

WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 33 | 9 to 15 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

2

Media highlights

3

Key opinion leaders

GOVERNANCE TAKEAWAY OF THE WEEK

If mammography AI is live in your department, split recall monitoring by pathway this month. A year-long alternating-month deployment doubled cancer detection without slowing reading, but abnormal interpretations rose only in diagnostic mammography, 18.7% versus 12.1%, while screening stayed flat. A single pooled recall metric would have hidden that shift. Add abnormal interpretation rate by pathway to your monthly monitoring file; the same file becomes the seed of your EU AI Act Article 72 post-market record.

SECTION 0

EXECUTIVE SUMMARY

Key trends, week of 9 to 15 August 2026

[REAL-WORLD MONITORING] MAMMOGRAPHY AI RAISED DETECTION, AND RECALLS, BUT ONLY IN DIAGNOSTIC WORK

Lee SE et al. (European Radiology) alternated a commercial mammography AI on and off monthly for a year across 4577 examinations read by four radiologists. Cancer detection roughly doubled (22.3 vs 11.5 per 1000), reading time was unchanged (65.0 vs 64.3 s), and the abnormal interpretation rate rose only in diagnostic mammography (18.7% vs 12.1%), not in screening. The on-off design itself is a reusable local monitoring template.

[OPERATIONS] DXA REPORT DRAFTING AUTOMATED ACROSS FOUR SITES, WITH THE AUDIT TO PROVE IT

Lakhani P et al. (JACR) deployed vendor-agnostic AI OCR report drafting for DXA at two academic and two community sites. Report creation time fell from 3.48 to 0.87 minutes, community turnaround dropped by roughly four days, and a 400-report audit against source images showed 99.9 to 100% numerical accuracy while exposing a 45% completeness baseline in the legacy process.

[BENCHMARK] THE PROMPT, NOT THE MODEL, DECIDED WHETHER AN LLM WAS USABLE

Zhu J et al. (Journal of Medical Internet Research) tested GPT-4o, Claude 3.7 Sonnet, Qwen2.5-VL 72B and Gemma 3 27B on femoral head osteonecrosis. Detection AUC went from 0.55 to 0.91 when a deidentified radiology description accompanied the image, grading reliability from ICC 0.51 to 0.97, and adding more images degraded performance. Commercial and open-source models did not differ overall (P = .83).

[FLAGSHIP] 3D PATHOLOGY BECOMES READABLE: TRIAGE KEEPS THE PATHOLOGIST IN CHARGE

Gao G et al. (Nature Biomedical Engineering) present TRICARE, which scores every 2D level of a 3D pathology volume using neighbouring-depth context and routes the highest-risk cross sections to the pathologist for final diagnosis. In prostate biopsy risk stratification and Barrett's esophagus screening it shows potential to improve on slide-based histopathology, which samples under 1% of the tissue.

[WATCHPOINT] MONITORING SPLITS ARE WHERE THE SIGNAL LIVES

Three of this week's results only exist because someone split what most departments pool. Mammography abnormal interpretation rate moved in diagnostic but not screening work. DXA completeness gaps sat at community sites, not academic ones. LLM grading reliability swung from ICC 0.49 to 0.97 on input configuration alone. Aggregate metrics would have hidden all three. Design your monitoring to disaggregate by pathway, by site and by configuration from the first day, or accept that your dashboard will average the problem away.

SECTION 1

PEER-REVIEWED PAPERS

Source: PubMed | Verified indexing | No preprints | 9 to 15 August 2026

Impact of AI assistance on reading time, cancer detection rate, and abnormal interpretation rate in screening and diagnostic mammography: a prospective alternating-month study.

Lee SE, et al. | *European radiology* | 2026-08-13

Before extending mammography AI from screening into the diagnostic worklist, check what it does to your abnormal interpretation rate. Four radiologists at one institution read 4577 consecutive mammograms over 12 months with a commercial AI system alternately displayed and hidden month by month, a prospective on-off design that doubles as a monitoring template. Cancer detection was higher with AI and reading time did not move, but abnormal interpretations rose only in the diagnostic pathway.

