplemental ultrasound from the mammography appointment, rather than at generic awareness campaigns. Treat the epidemiology cautiously: social media posters are not a representative screening population, absolute percentages are small, and cross-language sentiment comparison carries method risk even with negation handling, so this generates hypotheses about why adherence stalls rather than measuring it.

PMID 42537013 DOI 10.2196/97772

## Bridging the AI chasm in oncology: a standardized platform to enable the in silico clinical validation of AI models within the CHAIMELEON Project.

Galiana-Bordera A, et al. | [European radiology experimental](#) | 2026-07-31

A web-based platform built within the European CHAIMELEON project, aligned with EUCAIM, for in silico clinical validation of oncology AI models before real-world deployment. The architecture uses Kubernetes microservices combining a REST API, an ORTHANC PACS and a Keycloak/OAuth2 security layer, with a customized OHIF DICOM viewer. Clinicians assess each case in three sequential stages, namely standard clinical assessment, assessment with AI predictions, and final review against clinical endpoint ground truth, while the platform records perceptions of AI utility, trust, workflow integration and usability.

**Key metrics:** The platform provided access to multimodal data across five cancer types and clinicians completed structured validations. Survey responses were favourable: 93% found the platform intuitive, over 80% would recommend it, and agreement on AI utility exceeded 40% across most endpoints, which the authors read as reception of AI as a supportive second reader that prompts clinical re-evaluation rather than simple consensus. Registered as ClinicalTrials.gov NCT06950996. The abstract reports usability and perception outcomes only and publishes no diagnostic accuracy figures for the models validated on the platform.

**Clinical relevance:** The gap between a published AUC and a deployable product is a validation-infrastructure problem, and this is one of the few funded attempts to standardise that step rather than leave every site to improvise. The three-stage design is the substantive contribution, since measuring the clinician before AI, with AI, and against ground truth captures automation bias and appropriate reliance, which single-shot reader studies cannot. Its relevance to European practice is procedural: EUCAIM alignment means this is positioning itself as reusable regulatory-grade evidence infrastructure under the EU AI Act. Judge it accordingly, as the reported outcomes are usability and perception, agreement on AI utility above 40% is a modest bar, and no accuracy results are given.

PMID 42536277 DOI 10.1186/s41747-026-00770-7

### SECTION 2

## INDUSTRY & REGULATION

Sources: [Official registers](#) | [Press releases](#) | [FDA databases](#)

### [FDA] DeepHealth receives FDA clearance for AI-powered breast ultrasound

DeepHealth (RadNet) | 30 July 2026

FDA 510(k) clearance for DeepHealth Breast Ultrasound, the commercial name for the See-Mode SMART-B device cleared under K260303. The software automates lesion detection, characterization and reporting for sonographers and radiologists. Validation included a multi-reader multi-case study with 16 US board-certified radiologists plus live clinical use under regulated research protocols. The product is commercially available and eligible for reimbursement under an existing Category III CPT code for quantitative ultrasound tissue characterization, and RadNet plans to implement it across its network by the end of 2026, covering an estimated 700,000 breast ultrasound studies annually.

Source: [GlobeNewswire](#)

### [FDA] ThinkSono receives FDA clearance for AI ultrasound guidance in suspected DVT

ThinkSono | 29 July 2026

FDA 510(k) clearance for ThinkSono Guidance, described as the first AI-powered ultrasound guidance software cleared for the evaluation of suspected deep vein thrombosis. The mobile app connects to compatible handheld point-of-care scanners and guides non-ultrasound-trained staff in real time through probe positioning and venous compression cine-loop capture; recorded loops are sent to a dashboard where a remote qualified clinician grades image quality and assesses vein compressibility. Operators in clinical studies performed exams after 60 to 90 minutes of standardized training. The software does not interpret images. Breakthrough Device Designation was granted in February 2026.

Source: [AccessNewswire](#)

**[REGULATION] Caristo Diagnostics wins FDA De Novo authorization for CaRi-Heart coronary inflammation analysis***Caristo Diagnostics | 29 July 2026*

The FDA granted De Novo authorization for CaRi-Heart, described as the first and only technology authorized in the US to quantify coronary inflammation from routine coronary CT angiography. Developed from University of Oxford research, it detects inflammation-related changes in perivascular fat, reports them as a fat attenuation index score, and generates a personalized 10-year cardiovascular mortality risk estimate. Commercial launch in the US begins this quarter with expansion in Q4 2026. AMA Category III CPT codes 0992T and 0993T have applied since 1 January 2026.

