

WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 38 | 13 to 19 September 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

2

Industry / FDA items

2

Media highlights

3

Key opinion leaders

GOVERNANCE TAKEAWAY OF THE WEEK

This week makes one deployment lesson concrete: neither clearance nor good aggregate accuracy decides where a tool should run. A prospective shadow-mode study of an FDA-cleared aneurysm detector found operational value strongly favorable in inpatients, marginal in the emergency department and unfavorable in outpatients, with the same algorithm in the same hospital. An expert panel on AI ARIA detection reached the matching policy answer from the governance side: conditional implementation only, with a radiologist in the loop, quality assurance and prospective performance monitoring. The action for a radiology leader: write deployment scope per care setting into acceptance tests, and use shadow mode to measure incremental yield and gain-to-pain in each setting before the tool touches a live workload.

SECTION 0

EXECUTIVE SUMMARY

Key trends, week of 13 to 19 September 2026

[EVIDENCE] ONE CLEARED ALGORITHM, THREE DEPLOYMENT ANSWERS

Goldberg-Stein S et al. (JACR) ran Aidoc's FDA-cleared aneurysm detector in prospective shadow mode over 3,856 CTAs: gain-to-pain 2.57 in inpatients, 1.0 in the emergency department, 0.67 in outpatients. Same tool, same hospital, opposite deployment decisions by setting. Shadow mode delivered the answer before any workflow changed.

[OVERSIGHT] ARIA AI GETS CONDITIONS, NOT A GREEN LIGHT

Petrella JR et al. (AJR) publish a multidisciplinary expert position on AI ARIA detection under anti-amyloid therapy: likely to enhance safety, but only as decision support with a radiologist in the loop, ongoing QA and prospective performance monitoring, against a market of tools that vary substantially in regulatory status and validation evidence.

[VALIDATION] THE LEADERBOARD INVERTS ON NORMAL CASES

Lalinia M et al. (J Imaging Inform Med) release FAR-POLYP-SEG, 8,181 prospective colonoscopy frames from 455 patients: the models with the best internal Dice produced false-positive rates up to 59.5% on 7,749 normal-mucosa frames, an almost complete inversion of the ranking. Test sets without normal studies cannot predict alert burden.

[CLINICAL] NON-INVASIVE RCC GRADING HITS ITS EVIDENCE CEILING

Kiani I et al. (European Radiology) pool 20 studies of quantitative MRI against WHO/ISUP grade in renal cell carcinoma: ADC alone reaches AUC 0.71 with weak specificity, radiomics and deep learning reach AUC 0.90 but heterogeneously. External validation and decision-impact studies are still the missing step before biopsy triage.

[WATCHPOINT] LOCAL CONTEXT IS THE VARIABLE, NOT THE MODEL

Three papers this week make the same point from different angles: an aneurysm detector whose operational value flips sign between inpatients and outpatients, segmentation models whose internal ranking inverts on normal mucosa, and an ARIA panel that conditions deployment on local QA and monitoring. What varies is not the algorithm but the population, casemix and setting it meets. Acceptance tests that ignore setting are measuring a different deployment than the one you will run.

SECTION 1

PEER-REVIEWED PAPERS

Source: PubMed | Verified indexing | No preprints | 13 to 19 September 2026

Prospective Shadow-Mode Evaluation of an Artificial Intelligence Tool for Intracranial Aneurysm Detection on CT Angiography: Incremental Yield and Operational Impact.

Goldberg-Stein S, et al. | *Journal of the American College of Radiology* | 2026-09-15

Before extending a cleared detection tool beyond the setting it was bought for, run it in shadow mode and cut the results by care setting: this prospective study processed 3,856 consecutive brain CT angiograms through Aidoc's FDA-cleared aneurysm algorithm with radiologists blinded to the AI, adjudicating only the discordant cases. The design produced operational deployment metrics (incremental yield, gain-to-pain, number-needed-to-examine) before a single patient pathway was changed.

