

WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 39 | 20 to 26 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

4

Industry / FDA items

3

Media highlights

4

Key opinion leaders

GOVERNANCE TAKEAWAY OF THE WEEK

This week draws the line between accuracy evidence and deployment evidence unusually clearly. A Nature Medicine screening model arrives with 12-center validation on 80,612 patients, real-world calibration and a prospective cohort, while the week's meta-analyses and benchmarks, LVO triage on noncontrast CT, cerebral microbleed detection, LLM decision support, all stop at retrospective accuracy, and say so. The action for a radiology leader: add one line to every AI proposal your committee reviews, classifying its evidence as accuracy-based or workflow-based, and refuse to approve an unsupervised or workflow-changing role on accuracy-only evidence. A bounded human-in-the-loop role with a prospective monitoring plan is the correct ceiling until the workflow evidence exists.

SECTION 0

EXECUTIVE SUMMARY

Key trends, week of 20 to 26 September 2026

[EVIDENCE] WHAT DEPLOYMENT-GRADE VALIDATION LOOKS LIKE

Zhou J et al. (Nature Medicine) validate EAGLE, an esophageal cancer detector on plain chest CT, across 12 centers and 80,612 patients: external testing, low-dose CT, a 35,402-scan real-world calibration cohort that cut false positives 72.7%, and a prospective cohort reaching PPV 42.2%. The staged trajectory, not the headline sensitivity, is the lesson.

[OVERSIGHT] AI READS NCCT BETTER, BUT ONLY ON PAPER

Zhou H et al. (JMIR) pool 10 retrospective studies: NCCT-based AI is 16 to 18 points more sensitive than human readers for large vessel occlusion, at low GRADE certainty and with zero workflow evidence. The authors bound the defensible role to a human-in-the-loop escalation prompt, not autonomous triage.

[VALIDATION] COMPLEXITY IS WHERE LLM SAFETY BREAKS

Eroglu OO et al. (JMIR) benchmark three LLMs on ASCCP cervical screening management with complexity-oversampled scenarios: error rates climb from 3.1% on simple cases to 29% on complex ones, and the prompt package shifts safety as much as the model choice. Acceptance tests that sample only routine cases will miss this entirely.

[QUALITY] OPEN-SOURCE REPORT QA REACHES COMMERCIAL PARITY

Li J et al. (European Radiology Experimental) show a six-step radiological chain-of-thought framework lifting seven LLMs on 1,200 QA-repository reports (mean F1 0.77 to 0.85); Llama-3.3-70B with RadCoT matches GPT-4o standard, making on-premises, privacy-preserving report QA a realistic design.

[WATCHPOINT] ACCURACY EVIDENCE IS OUTRUNNING WORKFLOW EVIDENCE

Every comparative result this week, LVO triage, microbleed detection, LLM decision support, stops at accuracy measured retrospectively. The one paper that goes further, EAGLE, needed 12 centers, real-world calibration and a prospective cohort to do it. Read every vendor claim this week against that gap: what was measured, and in whose workflow.

SECTION 1

PEER-REVIEWED PAPERS

Source: PubMed | Verified indexing | No preprints | 20 to 26 September 2026

Large-scale esophageal cancer screening through noncontrast computed tomography and artificial intelligence.

Zhou J, et al. | *Nature medicine* | 2026-09-22

Before treating opportunistic screening claims as aspirational, study this as the reference case: EAGLE detects esophageal cancer and precancerous lesions on plain chest noncontrast CT, a task long considered impossible, and was validated the way a screening tool must be. Trained on 6,813 patients from two centers, it was then tested across 12 centers in three countries on 80,612 patients, spanning opportunistic reuse of existing CT, low-dose CT within lung screening, a real-world calibration cohort and a prospective hospital cohort.