Key metrics: 4577 examinations, mean age 51.7 years. Overall cancer detection rate 22.3 vs 11.5 per 1000 ($p = 0.005$). Abnormal interpretation rate in screening 9.5% vs 8.4% ($p = 0.293$), in diagnostic mammography 18.7% vs 12.1% ($p = 0.026$). Mean reading time 65.0 vs 64.3 s ($p = 0.723$) in the 2917-examination subset with times of 5 minutes or less.

Clinical relevance: A single pooled recall metric would have missed the diagnostic-pathway shift entirely; the effect appears only when screening and diagnostic examinations are monitored separately. For your AI committee this changes monitoring: add abnormal interpretation rate by pathway to the monthly post-market file (EU AI Act Article 72), and note the alternating-month on-off design as a low-cost local template for measuring what a tool does to your own readers under Article 14 human oversight.

PMID 42593499 DOI 10.1007/s00330-026-12793-0

Prompt Configurations for Multimodal Large Language Models in Diagnosing and Staging Osteonecrosis of the Femoral Head: Multimodel Retrospective Observational Diagnostic Study.

Zhu J, et al. | *Journal of medical Internet research* | 2026-08-10

Do not accept an evaluation of a multimodal LLM, from a vendor or your own team, that does not pin the exact prompt configuration. Four models (GPT-4o, Claude 3.7 Sonnet, Qwen2.5-VL 72B, Gemma 3 27B) diagnosed and staged osteonecrosis of the femoral head on radiographs and MRI under three input configurations. Adding a deidentified radiology description to the image moved detection from near chance to strong, while adding more images made models worse.

Key metrics: 159 radiograph patients (318 hip-level observations) and 170 MRI patients (340 observations), single center. Detection AUC 0.55 (SD 0.03) with single-image input vs 0.91 (SD 0.01) with image plus description. Staging accuracy 0.65, 0.78 and about 0.59 for single-image, image-plus-description and multi-image input. ARCO grading reliability ICC 0.51, 0.97 and 0.49 respectively, against surgeon interrater ICC 0.70. Commercial vs open-source models: no significant overall difference ($P = .83$).

Clinical relevance: Performance here is a property of the model plus its input configuration, not of the model alone, and grading reliability swings from unusable to near-perfect on the same cases. For your AI committee this changes acceptance testing: the test report must record the deployed prompt and input configuration, and any change to that configuration re-triggers the test (Article 15 accuracy and robustness). Open-source parity with commercial models also weakens any governance argument that relies on vendor reputation.

PMID 42573734 DOI 10.2196/92919

Vendor-Agnostic Multisite Automated Dual-Energy X-Ray Absorptiometry Reporting Using Artificial Intelligence-Based Optical Character Recognition: Impact on Workflow Efficiency and Accuracy.

Lakhani P, et al. | *Journal of the American College of Radiology* | 2026-08-13

If you are automating report drafting, this is the deployment file to copy. A vendor-agnostic AI optical character recognition pipeline extracted DXA measurements from DICOM images and drafted complete reports across four sites of one health system, two academic and two community, with a pre-post operational evaluation and a 400-report accuracy audit scored against the source images rather than against the prior report.

Key metrics: Median report creation time fell from 3.48 to 0.87 minutes (academic) and 1.42 to 0.63 minutes (community); median turnaround time from 2.40 to 0.96 hours (academic) and from 133.82 to 42.10 hours, roughly four days, at community sites (all $P < .0001$). AI drafts vs original reports: numerical accuracy 99.9% vs 99.4% ($P = .022$) and 100% vs 99.6% ($P = .031$); completeness at community sites 100% vs 45% ($P < .001$); diagnostic accuracy at least 99.5% in all cohorts.

Clinical relevance: The audit design is the transferable part: accuracy was defined against source images, which exposed that the human baseline was itself incomplete at community sites. For your AI committee this changes acceptance testing and monitoring: define draft-versus-source concordance as the acceptance metric for any report-drafting tool, then keep sampling it monthly as the post-market KPI (Article 72; ISO/IEC 42001 performance evaluation clause).

PMID 42595274 DOI 10.1016/j.jacr.2026.08.007

Adapted foundation models for breast MRI triaging in contrast-enhanced and non-contrast-enhanced protocols.