Key metrics: 3,856 CTAs, examination positive rate 5.1% (195 of 3,856); radiologist-AI concordance 96.3%. AI sensitivity 0.846 (0.787 to 0.894) vs radiologist 0.718 (0.649 to 0.779); specificity 0.987 vs 0.985. Relative enhanced detection rate 39% (55 AI-only true positives per 140 radiologist true positives); incremental detection ratio 1.83; gain-to-pain ratio 1.20; number-needed-to-examine 70.1. By setting: inpatient rEDR 78.3%, GPR 2.57, NNE 28.9; emergency GPR 1.0, NNE 85.3; outpatient rEDR 14.1%, GPR 0.67, NNE 130.3. 56.4% of AI-only aneurysms were under 3 mm.

Clinical relevance: The same cleared algorithm was operationally favorable for inpatients, marginal in the emergency department and unfavorable for outpatients, in one institution over one period. Clearance answers none of that; only local measurement does. For your AI committee this changes acceptance testing and monitoring: adopt shadow mode as the standard acceptance-test method for detection tools, require operational metrics per care setting, and write deployment scope by setting rather than hospital-wide; this is concrete Article 9 risk management and it seeds the Article 72 post-market monitoring series.

PMID 42742498 DOI 10.1016/j.jacr.2026.07.018

Quantitative MRI for World Health Organisation/International Society of Urological Pathology Grading of Renal Cell Carcinoma: a systematic review and diagnostic meta-analysis.

Kiani I, et al. | *European radiology* | 2026-09-18

Before treating quantitative MRI as a substitute for biopsy in grading renal cell carcinoma, look at what pooling the evidence actually shows: this systematic review and diagnostic meta-analysis synthesised 20 studies of quantitative MRI biomarkers (diffusion, relaxometry, chemical exchange saturation transfer, radiomics) against WHO/ISUP grade, with formal random-effects meta-analysis feasible for 12.

Key metrics: Seven ADC studies: pooled sensitivity 0.84 (95% CI 0.77 to 0.89), specificity 0.57 (95% CI 0.51 to 0.63), summary AUC 0.71 for high-grade disease; low-grade tumours showed higher ADC (mean difference 0.21×10^{-3} mm²/s; 95% CI 0.11 to 0.30). Five MRI-inclusive radiomics/deep learning studies: pooled sensitivity 0.79 (0.64 to 0.89), specificity 0.86 (0.74 to 0.93), AUC 0.90.

Clinical relevance: Diffusion metrics alone grade RCC with modest accuracy and weak specificity, which limits their value for biopsy triage on their own. Radiomics and deep learning models pool substantially better but with less consistency across studies, and the authors are explicit that standardised multiparametric protocols, external validation and decision-impact studies are still required before clinical implementation. Kept on clinical merit: this is the current evidence ceiling for non-invasive RCC grading.

PMID 42758284 DOI 10.1007/s00330-026-12837-5

Deep Learning Based on Magnetic Resonance Imaging for Preoperative Prediction of Pituitary Neuroendocrine Tumors Subtypes.

Yang Z, et al. | *Academic radiology* | 2026-09-14

Molecular subtyping of pituitary neuroendocrine tumours currently requires postoperative pathology; this two-centre study trained a YOLOv11 deep learning model to predict transcription-factor lineage (PIT1, TPIT, SF1, or no distinct lineage) preoperatively from routine MRI sequences in 888 patients, and tested it with internal, temporal external and multicenter validation.

Key metrics: 888 patients enrolled January 2021 to December 2025; mean age 54.32 years, 52.9% female. Best performance from sagittal contrast-enhanced T1WI (mAP50-95 0.901); weakest from coronal T2WI (mAP50-95 0.756). The abstract reports robust performance across internal, temporal external and multicenter cohorts without further per-cohort figures.

Clinical relevance: If lineage can be read from preoperative MRI, medical-first strategies for some PIT1 and SF1 tumours become plannable before surgery rather than after pathology. The validation design is stronger than most single-centre work in this space (temporal external plus multicenter), but the cohorts are retrospective and the authors position the model as warranting prospective validation before it aids personalised management. Presented on clinical merit.

PMID 42736093 DOI 10.1016/j.acra.2026.08.069

Automated three-dimensional radiomic body composition analysis enhances survival prediction in resectable non-small cell lung cancer.