Key metrics: Multicenter external test (8 centers, n = 11,466): specificity 98.5%, sensitivity 90.0% for cancer and 52.5% for precancerous lesions. LDCT validation (2 centers, n = 1,607) comparable. Real-world calibration (3 centers, n = 35,402): false positives reduced 72.7% with sensitivity preserved. Prospective hospital validation (n = 17,446): PPV 42.2%. Real-world low-dose screening (n = 10,959): specificity 99.94%. Paired CT-endoscopy cohorts (n = 702): sensitivity 65.0% for precancerous lesions and 78.4% for stage I cancer at a higher-sensitivity operating point. Registered: ChiCTR2300074806.

Clinical relevance: Clinically, this reframes the chest CT already flowing through every lung screening program as an esophageal screening opportunity, with endoscopy referral for the high-risk minority. The staged evidence, from external testing through real-world calibration to prospective PPV, is what a deployment-grade validation trajectory looks like in imaging AI.

PMID 42773211 DOI 10.1038/s41591-026-04656-4

MRI-based fetal gestational age estimation using a structure-aware self-supervised network.

Gan H, et al. | *European radiology* | 2026-09-21

When menstrual dating and first-trimester ultrasound are unavailable or unreliable in mid-to-late pregnancy, fetal brain MRI may offer a quantitative fallback: this structure-aware self-supervised network estimates gestational age from routine coronal T2-weighted images to within about 0.8 weeks.

Key metrics: 1,630 fetal brain MR images from 207 singleton pregnancies (GA range 22 to 38 weeks), single institution, retrospective, patient-level 80/20 split. Test set MAE 0.793 weeks (95% CI 0.574 to 0.860); R2 0.934 (95% CI 0.918 to 0.967); p < 0.001 for linear association with reference GA.

Clinical relevance: The reference standard itself (last menstrual period plus first-trimester ultrasound) is imperfect in exactly the cases where the tool would be used, and the study is single-center with no external test set. The authors are explicit that prospective, multicenter validation is required before clinical deployment; until then this is a promising complement to conventional dating in selected cases, not a replacement.

PMID 42768216 DOI 10.1007/s00330-026-12845-5

Biphasic CT clustering-based habitat radiomics predicts WHO/ISUP nuclear grade in clear cell renal cell carcinoma.

Luo S, et al. | *Abdominal radiology (New York)* | 2026-09-22

For noninvasive WHO/ISUP grading of clear cell renal cell carcinoma, this study asks a sharper question than most radiomics work: which part of the tumor carries the signal. Habitat clustering of biphasic CT split tumors into five subregions, and the viable-necrosis transition zone alone matched whole-tumor models.

Key metrics: 473 ccRCC patients (training n = 337, test n = 136), retrospective, nnU-Net segmentation, k-means habitats on corticomedullary and excretory phase attenuation and their difference. Transition-zone (subregion 4) random forest AUC 0.773 vs whole-tumor 0.819 in CMP (DeLong p = 0.152) and 0.810 vs 0.827 in EP (p = 0.593); necrotic-core model AUC 0.771 vs 0.802 in EP (p = 0.219); subregion 1+4 combination performed best among multi-subregion models.

Clinical relevance: The spatial mapping is the contribution: grading signal concentrates at the viable-necrosis interface, which is biologically plausible and gives the model a visualizable anatomical anchor rather than an opaque feature vector. Evidence remains retrospective and single-cohort; the authors call for prospective multicenter validation with histopathological correlation before clinical use.

PMID 42771041 DOI 10.1007/s00261-026-05793-7

Foundation Models in Ophthalmic Artificial Intelligence: Current Status and Future Directions.

Ren L, et al. | *IEEE journal of biomedical and health informatics* | 2026-09-22

For a structured view of where foundation models actually stand in ophthalmic imaging, this review organizes the field into vision foundation models, vision-language models and LLM-based agentic systems, maps the public multimodal datasets behind them (OCT, color fundus, slit lamp), and names the translation barriers plainly.

Key metrics: Narrative review; the abstract reports no quantitative performance figures. Scope covers structural, vascular and anterior-segment modalities, public dataset inventories, pretraining and cross-modal representation strategies.

Clinical relevance: The barrier list is the useful part for any imaging domain, not only ophthalmology: data imbalance across modalities, imperfect image-text semantic alignment, limited robustness and generalizability, hallucinations in generative agents, and workflow integration. The authors' proposed remedies, standardized multimodal data resources and rigorous multi-center validation, describe the gap between benchmark performance and clinical readiness.