Nguyen TT, et al. | *European radiology* | 2026-08-13

Foundation-model triage of abbreviated breast MRI is approaching a usable rule-out: at 97.5% sensitivity, roughly one in five examinations without immediately actionable findings could be deprioritized, with near-identical performance whether or not contrast was given. A DINOv2-based medical slice transformer was tested across four abbreviated protocols to rule out BI-RADS 4 or higher examinations.

Key metrics: 1847 single-breast MRI examinations (377 BI-RADS 4 or higher) from an in-house dataset plus 924 external examinations (Duke); 1448 patients, mean age 49 years. AUC 0.77 (SD 0.04) for T1 subtraction plus T2w and 0.74 (SD 0.04) for DWI1500 plus T2w, with no significant differences across protocols. Specificity at 97.5% sensitivity: 19% contrast-enhanced, 17% non-contrast. External validation of the T1sub protocol: AUC 0.77; 88% of attention maps rated good or moderate.

Clinical relevance: The clinically important detail is where the false negatives concentrate: predominantly lesions under 10 mm and non-mass enhancements, exactly the findings abbreviated protocols already handle worst. Near-parity of the non-contrast protocols matters for screening programs weighing gadolinium exposure, and the external AUC of 0.77 held at the level of the internal result.

PMID 42593493 DOI 10.1007/s00330-026-12782-3

Deep-learning triage of three-dimensional pathology datasets for comprehensive and efficient pathologist assessments.

Gao G, et al. | *Nature biomedical engineering* | 2026-08-12

Non-destructive 3D pathology images the whole biopsy instead of the under 1% of tissue sampled by a standard slide, but the resulting datasets are too large for manual review. TRICARE, a deep-learning triage framework, scores every 2D level in the volume using context from neighbouring depth levels and hands the highest-risk cross sections to the pathologist, who keeps the final diagnosis. It was evaluated for prostate cancer biopsy risk stratification and dysplasia screening in Barrett's esophagus.

Key metrics: The abstract reports the comparisons qualitatively, without numeric values: depth-context scoring outperformed models based on isolated 2D levels, and AI-triaged 3D pathology showed potential to improve detection of high-risk disease over slide-based 2D histopathology while optimizing pathologist workload. Each standard 2D section represents less than 1% of the biopsy volume.

Clinical relevance: Kept on scientific merit: this is the most credible route yet published for 3D pathology into practice, precisely because it does not remove the pathologist. Triageing a data format that cannot be read exhaustively is a different adoption problem from replacing a reader, and both use cases, prostate risk stratification and Barrett's surveillance, are settings where 2D undersampling has a known clinical cost.

PMID 42587049 DOI 10.1038/s41551-026-01760-1

Deep Learning-Based Enhancement of Already Diagnostic-Quality MRI for Alzheimer's Disease Classification: Effects on Model Performance and Training Data Requirements.

Zhang Z, et al. | *Journal of magnetic resonance imaging* | 2026-08-13

Image enhancement is usually bought to rescue poor images. Here an FDA-cleared deep learning enhancement tool (SubtleHD) was applied to already diagnostic-quality T1 brain MRI before training Alzheimer's disease classifiers, and it improved both classification performance and data efficiency, suggesting conventional definitions of image quality underestimate the information available to machine learning.

Key metrics: 2293 ADNI scans (1605 training, 229 validation, 459 internal test) and 270 external NACC scans. ResNet34 accuracy 85.2% to 88.7% and macro-AUC 0.951 to 0.968 with enhancement (both significant); DenseNet121 differences not significant ($p = 0.18$ and $p = 0.29$). Training on 70% of the enhanced dataset matched the full standard-of-care dataset. External test: accuracy 49.2% and macro-AUC 0.679 for the standard-trained model vs 63.0% and 0.819 for the enhancement-trained model, which kept an advantage on unenhanced external images (macro-AUC 0.772).

Clinical relevance: Two honest readings coexist: enhancement extracted signal that conventional image-quality definitions miss, and the external performance of both arms, 49% to 63% accuracy, shows how much of the internal result did not travel. The data-efficiency finding, matching full-dataset performance with 70% of the training data, is the practically interesting one for groups training on limited local cohorts.

PMID 42593101 DOI 10.1002/jmri.70507

Prediction of Central Lymph Node Metastasis in Papillary Thyroid Microcarcinoma Using a Deep Learning Radiomics Model Based on SAM3 Automatic Segmentation of Ultrasound Images: A Multicenter Cohort Study.