Huang Y, et al. | *European radiology experimental* | 2026-09-18

Standard staging of resectable non-small cell lung cancer ignores the host: this multicenter study built a fully automated deep learning pipeline for three-dimensional body composition segmentation, extracted radiomic features from both tumour and body composition, and integrated them with extreme gradient boosting to predict overall survival across training, internal and external validation cohorts.

Key metrics: 1,038 patients (mean age 61.8 +/- 10.7 years; 58.66% male); 293 deaths (28.2%) over median follow-up 3.31 years. Tumour score and body composition score independently associated with overall survival (hazard ratios 2.72 and 2.03; both $p < 0.001$). Adding body composition significantly improved discrimination over tumour-only models across cohorts (all $p < 0.05$); the comprehensive model reached AUCs above 0.80 for 1, 2, 3 and 5-year survival. SHAP-derived phenotypes stratified four prognostic groups (all log-rank $p < 0.05$).

Clinical relevance: Body composition is prognostic information the CT already contains and nobody bills for; automating its extraction turns it into a reproducible input for postoperative risk stratification and follow-up intensity decisions. External validation and interpretability analysis are in place, and the incremental value over tumour-only models is formally tested. Kept on clinical merit.

PMID 42758420 DOI 10.1186/s41747-026-00802-2

Radiomics-based prediction of pituitary adenoma consistency: a systematic review and meta-analysis.

Fusco G, et al. | *La Radiologia medica* | 2026-09-18

Pituitary adenoma consistency drives the surgical approach but conventional imaging cannot call it reliably; this systematic review and meta-analysis (pre-registered, PRISMA-DTA) pooled 14 radiomics studies predicting consistency from MRI to establish what the field can currently deliver for presurgical planning.

Key metrics: 14 studies; pooled AUC 0.86 (95% CI 0.76 to 0.92, $I^2 = 92.45\%$). HSROC sensitivity 0.74 (95% CI 0.56 to 0.87), specificity 0.79 (95% CI 0.73 to 0.84). No significant performance differences across algorithms, dimensionality or sequence; T2-weighted and combined-sequence models trended higher. Risk of bias low to moderate; heterogeneity high with small-study effects in the main analysis.

Clinical relevance: Good pooled discrimination, but the caveats are the finding: heterogeneity of 92%, small-study effects and limited external validation mean the pooled AUC is an optimistic ceiling, not an expected local performance. The authors restrict clinical generalizability accordingly and call for multicenter external validation and standardised radiomics workflows. Presented on clinical merit as the state of the evidence for a genuinely surgery-relevant prediction.

PMID 42758236 DOI 10.1007/s11547-026-02281-2

Artificial Intelligence-Based Detection of ARIA on MRI During Alzheimer Disease Therapy: Expert Opinion on Responsible Clinical Integration.

Petrella JR, et al. | *AJR. American journal of roentgenology* | 2026-09-16

Before buying an AI tool for ARIA surveillance under anti-amyloid therapy, adopt the conditions this multidisciplinary panel of neuroradiologists and Alzheimer disease clinicians attached to deployment: clinical decision support only, inside a radiologist-in-the-loop framework, with ongoing quality assurance and prospective monitoring of clinical performance. The panel reviewed the evidence base for AI-assisted ARIA detection as surveillance MRI volumes grow with routine anti-amyloid prescribing.

Key metrics: This is a consensus document; the abstract reports no pooled quantitative performance figures. Its substantive findings are qualitative: substantial variation among commercially available ARIA tools in regulatory status, technical capabilities and validation evidence, and an evidence gap on whether improved detection changes clinical outcomes.

Clinical relevance: The panel's position is a governance template: conditional implementation, named oversight, QA and prospective monitoring as preconditions rather than aspirations. For your AI committee this changes validation and monitoring: treat ARIA tools as a class where vendor regulatory status must be verified case by case, and where deployment approval should be written as conditional on a running local performance series; this maps directly to EU AI Act Article 14 human oversight and Article 72 post-market monitoring.

PMID 42747213 DOI 10.2214/AJR.26.35514

FAR-POLYP-SEG: A Prospective Single-Center Colonoscopy Dataset for Colorectal Polyp Segmentation with Patient-Level Metadata and Baseline Cross-Dataset Evaluation.