PMID 42771632 DOI 10.1109/JBHI.2026.3736832

AI Decision Support for Emergency Triage of Large Vessel Occlusion Using Noncontrast Computed Tomography: Systematic Review and Bayesian Diagnostic Test Accuracy Network Meta-Analysis.

Zhou H, et al. | *Journal of medical Internet research* | 2026-09-24

Before positioning noncontrast CT AI as an LVO triage layer, fix its role in writing: this Bayesian network meta-analysis finds AI more sensitive than human readers on NCCT, but on retrospective, accuracy-only evidence, and the authors themselves bound the appropriate use to a human-in-the-loop escalation prompt for CTA prioritization, tele-stroke consultation or transfer discussion.

Key metrics: 10 retrospective validation studies, 11 datasets, 28 study arms, 3,632 patients. Unimodal imaging AI: sensitivity 0.78 (95% CrI 0.68 to 0.86), specificity 0.88 (0.81 to 0.94), DOR 30.89. Expert readers: sensitivity 0.62 (0.48 to 0.75); nonexpert readers 0.60 (0.45 to 0.75); specificity approximately 0.86 for both. AI sensitivity advantage +0.16 over experts and +0.18 over nonexperts, no clear specificity separation. Multimodal AI: sensitivity 0.81, specificity 0.92, indirect evidence only. GRADE certainty: low.

Clinical relevance: The evidence is accuracy-based and establishes nothing about workflow effectiveness, reperfusion times or outcomes, and the certainty is low. For your AI committee this changes acceptance testing: an NCCT LVO tool should be accepted only into a bounded escalation role with a named human decision-maker downstream, which is precisely the Article 14 human-oversight design, and any autonomy claim should trigger a demand for prospective workflow evidence.

PMID 42785737 DOI 10.2196/105068

RadCoT: a Radiological Chain-of-Thought framework for enhanced error detection in radiology reports.

Li J, et al. | *European radiology experimental* | 2026-09-22

Before buying a commercial LLM layer for report quality assurance, test what structured prompting does to an open-source model on your own error repository: RadCoT, a six-step chain-of-thought framework mirroring section-based radiological review, lifted all seven tested models and brought Llama-3.3-70B to parity with GPT-4o under standard prompting.

Key metrics: 1,170 clinician-validated errors from a departmental QA repository (2021 to 2024) in 900 reports, plus 300 error-free controls; 1,200 reports balanced across radiography, ultrasound, CT and MRI; five error types. Mean micro-averaged F1 rose from 0.77 (SD 0.06) with standard prompting to 0.85 (SD 0.06) with RadCoT ($p = 0.003$). GPT-4o with RadCoT: F1 0.93. Llama-3.3-70B with RadCoT: F1 0.89 vs GPT-4o standard 0.88 (adjusted $p = 0.28$). Interpretation errors: F1 0.57 to 0.75.

Clinical relevance: An on-premises, privacy-preserving QA layer built on an open-source model is now a defensible design, which matters wherever report text cannot leave the institution. For your AI committee this changes monitoring and acceptance testing: a curated, clinician-validated local error set like this study's 1,170-error repository is the acceptance benchmark any report-QA tool should be run against before deployment, and the prompting framework must be version-controlled as part of the validated system under Article 15 accuracy and robustness.

PMID 42771118 DOI 10.1186/s41747-026-00804-0

Safety-Oriented Benchmarking of Large Language Models in Risk-Based Management of Abnormal Cervical Screening Results: Scenario-Based Benchmark Study.

Eroglu OO, et al. | *Journal of medical Internet research* | 2026-09-22

Before any LLM touches risk-based management decisions, benchmark it the way this group did: oversample the complex, history-dependent scenarios where models actually fail. Across three LLMs on ASCCP cervical screening management, safety varied sharply by model, prompt configuration and case complexity.