Song J, et al. | *Academic radiology* | 2026-08-12

For papillary thyroid microcarcinoma, the operative decision often hinges on occult central lymph node metastasis. A deep learning radiomics model on ultrasound, built on fully automatic SAM3 segmentation, predicted central node metastasis across four training centers and one external site, removing the manual contouring step that has kept such models out of routine use.

Key metrics: 1859 patients from four centers (1674 training, 185 validation) plus 140 from an external facility, 1999 evaluated in total. SAM3 segmentation Dice 0.882 (validation) and 0.859 (external testing). AUC of the combined deep learning radiomics model 0.901 (training), 0.875 (validation), 0.853 (external testing), against 0.779 to 0.825 for radiomics alone and 0.812 to 0.844 for deep learning alone. Decision curve analysis supported clinical utility.

Clinical relevance: Kept on clinical merit: the external AUC of 0.853 with fully automatic segmentation is the relevant number, since manual-ROI radiomics pipelines have repeatedly failed to survive contact with routine practice. Whether 0.85 discrimination is enough to change decisions on prophylactic central neck dissection is the clinical question the paper leaves open.

PMID 42586895 DOI 10.1016/j.acra.2026.07.064

From voxel discovery to regional interaction: A multi-level interpretable framework for Alzheimer's disease diagnosis.

Shu K, et al. | *Medical image analysis* | 2026-08-12

Interpretability built into the architecture rather than bolted on afterwards: a three-stage framework for Alzheimer's diagnosis on structural MRI that discovers disease-relevant voxels through hierarchical evidential masking, aggregates them into anatomically coherent regions, and models interactions between regions with a graph attention network, making the diagnostic reasoning visible at every level rather than reconstructed post hoc.

Key metrics: The abstract makes only comparative claims without numeric values: performance on the ADNI and AIBL datasets is described as comparable to state-of-the-art black-box models, while the framework autonomously reconstructed established neuropathological trajectories. Source code is publicly available.

Clinical relevance: Kept on scientific merit: post-hoc saliency maps have repeatedly failed to convince clinicians because of background leakage and anatomical incoherence, and this design addresses both structurally rather than cosmetically. If the comparable-performance claim survives independent testing, inherently interpretable architectures become a realistic alternative to explanation layers added after training.

PMID 42585730 DOI 10.1016/j.media.2026.104243

SECTION 2

INDUSTRY & REGULATION

Sources: Official registers | Press releases | FDA databases

[REGULATION] Medicare approves new technology add-on payment for Aidoc CARE CT Body Multi-Triage

Aidoc | 13 August 2026

What it covers: CMS finalized a New Technology Add-on Payment of up to \$137.53 per eligible inpatient case, about 65% of the average technology cost per the ACR summary, from 1 October 2026 for three years (procedure code XEZ5XKC). What the tool is cleared for: workflow triage that flags suspected urgent findings on chest, abdomen and pelvis CT; a triage clearance prioritizes the worklist, it does not diagnose. Evidence level: the payment came through the alternative pathway, which waives proof of substantial clinical improvement for breakthrough-designated devices, and CMS is repealing that pathway from FY2028. Questions before you rely on it: which of your inpatient cases actually qualify, how coding will be captured, and what your own override and concordance rates show, since reimbursement changes the economics but not the local validation burden.

Source: [PR Newswire](#)**[PARTNERSHIP] Twelve US health systems form Diagnostic AI Consortium with Aidoc**

Diagnostic AI Consortium | 11 August 2026

What it is: twelve systems, including Cedars-Sinai, Mount Sinai, Northwell, Northwestern Medicine and Houston Methodist, covering roughly 20 million patients a year, committing to shared standards for how diagnostic AI is evaluated, deployed and governed, with Aidoc supplying the technical infrastructure and first results promised for 2027. Evidence level: an announcement of intent, no data yet. Why it matters for governance: shared deployer-side evaluation and monitoring standards are exactly what the EU is regulating into existence; in the US they are emerging by consortium. The question to watch: whether standards built on a single vendor's platform prove portable, and whether outcome measurement is independent of the infrastructure provider being measured.