Lalinia M, et al. | *Journal of imaging informatics in medicine* | 2026-09-15

Before letting an internal leaderboard pick your segmentation model, test it on normal cases: FAR-POLYP-SEG is a prospective single-centre colonoscopy dataset built as a validation resource, with 8,181 frames from 455 patients (432 expert-segmented polyp frames plus 7,749 normal-mucosa frames) and patient-level clinical metadata, on which six architectures were benchmarked under one standardised patient-grouped protocol.

Key metrics: Patient-grouped five-fold cross-validation, no patient in both train and test of any fold. PraNet best internal Dice 0.755; nnU-Net best internal IoU 0.665. On the 7,749 normal-mucosa frames, gated false-positive rates ran from 24.9% (YOLOv11m-seg) to 59.5% (nnU-Net), described by the authors as an almost complete inversion of the internal Dice ranking. External Kvasir-SEG Dice 0.756 to 0.829 with ordering largely preserved. Dataset, metadata and code publicly released.

Clinical relevance: The model that wins on lesion overlap can be the one that alarms most on healthy tissue: a ranking built only on positive cases predicts neither alert burden nor specificity in service. For your AI committee this changes acceptance testing: require normal studies in every local test set, report false positives on them as a separate figure alongside the headline metric, and weight the choice of tool by both; this is the practical content of Article 15 accuracy and robustness and of Article 10 test-data governance.

PMID 42745112 DOI 10.1007/s10278-026-02268-5

Multi-contrast MRI acceleration via post-reconstruction fusion.

Nazarov A, et al. | *Medical image analysis* | 2026-09-14

Multi-contrast brain MRI is slow because each 3D contrast is acquired at full resolution independently; this work accelerates the protocol rather than the sequence, acquiring each contrast with a fourfold-reduced phase-encoding matrix along a contrast-specific axis on standard scanner settings, then fusing the complementary volumes with the FARD network (frequency attention residual denoising) to restore apparent isotropic 1 mm³ detail across contrasts.

Key metrics: Approximately 4x protocol acceleration on the prospective clinical-scanner cohort (10.4 to 2.6 minutes). FARD achieved the strongest or near-strongest PSNR and SSIM among evaluated reference-free methods and approached reference-guided performance without a high-resolution reference contrast. Complementary large-scale retrospective evidence on BraTS-GLI 2024, which does not model prospective scanner effects.

Clinical relevance: The practical point is deployability: no raw k-space access, no non-Cartesian trajectories, no custom sequences, so the approach can run on standard clinical acquisitions. Evidence includes a prospectively acquired clinical-scanner cohort under real vendor-reconstruction conditions, which most acceleration papers lack. Diagnostic-quality evaluation by readers on pathology remains the open step before clinical use. Presented on clinical merit.

PMID 42753300 DOI 10.1016/j.media.2026.104297

SECTION 2

INDUSTRY & REGULATION

Sources: [Official registers](#) | [Press releases](#) | [FDA databases](#)**[FDA] DeepHealth's foundation-model chest X-ray CAde clears the FDA***DeepHealth* | 16 September 2026

Cleared for what: adjunctive computer-aided detection and localization of suspected thoracic abnormalities on chest radiographs (lung nodules, pneumothorax, consolidation, mediastinal abnormalities), delivered through DeepHealth's Radiology AI Suite; the company presents Chest XRay as one of the first FDA-cleared chest X-ray products built on a proprietary foundation model. Evidence: company announcement and trade coverage; the 510(k) route means substantial equivalence to a predicate, and no peer-reviewed performance study is cited in the announcement. Questions for the vendor: which findings are inside the cleared claim and at what operating points; what data supported the 510(k) and how the foundation-model backbone was adapted and locked; whether a Predetermined Change Control Plan applies and what changes it lets the vendor push without a new clearance; performance by age group and scanner vendor. Foundation-model marketing does not change the deployer's job: the cleared claim, not the backbone, defines what you must locally validate, and a CAde adjunct still needs baseline reader metrics before go-live.