Key metrics: 60 synthetic scenarios oversampling complex decision nodes, 3 models, 2 prompt arms, 3 repetitions each, 1,080 observations, blinded specialist rating (weighted kappa 0.839). Unsafe major error-free rate under the guideline-directed prompt package: GPT-5.3 100% (95% CI 94 to 100), Gemini 3 Flash 98.9%, DeepSeek V3.2 75%. Guideline-directed prompting: OR 3.76 for unsafe-error-free performance, OR 8.27 for exact concordance (both $p < .001$). Error rate 3.1% in low-complexity vs 29% in high-complexity scenarios; dominant error types: undermanagement, genotype misinterpretation, history neglect.

Clinical relevance: A model that is flawless on simple cases can fail 29% of the time on complex ones, and the prompt package moved safety as much as the model choice did. For your AI committee this changes acceptance testing: benchmark LLM decision support on a complexity-stratified scenario set built from your own guideline edge cases, and pin the exact prompt configuration as a controlled configuration item, with any prompt change re-triggering validation under Article 9 risk management.

PMID 42772748 DOI 10.2196/98131

Accuracy of Deep Learning in Detecting Cerebral Microbleeds: Systematic Review and Meta-Analysis.

Feng Y, et al. | *Journal of medical Internet research* | 2026-09-21

For cerebral microbleed counting, a task that is accurate but expensive when done manually, the pooled evidence on deep learning is favorable but thin: lesion-level detection approaches expert performance, while patient-level evidence rests on five MRI studies.

Key metrics: Patient level (5 studies): sensitivity 0.89 (95% CI 0.76 to 0.96), specificity 0.86 (0.77 to 0.92), DOR 50 (12 to 212). Lesion level: sensitivity 0.96 (0.93 to 0.98), specificity 0.98 (0.94 to 0.99), DOR 1,061 (293 to 3,848). Direct-extraction subgroup: sensitivity 0.98, specificity 0.98, DOR 1,738. PROSPERO CRD42024628447; search updated July 2026.

Clinical relevance: The lesion-versus-patient gap is the clinically meaningful reading: models find individual microbleeds well, but per-patient classification, which drives anticoagulation and amyloid-therapy decisions, carries wider uncertainty on a small evidence base. The authors' call for multicenter studies is warranted before these tools inform treatment decisions.

PMID 42767630 DOI 10.2196/95041

SECTION 2

INDUSTRY & REGULATION

Sources: Official registers | Press releases | FDA databases

[FDA] a2z Radiology AI clears 10 more abdomen-pelvis CT triage findings

a2z Radiology AI | 21 September 2026

Cleared for what: triage and notification (CADt) of 10 additional suspected findings on abdominopelvic CT with or without contrast in adults 22 and older, including obstructing ureteral stone, acute appendicitis, retroperitoneal hematoma, hemoperitoneum, splenic, hepatic and renal traumatic injury, bringing the portfolio to 17 findings across two devices; the software flags and communicates suspected positives so teams can prioritize review, and radiologists still read the full scan. It also carries FDA Breakthrough Device designation. Evidence: company announcement; no peer-reviewed performance publication is cited for the new findings. Questions for the vendor: per-finding sensitivity and specificity with the test-set composition behind them; how the 10 new findings were validated and on which scanner fleets; expected alert volume per 1,000 studies in a mix like yours; what the system logs and whether you can export it without the vendor. Regardless of clearance, a CADt tool needs local alert-burden and miss-pattern monitoring per care setting from day one, and a triage clearance is not a detection claim.

Source: [Business Wire](#)**[MARKET] Merge launches MergelQ, an AI layer inside its enterprise imaging platform**

Merge (Merative) | 22 September 2026

What it is: an AI infrastructure layer embedded in Merge's enterprise imaging platform, launched with two workflow applications, RelevancyIQ (ranks prior reports by relevance to the active exam) and NavilQ (retrieves specific findings from historical reports), running inside Merge Workflow Orchestrator; the pitch is AI without standalone installations. Evidence: vendor launch materials only; no published performance data. These are workflow tools, not diagnostic devices, so the regulatory bar is lower but the governance questions remain: how relevance ranking was validated and what a wrong prior-selection does downstream; whether retrieval errors are logged and auditable; how the fine-tuned foundation models it embeds are versioned and updated; what happens to your data governance when third-party models run inside the archive. Worth monitoring: silent model updates inside platform layers are exactly the change class your AI inventory tends to miss.