Source: [Radiology Business](#)**[FDA] Leica Biosystems announces multiple FDA 510(k) clearances, including first AI-assisted quality control software for digital pathology**

Leica Biosystems | 11 August 2026

Cleared for: Aperio iQC DX, standalone AI-assisted quality control for clinical digital pathology that automatically detects six acquisition artifacts while slides are still on the scanner (air bubbles, pen marks, clipped tissue, missing tissue, out-of-focus regions, striping), alongside clearance of the Aperio GT 180 DX scanner. Evidence level: vendor announcement; no published performance study is cited. Questions for the vendor: per-artifact sensitivity and specificity, measured impact on rescan rates in a live laboratory, robustness across stains and scanner fleets, and whether QC flags are logged in an exportable form. Monitoring you need anyway: track rescan rate and pathologist-reported missed artifacts from day one; a QC tool that silently under-flags is worse than none.

Source: [Leica Biosystems](#)**[FDA] CliniComp receives FDA 510(k) clearance for EHR-integrated PACS viewer**

CliniComp | 10 August 2026

Cleared for: a PACS Viewer as a Medical Image Management and Processing System (MIMPS), providing diagnostic-quality viewing, manipulation and processing inside the company's New Era EHR. What it is not: a viewer clearance carries no detection or triage claim, whatever the surrounding language about native AI prioritization of high-acuity cases implies. Evidence level: vendor materials only. Questions for the vendor: which AI-driven prioritization functions sit inside the cleared intended use and which are unregulated workflow features, what performance data support the prioritization claims, and whether display fidelity has been validated per modality, since a viewer changes what every reader sees on every case.

Source: [HIT Consultant](#)

[PARTNERSHIP] Coreline Soft and Optellum announce intent to collaborate across the lung cancer pathway

Optellum and Coreline Soft | 13 August 2026

What it is: a declared intent to collaborate, not a signed distribution agreement, to make Coreline Soft's FDA-cleared AVIEW Lung Nodule CAD detection software available through Optellum's LungOS platform, which carries its own cleared CADx risk assessment and nodule pathway management. The regulatory point buyers should hold onto: detection (CADE) and risk scoring (CADx) are separate claims cleared on separate evidence; wiring them into one workflow does not create a cleared end-to-end pathway. Questions for the vendors: which jurisdiction-specific claims will apply to the combined workflow, and what evidence covers the handoff from one system's detection output to the other's risk input.

Source: [PR Newswire](#)

SECTION 3

POST-MARKET SIGNAL

Drift | Recalls and MAUDE | Real-world monitoring | EU AI Act Article 72

SIIM publishes practical strategies for post-production monitoring of radiology AI

SIIM | 13 August 2026

SIIM published a recap of its webinar on monitoring deployed AI, and it reads as a field report on what Article 72-style obligations look like in practice. Penn Medicine runs a governance committee spanning demo approval through post-deployment monitoring. Dasa tracks drift on an in-house MRI denoising model by watching input voxel distributions and input-output similarity over time. Greensboro Radiology has declined FDA-cleared tools over poor real-world performance and monitors deployed algorithms by comparing predictions against radiologist impressions extracted from reports. The shared message: clearance is the start of the evidence obligation, not the end.

Link: <https://siim.org/artificial-intelligence/effective-strategies-for-post-productio...>

Radiologists report relatively few meaningful benefits from breast imaging AI

Radiology Business | 11 August 2026

A Clinical Imaging survey of 215 Society of Breast Imaging members, covered 11 August, is a real-world adoption signal worth filing: about half worked in organizations with breast cancer detection AI deployed, yet only a minority reported measurable improvement in callback rate, biopsy rate or burnout. Reported recall rates fell around 35% among current users against non-user expectations, unnecessary biopsies only about 9%, and 29.4% of users saw a notable burnout effect. Adoption is running ahead of measured benefit; departments should be tracking their own callback and biopsy deltas rather than assuming projected ones.

Link: <https://radiologybusiness.com/topics/artificial-intelligence/radiologists-report...>

SECTION 4

PLAYBOOK SNIPPET

One concrete step your AI committee can take this month

EU AI Act Article 72 (post-market monitoring); ISO/IEC 42001 performance evaluation clause

Step: Sample 20 AI outputs per month against the final signed report and track one number: the concordance rate. Plot it monthly and act on the trend, not on any single dip.