Source: [Diagnostic Imaging](#)**[FDA] MICS-PET: MR-guided PET enhancement with automated amyloid and tau quantification cleared***Microstructure Imaging (MICS)* | 15 September 2026

Cleared for what: a neurological PET software platform combining MR-guided PET denoising and super-resolution with automated quantification, including Centiloid for amyloid PET, CenTauR for tau PET, regional SUVr and z-scores against normative controls, supporting amyloid, tau and FDG tracers; MICS, an NYU Langone spinout, calls it the first FDA-cleared MR-guided PET enhancement software. Evidence: company press release citing development and validation across major-vendor PET and MRI systems; no peer-reviewed clinical accuracy study is referenced in the release. Questions for the vendor: what reference standard anchored the quantification validation and on which scanner generations; how enhancement affects quantitative fidelity in borderline amyloid cases, since the same processing that sharpens images can shift SUVr; which normative database generates the z-scores and how it matches your population. Relevant to any centre building anti-amyloid therapy workflows, where quantification feeds ARIA-adjacent treatment decisions; set up concordance checks against your current quantification pipeline before switching.

Source: [Business Wire](#) via [BioSpace](#)

SECTION 3

POST-MARKET SIGNAL

Drift | *Recalls and MAUDE* | *Real-world monitoring* | *EU AI Act Article 72***COMMENTARY, NOT REPORTED NEWS****Commentary: this week's evidence is an argument for setting-stratified monitoring**

No new post-market surveillance event (drift report, recall, MAUDE signal or Article 72 guidance) surfaced in the window, so this is commentary on this week's evidence rather than reported news. Two of this week's papers show why a single monitoring line per tool is not enough. The JACR shadow-mode study found the same cleared aneurysm algorithm operationally favorable in inpatients (gain-to-pain 2.57) and unfavorable in outpatients (0.67): a pooled concordance metric would have averaged those into a false reassurance. The FAR-POLYP-SEG benchmark found false-positive rates on normal tissue inverting the internal accuracy ranking of six models. The monitoring consequence of both: stratify your monthly concordance series by care setting, and track false positives on studies the final report calls normal as its own line. Aggregate concordance can hold steady while one setting quietly degrades.

SECTION 4

PLAYBOOK SNIPPET

*One concrete step your AI committee can take this month***MDR and FDA intended-use statements; conformity assessment**

Step: For every cleared tool you run or plan to buy, copy the intended-use sentence verbatim to the top of its deployment file. Under it, list every way your intended use differs: population, modality, scanner fleet, care setting, reader mix. Treat each listed difference as one named local validation task with an owner and a date.

Why: Clearance is a statement about a specific claim on specific evidence, not about your deployment. A viewer clearance is not a detection claim, and a De Novo authorization means no predicate device exists, so the local validation burden is higher, not lower.

Worked example: Goldberg-Stein S et al. show what one such difference is worth: the same FDA-cleared aneurysm algorithm, measured prospectively in shadow mode over 3,856 CTAs, returned a gain-to-pain ratio of 2.57 in inpatients and 0.67 in outpatients. The care-setting line in the deployment file flips the decision by itself. JACR, PMID 42742498, DOI 10.1016/j.jacr.2026.07.018.

SECTION 5

MEDIA HIGHLIGHTS

*AuntMinnie | Radiology Business | The Imaging Wire | Diagnostic Imaging | ITN***AuntMinnie: AI boosts brain aneurysm detection on CT angiography***AuntMinnie.com | 16 September 2026*

AuntMinnie's coverage of the Northwell Health JACR shadow-mode study (this week's lead paper) pushed it into the site's most-read list of the week. The trade angle matches the study's own caution: the tool surfaced 55 aneurysms radiologists had not reported, most under 3 mm, but operational value varied sharply by care setting. The story reaching the top of the most-read list signals that operational, setting-level evaluation is becoming the frame trade media apply to detection AI, not just standalone accuracy.

Link: <https://www.auntminnie.com/clinical-news/ct/article/15835193/ai-boosts-brain-ane...>

Diagnostic Imaging podcast: AI reimbursement remains radiology's unsolved business case*Diagnostic Imaging | 17 September 2026*

Part 2 of the Reading Room podcast on AI and reimbursement, with Lauren Nicola, MD, Christoph Wald, MD, PhD, MBA, and Peter Shen, examines why evaluating value and return on investment for AI software remains hard, the physician time it consumes, and possible paths to more consistent AI reimbursement. For deployers the practical point is that without reimbursement clarity, the business case for most tools still rests on efficiency and risk reduction you must measure yourself, which is one more reason acceptance testing and monitoring data double as financial evidence.