Source: [Axis Imaging News](#)**[MARKET] NVIDIA releases NV-Reason-CT, an open vision-language model for 3D CT reasoning**

NVIDIA | 23 September 2026

What it is: an open-weights vision-language model for volumetric chest and abdomen CT that generates structured draft reports with explicit chain-of-thought reasoning and supports multi-turn follow-up, released for developers rather than clinical use; it extends NV-Reason-CXR, whose multireader study is accepted at RSNA 2026. Evidence: vendor technical blog; no regulatory status, and any clinical claim would require a downstream developer to seek clearance on their own evidence. Why a department should care anyway: open CT reasoning models lower the barrier for vendors and academic groups building report-drafting tools, so procurement will increasingly see products wrapping such models. Questions to ask any vendor building on it: what was changed from the open base model, on what data it was fine-tuned, and which of the claims were validated on the product rather than inherited from the base model's publications.

Source: [NVIDIA developer blog](#)**[PARTNERSHIP] Leica Biosystems adds Tempus Paige AI applications to the Aperio AI Store**

Leica Biosystems / Tempus | 24 September 2026

What it is: two Paige computational pathology products from Tempus, PanCancer Detect (flags suspicious regions across more than 40 cancer types, built on a pathology foundation model) and the Paige Prostate suite, added to Leica's Aperio AI Store inside the HALO AP image management system. Evidence and status: both products are explicitly research use only in this channel, not for diagnostic procedures; HALO AP itself is CE-IVDR marked in the EU but RUO in the US. The label is the point for governance: an RUO tool inside a clinical-adjacent platform is a boundary that needs policing. Questions before adoption: what keeps RUO outputs out of clinical decision paths; who audits actual usage; and what changes contractually and regulatorily if a diagnostic version later ships under the same interface.

Source: [Leica Biosystems announcement](#)

SECTION 3

POST-MARKET SIGNAL

*Drift | Recalls and MAUDE | Real-world monitoring | EU AI Act Article 72***CADeNCE: the largest real-world test of a deployed imaging AI to date****21 September 2026**

Cosmo announced results of the CADeNCE study (Dominitz et al., Gastroenterology): a cluster-randomized quality-improvement evaluation of the GI Genius colonoscopy CADe system across more than 334,000 procedures by 816 endoscopists at 139 US Veterans Health Administration facilities, 42 sites with CADe against 97 controls, under routine clinical conditions. Availability of CADe was associated with 22 percent higher odds of adenoma detection (adjusted OR 1.22, 95% CI 1.15 to 1.28; ADR 50.7 to 54.9 percent), with benefit across all baseline detection levels and experience. This is what post-market evidence should look like: a randomized real-world design at national scale, run on the deployed system rather than a curated test set, answering whether trial effects survive routine practice. Notable detail for monitoring plans: the study was run on an earlier software generation than the currently marketed version, a reminder that post-market evidence attaches to a version, not a product name.

Link: <https://healthtechnologynet.com/2026/09/21/largest-real-world-study-of-the-gi-ge...>

Practice commentary: model retirement is the post-market phase nobody plans**25 September 2026**

A critical-care practice newsletter pulled together the emerging literature on decommissioning deployed clinical AI, including the JACR off-boarding framework from radiology, a scoping review finding post-deployment monitoring universally recommended but operationally immature, and an NEJM AI analysis of the measurement trap in which a model that changes outcomes corrupts the statistics used to monitor it. Labelled here as commentary rather than an imaging post-market event: the transferable points for radiology AI programmes are to require a written exit and transition plan before any tool goes live, to track vendor end-of-support dates by name, and to treat silent version updates as the most common and least visible of the three exit modes.