Why: Drift is invisible in aggregate accuracy and visible in a monthly line; a series started now is history you will already own when the monitoring obligation bites.

Worked example: Lakhani P et al. audited 400 AI-drafted DXA reports against the source images, not the prior reports, across four sites: 99.9 to 100% numerical accuracy, and the audit surfaced a 45% completeness rate in the human baseline at community sites. JACR, PMID 42595274, DOI 10.1016/j.jacr.2026.08.007.

SECTION 5

MEDIA HIGHLIGHTS

*AuntMinnie | Radiology Business | The Imaging Wire | Diagnostic Imaging | ITN***Factors that make radiologists less likely to be fooled by large language models***Radiology Business | 12 August 2026*

Coverage of a Radiology study (published 11 August) in which 10 readers interpreted 100 curated chest imaging cases with and without LLM assistance, randomized to a high-accuracy (GPT-5 based, 76%) or low-accuracy (GPT-4o based, 27%) model. Model confidence (OR 3.82) and reader expertise (OR 2.06) were independently associated with successful collaboration, and expertise protected against accepting incorrect suggestions (OR 0.54). The uncomfortable finding: higher rationale quality increased acceptance of incorrect suggestions (OR 1.71), meaning persuasive explanations cut both ways.

Link: <https://radiologybusiness.com/topics/artificial-intelligence/factors-make-radiol...>

American Board of Radiology taking cautious approach to AI*Radiology Business | 13 August 2026*

The ABR set out its AI position in a 12 August update: no AI is currently used in certification decisions or exam content creation, and scoring determinations remain with qualified human experts. The board has stood up an AI advisory group for governance and ethics and a trustee-level committee on assessment use cases, with guardrails established before any adoption. A certification body publishing its human-oversight boundary is itself a governance signal for the specialty.

Link: <https://radiologybusiness.com/topics/artificial-intelligence/american-board-radi...>

SECTION 6

PROFESSIONAL SOCIETIES

*SIIM | ACR | RCR | Official publications and event coverage***ACR releases detailed summary of the FY2027 IPPS final rule***13 August 2026*

The ACR published its detailed member summary of the CMS FY2027 Inpatient Prospective Payment System final rule on 13 August: a 2.3% hospital payment update, roughly \$779 million in new-technology case payments, a new technology add-on payment for an imaging AI triage device, and the end of the alternative NTAP pathway from FY2028, after which every technology seeking add-on payment must demonstrate substantial clinical improvement. For AI vendors and buyers alike, the evidence bar for US inpatient reimbursement just moved up.

Source: <https://www.acr.org/News-and-Publications/2026/fy2027-ippes-final-rule-detailed-summary>

SECTION 7

KEY OPINION LEADERS

*LinkedIn | Public statements | Week of 9 to 15 August 2026***Curt Langlotz***Professor of Radiology, Medicine, and Biomedical Data Science, Stanford; Director, Stanford AIMI*

Active week of curation (posts of 12 and 15 August): amplified Eric Horvitz's work on confidence-based control of clinical report generation, a threshold-controlled approach where an operating point on the false-positive curve gates what an LLM writes into the draft report; reshared Stanford AIMI's new pediatric echocardiography AI work on congenital heart disease published in *Circulation*; and relayed the *Nature Health* analysis of 617,827 health conversations with Microsoft Copilot, noting that Copilot is the small platform next to ChatGPT's reported 40 million health queries.

Source: [LinkedIn](#)

Amine Korchi*Radiologist and HealthTech investor, Geneva*

Post of 13 August: flagged Vocca as a startup to watch as it expands from San Francisco to New York, part of his continuing thread on voice AI agents taking over administrative load in medical practices rather than diagnostic tasks, the segment of medical AI where deployment friction and regulatory burden are lowest.

Source: [LinkedIn](#)

Daniel Pinto dos Santos*Radiologist, University Hospitals of Cologne and Frankfurt; EuSoMII board*

Repost of 12 August: amplified Marius George Linguraru's announcement of the joint EuSoMII and MICCAI session on the future of AI in medical imaging at the EuSoMII Annual Meeting in Heraklion this October, with safe integration of AI tools in clinical practice as the stated theme, a signal of the clinical and technical societies converging on deployment safety rather than model development.

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 9 to 15 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.