Source: [PR Newswire](#)**[PARTNERSHIP] Median Technologies and Olea Medical, a Canon Medical Systems company, partner on commercial deployment of eyonis LCS***Median Technologies and Olea Medical | 30 July 2026*

A global strategic commercial agreement under which Olea Medical, a Canon Medical Systems subsidiary, will support commercialization of the AI lung cancer screening software eyonis LCS across the United States and Europe, with other geographies possible over time. Canon is a major supplier of the CT systems that eyonis LCS runs on. The agreement follows FDA 510(k) clearance in February 2026 and CE marking in July 2026, and complements Median's earlier partnership with Tempus AI, forming part of a non-exclusive multi-channel commercialization strategy.

Source: [Business Wire](#)**[PARTNERSHIP] Vanderbilt Health and Siemens Healthineers launch \$87M value partnership covering imaging and AI***Vanderbilt Health and Siemens Healthineers | 27 July 2026*

A multi-year, \$87 million Value Partnership to modernize and standardize diagnostic imaging, radiation oncology and related workforce development across Vanderbilt Health hospitals in the Nashville area. Siemens Healthineers becomes the primary provider of MRI, CT, molecular imaging, interventional radiology and Varian radiation oncology systems for sites including Vanderbilt University Hospital and Monroe Carell Jr. Children's Hospital. The parties will also explore data infrastructure, AI-enabled technologies and workflow innovation in radiology and personalized medicine.

Source: [Siemens Healthineers](#)

## SECTION 3

**MEDIA HIGHLIGHTS***AuntMinnie | Radiology Business | The Imaging Wire | Diagnostic Imaging | ITN***AI closes the mammography gap between generalists and specialists***The Imaging Wire | 27 July 2026*

Coverage of a new Radiology study in which DeepHealth's ProFound Pro 2.x was deployed as a safeguard review: after a radiologist's initial read, AI re-analysed non-recalled mammograms and routed those above a risk threshold to an expert reviewer. Across 109 breast imaging facilities and 578,000 DBT screening exams from 2021 to 2022, the cancer detection rate of general radiologists rose 25%, from 3.76 to 4.99 per 1,000 exams, with positive predictive value up 15%. Specialists showed no statistically significant change, and the gap between AI-aided generalists and specialists was no longer significant. Recall rates for generalists did rise, from 9.06% to 10.4%.

Link: <https://theimagingwire.com/2026/07/26/ai-based-workflow-for-mammography-screenin...>

## How AI will reshape radiology without replacing radiologists

*Smithsonian Magazine* | 29 July 2026

A mainstream feature revisiting Geoffrey Hinton's 2016 prediction that radiologists would be replaced within five years. It notes that radiology's ranks are instead growing, with the number of practitioners expected to expand by 26% or more over the next three decades, while accepting the part of the prediction that held, namely that radiologists now work alongside systems that match or exceed human performance on specific tasks. Useful as a barometer of how the specialty is being described to a general audience.

Link: <https://www.smithsonianmag.com/innovation/how-ai-will-reshape-radiologywithout-r...>

## AI revolution in radiology is about education, mature data, guard rails and consent, not job loss

*Croakey Health Media* | 29 July 2026

Conference reporting published this week from Intelligence26: AI and the Future of Practice, hosted by the Royal Australian and New Zealand College of Radiologists in Sydney on 24 and 25 July, which drew more than 250 local and international professionals across medical AI, diagnostic imaging, radiation oncology and data science. The reporting frames the practical agenda as education, data maturity, guard rails and patient consent rather than workforce displacement.

Link: <https://www.croakey.org/ai-revolution-in-radiology-is-about-education-mature-dat...>

## Can radiologists distinguish AI-generated radiological images from real ones?

*Radiology News* | 27 July 2026

Summary of a Clinical Radiology paper testing whether radiologists can tell synthetic images, produced with a DreamBooth fine-tuning implementation, from real ones, and which factors predict correct classification. Relevant to data provenance and to the integrity of training and assessment material as generative models become easier to apply to medical images.

Link: <https://radiologynews.com/real-or-not-real-can-radiologists-distinguish-artifici...>

### SECTION 4

## PROFESSIONAL SOCIETIES

*SIIM | ACR | RCR | Official publications and event coverage*

## ESR publishes the report from its European Parliament event on imaging and cardiovascular health

27 July 2026

The European Society of Radiology released the full report from Imaging Saves Lives: The Role of Imaging in Improving Cardiovascular Health, the event it co-organised with the European Association of Nuclear Medicine at the European Parliament on 2 July. The report sets out the key discussions, main messages and recommendations on how imaging supports implementation of the EU Safe Hearts Plan across prevention, early diagnosis, treatment and follow-up. The accompanying ESR and ESCR policy briefing names responsible integration of artificial intelligence, alongside digital innovation, research and data capacity and reduced health inequalities, as cross-cutting priorities for delivering the plan.