Link: <https://iccn.substack.com/p/every-model-your-unit-uses-will-be>

SECTION 4

PLAYBOOK SNIPPET

*One concrete step your AI committee can take this month***EU AI Act Article 12 (record-keeping) and Article 19 (automatically generated logs)**

Step: At procurement, ask the vendor in writing what the system logs, for how long, in what format, and whether you can export it without their involvement. If the answer is a dashboard, the answer is no. Add the log-export clause to the contract before signature, not after go-live.

Why: Every post-market obligation downstream assumes you can reconstruct what the model saw and said on a given date. Retrofitting logging after deployment is a contract renegotiation, not a configuration change.

Worked example: Zhou J et al. show what retained case-level outputs buy you: EAGLE's real-world calibration on 35,402 scans cut false positives 72.7% while preserving sensitivity, an adjustment only possible because outputs across three centers could be re-analyzed after the fact. Demand the same reconstruction capability from any tool you deploy. Nature Medicine, PMID 42773211, DOI 10.1038/s41591-026-04656-4.

SECTION 5

MEDIA HIGHLIGHTS

*AuntMinnie | Radiology Business | The Imaging Wire | Diagnostic Imaging | ITN***Mainstream press picks up the Northwell aneurysm shadow-mode study***Medical News Today | 25 September 2026*

Medical News Today ran a lay explainer on the JACR study featured in last week's issue: the FDA-cleared aneurysm detector that found 55 aneurysms radiologists missed across 3,856 CTAs, with benefit concentrated in inpatient and emergency settings. The coverage is notably faithful to the study's caveats, quoting that further research must show whether increased detection improves outcomes. When a setting-stratified deployment study crosses into consumer health media, expect patients and referrers to arrive with questions about whether your institution uses AI for aneurysm detection.

Link: <https://www.medicalnewstoday.com/articles/ai-complements-radiologists-improve-br...>

Legal analysis: what the first NTAP for imaging AI changes on 1 October*Mondaq (law firm analysis) | 24 September 2026*

A healthcare law analysis walks through CMS's first-ever New Technology Add-On Payment for an AI diagnostic tool, Aidoc's CARE Body CT Multi-Triage (approved 13 August, supplemental hospital payments start 1 October 2026 and run three years). The authors' framing matters for buyers: NTAP reduces the financial risk of early adoption but is time-limited, so a purchase justified by the add-on payment needs a plan for the economics after it expires. Reimbursement is becoming an active driver of imaging AI procurement decisions, which raises rather than lowers the burden on local validation: a subsidized deployment is still your deployment.

Link: <https://www.mondaq.com/unitedstates/healthcare/1846642/cms-greenlights-enhanced-...>

Radiology societies push back on the 2027 Medicare Physician Fee Schedule*Radiology Business | 24 September 2026*

Radiology Business summarizes comment letters from ACR, SIR, ASNR and ASTRO on the proposed 2027 MPFS: opposition to the blanket 2.5 percent efficiency adjustment on non-time-based procedures, concern over practice-expense and site-of-service changes, modifier-25 payment cuts, and the forced transition from MIPS to MVPs. The ACR's core argument is directly relevant to AI economics: advancements in imaging technology do not necessarily translate into efficiency, so payment policy that assumes AI-driven productivity gains ahead of evidence squeezes the same departments that are being asked to fund AI governance.

Link: <https://radiologybusiness.com/topics/healthcare-management/healthcare-policy/rad...>

SECTION 6

PROFESSIONAL SOCIETIES

*SIIM | ACR | RCR | Official publications and event coverage***RCR responds to the UK National AI Commission recommendations***22 September 2026*

The Royal College of Radiologists published its reading of the National AI Commission's 10 September recommendations on the future of AI regulation in healthcare, noting that radiology and oncology lead NHS AI deployment and that the College is working on proposals for a national post-deployment monitoring system, building on its existing deployment and monitoring guidance. The government has not yet responded to the Commission; the RCR positions itself to shape implementation, including an NHS England-commissioned national lecture series on AI.