Link: <https://www.diagnosticimaging.com/view/reading-room-podcast-current-landscape-ai...>

SECTION 6

PROFESSIONAL SOCIETIES

SIIM | ACR | RCR | Official publications and event coverage

Von der Leyen makes AI-supported mammography the flagship of "AI for Health"; ESR responds

16 September 2026

In her 2026 State of the Union address to the European Parliament on 16 September, European Commission President Ursula von der Leyen chose AI-supported breast cancer screening as the leading example of how AI can transform healthcare in Europe, citing screening's 30 to 40 percent mortality reduction among participating women. The ESR welcomed the recognition of radiology's role. For imaging leaders this is political capital: screening AI now has explicit top-level EU backing, which will accelerate funding calls and deployment pressure ahead of the first EU Screening Week (28 September to 4 October), and makes disciplined local validation the differentiator between programmes that can show governance and those that only show adoption.

Source: <https://www.myesr.org/ai-powered-radiology-saves-lives/>**First radiologist appointed to the USPSTF; ACR welcomes the appointment**

17 September 2026

HHS announced eight new members of the US Preventive Services Task Force on 17 September, including Dennis Wulfeck, MD, MBA, DHA, a diagnostic radiologist with RadPartners MBB Radiology, the first radiologist to serve on the 16-member panel that sets US screening recommendations. The ACR, which has long advocated for specialty representation on the USPSTF, issued a statement welcoming the appointment. Task force composition shapes which imaging-based screening services get graded, and therefore covered; a radiologist voice inside that process matters for how future AI-supported screening evidence will be judged.

Source: <https://www.auntminnie.com/practice-management/careers/imaging-leaders/article/15835216/uspstf-names-radiologist-to-task-force-acr-responds>**EuSoMII joins the Brazilian Congress of Radiology with a live joint session**

18 September 2026

EuSoMII held a joint "CBR meets EuSoMII" session on 18 September at the 55th Brazilian Congress of Radiology (CBR26, Recife, 17 to 19 September), livestreamed for remote attendees, extending the European imaging-informatics society's collaboration into Latin America. Cross-regional society programming of this kind is where imaging-AI governance practice spreads between markets with very different regulatory regimes, and it signals EuSoMII's ambition beyond Europe ahead of its 2026 annual meeting in Crete.

Source: <https://www.linkedin.com/company/eusomii/posts/>

SECTION 7

KEY OPINION LEADERS

LinkedIn | Public statements | Week of 13 to 19 September 2026

Curt Langlotz

Professor of Radiology, Medicine, and Biomedical Data Science, Stanford University; Stanford AIMI

Amplified (17 September) Stanford AIMI's announcement of MedVAL, published in npj Digital Medicine (Asad Aali, Vasiliki Bikia et al.): 'Language models can generate clinical text in seconds. But verifying that text is factually accurate still requires valuable' expert time. MedVAL is a framework for expert-level validation of generated medical text, benchmarking model-as-judge validation against physician judges with quantitative evaluation. The relevance for deployers of report-drafting AI: the bottleneck is shifting from generation to verification, and validation tooling is becoming a research object in its own right.

Source: [LinkedIn](#)

Amine Korchi

Neuroradiologist, Geneva; HealthTech innovation and ventures

Posted 16 September on a newly published comparative evaluation in Chest of generative AI models for chest radiograph report generation in the emergency department, calling it an important study on vision-language models for CXR reporting and sharing side-by-side model outputs against the radiologist report. His framing echoes the week's verification theme: draft quality varies widely across models, and the examples show confident, detailed but divergent narratives from the same image.

Source: [LinkedIn](#)

Daniel Pinto dos Santos

Radiologist and imaging informatics researcher, Cologne/Frankfurt; EuSoMII

Amplified (18 September) the live 'CBR meets EuSoMII' joint session at the 55th Brazilian Congress of Radiology in Recife, and (17 September) ESR Journals' feature of a European Radiology paper on topogram-based anatomical labelling of CT series (RAPID), which tackles inconsistent DICOM metadata, the unglamorous data-governance layer any CT AI pipeline depends on before a single model can run reliably.

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 13 to 19 September 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.