Source: [https://www.linkedin.com/posts/esr-european-society-of-radiology\\_the-report-from-the-european-parliament-event-activity-7487507113272082432-TmZ1](https://www.linkedin.com/posts/esr-european-society-of-radiology_the-report-from-the-european-parliament-event-activity-7487507113272082432-TmZ1)

## European Board of Radiology opens EDiR registration for ECR 2027 and launches the EDiR Academy

31 July 2026

The European Board of Radiology began accepting registrations for the European Diploma in Radiology examination to be held at ECR 2027 on 2 March 2027, and launched the EDiR Academy, which consolidates its official preparation resources, namely self-assessments, courses and exam simulations, into a single structured pathway. A dedicated EDiR Preparation Pathway package for ECR 2027 candidates combines knowledge evaluation, topic-focused self-assessment, case-based training and a mock exam across four stages, priced at 200 euro.

Source: <https://www.auntminnieeurope.com/radiology-education/article/15831408/ebr-accepts-registrations-for-edir-launches-edir-academy>

### SECTION 5

## KEY OPINION LEADERS

LinkedIn | Public statements | Week of 26 July to 1 August 2026

### Curt Langlotz

*Professor of Radiology, Medicine and Biomedical Data Science, Stanford University School of Medicine*

Posted on 30 July about co-writing, with Fatima Rodriguez, the Circulation editorial accompanying a paper on coronary artery calcium detection on non-gated CT, quoting the editorial's conclusion that the time for clinical implementation is now. The post links to trade coverage reporting that coronary calcium scores derived from nongated CT scans predicted cardiovascular risk. It is a direct statement from one of the field's most cited voices that opportunistic CAC scoring has passed from research question to implementation question, and it lands in the same week as new registry evidence that AI-detected proximal calcium on attenuation maps adds prognostic value.

Source: [LinkedIn](#)

### Curt Langlotz

*Professor of Radiology, Medicine and Biomedical Data Science, Stanford University School of Medicine*

Posted on 30 July about a recorded conversation on the effect of AI on the radiologist workforce, the need for patient consent and other current trends, filmed at the Royal Australian and New Zealand College of Radiologists meeting and published by Croakey Health Media. The framing is notable for what it puts at the centre: consent and workforce structure rather than model accuracy.

Source: [LinkedIn](#)

### Pranav Rajpurkar

*Associate Professor, Harvard University; co-founder of a2z Radiology AI*

Posted on 1 August congratulating the teams competing in the ReXGrounding CT challenge at MICCAI 2026, which asks models to ground free-text findings to 3D CT segmentations, and encouraging new entrants. The shared leaderboard shows twelve submissions ranked by mean Dice on a public 50% subset of the held-out test set, with the top entries clustered near 0.33 and final results to be revealed on the full test set at MICCAI 2026. The modest Dice scores are the useful signal: linking narrative report language to voxel-level anatomy remains substantially unsolved.

Source: [LinkedIn](#)

### Daniel Pinto dos Santos

*Radiologist and medical imaging informatics researcher, Germany*

Shared on 31 July a EuSoMI announcement on how artificial intelligence is reshaping musculoskeletal radiology, spanning current clinical applications through to the technologies that may define future diagnostic workflows. Musculoskeletal imaging has been one of the earliest and densest areas of commercial AI deployment, which makes society-level stock-taking of what is actually in clinical use there a useful reference point.

Source: [LinkedIn](#)

**QUALITY ASSURANCE NOTE**

*Compiled by Dr. Sergey Morozov from publicly available sources: peer-reviewed papers indexed in PubMed (PMIDs and DOIs verified, no preprints or arXiv), official press releases, regulatory databases and specialist media. It does not constitute medical, regulatory or financial advice. All papers have publication dates verified within 26 July to 1 August 2026. Paper selection is deterministic (pinned PubMed query, Q1/flagship journal allow-list, exclusion of pure reviews [meta-analyses kept], radiomics-only and interventional-radiology work, then a transparent ranking score). Industry items come from official press releases / regulatory notices; KOL items from public LinkedIn activity within the window.*