Source: <https://www.rcr.ac.uk/news-policy/latest-updates/national-ai-commission-recommendations-what-they-mean-for-clinical-radiology-and-clinical-oncology/>

RSNA extends RadLex to standardize imaging series naming

23 September 2026

RSNA announced a RadLex initiative to develop standardized imaging series naming conventions, targeting the variability across vendors and institutions that breaks hanging protocols and complicates automated series selection for AI pipelines. It builds on the RadLex Playbook and LOINC harmonization. Unglamorous but foundational: series-level metadata inconsistency is a common silent failure mode for deployed imaging AI, and a shared naming standard is upstream data governance in the Article 10 sense.

Source: <https://www.rsna.org/news/2026/september/radlex-efforts-targets-imaging-data-variability>

ACR and partner societies request a CPT code for an AI breast malignancy risk score

24 September 2026

At the September AMA CPT Editorial Panel, ACR, RSNA, ARRS, AAR, SIR and ASNR jointly requested new codes, including a Category III code for a five-year AI-derived breast malignancy risk score based on screening mammography. If approved, Category III codes take effect 1 January 2027. A dedicated code is the first step toward reimbursement for AI risk stratification, and its existence will shape which breast AI products vendors bring forward.

Source: <https://www.acr.org/News-and-Publications/2026/acr-seeks-new-cpt-codes>

ESR Patient Advisory Group: imaging at the centre of the EU Safe Hearts Plan

23 September 2026

The ESR Patient Advisory Group issued a statement welcoming the EU Safe Hearts Plan for cardiovascular health and listing priorities for its implementation: equitable access to cardiovascular imaging, prevention and early detection including presymptomatic identification in at-risk individuals, integration of imaging into care pathways, and investment in innovation and capacity. Cardiovascular imaging, including AI-supported CT, is positioned as a cornerstone of an EU-level prevention agenda.

Source: <https://www.myesr.org/radiology-saves-lives-esr-pag-statement/>

SECTION 7

KEY OPINION LEADERS

LinkedIn | Public statements | Week of 20 to 26 September 2026

Nina Kottler

Chief Medical AI Officer, Mosaic Clinical Technologies

Posted 21 September a long reflection on building a future in which AI improves care while strengthening the human clinical role, organized around three theses: responsible AI requires lifecycle governance ("Validation at deployment is not enough. We need continuous monitoring of both AI performance and human-AI interaction, clinical review of meaningful safety signals, and clear pathways from monitoring to action"); scaled governance requires resources, so payment models must fund the monitoring infrastructure and physician expertise rather than only the AI product; and AI can expand radiology's role via imaging biomarkers only if the specialty first addresses capacity and then owns validation and clinical meaning.

Source: [LinkedIn](#)

Curt Langlotz

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

Amplified (23 September) Zongwei Zhou's announcement that his Johns Hopkins team received a four-year, 2.4 million dollar NIH R01 to develop AI detecting three types of abdominal cancers on CT, building on work in which their model identified signs of pancreatic cancer on CT a median of 11 months before diagnosis, in some cases up to three years earlier. Early-detection AI on routine CT is consolidating into a funded research programme, the same territory as this week's Nature Medicine esophageal screening paper.

Source: [LinkedIn](#)

Woojin Kim

Chief Strategy Officer and CMIO at HOPPR; CMO, ACR Data Science Institute; MSK radiologist

Posted 22 September ahead of his invited lecture at Brown University's Warren Alpert Medical School Alumni Connections series on 23 September, 'Navigating the New Frontier: Generative and Agentic AI in Medical Imaging'. Generative and agentic AI in imaging is now standing curriculum for medical audiences, delivered by people who sit at the vendor-society interface.

Source: [LinkedIn](#)

Amine Korchi

Neuroradiologist, Geneva; HealthTech innovation and ventures

Posted 22 September on the invisibility of radiology workload: an ER physician's pressure is visible in the waiting room, while the radiologist's is a 50-to-100-line worklist on a screen, each line a waiting patient, read under constant interruption. His question, whether perception of radiologist workload would change if those patients were physically queued outside the reading room, speaks to the capacity strain that drives both AI adoption and the burnout debate.

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 20 to 26 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.